

- c) w przypadku większej liczby wyładowań, np. przy „p” powtarzających się wyładowaniach w okresie zmian napięcia sinusoidalnego, energia ΔW dodana przez źródło wyraża się wzorem:

$$\Delta W = \sum_{i=1}^{i=p} U_i \cdot \Delta q;$$

U_i — napięcie na modelu układu w chwili i - tego wyładowania we wtrącinie gazowej,

Δq — ładunek, jaki jest dostarczony od pojemności C_3 do pojemności C_2 zaraz po wyładowaniu.

Podany wyżej wzór, określający wartość energii dostarczonej ze źródła wskutek wyładowań we wtrącinie gazowej dielektryku, jest słuszny, o ile zachodzi warunek:

$$C_2 \ll C_1 \ll C_3.$$

W przypadku daleko posuniętej erozji dielektryku (wtrącina wydłużyła się znacznie w kierunku pola) warunek $C_2 \ll C_1$ może nie być spełniony i wówczas podana zależność nie jest słuszna;

- d) gdy liczba wyładowań jest większa od jedności, to ubytek energii wywołany przez każde wyładowanie, niezależnie od chwilowej wartości napięcia panującego na układzie, jest zawsze ten sam;
- e) ubytek energii układu, po uwzględnieniu oporności rozładowującej R (zgodnie z rys. 8), równa się energii zużytej w tym oporze, to znaczy:

$$\Delta W = \int_0^{\infty} i^2 R dt.$$

A. Bartak wykazał [32], że powyższe wnioski są również słuszne, gdy napięcie gaszenia wyładowania nie równa się zeru. Uogólnił on jeszcze bardziej zagadnienie przyjmując założenie, że wartości napięcia zapłonu i wartości napięcia gaszenia wyładowania są zależne od biegunowości napięcia.

Z przypadkiem zależności wartości napięcia zapłonu i napięcia gaszenia wyładowania we wtrącinie gazowej od biegunowości napięcia sinusoidalnie zmiennego zetknął się np. A. E. W. Austen [33] w czasie przeprowadzanych przez siebie obserwacji wyładowań wewnętrznych w izolacji stałej.

A. Bartak [32] udowodnił, że i przy powyższym założeniu wzory określające energię straconą w układzie i dostarczoną ze źródła do układu są słuszne i że zachodzi równość pomiędzy wartościami tych energii. Trzeba jednak zaznaczyć, że wspomniana równowaga energii jest tylko słuszna w odniesieniu do całego cyklu wyładowań. Przez cykl wyładowań rozumiany jest tutaj taki okres, po którym fazy wyładowań w stosunku do napięcia panującego na szczelince gazowej zaczynają się powtarzać (są takie

same i występują w tej kolejności jak w cyklu poprzednim). Cykl taki może obejmować pewną liczbę okresów zmian napięcia sinusoidalnego. Energię zużytą w układzie i dostarczoną ze źródła wskutek wyładowań w szczelince gazowej dielektryku w odniesieniu do pełnego cyklu wspomnianego wyżej, można wyrazić wzorem:

$$\Delta W = \sum_{i=1}^{i=n} U_i \cdot \Delta q;$$

n — liczba wyładowań we wtrącinie gazowej, które to wyładowania zachodzą w czasie trwania jednego cyklu,

U_i — napięcie panujące na układzie (dielektryku) w chwili i - tego wyładowania we wtrącinie gazowej,

$\Delta q = C_3 \cdot \Delta U$ — ładunek dostarczony ze źródła do układu, wywołany wyładowaniem we wtrącinie gazowej, przy czym

ΔU — spadek napięcia na układzie, wywołany wyładowaniem we wtrącinie gazowej,

C_3 — stanowi praktycznie całkowitą pojemność układu izolacyjnego. (Model układu — rys. 8; $C_3 \gg C_1, C_2$).

Podany wyżej wzór na traconą energię w układzie wskutek wyładowań we wtrącinie gazowej prowadzi do wniosku o dużym znaczeniu praktycznym. Wniosek ten jest następujący: wartość energii traconej wskutek wyładowania we wtrącinie gazowej, znajdującej się wewnątrz dielektryku, może być określona przy pomocy parametrów U_i i Δq . Parametry te można pomierzyć w obwodzie elektrycznym, w który włączony jest układ zawierający rozpatrywany dielektryk. Duże znaczenie podanego wyżej wniosku wynika z tego, że energia ΔW jest praktycznie miarą energii traconej na erozję wtrąciny gazowej. Jeżeli chodzi zaś o wymiary wtrąciny gazowej, na podstawie których można by wnioskować o tym, jak dalece została posunięta szkodliwa erozja dielektryku, to B. Heller i J. Chladek [10] wykazali, że nie mogą być one określone drogą pomiarów w sposób podobny jak energia ΔW .

4.4. Intensywność jonizacji

W celu ustalenia szybkości starzenia układu izolacyjnego [34], które to starzenie jest spowodowane wyładowaniami niezupełnymi, dokonuje się pomiarów wielkości zwanej najczęściej „intensywnością jonizacji”. W przypadku gdy wyładowanie ma miejsce w pojedynczej wtrącinie gazowej, za punkt wyjścia może posłużyć wzór na wartość ładunku, który to ładunek w wyniku wyładowania we wspomnianej wtrącinie zostaje dostarczony ze źródła do układu:

$$\Delta q = C_3 \Delta U.$$

(Oznaczenia podane w poprzednim rozdziale).

Należy podkreślić, że mierzone wartości ΔU są bezpośrednio związane z „intensywnością jonizacji”, natomiast wartości mierzonych impulsów napięcia, występujących na impedancji włączonej w szereg z badanym układem, zależą od własności elektrycznych zewnętrznego obwodu zasilającego układ [30]. Miara „intensywności jonizacji”, bywa także wartość traconej energii wskutek wyładowania we wtrącinie gazowej lub tracona moc, gdy ta energia zostanie odniesiona do czasu. Traconą energię, którą dostarcza układ, wyraża wzór:

$$\Delta W = \frac{1}{2} U_j \Delta q;$$

U_j — napięcie na układzie, przy którym rozpoczyna się wyładowanie we wtrącinie gazowej,

ΔW — strata energii we wtrącinie (zgodnie z poprzednim rozdziałem, gdy wtrącina jest bardzo mała i zachodzi warunek $C_1 \gg C_2$, wówczas $W_c = \Delta W$).

Pomiar strat energii lub mocy traconej w czasie wyładowań wewnętrznych daje zgodnie z [10] bezpośrednie informacje o intensywności erozji, wywołanej tymi wyładowaniami. Zdaniem autora określenia „intensywności jonizacji” przez pomiar impulsu napięciowego ΔU i pomiar strat energii ΔW nie są równoznaczne, ponieważ spadek napięcia ΔU jest zależny między innymi od pojemności układu C_3 , natomiast strata energii ΔW jest od tej pojemności niezależna.

Często „intensywność jonizacji”¹⁾ określa się wartością prądu [30], [9], [34], który jest dostarczany ze źródła układu w wyniku wyładowań we wtrącinie gazowej.

W praktycznych układach izolacyjnych występuje duża liczba wtrącin gazowych i pomiar „intensywności jonizacji” odnosi się zwykle do całego układu. W takim przypadku „intensywność jonizacji” może być określona wartością sumarycznego ładunku, który dopłynie ze źródła do układu w ciągu jednej czwartej okresu zmian napięcia zasilającego [34].

„Intensywność jonizacji” określa się także wartością sumarycznej mocy traconej w układzie w związku z wyładowaniami we wtrącinach gazowych dielektryku.

A. Veverka i J. Chladek [9], opierając się na wzorze określającym straty energii dla pojedynczej wtrąciny (podanym w poprzednim rozdziale), opracowali metodę watomierzową pomiaru „intensywności jonizacji”.

R. Gotszałk [35] opracował metodę wyznaczania strat jonizacji przy pomocy specjalnie zbudowanego w tym celu kompensatora prądu zmiennego.

¹⁾ Niekiedy zamiast terminu „intensywność jonizacji” używany jest termin „intensywność ulotu” (intensity of corona) [1] lub termin „intensywność wyładowań” [30].

Istnieje szereg innych metod oceny „intensywności jonizacji”. Do najstarszych należy pomiar współczynnika stratności ($\text{tg } \delta$) [30]. Należy zwrócić uwagę, że wartość $\text{tg } \delta$ wynika z sumarycznych strat zachodzących w izolacji badanego obiektu. Ten sposób określania „intensywności jonizacji” ma szereg wad i wskutek tego nie znalazł dotąd szerokiego zastosowania. Należą do nich stosunkowo mała czułość pomiaru oraz niepewność, czy nagły wzrost charakterystyki $\text{tg } \delta = f(u)$ jest wynikiem wyładowań wewnętrznych w dielektryku [34], [7], [4].

Do przyczyn, które mogą wpłynąć na przebieg charakterystyki $\text{tg } \delta = f(u)$, między innymi należą: zawartość w dielektryku zanieczyszczeń (wilgoć, grafit), których współczynnik stratności zmienia się nieliniowo ze zmianą napięcia, oraz wpływ temperatury [30].

Ostatnio rozpowszechnia się metoda pomiaru „intensywności jonizacji” przez pomiar przypadającej na jednostkę czasu liczby impulsów wywołanych wyładowaniami we wtrącinie gazowej. W danym więc przypadku „intensywność jonizacji” jest określona liczbą impulsów na jednostkę czasu, a zatem odbiega zupełnie od określeń „intensywności jonizacji” podanych wyżej. Ten sposób pomiaru „intensywności jonizacji” przy pomocy tzw. „przelicznika impulsów” może być szczególnie pożyteczny przy prądzie stałym i przy udarach [36], gdzie inne metody pomiaru są bardzo trudne.

To krótkie zestawienie różnych określeń „intensywności jonizacji” (wynikające ze sposobu jej pomiaru) miało na celu wskazanie, w jaki sposób w praktyce ocenia się ilościowo intensywność tego szkodliwego zjawiska towarzyszącego wyładowaniom niezupełnym w izolacji oraz wskazanie na teoretyczną i praktyczną niejednoznaczność tych określeń. Mimo to każdy pomiar „intensywności jonizacji” daje w praktyce duże usługi dzięki temu, że pozwala mniej lub więcej dokładnie ocenić szybkość starzenia izolacji wskutek wyładowań wewnętrznych i ewentualnie przewidywać, jak długi jeszcze okres czasu wspomniana izolacja może wytrzymać przykładane do niej napięcie.

Zgodnie z [37] pomiar wielkości ładunku związanego z pojedynczym wyładowaniem może dać pogląd na długość wtrąciny, w której to wyładowanie się odbywa. Jeśli wtrącin jest wiele, może to być pomiar największego ładunku spośród wszystkich ładunków związanych z wyładowaniami w różnych wtrącinach gazowych. Pomiar liczby impulsów przypadających na jednostkę czasu, rzuca światło na liczbę miejsc, w których skupiają się wyładowania.

Z pojęciem „intensywności jonizacji” wiąże się pojęcie „progu jonizacji niebezpiecznej”.

Jak podaje J. Coquillon [38], istnieje pewna grupa badaczy, która zgadza się, aby „próg jonizacji niebezpiecznej” oznaczał taką wartość napięcia przyłożonego do izolacji, której zwiększenie pociągałoby nagły wzrost intensywności zachodzącej w niej jonizacji. Zgodnie z [38] pewne wyłado-

wania o słabej intensywności, pojawiające się przy stosunkowo niskim napięciu, mogą nie być szkodliwe dla izolacji.

Samo określenie „napięcie progu jonizacji” (w literaturze francuskiej nazwane „la tension de seuil d'ionisation”) oznacza wartość napięcia, przy którym rozpoczyna się jonizacja towarzysząca wyładowaniom niezupełnym występującym w izolacji [27], [26], [37].

Należy jednak nadmienić, że wspomniany nagły wzrost „intensywności jonizacji” jest charakterystyczny dla izolacji, którą stanowi papier impregnowany olejem mineralnym. W przypadku gdy syciwem jest dielektryk chlorowany, zmiana „intensywności jonizacji” ma charakter stopniowy i „próg jonizacji” tutaj nie występuje tak wyraźnie jak w przypadku, gdy syciwo stanowi olej. Nagły więc wzrost „intensywności jonizacji” nie może stanowić ogólnego kryterium jonizacji szkodliwej [7], [39].

Według [7], [39] ważnym parametrem, którego znajomość może być bardzo pożyteczna przy ocenie stopnia możliwości uszkodzenia izolacji kondensatora wskutek przepięć, jest wartość napięcia, przy którym „intensywność jonizacji” zaczyna szybko wzrastać ze wzrostem czasu (wartość szczytowa napięcia nie jest zmieniana). Jednakże wspomniana wartość napięcia zależy dość silnie od wielu czynników, w związku z czym musi być jeszcze ona przedmiotem skrupulatnych badań.

4.5. Mechanizm uszkodzenia izolacji wskutek wyładowań wewnętrznych

Mała przenikalność dielektryczna i niska wytrzymałość dielektryczna gazów przyczyniają się do stosunkowo łatwego powstawania wyładowań we wtrącinach gazowych, znajdujących się w masie dielektryku.

J. H. Mason [5] stwierdził, że po pewnym czasie od chwili przyłożonego napięcia, wyładowania we wtrącinie przyległej do elektrody koncentrują się na brzegu tej wtrąciny. Nie znając dokładniej przyczyny tego zjawiska autor ten przypuszcza tylko, że skupienie się wyładowań na brzegu wtrąciny wynika z obniżenia się wytrzymałości powietrza w przypadku, gdy stanowi ono równoległe połączenie z dielektrykiem stałym.

Wyładowania drażą w dielektryku wgłębienia (jamy o średnicy rzędu 0,1 mm), na dnie których, w przypadku np. polietylenu, J. H. Mason stwierdził obecność przezroczystej żywicy. Napięcie początkowe wyładowania zależy od średnicy jamy i jest większe, gdy średnica jej jest mniejsza.

Wyładowania te, jeżeli się powtarzają, powodują stopniowo coraz dalsze niszczenie dielektryku.

B. Heller i J. Chladek [10] wykazali na drodze analitycznej, że zamieniona na ciepło energia wyładowania (oraz towarzyszący temu zjawisku lokalny wzrost temperatury) rośnie ze wzrostem głębokości wtrąciny

gazowej (wspomniani autorzy rozpatrywali wtrącinę o kształcie walca). To samo stwierdzenie można znaleźć także u innych autorów [16].

B. Heller i J. Chladek [10] wykazali również, że ilość zamienionej na ciepło energii w czasie wyładowań we wtrącinie gazowej jest proporcjonalna do strumienia przesunięcia ψ , określonego wzorem:

$$\psi = \varepsilon KS;$$

S — przekrój wtrąciny prostopadły do kierunku pola,
 K, ε — naprężenie i przenikalność dielektryczna izolacji, w której znajduje się dana wtrącina.

Rozpatrując erozję dielektryku (erozja o charakterze mechanicznym), polegającą na rozdzielaniu drobin materiału izolacyjnego, które to rozdzielanie jest wywołane energią kinetyczną elektronów i jonów, powstałych w wyniku jonizacji we wtrącinie gazowej, B. Heller i J. Chladek [10] stwierdzili, że energia tracona na ten rodzaj erozji jest także proporcjonalna do ψ . Należy jednak dodać, że energia tracona na erozję mechaniczną stanowi zaledwie tylko małą część energii zamienionej na ciepło [10]. Ponadto wykazali oni, że działanie erozyjne o charakterze mechanicznym przez bombardujące powierzchnię wtrąciny jony i elektrony, jest tym większe, im mniejsza jest głębokość wtrąciny gazowej.

Na podstawie tych badań, przeprowadzonych na próbkach polietylenu, zostało stwierdzone [5], że wyładowania koncentrują się w kilku miejscach tuż pod elektrodą i powodują w tych miejscach stopniowe drążenie dielektryku. Intensywność wspomnianej erozji rośnie szybko ze wzrostem naprężenia panującego w dielektryku [16].

Proces żłobienia jamy odbywa się tak długo, dopóki nie osiągnie ona pewnej krytycznej długości, a następnie, wskutek w dalszym ciągu powtarzających się wyładowań, od jej końca wgłąb dielektryku zaczynają rozwijać się cienkie (półzwęglone) kanaliki. Tworzenie się tych ostatnich przyspiesza przebicie dielektryku [22], [16].

J. H. Mason [5] tłumaczy ostatnią fazę przebicia pojawieniem się na czole tych półzwęglonych wąskich kanalików, naprężeń przekraczających wytrzymałość dielektryku. Naprężenia te panujące na przestrzeni 10^{-4} cm od czoła kanaliku powodują częściowe przebicie materiału. Analogiczny proces destrukcji obserwować można także np. w nylonie i polistyrenie [16].

Jako główną przyczynę omawianej erozji podaje się termiczny rozkład dielektryku [16]. Przy przejściu nawet niedużej części energii wyładowania w ciepło, może nastąpić lokalny wzrost temperatury, dochodzący do kilkuset stopni [16].

J. Fabre [11] podaje, że wskutek bombardowania jonowego, lokalne wzrosty temperatury mogą dochodzić do 2000°C . W takich temperaturach, wprawdzie panujących w bardzo małych strefach, następuje rozkład termiczny dielektryku. Na przykład mika, którą cechuje duża stabil-

ność chemiczna, w temperaturze 1200 °C ulega wspomnianemu rozkładowi. Warto dodać, że rozkład ten jest jedną z głównych przyczyn uszkodzenia izolacji stojanów maszyn wysokonapięciowych, w przypadku których, mika jest podstawowym materiałem izolacyjnym. Lepiszcza tej izolacji mikowej (asfalty, żywice syntetyczne) są jeszcze bardziej podatne na wspomniany rozkład.

Celuloza jest specjalnie czuła na działanie rozważanych wyładowań. J. B. Whithead i M. J. Kopper [40] badając wpływ ciężaru objętościowego papieru na wytrzymałość dielektryku i trwałość izolacji papieru impregnowanego olejem stwierdzili, że początek jej uszkodzenia występuje w kanałach olejowych. Jeżeli wartość napięcia przypadająca na kanał olejowy przekroczy wartość napięcia przebicia oleju, wówczas powstaje w tym kanale wyładowanie. W wyniku przebicia oleju w kanale zaczynają tworzyć się małe pęcherzyki gazowe. Dalsze wyładowania we wtrącinach gazowych powodują powstawanie produktów polimeryzacji oleju, tzw. wosków — „X” oraz wydzielanie się gazów [7], [22], głównie wodoru [20]. Powstawanie gazu powoduje zwiększenie się wymiarów wtrąciny. Woski „X” atakują celulozę [38]. Końcową fazą uszkodzenia izolacji papierowej jest uszkodzenie celulozy.

Szybkość starzenia, którego przyczyną są wyładowania niezupełne w izolacji kondensatorów, zależy od rodzaju syciwa. Syciwa chlorowane są bardziej odporne niż olej mineralny na starzenie wywołane wyładowaniami niezupełnymi [41], [21], [42].

Według [21] zawarte w papierze substancje mineralne, a także i folia aluminiowa, potrafią związać ze sobą pewną ilość drobin HCl powstałych wskutek rozkładu chlorowanych dwufenyli przy wyładowaniach niezupełnych w izolacji kondensatora. W związku z tym, w wyniku wspomnianych wyładowań, musi tworzyć się dostatecznie dużo drobin HCl, aby zniszczyć powoli celulozę.

R. J. Hopkins, T. R. Walters, M. S. Scoville [1] zajmując się zagadnieniem związku pomiędzy napięciem początkowym wyładowań niezupełnych w izolacji kondensatorów papierowo-olejowych a ich wytrzymałością dielektryczną stwierdzili, że przed rozpoczęciem próby napięcie, przy którym rozpoczynała się intensywna jonizacja (napięcie, przy którym występował nagły wzrost intensywności jonizacji), było znacznie wyższe dla kondensatorów o oleju mineralnym niż dla kondensatorów o oleju syntetycznym (pięciochloro-dwufenyl)¹⁾. Jednak już po okresie jednego tygodnia starzenia izolacji wspomnianych kondensatorów napięcie początkowe „intensywnej jonizacji” spadło do wartości $12 \div 50\%$ wartości początkowych w przypadku oleju mineralnego, podczas gdy dla pięciochloro-dwufenyli nawet po okresie starzenia jednorocznego napię-

¹⁾ W wielu krajach syciwo to stosuje się szeroko, nadając mu różne nazwy: clophen, intertin, piranol, sowol, askarel, pyralene [43].

cie początkowe „intensywnej jonizacji” nie uległo zmniejszeniu. Należy dodać, że gdy syciwo stanowił pięciochloro-dwufenyl, wspomniani autorzy nie zanotowali żadnych uszkodzeń izolacji kondensatorów po jednorocznym czasie starzenia, tymczasem po okresie jednego tygodnia sztucznego starzenia, izolacja niektórych kondensatorów o oleju mineralnym została przebita. Jak okazało się, przebiciu uległa izolacja tych kondensatorów, dla których po okresie jednego tygodnia starzenia napięcie początkowe „intensywnej jonizacji” spadło poniżej napięcia, przy którym kondensatory starzono.

J. Coquillon [28] podaje na podstawie poczynionych obserwacji, że „próg jonizacji” pozostawał przez długi czas bez zmiany w przypadku kondensatorów, których izolację stanowił papier z syciwem chlorowanym. W przypadku kondensatorów olejowych „próg jonizacji” szybko maleje.

Obserwacje J. Coquillona potwierdzają wyżej podane wyniki otrzymane przez amerykańskich autorów [1].

Ogólnie biorąc, dielektryki organiczne są bardzo podatne na starzenie wywołane wyładowaniami we wtrącinach gazowych. Wyładowania nawet w bardzo małych wtrącinach gazowych, lecz działające długotrwale, mogą być groźne dla izolacji organicznej [15].

Po bliższej analizie zjawiska można ustalić następujące mechanizmy degradacji dielektryku, wywołanej wspomnianymi wyładowaniami, a mianowicie [11]:

- a) stopniowa erozja dielektryku wskutek bombardowania jonami i elektronami ścianek dielektryku zamykających wtrącinę gazową,
- b) działanie chemiczne gazów wytwarzanych podczas wyładowania — działanie to jest zwiększone wskutek obecności wilgoci,
- c) tworzenie przewodzących dróg na powierzchni dielektryku (zwęglanie),
- d) tworzenie kanałów przewodzących, penetrujących wzdłuż grubości dielektryku, któremu sprzyja obecność dużych naprężeń na końcach kanału,
- e) ogrzewanie kumulatywne, wywołane powtrzymującymi się wyładowaniami we wtrącinie gazowej.

4.6. Częstotliwość występowania wyładowań we wtrącinach gazowych dielektryku przy napięciu stałym i zmiennym

Przy napięciu zmiennym wyładowania we wtrącinie gazowej dielektryku występują ze znacznie większą częstotliwością niż przy napięciu stałym. Stąd niszczenie dielektryku wywołane wyładowaniami w jego wtrącinach gazowych jest znacznie szybsze przy napięciu zmiennym niż przy napięciu stałym. Przy napięciu stałym, o ile wartość napięcia przy-

łożonego do układu nie zmienia się, wyładowania we wtrącinie gazowej mogą powtarzać się, lecz przerwy pomiędzy kolejnymi wyładowaniami zachodzącymi w danej wtrącinie są stosunkowo duże.

Pierwsze wyładowanie przy napięciu stałym występuje przy wartości napięcia około $10 \div 20\%$ niższej od wartości, która wynika z obliczeń napięcia początkowego U_{pw} zgodnie z wyżej podanym wzorem. Ładunki powstałe w czasie jonizacji towarzyszącej wyładowaniu osiadając na ściankach wtrąciny stwarzają pole kompensujące pole przyłożone [44], [16]. Powtórne wyładowanie we wtrącinie gazowej może wystąpić po neutralizacji wspomnianych ładunków, zachodzącej dzięki upływności skrośnej dielektryku i powierzchniowej ścianek ograniczających wtrącinę [15], [16]. Wartość odstępów czasowych pomiędzy wyładowaniami można oszacować przy pomocy stałej czasowej „samorozładowania” τ , którą w przybliżeniu określa wzór [16]:

$$\tau = \rho \cdot \varepsilon;$$

ρ, ε — oporność właściwa i przenikalność dielektryczna dielektryku kondensatora.

Dla ścisłości należy dodać, że w rzeczywistości czas samorozładowania kondensatora zależy od wartości napięcia, do którego został on naładowany [15], [45] (krzywa rozładowania odbiega od krzywej wykładniczej).

Na ogół stała czasowa τ jest stosunkowo duża i wyładowania we wtrącinie gazowej przy napięciu stałym powtarzają się z tak małą częstotliwością, iż nie odgrywają istotnej roli w procesie starzenia izolacji. Jeżeli τ jest stosunkowo małe, co może mieć miejsce przy znacznym nagrzaniu lub zawilgoceniu izolacji, wówczas liczba wyładowań w jednostce czasu może być dość znaczna i wpływ ich na starzenie dielektryku może okazać się nie do pominięcia. Przy napięciu zmiennym stała τ praktycznie nie wpływa na odstępy czasu pomiędzy kolejnymi wyładowaniami w danej wtrącinie gazowej. Przy stałej wartości szczytowej napięcia, częstotliwość wyładowań we wtrącinie gazowej zależy od częstotliwości przyłożonego napięcia. Liczba wyładowań (we wtrącinie gazowej) na jednostkę czasu rośnie ze wzrostem częstotliwości napięcia zasilającego.

Częstotliwość zmian przyłożonego napięcia będzie miała istotny wpływ na szybkość erozji dielektryku powstałej wskutek wyładowań wewnętrznych. Zgodnie z [2] „czas życia” izolacji polietylenowej kabla jest odwrotnie proporcjonalny do częstotliwości przykładanego do niej napięcia. Zwiększenie częstotliwości przykładanego napięcia stanowi więc często stosowaną metodę przyspieszania starzenia izolacji, wywołanego wyładowaniami w jej wtrącinach gazowych.

Przy napięciu zmiennym początkowe wyładowanie występuje, gdy wartość szczytowa napięcia wynosi U_{pw} . Po pierwszych wyładowaniach we wtrącinie gazowej pole panujące w niej jest superpozycją pola przy-

łożonego i pola wywołanego ładunkiem, który powstał wskutek jonizacji gazu we wtrącinie i osiadł na ściankach ją ograniczających.

Stwierdzono, że chociaż napięcie początkowe wyładowania we wtrącinie gazowej praktycznie nie zależy od przewodności ścianek, pomiędzy którymi to wyładowanie występuje, to jednak wspomniana przewodność ma wpływ na intensywność i liczbę wyładowań, występujących w jednym półokresie zmian przyłożonego napięcia.

Jeżeli ścianki wtrąciny są metalizowane, wówczas wyładowania występują z większą częstotliwością i cechuje je większa intensywność towarzyszącej im jonizacji [16]. W przypadku gdy tylko jedna ze ścianek wtrąciny jest metalowa, liczba wyładowań w jednym półokresie zmian napięcia i intensywność jonizacji związanej z każdym wyładowaniem zależą od biegunowości napięcia.

J. H. Mason [5] wykazał, że bardziej intensywna jonizacja występuje wtedy, gdy metalowa ścianka wtrąciny jest biegunowości ujemnej. Autor artykułu stwierdził, że w przypadku napięć udarowych jonizacja związana z wyładowaniem we wtrącinie gazowej (z jedną ścianką metalową)



Rys. 11. Potencjał metalowej ścianki wtrąciny niższy od potencjału przeciwległej ścianki utworzonej przez dielektryk

jest również bardziej intensywna wtedy, gdy ścianka metalowa ma potencjał niższy niż ścianka przeciwna. Autor przeprowadził próby na modelu, który stanowiła płytka pleksiglasu, zawierająca cylindryczną wtrącinę gazową, przylegającą z jednej strony do elektrody. Otrzymane impulsy napięcia przy wyładowaniu we wspomnianej wtrącinie podają



Rys. 12. Potencjał metalowej ścianki wtrąciny wyższy od potencjału przeciwległej ścianki utworzonej przez dielektryk

oscyllogramy na rys. 11 i rys. 12. (Długość czoła udaru — $1\text{ }\mu\text{s}$, długość osi czasowej — $3\text{ }\mu\text{s}$. Napięcie udarowe zostało odfiltrowane. Impuls napięcia odpowiadający wyładowaniu we wtrącinie gazowej jest oznaczony literą A).

4.7. Długość życia izolacji uwarunkowana wyładowaniami

Warto zaznaczyć, że oba rodzaje wyładowań — wyładowania samodzielne i niesamodzielne — prowadzą do podobnych zależności czasu „życia” izolacji w funkcji przyłożonego do niej napięcia oraz że charakter tych zależności jest zgodny z doświadczeniem [25], [46]. Zależność czasu „życia” izolacji w funkcji napięcia do niej przyłożonego otrzymuje się w sposób na przykład podany przez B. Hellera i A. Veverkę [25]. Przyjmują oni założenie, że gdy pewien parametr, mający wpływ na wytrzymałość elektryczną izolacji, jak na przykład głębokość wtrąciny gazowej, która zwiększa się wskutek erozji, osiągnie pewną wartość krytyczną, wówczas nastąpi przebicie izolacji. Zakładają oni także, że szybkość zmian tego parametru jest proporcjonalna do energii doprowadzonej do ścianek wtrąciny gazowej przez bombardujące ją elektrony i jony dodatnie. W rezultacie pewnych przekształceń i dopuszczalnych uproszczeń dla przypadku starzenia izolacji wskutek wyładowań elektrycznych w jej wtrącinach gazowych, wspomniani autorzy otrzymali ostatecznie zależność o następującej postaci:

$$U = -k_1 \lg t + k_2,$$

gdzie:

U — oznacza napięcie przyłożone do izolacji,

t — stanowi zaś „czas życia” tej izolacji przy przyłożonym do niej napięciu U .

Trzeba jednak zaznaczyć, że chociaż oba mechanizmy wyładowań w szczelinie gazowej prowadzą do tego samego charakteru „krzywej życia”, to jednak nie oznacza to, że prowadzą one do tych samych wyników liczbowych określających długość „życia” izolacji w zależności od wartości napięcia do niej przyłożonego.

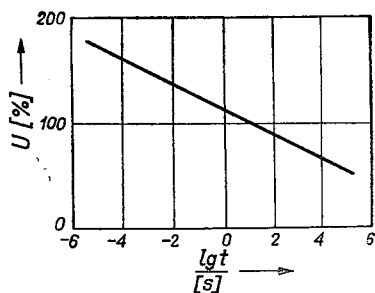
Zależność wyprowadzona przez B. Hellera i A. Veverkę określa ten sam charakter zmian napięcia w funkcji czasu t (długość „życia”), jaki C. Moses [46] otrzymał doświadczalnie dla izolacji mikowej, stosowanej w maszynach wirujących rys. 13. Przeprowadził on badania przy napięciu o częstotliwości 60 Hz. Do wspomnianych badań użył próbek izolacji nowej, którą stanowiła taśma mikowa.

Dla organicznych dielektryków [15], w szczególności dla papieru przesyconego dielektrykiem płynnym, przebieg zależności naprężenia panującego w izolacji (K_p) w funkcji czasu życia (t) podaje wzór empiryczny (analogiczny jak w przypadku starzenia elektrochemicznego):

$$K_p = A \cdot e^{-\frac{1}{m}}.$$

Dla kondensatorów papierowych $m = 7 \div 8$.

Wzór podany jest słuszny przy naprężeniach $K_p > K_{pw}$ (K_{pw} — naprężenie początkowe wyładowań). Dla wartości naprężeń stosunkowo nieznacznie przekraczających wartości naprężeń, przy których rozpoczynają się wyładowania we wtrącinach gazowych izolacji, krzywe „życia” uwzględniające jedynie starzenie jonizacyjne będą przebiegały inaczej,



Rys. 13. U — wytrzymałość dielektryczna izolacji mikowej w procentach jej wytrzymałości jednonominutowej, t — czas przebiecia

niż wynika to z podanych wyżej wzorów. Krzywe „życia” w danym przypadku będą zbliżały się asymptotycznie do prostej poziomej wyrażonej równaniem:

$$K_p = K_{pw}.$$

W rzeczywistości najczęściej na długość „życia” izolacji kondensatora wpływa wiele czynników i to w bardzo złożony sposób.

M. D. Zanobetti, Ph. R. Coursey, A. Dejou, G. G. Garton, P. Gaussens i G. Soulages [7] omawiając wpływ jonizacji na długość „życia” kondensatorów powątpiewają, aby można było podać ogólne prawo, wiążące długość „życia” kondensatora z naprężeniem panującym w jego izolacji.

4.7.1. Wyładowania powierzchniowe

Intensywność jonizacji, związanej z wyładowaniami niezupełnymi występującymi w izolacji kondensatorów, zależy od wartości napięcia przyłożonego do tej izolacji.

R. J. Hopkins, T. R. Walters i M. E. Scoville [1], badając wpływ wyładowań niezupełnych na długość „życia” izolacji papierowej impregnowanej stwierdzili, że począwszy od pewnej wartości napięcia (napięcie zmiennej częstotliwości technicznej) następuje gwałtowny wzrost intensywności jonizacji. Napięcie, przy którym powstaje gwałtowny skok intensywności jonizacji, podani wyżej autorzy nazywają „corona starting voltage”, co zgodnie z terminologią przyjętą w krajowej literaturze oznacza napięcie początkowe ulotu (świecenia) [47], [48].

Początkowe powolne zmiany intensywności jonizacji, występujące przy stosunkowo niewysokim napięciu przyłożonym do izolacji, odnoszą się do wyładowań zachodzących w odosobnionych i rozrzuconych po ca-

łym dielektryku wtrącinach gazowych. Wyładowania te były omawiane wyżej. J. R. Nye, W. R. Wilson [49] nazywają je wyładowaniami „przypadkowymi” (random corona), gdyż charakter zmian intensywności towarzyszącej im jonizacji ze wzrostem przyłożonego napięcia może być dość różny — przypadkowy. Intensywność jonizacji towarzyszącej tym wyładowaniom „przypadkowym” może rosnąć lub maleć ze wzrostem przyłożonego napięcia. Jonizacja ta może nawet zupełnie zaniknąć. Wspomniany nagły skok intensywności jonizacji jest wynikiem jonizacji towarzyszącej wyładowaniom powstałym na powierzchni dielektryku tuż przy krawędzi folii. Stąd też wyładowania te zwane są powszechnie „wyładowaniami powierzchniowymi”. Ze względu na to, że wyładowaniom powierzchniowym towarzyszy intensywna jonizacja, obecność ich jest bardzo szkodliwa dla izolacji kondensatorów. Napięcie pracy kondensatora zainstalowanego w sieci prądu zmiennego powinno być znacznie niższe od napięcia początkowego światlenia [15], a nawet więcej, wspomniane napięcie początkowe światlenia, charakterystyczne dla danego kondensatora, powinno być tak duże, aby zapobiec uszkodzeniu jego izolacji podczas prób napięciem probierczym oraz podczas przebieg [49]. Nie oznacza to jednak, aby dla izolacji kondensatora nie były dopuszczalne przebiegi o wartościach przekraczających wartość jego napięcia początkowego światlenia. Można dopuścić, aby krótkotrwałe przebiegi przekraczały napięcie początkowe światlenia, pod warunkiem jednak, że zostanie udowodnione przy pomocy odpowiednich prób, że przebiegi te w praktyce nie wpłyną na długość „życia” kondensatorów [21].

Pojawienie się na pewien okres czasu wyładowań świetlających w izolacji kondensatora może spowodować obniżenie się wartości napięcia początkowego ich powstawania. Zagadnieniem tym, jako wpływem wspomnianych przebiegów na uszkodzenie izolacji kondensatorów, zajmowało się wielu badaczy. Między innymi R. J. Hopkins, T. R. Walters i M. E. Scoville [1] stwierdzają, że jeżeli napięcie początkowe ulotu nie obniży się wskutek przebiegów więcej niż o 15% wartości pierwotnej, to dielektryk po pewnym czasie „odpoczynku” (wyłączony spod napięcia na okres ok. 1 tygodnia) odzyska z powrotem swe własności. Rozumie się przez to, że wartość napięcia początkowego ulotu powraca do swego pierwotnego poziomu. Obniżenie się wartości napięcia początkowego ulotu o $30 \div 40\%$ świadczy już o trwałym uszkodzeniu dielektryku, gdyż w tym przypadku napięcie początkowe ulotu nie odzyskuje swej pierwotnej wartości pomimo, że tak jak wyżej dielektryk zostaje na pewien okres czasu wyłączony spod napięcia. Podobne wnioski wyciągają również inni badacze zajmujący się wpływem wyładowań niezupełnych na starzenie izolacji kondensatorów statycznych.

Zgodnie z [7] dielektryk kondensatorowy „zapomni” o skutkach wyładowań pod warunkiem, że skutki te nie będą zbyt poważne, przy czym

wymagany okres „odpoczynku”, aby napięcie początkowe wyładowań świetlających osiągnęło swą pierwotną wartość, zależy od wielu czynników.

Wartość napięcia początkowego wyładowań świetlających zależy od przenikalności dielektrycznych dielektryku stałego i otaczającego go ośrodka (dielektryk płynny lub gazowy), wytrzymałości dielektrycznej tegoż ośrodka, grubości dielektryku stałego, kształtu i stopnia symetrii elektrod [16].

Wyładowania powierzchniowe są odtwarzalne w sposób regularny na próbkach o tych samych wymiarach i wykonanych z tego samego dielektryku. Z tego powodu wyładowania powierzchniowe należą do grupy wyładowań zwanych wyładowaniami „charakterystycznymi” (characteristic corona) [49].

Mimo że w kondensatorach wysokonapięciowych dielektryk stały jest najczęściej otoczony dielektrykiem płynnym, to jednak zawsze należy liczyć się z obecnością wtrącin gazowych, występujących zwłaszcza przy krawędziach folii. Przyczyny powstawania wtrącin gazowych zostały podane wyżej.

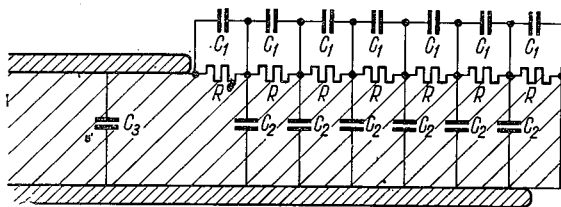
Wyładowania powierzchniowe mogą łatwiej rozwijać się we wtrącinach gazowych niż w płynnym dielektryku. W miarę wzrostu liczby wyładowań zachodzących w danej wtrącinie jej wymiary stopniowo wzrastają.

Chociaż mechanizm rozwoju wyładowań powierzchniowych jest bardziej złożony niż mechanizm rozwoju wyładowań wewnętrznych, to jednak oba te rodzaje wyładowań mają wiele cech wspólnych [16]. Jak poucza doświadczenie, w przypadku wyładowań powierzchniowych może być stosowana ta sama technika ich detekcji i te same kryteria z nią związane co i w przypadku wyładowań wewnętrznych [18]. Podobnie jak w przypadku wyładowania wewnętrznego, tak i w przypadku pojedynczego wyładowania zewnętrznego impuls napięcia może być kojarzony z rozładowaniem pewnej pojemności. Przy uproszczonym sposobie rozpatrywania mechanizmu wyładowania powierzchniowego można posłużyć się także schematem układu podanym na rys. 8. W tym jednak przypadku C_1 oznacza pojemność między elementami powierzchni objętej wyładowaniem a elektrodą, od której wspomniane wyładowanie rozwija się, C_2 — oznacza zaś pojemność dielektryku stałego połączoną szeregowo z pojemnością, C_1 , C_3 — przedstawia ogólną pojemność układu włączoną równolegle z układem szeregowo połączonych pojemności C_1 i C_2 [18].

W wielu przypadkach okazuje się jednak konieczne uwzględnienie oporności powierzchniowej, a pojemności C_1 i C_2 traktuje się jako rozłożone. W związku z tym schemat zastępczy dielektryku dla przypadku wyładowań powierzchniowych będzie wyglądał jak na rys. 14 [15], [49].

Jedną z istotnych różnic pomiędzy wyładowaniami zewnętrznymi a wewnętrznymi jest to, że wyładowania zewnętrzne mogą obejmować

odcinki dielektryku stosunkowo znacznie oddalone od elektrody, podczas gdy wyładowania wewnętrzne zamykają się w maleńkich wtrącinach gazowych [16]. Wskutek tego intensywność jonizacji pojedynczego wyładowania wewnętrznego jest znacznie mniejsza od intensywności jonizacji pojedynczego wyładowania powierzchniowego.



Rys. 14. Schemat zastępczy: C_1 , R — rozłożone pojemności i oporności powierzchniowe, C_2 — rozłożone pojemności poprzeczne (w układach praktycznych oporności poprzeczne ze względu na swe duże wartości mogą być pominięte)

Mechanizm wyładowań powierzchniowych

Powstawanie wyładowania wzdłuż powierzchni dielektryku jest spowodowane zmianami w czasie napięcia przyłożonego do układu. Pojemności C_2 ładowane są poprzez opory R . Jeżeli napięcie przyłożone do układu wzrasta znacznie szybciej niż napięcie, do którego ładowane są pojemności C_2 , wówczas na pojemnościach C_1 (stosunkowo małych) mogą pojawiać się spadki napięcia wystarczające do wywołania jonizacji powierzchniowej.

Schemat przedstawiony na rys. 14 może być stosowany zarówno dla izolacji, którą stanowią dielektryki impregnowane (np. papier impregnowany chloro-dwufenylem lub olejem mineralnym), jak i izolacji, którą stanowią dielektryki suche. W tym ostatnim przypadku schemat zastępczy dielektryku można znacznie uprościć uważając, że wartości oporności R są nieskończenie duże (uwzględniając w schemacie zastępczym tylko pojemności C_1 , C_2).

Mechanizm wyładowań powierzchniowych został stosunkowo szeroko omówiony przez J. R. Nye'a i W. R. Wilsona [49]. W oparciu o wymienioną pracę wspomnianych autorów zostaną niżej podane niektóre ważniejsze cechy wyładowań powierzchniowych.

Bardzo istotną rolę w mechanizmie wyładowań powierzchniowych odgrywa ładunek przestrzenny, który powstaje w wyniku jonizacji towarzyszącej tym wyładowaniom i gromadzi się w pobliżu elektrody. Ładunek ten stwarza własne pole nakładające się na pole wywołane napięciem przyłożonym do układu.

Na przebieg wyładowań ma także duży wpływ oporność powierzchni dielektryku. W związku z tym własności wyładowań powierzchniowych

w układach z dielektrykami impregnowanymi będą nieco odmienne od własności tych wyładowań w układach z dielektrykami suchymi. Jak już wyżej wspomniano, oporność powierzchniowa dielektryków impregnowanych płynami w przeciwieństwie do dielektryków suchych nie może być pominięta i stąd wynika różnica wspomnianych własności.

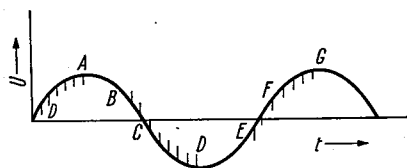
a) Kondensatory z dielektrykami impregnowanymi płynami

Do typowych dielektryków impregnowanych płynami, stosowanych szeroko przy produkcji kondensatorów, należy papier impregnowany olejem mineralnym lub syntetycznym.

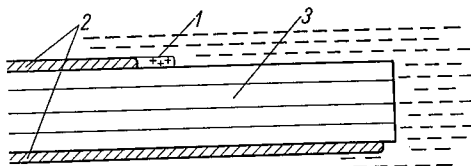
Napięcie początkowe wyładowań powierzchniowych w przypadku wspomnianych dielektryków zależy od częstotliwości. Ze wzrostem częstotliwości napięcie początkowe wyładowań powierzchniowych maleje.

Przy napięciu sinusoidalnie zmiennym wyładowania powierzchniowe występują zarówno w czasie jego wzrastania, jak i malenia — rys. 15. Godne jest uwagi, że występują one także przy wartości napięcia bliskiej zeru.

Pole elektryczne w pobliżu folii, gdzie rozpoczynają się wyładowania, nie jest wywołane jedynie przyłożonym napięciem. W czasie wyładowań powstaje bowiem ładunek przestrzenny, który ma wpływ na dalszy ich rozwój (rys. 16).



Rys. 15. Fazy wyładowań „charakterystycznych” w stosunku do napięcia okresowego częstotliwości technicznej dla typowego dielektryku kondensatorowego (papier impregnowany olejem). U — napięcie, t — czas



Rys. 16. Gromadzenie się ładunku przestrzennego we wtrącinie gazowej naprzeciw folii metalowej. 1 — wtrącina gazowa, 2 — folia aluminiowa, 3 — papier

Pole wywołane ładunkiem przestrzennym może doprowadzić do wyładowania wtedy, gdy napięcie przyłożone ma wartość zbliżoną do zera. Ładunki dodatnie, które nagromadzą się w czasie wyładowań w przedziale zmian napięcia $O-A$, nie mogą w okresie $A-B$ przepłynąć z powrotem do elektrody. Elektroda dodatnia ma w tym czasie potencjał wyższy lub nieznacznie różniący się od potencjału wywołanego przez dodatni ładunek przestrzenny.

Gdy potencjał elektrody spadnie do pewnej wartości (punkt — B), naprężenie w pobliżu folii wywołane ładunkiem przestrzennym jest tak duże, że wyładowanie rozpoczyna się ponownie i ładunek nagromadzony naprzeciw folii zaczyna stopniowo zmniejszać się. To zmniejszanie się

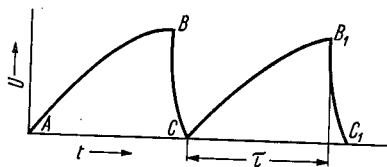
ładunku dodatniego aż do zupełnej jego neutralizacji, po której odbywa się gromadzenie ujemnego ładunku przestrzennego, zachodzi w przedziale zmian napięcia $B-D$. W dalszym ciągu opisane zjawiska powtarzają się. Wyładowanie w punkcie B zachodzi pod wpływem nagromadzonego ładunku przestrzennego.

Na podstawie powyższego opisu można powiedzieć, że ładunek przestrzenny powstały w jednym półokresie podtrzymuje wyładowanie w drugim półokresie. J. R. Nye i R. W. Wilson nazywają ten rodzaj wyładowań wyładowaniami „podtrzymywanymi” (sustained corona). Stąd wynika, że wyładowania powierzchniowe przy napięciu sinusoidalnym o częstotliwości technicznej są jednocześnie wyładowaniami „podtrzymywanymi”.

Przy napięciach bardzo powoli zmieniających się, jak to ma na przykład miejsce w czasie ładowania kondensatorów głównych generatora udarów napięciowych lub prądowych, wyładowania o charakterze powierzchniowym praktycznie nie występują, lecz pojawiają się one natomiast, gdy następuje szybkie rozładowanie kondensatora. Wyładowania te jednak nie należą do wyładowań „podtrzymywanych”.

Przy napięciach powtarzających się okresowo (rys. 17), wyładowania „podtrzymywane” nie zachodzą, jeśli pomiędzy dwoma cyklami napięciowymi upływa dostatecznie długi czas lub gdy okres wyładowań poprzedzony jest stosunkowo długim powolnym wzrostem napięcia (wzrost napięcia jest tak powolny, że nie wywołuje wyładowań).

Rys. 17. $A - B$ powolny wzrost napięcia, $B - C$ — maleńie napięcia, w czasie którego powstają wyładowania, τ — czas odprężenia



O ile pomiędzy wyładowaniami upływa dostatecznie długi okres czasu, wówczas ładunek przestrzenny, który odgrywa istotną rolę w mechanizmie tego typu wyładowań, rozplywa się po powierzchni dielektryku. Zjawisko to może być także wyjaśnione przy pomocy obrazu fizycznego dielektryku zgodnego z rys. 14. Taki rodzaj wyładowań J. R. Nye i R. W. Wilson nazywają wyładowaniami „przejściowymi” (transient corona).

Ze zjawiskiem eliminacji wpływu ładunku przestrzennego łączy się pojęcie tzw. „czasu odprężenia” (relaxation time). W przypadku napięć okresowych symetrycznych¹⁾ jest to czas pomiędzy dwoma kolejnymi szczytami napięcia.

W przypadku napięć okresowych niesymetrycznych, czas „odprężenia” mierzy się między dwoma kolejnymi odcinkami krzywej napięcia, dla których występują wyładowania (rys. 17).

¹⁾ tj. cechujących się tym, że dodatnie i ujemne przebiegi czasowe napięcia są tego samego kształtu.

Gdy czas „odprężenia” jest dostatecznie duży, wyładowania „podtrzymywane” nie pojawiają się. Wynika z tego, że napięcie początkowe wyładowań „podtrzymywanych” powinno zależeć od częstotliwości zmian przyłożonego napięcia. Zależność ta została potwierdzona doświadczalnie.

b) *Kondensatory, w których izolację stanowią dielektryki suche*

W przypadku gdy izolację kondensatora stanowi dielektryk suchy (dielektryk o dużej oporności powierzchniowej), wówczas, jak to już wyżej wspomniano, jego schemat zastępczy (rys. 14) upraszcza się do układu utworzonego jedynie z samych pojemności.

Rozkład napięcia na powierzchni dielektryku w takim układzie zależny jest jedynie od stosunku pojemności C_1/C_2 . Wyjaśnia to całkowicie, dlaczego napięcie początkowe wyładowań powierzchniowych w przypadku dielektryków suchych, jak to stwierdzono doświadczalnie, nie zależy od częstotliwości.

Jeżeli do kondensatora z dielektrykiem suchym zostanie przyłożone napięcie stałe, którego wartość jest powoli podnoszona, wówczas, w przeciwieństwie do przypadku kondensatorów z izolacją papierową impregnowaną, wyładowania powierzchniowe pojawiają się przy pewnej wartości przyłożonego napięcia. Wyładowania te jednak po pewnym czasie zanikają, jeżeli po ich ukazaniu się wartość przyłożonego napięcia nie była nadal zwiększana.

Tę wartość napięcia J. R. Nye i W. R. Wilson nazywają „wstępnym napięciem początkowym” wyładowań (initial corona — starting voltage). Odpowiadające temu napięciu naprężenie krytyczne E_k może być nazwane wstępnym „naprężeniem początkowym”.

Wspomniana wyżej wartość napięcia nie różni się praktycznie od wartości szczytowej sinusoidalnego napięcia początkowego wyładowań powierzchniowych.

Gdy napięcie stałe przyłożone do kondensatora zwiększone jest w dalszym ciągu, wówczas przy wzroście jego wynoszącym około 50% wartości „wstępnego napięcia początkowego” wyładowania pojawiają się ponownie, lecz także podobnie jak przed tym zanikają one, jeżeli wartość przyłożonego napięcia dalej nie wzrasta. Takie pojawianie się i zanikanie wyładowań może mieć miejsce aż do 6 ÷ 7-krotnej wartości „wstępnego napięcia początkowego”, po czym wyładowania stają się już trwałe.

Przy przebiegach napięciowych powtarzalnych, na przykład takich, które przypominają kształtem przebieg czasowy napięcia ładowania i rozładowania¹⁾ kondensatora (rys. 17), wyładowania pojawiają się w pobliżu szczytu napięcia przy wartości zbliżonej do wartości „wstępnego napięcia początkowego”.

¹⁾ Doświadczenia zostały przeprowadzone [49] dla czasów ładowania (krzywa A-B) zmieniających się w granicach 4—600 sek i czasów rozładowania (krzywa B-C) w granicach 0,002—2 sek.

Jeżeli wartość szczytowa przyłożonego napięcia nie jest nadal zwiększana, to wyładowania po pewnym czasie zanikają, podobnie jak to miało miejsce przy napięciu stałym. Wspomniany proces pojawiania się i zanikania wyładowań powtarza się aż dotąd, dopóki wartość szczytowa przyłożonego napięcia nie równa się podwójnej wartości „wstępnego napięcia początkowego” wyładowań. Wyładowania występują wtedy w pobliżu szczytu napięcia (punkt B) i w końcowej części krzywej rozładowania (punkt C).

Wspomniane wyżej zanikanie wyładowań tłumaczy się kompensacją pola przyłożonego polem wytworzonym przez ładunek przestrzenny, powstały w wyniku jonizacji towarzyszącej tym wyładowaniom. Od chwili pojawienia się ładunku przestrzennego, naprężenie przy krawędzi folii jest superpozycją naprężeń, z których jedno wywołane jest przyłożonym napięciem oraz drugie wywołane w tym samym miejscu wspomnianym ładunkiem przestrzennym.

Jeżeli przez E zostanie oznaczone naprężenie wypadkowe w pobliżu folii, to można napisać, że:

$$E = E_n - E_p,$$

gdzie E_n — oznacza naprężenie związane z przyłożonym napięciem,

E_p — wyraża naprężenie wywołane ładunkiem przestrzennym.

Gdy naprężenie E osiągnie wartość naprężenia krytycznego E_k , wówczas następuje wyładowanie.

Przy mniejszych wartościach szczytowych napięcia, wyładowania mogą pojawiać się jedynie, gdy napięcie przechodzi przez punkt B (rys. 17). Jeżeli, jak to już wyżej wspomniano, przyłożone napięcie jest stopniowo zwiększane, to przy pewnej jego wartości szczytowej wyładowania zaczynają pojawiać się również, gdy napięcie przechodzi przez punkt C. Jeśli wypadkowe naprężenie przy krawędzi folii, odpowiadające punktowi B, równa się E_k , to naprężenie odpowiadające punktowi C (wywołujące także wyładowania) wynosi $-E_k$ (ponieważ wywołane jest polem ładunku przestrzennego, które to pole jest przeciwnie skierowane do pola wywołanego napięciem przyłożonym).

Wyładowanie odpowiadające punktowi C zachodzi jedynie pod wpływem pola wywołanego ładunkiem przestrzennym, gdyż napięcie przyłożone do układu jest w tym momencie równe zero. A zatem

$$-E_k = 0 - E_p.$$

Gdy jednak w czasie swego kolejnego przebiegu napięcie ponownie wzrasta (krzywa C-B₁), to wyładowanie, które powstało, gdy napięcie przechodziło przez punkt C, zanika. Pojawia się ono dopiero ponownie, gdy napięcie osiąga wartość szczytową (punkt B). Naprężenie wypadkowe na krawędzi folii będzie w tym przypadku miało wartość równą naprężeniu początkowemu wyładowań E_k , przy czym

$$E_k = E_n - E_p.$$

Jeżeli zostanie przyjęte założenie, że naprężenie $-E_p$ (wywołane ładunkiem przestrzennym) wskutek dużej oporności powierzchniowej dielektryku nie ulega zmianie w czasie wzrostu napięcia od punktu C do punktu B_1 , to z ostatnich dwóch równań wynika, że warunkiem wystąpienia wyładowań niezanikających jest spełnienie zależności: $E_n = 2E_k$, gdy napięcie osiąga wartość maksymalną. Oznacza to, że aby pojawiły się wyładowania niezanikające, naprężenie wywołane przyłożonym napięciem powinno być równe podwójnej wartości „wstępnego naprężenia początkowego wyładowań”.

Jak okazuje się, ostatni wniosek jest słuszny również w przypadku innych przebiegów okresowych powtarzalnych. Wniosek ten da się streścić następująco: aby występowały wyładowania przy największych i najmniejszych wartościach napięcia okresowo zmiennego, różnica pomiędzy tymi wartościami powinna być równa podwójnej wartości „wstępnego napięcia początkowego” wyładowań. Gdy podany warunek jest spełniony, wyładowania mają wówczas również charakter trwałe. W przypadku napięcia sinusoidalnie zmiennego ma to miejsce wtedy, gdy wartość szczytowa napięcia równa się wartości wstępnego napięcia początkowego wyładowań. Bezwzględny przyrost napięcia pomiędzy dwoma kolejnymi szczytami napięcia (różnego znaku) równa się wówczas podwójnej wartości wstępnego napięcia początkowego wyładowań (wstępne napięcie początkowe wyładowań pojawia się w momencie, gdy napięcie przechodzi przez wartość maksymalną lub minimalną).

Z powyższego wynika również, że pojawienie się wstępnego napięcia początkowego wyładowań przy napięciu sinusoidalnie zmiennym zapoczątkowuje wyładowania trwałe.

Wpływem kształtu napięcia oraz długością przerw pomiędzy impulsami napięciowymi na napięcie początkowe wyładowań zajmowali się również R. J. Hopkins, T. R. Walters i M. E. Scoville [1]. Stwierdzili oni między innymi, że długości przerw pomiędzy impulsami napięciowymi mają wpływ na napięcie początkowe wyładowań. Wyrazili oni także pogląd, że wartość napięcia początkowego wyładowań powierzchniowych zależy od całkowitej zmiany napięcia przyłożonego do układu (tj. różnicy pomiędzy jego maksymalną i minimalną wartością).

4.7.2. Starzenie izolacji wywołane wyładowaniami zewnętrznymi

Wyładowania zewnętrzne pod wieloma względami w swym działaniu na dielektryki przypominają wyładowania wewnętrzne. Do procesów, które powodują stopniową degradację dielektryku a związane są z wyładowaniami zewnętrznymi, należą:

- 1) erozja wywołana wyładowaniem jonowym,
- 2) chemiczny rozkład dielektryku wywołany przez chemicznie aktywne gazy lub parę wodną, które wydzielają się wskutek wyładowań. Dal-

szym skutkiem tego działania wyładowań okazuje się mechaniczne osłabienie dielektryku i zwiększenie jego strat dielektrycznych,

3) nagrzewanie kumulatywne dielektryku wskutek powtarzających się wyładowań.

Przy stosunkowo szybkich zmianach napięcia, wyładowania mogą rozwijać się w głąb dielektryku.

Przy niższych napięciach działanie wyładowań ma charakter bardziej dyskretny, jednak gdy występują one przez stosunkowo długi okres czasu, dielektryk może zostać tak osłabiony, iż szereg powtarzających się przebiegów może go całkowicie uszkodzić.

Nie wnikając bliżej w mechanizm przebicia dielektryku wywołanego wyładowaniami zewnętrznymi można stwierdzić ogólnie, że w zależności od „długości życia” dominują różne procesy stopniowej degradacji dielektryku. Po czasach względnie długich (tygodnie, miesiące, lata), przy stosunkowo niskich naprężeniach, przebicie izolacji związane jest z uszkodzeniem chemicznym. „Długość życia” jest zależna od chemicznej odporności dielektryku na działanie aktywnych produktów tworzących się wskutek wyładowań [18]. Przebicie nie może być w tym przypadku rozpatrywane ogólnie, ponieważ możliwe są różne jego mechanizmy. W przypadku materiałów celulozowych chemiczne procesy wywołane wyładowaniami powierzchniowymi prowadzą do mechanicznego uszkodzenia izolacji. W miarę czasu wspomniane niszczenie powoli przechodzi od powierzchni w głąb dielektryku.

Przy większych naprężeniach, dla „długości życia” zmieniającej się od minut do setek godzin, wytrzymałość układu zależy głównie od własności otaczającego dielektryk stały ośrodka.

Przy stosunkowo wysokich wartościach naprężeń, napięcie przebicia uwarunkowane jest zarówno wytrzymałością dielektryku stałego, jak i wytrzymałością otaczającego ten dielektryk ośrodka.

Warto dodać, że typowym przykładem, gdzie wyładowania powierzchniowe odgrywają istotną rolę w procesie starzenia dielektryków, są kondensatory przeznaczone do pracy w układach generatorów udarowych [3]. Do takich samych wniosków doszedł autor przeprowadzając próby życia kondensatorów papierowo-wazelinowych, pracujących w układzie generatora udarów napięciowych.

Na zakończenie należy stwierdzić, iż zarówno dla przypadku napięć sinusoidalnie zmiennych, jak i dla przypadku napięć udarowych pozostaje do rozwiązania jeszcze stosunkowo dużo ważnych problemów związanych z wpływem „jonizacji” na „starzenie się” izolacji kondensatorów.

WYKAZ LITERATURY

1. Hopkins R. J., Walters T. R., Scoville M. E.: Development of Measurements and their Relation to the Dielectric Strength of Capacitors. Tr. AIEE, Vol. 70, 1951, s. 1643.
2. Dakin T. W., Philofsky H. M., Divens W. C.: Effect of Electric Discharges on the Breakdown of Solid Insulation. Comm. and. Electr., 1954, No 12, s. 155.
3. Chu S. C.: Capacitors for Impulse-Current Generators., Electr. Times, 1961, 139, s. 509.
4. Austen A. E. W., Hackett W.: Internal Discharges in Dielectrics. Their Observation and Analysis. Journal JEE, 1944, s. 298.
5. Mason J. H.: The Deterioration and Breakdown of Dielectrics Resulting from Internal Discharges. J. E. E., 1951, cz. I, 98, s. 44.
6. Renne W. T., Fainskij W. M., Warszawskij D. S.: Woskoobrazowanie w izolacji bumażnomasliannykh kondensatorow. Elektrichestwo, 1953, No 12.
7. Zanobetti D., Coursey Ph., Garton C. G., Dejou A., Gaoussens P., Soulages G.: L'Ionisation dans les Condensateurs Industriels. CIGRE, 1958, Rap. 141.
8. Heller B., Veverka A.: Analýsa jonisacnich pochodu w tuchych isolantach pri prumyslovem kmitoctu. Prace Ustavu Pro Elektrotechniku, 1954, No 1, s. 33.
9. Veverka A., Chladek J.: Pístroj na meřeni ionizace. Prace Ustavu Pro Elektrotechniku, 1937, No 7, s. 21.
10. Heller B., Chladek J.: Zur Problematik der Korona-Entladung im festen Dielektrikum. Acta Technica, 1962, No 5, s. 391.
11. Fabre J.: Les isolants solides: Théories Modernes. 1958, D. 181.
12. Fabre J.: Les Critères Chimiques de dégradation du papier imprégné d'huile dans les appareils électriques. Rev. Gen. E., 1957, No 12, s. 19.
13. Fabre J.: Les Lois de dégradation du papier imprégné d'huile dans les transformateurs. Bull. S. F. E., 1959, tom IX, s. 409.
14. Church H.: Factors Affecting the Life of Impregnated-Paper Capacitors. P. I. E. E., cz. III, 98, 1952.
15. Renne T. W.: Elektriczeskie kondensatory. Moskwa 1959.
16. Birks J. B., Schulman J.: Progress in Dielectrics. Tom I. London 1959.
17. Jasicki Z.: Zagadnienie kondensatorów elektroenergetycznych (M. K. W. S. E. 1962). Energetyka, 1963, No 1, s. 32.
18. Whitehead S.: Dielectric Breakdown of Solids. Oxford 1951.
19. Skowroński J. I. i inni: Poradnik materiałoznawstwa elektrycznego. Warszawa 1959, s. 50.
20. Bagaliej I. W.: O processach rozruszenja dielektrika bumažno-masliannykh kondensatorow dla impulsnych schem. Elektrichestwo, 1955, No 3.
21. Held W., Kunze R. C.: Glimmentladungen und Lebensdauer von Starkstromkondensatoren. ETZ, Z-11, 1961, s. 333.
22. Skanawi G. J.: Fizika dielektrikow. Moskwa 1958.
23. Kuczinskij G. S., Renne W., Foinickij W. M.: Wybor tołszcziny dielektrika dla bumażnykh siłowych kondensatorow. Elektrichestwo, 1954, No 6.
24. Schmidt J.: Prubeň elektrickich vyboju v uzavřených plynem naplnených dutinách papirovehu, olejem napusteneho dielektrika. Elektr. Obz., 1963, No 1, s. 8.
25. Heller B., Veverka A.: Vliv ionizacnich pochodu v mezerach dielektrika na život isolace. Prace Ustavu Pro Elektrotechniku, 1954, No 1, s. 47.

26. Renaudin D.: Caractère et mesure de l'ionisation aigrette dans les diélectriques. Bull. Soc. Franc. Electr., 1952, s. 431.
27. Langlois-Berthelot R.: Comment utiliser la mesure des phénomènes d'ionisation pour connaître la fatigue des isolants. CIGRE, 1952, Rap. 132.
28. Dakin W., Lim J.: Corona Measurement and Interpretation. Trans. AIEE, P. III, 1957, s. 1059.
29. Oudin J., Thevenon H.: Influence de l'épaisseur des papiers sur l'ionisation des câbles à haute tension. Rev. Gén. d. Electr., T. 62, No 12, 1953, s. 581.
30. Goliński J., Lech W., Żołędziowski A.: Jonizacja w transformatorach i maszynach elektrycznych. Rozprawy Elektrotechniczne T. III, 1957, z-3, s. 403.
31. Veverka A., Chládek J.: Jonizační pochody v pevném dielektriku a jejich energetické pomery. Prace Ustavu Pro Elektrotechniku 1955, No 3, s. 37.
32. Bartak A.: Energetika wyboju w pevném dielektriku. Prace Ustavu Pro Elektrotechniku CSAV 1956, No 5.
33. Austen A. E. W.: Some Observations of Discharges in Insulation under A. C. Stress. Ref. L/T 141, 1943.
34. Adamec V.: Zariadenie na meranie jonizacie v izolatorach. El. Obz. 47, 1958, C. 38, s. 418.
35. Gotszalk R.: Pomiar strat jonizacji w izolacji generatorów wysokiego napięcia. Rozprawa doktorska 1962.
36. Awrutin A. D., Dawidowa L., Ławrowa D. S., Renne W. T.: Исследование некоторых факторов влияющих на развитие ионизационных процессов в диэлектрике бумажномасляных конденсаторов. Передача Энергии Постоянным и переменным током 1961, No 7, s. 231.
37. Ross C. W., Courdts E. B.: Considerations in Specifying Corona Tests. Tr. AIEE 1956, vol. 76, cz. III, s. 63.
38. Coquillion J.: Progrès des diélectriques chlorés. Incidences de leurs caractéristiques sur les propriétés des condensateurs au papier imprégné pour tensions alternatives et de certains condensateurs spéciaux. Bull. Soc. Fr. Electr. 1957, No 81, s. 534.
39. Soulages G.: Evolution dans la technique des condensateurs. Bull. Soc. Fr. El. 1959, t. IX, No 104, s. 449.
40. Whitehead J. B., Kopper J.: The Dielectric Strength and Life of Impregnated-Paper Insulation — IV. Tr. AIEE 1945, s. 171.
41. Meier K.: Elektrische Eigenschaften von Starkstromkondensatoren. Bull. de l'Ass. Suisse d. Electr. 1958, No 2, s. 37.
43. Morozow M. M., Miedwiediew S. K.: Kondensatory dla siłowych установок. Elektrizestwo, 1955, No 7, s. 123.
44. Sirotinskij L. J.: Technika vysokich napięć, Moskwa-Leningrad I, II, 1953, s. 60.
45. Renne W. T.: Skorość samorozrządu jak kriterii dla oceny jakości bumażnych kondensatorów postojennogo naprjażenia. Żurnal Techn. Fiziki, 1947, t. XVII, No 1, s. 49.
46. Moses G. L.: Alternating and Direct Voltage Endurance Studies on Mica Insulation for Electric Machinery. Tr. AIEE, 1951, s. 763.
47. Jakubowski J. L.: Technika wysokich napięć. Warszawa 1951.
48. Szpor S.: Wytrzymałość dielektryczna i technika izolacyjna. Warszawa 1959.
49. Nye I. R., Wilson W. R.: Physical Concepts of Corona in Capacitors. Proc. Apparatus Syst., 1953, No 7, str. 781.
50. Hall F., Ruseck R.: Discharge Inception and Exinction in Dielectric Voids. Proceedings IEE, 1953, s. 47.

Z. SZCZEPAŃSKI

DÉCHARGES INCOMPLÈTES COMME LA CAUSE DU VIEILLISSEMENT
DES DIÉLECTRIQUES DES CONDENSATEURS ÉNERGÉTIQUES
HAUTE TENSION

Résumé

On a discuté les mécanismes du vieillissement des isolants des condensateurs. Le plus de place on a consacré au vieillissement causé par l'ionisation. En cas du mécanisme thermique du vieillissement on a donné le critère du vieillissement du papier isolant en considérant le degré de polymérisation de la cellulose. On a discuté les facteurs accélérant le vieillissement thermique des isolants de papier.

Concernant le vieillissement électrochimique on a donné les facteurs influant sur la valeur du courant de fuite. On a décrit les mécanismes du vieillissement des isolants de papier sursaturé par les diélectriques chlorés. On a discuté l'influence des stabilisateurs sur le retard du vieillissement électrochimique.

En cas du vieillissement causé par l'ionisation on a considéré l'influence des décharges dans les occlusions gazeuses ainsi que les décharges extérieures sur la dégradation des isolants des condensateurs. On a discuté: causes de la naissance des occlusions gazeuses dans les diélectriques, mécanisme des décharges dans les occlusions gazeuses ainsi que valeur de la tension initiale de leur naissance. On a présenté le bilan énergétique des décharges dans les occlusions gazeuses et, sortant de ce bilan, on a donné de différentes définitions de l'intensité d'„ionisation”.

On a discuté le mécanisme de l'endommagement des diélectriques causé par les décharges dans les occlusions gazeuses. On a montré que la décomposition thermique des diélectriques constitue la cause principale de l'érosion des diélectriques solides. On a discuté les processus liés avec l'endommagement des isolants de papier sursaturés par huile et diélectriques chlorés. On a présenté les résultats des expériences, ayant pour le but la comparaison de la vitesse du vieillissement d'„ionisation” des isolants de papier sursaturés par huile et sursaturés par les diélectriques chlorés.

On a étudié l'influence de la vitesse des changements de la tension appliquée sur la fréquence des décharges dans les vacuoles gazeuses ainsi que l'influence de ces décharges sur la durée de vie des isolants.

En discutant les problèmes liés avec les décharges extérieures, on a touché le problème de l'influence des surtensions et de la valeur de la tension initiale de ces décharges sur la durée de vie des isolants des condensateurs. On a décrit le mécanisme de développement des décharges extérieures en cas des diélectriques secs et des diélectriques imprégnés (papier imprégné) par huile ou par diélectriques chlorés.

On a décrit l'influence de la forme et de la fréquence de la tension appliquée sur la valeur de la tension initiale des décharges extérieures.

On a présenté le modèle physique du diélectrique en cas des décharges extérieures et on a discuté l'influence de la charge spatiale sur le mécanisme de ces décharges.

En se basant sur l'oeuvre de J. R. Nye, W. R. Wilson on a donné quelques définitions concernant les décharges extérieures. Enfin on a discuté le vieillissement des isolants causé par les décharges extérieures.

Z. SZCZEPAŃSKI

UNVOLLSTÄNDIGE ENTLADUNGEN ALS URSACHE VON ALTERUNG
DES DIELEKTRIKUMS IN ENERGETISCHEN
HOCHSPANNUNGSKONDENSATOREN

Zusammenfassung

In diesen Artikel wurden die Mechanismen der Isolationsalterung in Kondensatoren besprochen. Der grösste Teil dieses Aufsatzes wurde den Alterungserscheinungen, welche durch unvollständige Entladungen hervorgerufen werden, gewidmet. Im Falle des thermischen Alterungsmechanismus wurde das Kriterium der Alterung von Isolationspapier unter Berücksichtigung des Polymerisationsgrades des Zellstoffs angegeben. Weiter wurden Faktoren, die eine Beschleunigung der Alterung von Papierisolation durch Wärme verursachen, besprochen.

Bezüglich der elektrochemischen Alterung wurden Faktoren, die den Ableitungswert des Stromes beeinflussen, angegeben. Hier wurde auch der Mechanismus der elektrochemischen Alterung des Isolationspapiers, das mit einem übersättigt chlorierten Dielektrikum präpariert wurde, beschrieben. Es wurde auch der Einfluss von Stabilisatoren auf die Verzögerung der elektrochemischen Alterung besprochen.

Im Falle einer Alterung, die durch unvollständige Entladung hervorgerufen wurde, wird der Einfluss der Innen- und Oberflächenentladungen auf die Degradation der Isolierung in Kondensatoren erwogen. Es wurden auch die Entstehungsursachen der Gaseinschlüsse im Dielektrikum, der Mechanismus der Innenentladungen sowie der Wert der Anfangsspannung bei ihrem Entstehen besprochen. Ferner wurde die Energiebilanz der Innenentladungen, auf welche sich die verschieden angegebenen Definitionen der Ionisationsintensität beziehen, dargestellt. Weiter wurde der Beschädigungsmechanismus im Dielektrikum infolge der in ihm auftretenden Innenentladungen besprochen.

Ebenfalls wurde auch auf die thermische Zersetzung des Dielektrikums als Hauptursache der Erosion im konstanten Dielektrikum hingewiesen. Weiterhin wurden die Prozesse, die im Zusammenhange mit der Beschädigung des Isolationspapiers, das mit öl und chloriertem Dielektrikum durchsättigt ist, besprochen. Es wurden Ergebnisse von Versuchen, als Vergleich der Beschleunigung der Ionisationsalterung des Isolationspapiers, welches mit öl und chloriertem Dielektrikum durchsättigt ist, dargestellt.

Auch wurde der Einfluss der Beschleunigungsänderung der angelegten Spannung auf die Entladungsfrequenz in den Gaseinschlüssen sowie der Einfluss dieser Entladungen auf die Lebensdauer der Isolation behandelt.

Bei der Besprechung von Fragen, die im Zusammenhange mit den Oberflächenentladungen stehen, wurde auch die Frage des Einflusses von Überspannungen und des Anfangsspannungswertes dieser Entladungen auf die Lebensdauer der Kondensatorenisolation erwogen. Es wurde auch der Mechanismus der Entwicklung der Oberflächenentladungen in Fällen von angewandtem trockenem Dielektrikum oder imprägniertem Dielektrikum und Papier behandelt. Besprochen wurde ferner der Einfluss der Form sowie der Frequenz der angelegten Spannung auf den Wert der anfänglichen Oberflächenentladespannungen.

Es wurde auch ein physisches Modell des Dielektrikums für den Fall der Oberflächenentladungen dargestellt sowie der Einfluss der Raumladung auf den Mechanismus dieser Entladungen.

In Anlehnung auf die Arbeiten von J. R. Nye und W. R. Wilson wurden einige Bezeichnungen betr. der Oberflächenentladungen angegeben.

Zum Schluss wurde die Alterung von Isolation, die durch Oberflächenentladungen hervorgerufen wurde, besprochen.

3. ЩЕПАНСКИ

НЕПОЛНЫЕ РАЗРЯДЫ В КАЧЕСТВЕ ПРИЧИНЫ СТАРЕНИЯ ДИЭЛЕКТРИКА ЭНЕРГЕТИЧЕСКИХ КОНДЕНСАТОРОВ ВЫСОКОГО НАПРЯЖЕНИЯ

Резюме

Описано механизмы старения изоляции конденсаторов. Больше всего места посвящено старению, вызванному неполными разрядами. В случае теплового механизма старения был приведен критерий старения изоляционной бумаги, учитывающий степень полимеризации целлюлозы. Рассмотрены факторы ускоряющие тепловое старение бумажной изоляции.

Относительно электрохимического старения, приведены факторы влияющие на значение тока утечки. Описаны механизмы электрохимического старения бумажной изоляции, насыщенной хлорированным диэлектриком. Представлено влияние стабилизаторов на замедление электрохимического старения.

В случае старения, вызванного неполными разрядами, рассмотрено влияние внутренних и поверхностных разрядов на ухудшение изоляции конденсаторов. Проанализировано причины возникновения воздушных включений в диэлектрике, механизм внутренних разрядов и значение начального напряжения их появления. Представлен энергетический баланс внутреннего разряда, на основании которого были приведены разные определения интенсивности „ионизации”. Описан механизм повреждения диэлектрика из-за происходящих в нем внутренних разрядов. Термическое разложение диэлектрика указано как главная причина эрозии твердых диэлектриков. Описаны процессы, связанные с повреждением бумажной изоляции, насыщенной маслом и хлорированными диэлектриками. Представлены результаты экспериментов, имеющих целью сравнение скорости „ионизационного” старения изоляции, насыщенной маслом и насыщенной хлорированными диэлектриками.

Рассмотрено влияние скорости изменений приложенного напряжения на частоту разрядов в воздушных включениях и влияние этих включений на время жизни изоляции.

При описании вопросов, связанных с поверхностными разрядами, затронут вопрос влияния перенапряжений и значения начального напряжения этих разрядов на продолжительность жизни изоляции конденсаторов. Описан механизм развития поверхностных разрядов в случае сухих диэлектриков и диэлектриков импрегнированных (импрегнированная бумага) маслом или хлорированными диэлектриками.

Описано влияние форм и частоты приложенного напряжения на значение начального напряжения поверхностных разрядов.

Представлена физическая модель диэлектрика для случая поверхностных разрядов и описано влияние пространственного заряда на механизм этих разрядов.

На основании труда Дж. Р. Найна и У. Р. Уильсона приведены некоторые определения, касающиеся поверхностных разрядов. Наконец описано старение изоляции вызванное поверхностными разрядами.

JULIUSZ LECH KULIKOWSKI

O strukturze i własnościach asymptotycznych pewnej klasy kwantowych układów decyzyjnych

Rękopis dostarczono 20.5.1963

W artykule rozważono zagadnienie asymptotycznej oceny własności układu decyzyjnego z kwantowaniem sygnału przeznaczonego do wykrywania sygnałów radiolokacyjnych na tle korelowanych zakłóceń. Ocenę układu oparto o pewne własności graniczne złożonych łańcuchów Markowa. Rozważania uzupełniono uwagami dotyczącymi optymalnej struktury logicznej cyfrowego układu liczącego realizującego optymalny algorytm obróbki sygnału.

1. WSTĘP

Rozwój metod statystycznych w teorii odbioru sygnałów na tle zakłóceń doprowadził do znacznej dysproporcji między teoretycznymi możliwościami wyboru optymalnego (w sensie pewnych kryteriów) algorytmu odbioru a możliwościami jego technicznej realizacji. Dotychczasowy rozwój techniki odbiorczej szedł głównie w kierunku realizacji optymalnych algorytmów przy pomocy urządzeń elektronicznych pracujących na zasadzie analogowej, a więc posiadających nieprzeliczalny zbiór stanów wejściowych. Urządzeniom tym właściwe są: z jednej strony — względna prostota konstrukcji i stosunkowo niewielka ilość elementów potrzebnych do zrealizowania typowych operacji matematycznych występujących w technice odbioru (np. wyliczenia korelacji wzajemnej sygnałów, sumowania koherentnego itp.), z drugiej strony jednak — właściwa wszelkim układom analogowym mała dokładność wykonania tych operacji, uzależniona od takich zjawisk ubocznych, jak np. ograniczenie dynamiki sygnału spowodowane nieliniowością charakterystyk urządzeń elektronicznych, ograniczenie pamięci sumatorów spowodowane nieuniknionym tłumieniem sygnału w elementach opóźniających itp. Trudności te spowodowały w ostatnich latach wzrost zainteresowania techniką cyfrowej realizacji optymalnych algorytmów detekcji sygnałów [1], [2], [3], [4]. Ściśle biorąc, opisane w literaturze układy tego typu mają charakter analogowo-cyfrowy, przy czym w części analogowej dokonuje się wstępnych operacji matematycznych nad sygnałem wejściowym wraz z operacją kwantowania wartości chwilowej sygnału, a w części cyfrowej — operacji

zliczania impulsów skwantowanych lub pewnych prostych operacji logicznych nad ciągami impulsów, jak np. znalezienia środka serii impulsów w pewnych układach estymacyjnych. Tego rodzaju układy nie wyczerpują oczywiście wszystkich możliwości zastąpienia operacji analogowych przez cyfrowe w układach detekcyjnych. Wydaje się, że w perspektywie należy dążyć do pełnego wyeliminowania operacji analogowych w części decyzyjnej urządzeń. Mając powyższe na celu rozważymy dalej niektóre teoretyczne aspekty zagadnienia cyfrowej realizacji optymalnego algorytmu odbioru sygnałów radiolokacyjnych na tle zakłóceń korelowanych. Uwzględnienie korelacji między kolejnymi próbkami sygnału odbieranego stanowi dość istotny warunek zbliżenia matematycznego modelu do rzeczywistej sytuacji, z jaką mamy do czynienia w radiolokacji i niektórych innych dziedzinach łączności. Z założeniem tym wiążą się jednak dość istotne trudności, gdyż optymalizacji musi podlegać nie tylko wartość progu kwantowania sygnału odbieranego, lecz także sama struktura układu decyzyjnego, tj. sposób wykonania operacji logicznych na skwantowanych próbkach. W sumie jest to zatem typowo nieanalityczne zagadnienie optymalizacyjne, być może dlatego właśnie dotychczasowe publikacje poświęcone są głównie układom kwantowym pracującym przy założeniu statycznej niezależności sygnałów wejściowych. W niniejszej pracy, kosztem pewnych ograniczeń nałożonych na sygnał wejściowy, trudność tę udało się ominąć, czego wynikiem jest praktycznie przydatna, choć wiążąca się z pewnymi trudnościami rachunkowymi, metoda łącznej optymalizacji struktury układu oraz dobierania wartości progów kwantowania. Co więcej, metoda ta nie wymaga konieczności wprowadzania pośrednich kryteriów optymalności (na przykład informacyjnych), lecz bazuje wprost na kryterium minimum strat detekcji. Ograniczenia, którymi przychodzi zapłacić za wymienione zalety metody, to: a) założenie, że sygnały wejściowe są procesami Markowa (dowolnie wysokiego, lecz skończonego rzędu), b) założenie, że ciąg sygnałów wejściowych jest dostatecznie długi, aby błędy wynikłe ze stosowania twierdzeń granicznych nie odgrywały praktycznej roli. Zauważmy, że w radiolokacji oba powyższe założenia nie są nowe: powszechnie stosowane układy eliminacji zakłóceń pasywnych z jednokrotnym odejmowaniem są optymalne w sensie minimum strat detekcji tylko dla zakłóceń o wykładniczej funkcji korelacji (proces Markowa pierwszego rzędu) i dla nieskończenie długich ciągów sygnałów odbieranych. Rezygnacja z tych założeń, które wydają się dość nierealne, prowadzi jednak do komplikacji układowych, zwykle nieuzasadnionych z technicznego punktu widzenia. W praktyce zatem często znajduje się układ optymalny dla prostszych, choć mniej realnych założeń, a następnie bada się przydatność tak otrzymanego układu w warunkach zbliżonych do rzeczywistości. Taką drogę postępowania sugeruje się także i w naszym przypadku.

2. OGÓLNE ZASADY PROJEKTOWANIA CYFROWYCH UKŁADÓW DECYZYJNYCH

Oznaczmy przez $X(t)$ proces przypadkowy obserwowany na wejściu urządzenia odbiorczego. Załóżmy, że zarówno w przypadku występowania sygnału i zakłóceń (hipoteza H_1), jak i samych zakłóceń (hipoteza H_0) proces $X(t)$ może być aproksymowany przy pomocy procesu Markowa k -go rzędu. Jest to uproszczenie polegające na traktowaniu sygnału użytecznego jako stochastycznego procesu gaussowskiego. Suma funkcji korelacji sygnału i statystycznie od niego niezależnych zakłóceń addytywnych może być aproksymowana przy pomocy funkcji korelacji odpowiadającej procesowi Markowa odpowiednio wysokiego rzędu (zob. na przykład [5] rozdz. I). Ponieważ proces Markowa rzędu k' zawsze można przedstawić formalnie jako proces rzędu $k > k'$, więc wygodnie będzie przez rząd k dwóch procesów rozumieć tę spośród liczb określających rząd procesu, jaki odpowiada zakłóceniom oraz sygnałowi i zakłóceniom, która jest większa.

W urządzeniach radiolokacyjnych typu impulsowego sygnał odbierany odpowiadający ustalonemu wycinkowi analizowanej przestrzeni może być (przy pominięciu nieistotnego dla nas kształtu impulsów sondujących) przedstawiony w postaci skończonego lub przeliczalnego ciągu próbek o nieprzeliczalnym zbiorze wartości. Ponieważ z założenia są to próbki odpowiadające procesom Markowa k -go rzędu, więc mają one własności łańcuchów Markowa k -go rzędu o nieprzeliczalnym zbiorze wartości. Jeżeli przez x_k ($k = 1, 2, \dots, k$) oznaczmy wartość próbki sygnału wejściowego odpowiadającą k -mu momentowi próbkowania, to gęstości prawdopodobieństwa sygnału wejściowego X będą miały następujące własności:

$$w_{SZ}(x_k | x_{k-1}, x_{k-2}, \dots, x_1) = w_{SZ}(x_k | x_{k-1}, x_{k-2}, \dots, x_{k-k'}), \quad (1a)$$

$$\begin{aligned} w_{SZ}(x_k, x_{k-1}, \dots, x_1) &= w_{SZ}(x_k | x_{k-1}, \dots, x_{k-k}) w_{SZ}(x_{k-1}, x_{k-2}, \dots, x_1) = \dots \\ &= w_{SZ}(x_k | x_{k-1}, \dots, x_{k-k}) w_{SZ}(x_{k-1} | x_{k-2}, \dots, x_{k-k-1}) \dots \\ &\dots w_{SZ}(x_{k+1} | x_k, \dots, x_1) w_{SZ}(x_k, x_{k-1}, \dots, x_1), \end{aligned} \quad (1b)$$

$$w_Z(x_k | x_{k-1}, x_{k-2}, \dots, x_1) = w_Z(x_k | x_{k-1}, x_{k-2}, \dots, x_{k-k''}), \quad (1c)$$

$$\begin{aligned} w_Z(x_k, x_{k-1}, \dots, x_1) &= w_Z(x_k | x_{k-1}, \dots, x_{k-k}) w_Z(x_{k-1}, x_{k-2}, \dots, x_1) = \dots \\ &= w_Z(x_k | x_{k-1}, \dots, x_{k-k}) w_Z(x_{k-1} | x_{k-2}, \dots, x_{k-k-1}) \dots \\ &\dots w_Z(x_{k+1} | x_k, \dots, x_1) w_Z(x_k, x_{k-1}, \dots, x_1), \end{aligned} \quad (1d)$$

$$k = \max(k', k'') \quad (1e)$$

Z ogólnej teorii detekcji wiadomo, że minimalne prawdopodobieństwo błędu detekcji otrzymuje się w przypadku przyjęcia następującego algorytmu odbioru:

$$\begin{aligned} H &= H_1 \text{ dla } \Lambda(\bar{x}) > \Lambda_0, \bar{x} = (x_1, x_2, \dots, x_K), \\ H &= H_0 \text{ dla } \Lambda(\bar{x}) \leq \Lambda_0, \end{aligned} \quad (2)$$

gdzie:

$$\Lambda(\bar{x}) = \frac{w_{SZ}(\bar{x})}{w_Z(\bar{x})} \text{ jest tzw. funkcją wiarygodności, a}$$

$\Lambda_0 = \frac{\alpha_{10}(1-p_1)}{\alpha_{01}p_1}$ jest pewną wartością progową zależną od prawdopodobieństwa a priori pojawienia się celu p_1 oraz od wag przypisywanych różnym rodzajom błędów detekcji (błędem I i II rodzaju) — α_{10} i α_{01} .

Na podstawie zależności (1) funkcję wiarygodności $\Lambda(\bar{x})$ można przedstawić w postaci

$$\Lambda(\bar{x}) = \frac{w_{SZ}(x_K, \dots, x_1)}{w_Z(x_K, \dots, x_1)} = \frac{w_{SZ}(x_K | x_{K-1}, \dots, x_{K-k})}{w_Z(x_K | x_{K-1}, \dots, x_{K-k})} \frac{w_{SZ}(x_{K-1}, \dots, x_1)}{w_Z(x_{K-1}, \dots, x_1)}. \quad (3)$$

Zależność powyższa pozwala zrealizować funkcję wiarygodności w postaci zależności rekurencyjnej

$$\Lambda_K(\bar{x}) = \lambda(x_K; x_{K-1}, \dots, x_{K-k}) \Lambda_{K-1}(\bar{x}), \quad (4)$$

gdzie:

$$\Lambda_K(\bar{x}) = \Lambda(x_K, x_{K-1}, \dots, x_1) = \frac{w_{SZ}(x_K, \dots, x_1)}{w_Z(x_K, \dots, x_1)}, \quad (4a)$$

$$\lambda(x_K, x_{K-1}, \dots, x_{K-k}) = \frac{w_{SZ}(x_K | x_{K-1}, \dots, x_{K-k})}{w_Z(x_K | x_{K-1}, \dots, x_{K-k})}, \quad (4b)$$

lub też w postaci

$$\Phi_K(\bar{x}) = \varphi(x_K, x_{K-1}, \dots, x_{K-k}) + \Phi_{K-1}(\bar{x}), \quad (5)$$

gdzie:

$$\Phi_K(\bar{x}) = \log \Lambda_K(\bar{x}), \quad (5a)$$

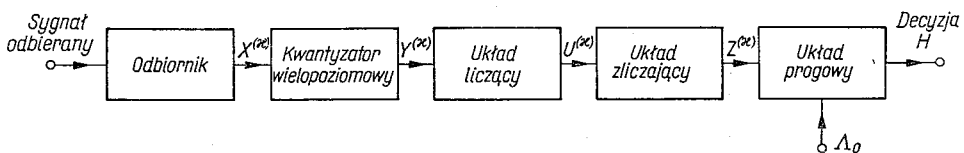
$$\varphi(x_K, x_{K-1}, \dots, x_{K-k}) = \log \lambda(x_K; x_{K-1}, \dots, x_{K-k}). \quad (5b)$$

Z punktu widzenia techniki realizacji optymalnego algorytmu zarówno postać (4), jak i (5) nadają się do zrealizowania w postaci układów liczących pracujących np. w systemie binarnym. Z wyrażen tych widać, że optymalny układ decyzyjny powinien mieć postać przedstawioną na rys. 1. Sposób działania układu jest następujący. Impulsowy sygnał wejściowy $X(t)$ ($t = t_1, t_2, \dots, t_k$) o nieprzeliczalnym zbiorze war-

tości jest podawany po wzmacnieniu i detekcji obwiedni na urządzenie kwantujące, transformujące wartości chwilowe sygnałów z przestrzeni ciągłej Ω_X do przestrzeni $(n+1)$ -elementowej Ω_n (gdzie $n+1$ — ilość dopuszczalnych wartości sygnału skwantowanego), przy czym funkcja opisująca układu kwantującego

$$y(t) = T\{x(t); t\}, \quad t = t_1, t_2, \dots, t_K, \quad (6)$$

w ogólnym przypadku (na przykład sygnałów niestacjonarnych) może zależeć od czasu t , tzn. że wartości kwantów mogą być zmienne w czasie, natomiast ich ogólna ilość narzucona przez warunki konstrukcyjne niech będzie ustalona. Tak otrzymany skwantowany w poziomie sygnał $Y(t)$



Rys. 1

jest podawany na wejście urządzenia liczącego, wyliczającego bieżące wartości funkcji λ lub φ , które następnie są zliczane dając wartość funkcji Δ_K (lub Φ_K). Ta ostatnia jest wreszcie porównywana z progiem decyzji Δ_0 w urządzeniu progowym pracującym na zasadzie analogowej lub cyfrowej stanowiąc podstawę podjęcia odpowiedniej decyzji. Jak widać z powyższego opisu, o jakości podejmowanych decyzji w pierwszym rzędzie decyduje sposób zaprojektowania układu kwantującego i liczącego, gdyż jest on związany z przyjętą metodą cyfrowej aproksymacji optymalnej funkcji decyzyjnej $\Delta(\bar{x})$. Sposobowi doboru głównych parametrów tych układów poświęcimy dalszy ciąg naszych rozważań przyjmując jednocześnie, że pozostałe bloki wchodzące w skład urządzenia decyzyjnego mają charakter konwencjonalny.

Założmy początkowo, że na wejściu urządzenia odbiorczego dany jest bezpośrednio sygnał skwantowany $Y(t)$ o wartościach należących do

$K(n+1)$ -elementowego zbioru $E_{K,n} = \bigcup_{n=1}^K \Omega_n^{(n)}$, gdzie $\Omega_n^{(n)} = \{\omega_0^{(n)}, \omega_1^{(n)}, \dots, \omega_n^{(n)}\}$

jest zbiorem poziomów kwantowych dla momentu próbkowania $t = t_n$, a symbol X oznacza iloczyn kartezyjski zbiorów.

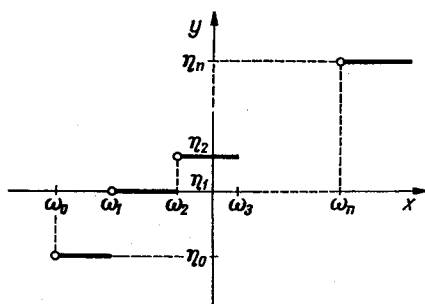
Z formalnego punktu widzenia proces $Y(t)$ może być scharakteryzowany w przybliżeniu jako proces Markowa k -go rzędu, przybierający „wartości” należące do $(n+1)^{k+1}$ -elementowego zbioru $S_{k,n}$ $(k+1)$ -elementowych ciągów postaci $\{y_{a_1}, y_{a_2}, \dots, y_{a_{k+1}}\}$, gdzie $y_{a_n} \in \Omega_n^{(n)}$ ($n = a_1, a_2, \dots, a_{k+1}$). Własności statystyczne takiego procesu określają w pełni prawdopodobieństwa początkowe oraz $(n+1)^{k+1} \cdot (n+1)^{k+1}$ -elementowe macierze prawdopodobieństw przejść:

$$\mathbf{P}_{SZ}^{(\infty)} = \begin{bmatrix} p_{11}^{(\infty)} & \dots & p_{1N}^{(\infty)} \\ \vdots & \ddots & \vdots \\ p_{N1}^{(\infty)} & \dots & p_{NN}^{(\infty)} \end{bmatrix}, \quad \mathbf{P}_Z^{(\infty)} = \begin{bmatrix} q_{11}^{(\infty)} & \dots & q_{1N}^{(\infty)} \\ \vdots & \ddots & \vdots \\ q_{N1}^{(\infty)} & \dots & q_{NN}^{(\infty)} \end{bmatrix}, \quad (7)$$

$$N = (n+1)^{k+1}, \quad k = \max(k', k''), \quad n = 1, 2, 3, \dots, \quad (7a)$$

$$\sum_{j=1}^N p_{ij}^{(\infty)} = \sum_{j=1}^N q_{ij}^{(\infty)} = 1 \quad \text{dla } i = 1, 2, \dots, N, \quad \kappa = 1, 2, \dots, K, \quad (7b)$$

przez $p_{ij}^{(\infty)}$ (odpowiednio $q_{ij}^{(\infty)}$) oznaczono tu prawdopodobieństwo zaobserwowania j -tej wartości wycinkowej procesu $Y(t)$, jeżeli poprzednio (w chwili t_*) zaobserwowano i -tą wartość wycinkową tego procesu, a na wejście przyłożono sygnał użyteczny wraz z zakłóceniem (odpowiednio — samo zakłócenie). Przez „wartość wycinkową” rozumiemy tu ciąg wartości chwilowych procesu $Y(t)$ obserwowanych przez $k+1$ kolejnych momentów czasu aż do momentu bieżącego włącznie. W przypadku przejścia od ciągu zaobserwowanych wartości postaci $\{y_{\alpha_0}, y_{\alpha_1}, \dots, y_{\alpha_k}\}$ do ciągu $\{y_{\beta_0}, y_{\beta_1}, \dots, y_{\beta_k}\}$ prawdopodobieństwo przejścia jest równe tożsamościowo zero, o ile nie są spełnione warunki identyczności $y_{\alpha_i} = y_{\beta_{i-1}}$ ($i = 1, 2, \dots, k$), zaś w przypadku spełnienia tych warunków, jest równe prawdopodobieństwu zaistnienia elementu y_{β_k} , jeżeli poprzednio zaobserwowano ciąg $\{y_{\alpha_1}, \dots, y_{\alpha_k}\}$. Wyliczenie niezerowych prawdopodobieństw przejścia $p_{ij}^{(\infty)}$ i $q_{ij}^{(\infty)}$ w przypadku znajomości funkcji gęstości prawdopodobieństwa sygnału ciągłego $w_{SZ}(x)$, $w_Z(x)$ i zbioru $E_{K,n}$ nie nastrocza większych trudności. Jeśli mianowicie przyjmiemy, że dla



Rys. 2

stacjonarnego sygnału wejściowego element kwantujący ma charakterystykę (zob. rys. 2)

$$y = \omega_i \quad \text{dla } x \in \eta_i, \quad (8)$$

gdzie

$$\eta_i = (\omega_i, \omega_{i+1}), \quad i = 0, 1, 2, \dots, n, \quad (8a)$$

to oznaczając przez $H_{\kappa}(\omega_i)$ obszar cylindryczny postaci

$$H_{\kappa}(\omega_i) = \prod_{\gamma=1}^{\kappa-k-1} R_1^{(\gamma)} \times \prod_{\delta=\kappa-k}^{\kappa} \eta_{i_{\delta}} \times \prod_{\varepsilon=\kappa+1}^K R_1^{(\varepsilon)}, \quad (9)$$

gdzie $R_1^{(\gamma)}$ oznacza oś liczb rzeczywistych będącą zbiorem wartości zmiennej $X(t_{\gamma})$, a wszystkie iloczyny są rozumiane w sensie kartezjańskim, tj. oznaczając przez H_{κ} zbiór wszystkich realizacji procesu $X(t)$ ($t = t_1, t_2, \dots, t_K$), które w momencie $t_{\kappa-k}$ przybierają wartości należące do $\eta_{i_{\kappa-k}}$ w momencie $t_{\kappa+1}$ — wartości należące do $\eta_{i_{\kappa-k+1}}$ itd. aż do t_{κ} , to szukane prawdopodobieństwa wyrażą się wzorami (wyrażenia w nawiasach $\{ \}$ określają przedziały całkowania):

$$p_{\alpha\beta} = \frac{\int w_{SZ}(x) dx}{\int w_{SZ}(x) dx} \cdot \frac{\{H_{\kappa-1}(\omega_{\alpha}) \cap H_{\kappa}(\omega_{\beta})\}}{\{H_{\kappa-1}(\omega_{\alpha})\}}, \quad \kappa > k, \quad (10a)$$

$$q_{\alpha\beta} = \frac{\int w_Z(x) dx}{\int w_Z(x) dx} \cdot \frac{\{H_{\kappa-1}(\omega_{\alpha}) \cap H_{\kappa}(\omega_{\beta})\}}{\{H_{\kappa-1}(\omega_{\alpha})\}}. \quad (10b)$$

Znając macierze prawdopodobieństw P_{SZ} i P_Z można zaprojektować układ liczący wyliczający funkcję $\lambda(y)$ (lub $\varphi(y)$). Ponieważ sygnał na wejściu tego układu zgodnie z założeniem jest skwantowany w poziomie, więc układ ten można traktować jako pewnego rodzaju automat skończony. Będziemy tak nazywali układ o skończonym zbiorze stanów wejściowych $\Omega_n^{(\kappa)}$, skończonym zbiorze stanów wyjściowych $U^{(\kappa)}$ i skończonej pamięci obejmującej k dyskretnych jednostek czasu. Reakcja układu $u^{(\kappa)}$ w chwili t_{κ} ($\kappa > k$) niech będzie jednoznaczną funkcją wymuszenia w chwili $t_{\kappa} - \omega_{i_{\kappa}}$ i stanu wewnętrznego układu $v^{(\kappa)}$ w chwili t_{κ} . Ten ostatni niech będzie z kolei funkcją ciągu $\{\omega_{i_{\kappa-k}}, \omega_{i_{\kappa-k+1}}, \dots, \omega_{i_{\kappa-1}}\}$ wymuszeń w k poprzedzających momentach czasu. Można to zapisać jak następuje.

$$u^{(\kappa)} = F(\omega_{i_{\kappa}}; v^{(\kappa)}; t_{\kappa}), \quad (11)$$

$$v^{(\kappa)} = G(\omega_{i_{\kappa-k}}, \omega_{i_{\kappa-k+1}}, \dots, \omega_{i_{\kappa-1}}; t_{\kappa}), \quad (12)$$

przy czym zakładamy, że funkcje F i G mogą zależeć od bieżącej chwili czasu t_{κ} (a nie tylko od numeru okresu próbkowego κ). Wyjście układu dla $\kappa \leq k$ można przyjąć za równe tożsamościowo zeru. W wyrażeniach (11) i (12) celowo traktujemy odrębnie wpływ k przedostatnich oraz ostatniego wymuszenia wejściowego, gdyż ułatwia to syntezę układu, w którym k przedostatnich wymuszeń wejściowych jest zapamiętywane w postaci pewnego zbioru wewnętrznych stanów energetycznych determinujących odpowiedź układu na ostatnie, $(k+1)$ -sze wymuszenie. Automat taki może pracować jako układ liczący, jeżeli zbiór stanów

wyjściowych $U^{(\infty)}$ utożsamimy z liczbami $\frac{p_{\alpha\beta}}{q_{\alpha\beta}}$ (albo $\log \frac{p_{\alpha\beta}}{q_{\alpha\beta}}$) wyrażonymi w pewnym kodzie, dogodnym dla przeprowadzenia operacji mnożenia (lub sumowania) w układzie zliczającym, o dowolnej, lecz skończonej liczbie symboli. Ten ostatni warunek narzuca pewne ograniczenia na dokładność wyliczenia funkcji λ lub φ , które można interpretować jako wynik operacji wtórnego kwantowania sygnału na wyjściu układu liczącego; w dobrze skonstruowanym urządzeniu efekt ten jest jednak do pominięcia. Zaletą powyższej interpretacji działania układu liczącego jest możliwość zastosowania szeroko w literaturze specjalnej (zob. na przykład [7]) omówionych metod syntezy automatów skończonych do projektowania kwantowych urządzeń decyzyjnych.

Jeżeli zbiór $S_{K,n}$ wszystkich K -elementowych ciągów $\bar{y} = \{y_n\}$, $y_n \in \Omega_n^{(\infty)}$ rozbijemy na dwa dopełniające się podzbiory $S_{K,n}^{(0)}$, $S_{K,n}^{(1)}$ według zasady

$$\begin{aligned} \bar{y} \in S_{K,n}^{(0)}, & \text{ jeśli } \Lambda_K(\bar{y}) \leq \Lambda_0, \\ \bar{y} \in S_{K,n}^{(1)}, & \text{ jeśli } \Lambda_K(\bar{y}) > \Lambda_0. \end{aligned} \quad (13)$$

analogicznej do zasady (2) dla sygnałów nieskwantowanych w poziomie, to średnie straty detekcji wyrażą się wzorem:

$$\bar{r} = p_0 a_{10} \cdot \sum_{\{\bar{y} \in S_{K,n}^{(1)}\}} w_Z(\bar{y}) + p_1 a_{01} \sum_{\{\bar{y} \in S_{K,n}^{(0)}\}} w_{SZ}(\bar{y}), \quad (14)$$

gdzie $w_{SZ}(\bar{y})$, $w_Z(\bar{y})$ oznaczają warunkowe prawdopodobieństwa ciągu $\bar{y} = \{y_n\}$, $n = 1, 2, \dots, K$, obliczone według wzorów analogicznych do (1a) ... (1e), w których prawdopodobieństwa warunkowe $w_{SZ}(y_n | y_{n-k}, \dots, y_{n-1})$ są tożsame z macierzami przejścia P_{SZ} , P_Z zadanymi przez wzory (7).

3. DOBÓR OPTYMALNYCH POZIOMÓW KWANTOWANIA

W poprzednich rozważaniach założyliśmy, że sygnał dany jest na wejściu urządzenia odbiorczego od razu w postaci skwantowanej. Przypuśćmy teraz, że warunek ten nie jest spełniony i urządzenie liczące jest poprzedzone przez układ kwantujący o charakterystyce (8). Powstaje problem, w jaki sposób należy dobrać wartości ω_i , $i = 0, 1, 2, \dots, n$, znając warunkowe funkcje gęstości $w_{SZ}(\bar{x})$ i $w_Z(\bar{x})$ oraz liczby a_{10} , a_{01} , p_0 i p_1 , aby straty detekcji były minimalne. Problem ten rozwiązaliśmy przy pewnych dodatkowych założeniach upraszczających.

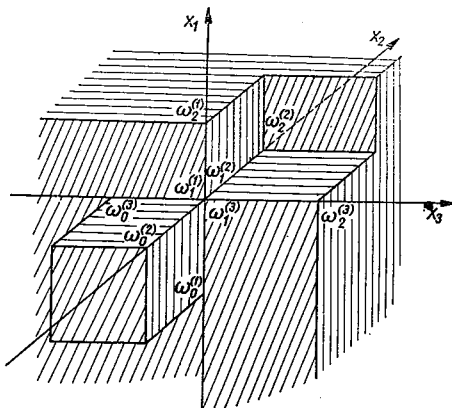
Zastąpienie ciągłego układu decyzyjnego układem kwantowym może być interpretowane geometrycznie jako przybliżenie pewnej ciągłej hiperpowierzchni w K -wymiarowej przestrzeni Euklidesa $(\Omega_X)^K$, zadanej równaniem

$$\Lambda(\bar{x}) - \Lambda_0 = 0 \quad (15a)$$

przez pewną hiperpowierzchnię schodkową (zob. rys. 3):

$$\Lambda(T(\bar{x})) - \Lambda_0 = 0, \quad (15b)$$

gdzie $T(\bar{x})$ jest charakterystyką elementu kwantującego.



Rys. 3

Srednie straty detekcji w układzie kwantującym można wyrazić następująco:

$$\begin{aligned} \bar{r} &= p_0 a_{10} \cdot \int_{\{\bar{x}: T(\bar{x}) \in S_{K,n}^{(1)}\}} w_Z(\bar{x}) d\bar{x} + p_1 a_{01} \cdot \int_{\{x: T(x) \in S_{K,n}^{(0)}\}} w_{SZ}(x) dx = \\ &= p_0 a_{10} \cdot \sum_{\{\bar{y} \in S_{K,n}^{(1)}\}} \int_{\{T(\bar{x}) = \bar{y}\}} w_Z(\bar{x}) d\bar{x} + p_1 a_{01} \cdot \sum_{\{\bar{y} \in S_{K,n}^{(0)}\}} \int_{\{T(\bar{x}) = \bar{y}\}} w_{SZ}(\bar{x}) d\bar{x}. \end{aligned} \quad (16)$$

Optymalny układ kwantujący będzie zatem układem, którego charakterystyka $\bar{y} = T(\bar{x})$ minimalizuje wielkość (16). Ponieważ jednak związek między kształtem hiperpowierzchni schodkowej $\bar{y} = T(\bar{x})$ i wartościami ω_i jest dosyć złożony, wygodniej będzie to samo zadanie sformułować nieco inaczej.

W wyniku skwantowania sygnału wejściowego ciągly proces Markowa k -go rzędu $X(t)$ zostaje zastąpiony przez dyskretny łańcuch Markowa (zob. Dodatek I) $Y(t)$ tegoż rzędu k . Własności statystyczne łańcucha Markowa są w pełni określone przez rozkład prawdopodobieństwa stanów początkowych i macierze przejścia $\mathbf{P}_{SZ}^{(n)}$ i $\mathbf{P}_Z^{(n)}$, które w przypadku procesów stacjonarnych są niezależne od numeru bieżącej chwili czasowej n . Elementy tych macierzy p_{ij} i q_{ij} , $i, j = 1, 2, \dots, (n+1)^k$ zależą od aktualnego wyboru progów kwantowania; muszą one zatem być dobrane w sposób optymalny z punktu widzenia rozróżnialności sygnału i zakłóceń. Optymalny kwantowy układ decyzyjny powinien i w tym przypadku realizować algorytm określony wyrażeniem (13); progi kwantowania należy zatem wybrać tak, aby wyrażenie (14) uczynić minimalnym.

W przypadku krótkich serii sygnałów odbieranych (małych K) średnie straty detekcji można obliczyć na drodze bezpośredniej wyliczając warunkowe (ze względu na obecność lub brak sygnału użytecznego) prawdopodobieństwa ciągów należących do klas $S_{K,n}^{(0)}$ i $S_{K,n}^{(1)}$. Elementy kolejnych potęg macierzy przejścia można wówczas obliczać w oparciu o znany z teorii macierzy wzór Perrona [8]. W praktyce procedura taka staje się zbyt kłopotliwa już przy $K > 5$, w związku z czym dla celów praktycznych obliczeń wygodnie jest posłużyć się zależnościami asymptotycznymi, słusznymi dla dużych K , a zwłaszcza twierdzeniem I (zob. Dodatek I). Powstaje w związku z tym problem oceny błędu popełnianego przy zastosowaniu powyższego twierdzenia przy skończonym K . Jak wynika z dowodu twierdzenia (zob. Sarymsakow [8], [14]), błąd ten jest rzędu $K^{-1/2}$. Pozwala to ocenić względny wpływ długości próbki na dokładność przybliżenia; ocena bezwzględnej wartości błędu w każdym indywidualnym przypadku wiąże się jednak z dużymi trudnościami, ponieważ błąd indywidualny zależy explicite od elementów macierzy przejścia i zbioru wartości procesu skwantowanego $Y(t)$. W praktyce należy liczyć się z wartościami $K \approx 20$, przy których rozkład graniczny daje oszacowanie dość dobre.

Wyliczenie średnich strat detekcji $\tilde{r}$ wymaga zatem oszacowania asymptotycznych własności zmiennej przypadkowej na wyjściu układu zliczającego, $Z = A_K(Y)$ przy dużych K . W tym celu wykorzystamy pewne własności macierzy przejścia $\mathbf{P}_{SZ}$ i $\mathbf{P}_Z$. Elementy tych macierzy, p_{ij} i q_{ij} , oznaczają warunkowe prawdopodobieństwa przejścia od i -tej do j -tej wartości wycinkowej procesu skwantowanego, przy czym wartości te są brane w dwóch momentach czasu odległych o okres próbkowania T . Jeżeli będziemy rozważać bardzo długie ciągi ($K \rightarrow \infty$) próbek, to sygnał na wyjściu układu zliczającego będzie sumował zmienne przypadkowe $U^{(\kappa)} = \log \frac{p_{\alpha\beta}}{q_{\alpha\beta}}$, $\kappa = 1, 2, \dots, K$ tworzące łańcuch Markowa o jednej z dwóch alternatywnych macierzy przejścia $\mathbf{P}_{SZ}$ lub $\mathbf{P}_Z$. Zgodnie z twierdzeniem I (Dodatek I), przy pewnych założeniach odnośnie macierzy $\mathbf{P}_{SZ}$ i $\mathbf{P}_Z$, sygnał na wyjściu logarytmicznego układu zliczającego [tj. wyliczającego wyrażenie (5)] przy $K \rightarrow \infty$ podlega asymptotycznie rozkładowi normalnemu o parametrach $m = (K+1)a$ i $\sigma^2 = -(K+1)\lambda_0''(0)$, których sens podają wyrażenia (I.8) — (I.10). W tych warunkach za optymalny układ decyzyjny mamy prawo uważać układ, który dla pewnego ustalonego i dostatecznie dużego K charakteryzuje się minimalnymi średnimi stratami detekcji określonymi zależnością

$$\tilde{r} = p_0 a_{10} \int w_Z^{(K)}(z) dz + p_1 a_{01} \int w_{SZ}^{(K)}(z) dz, \quad (17)$$

$$\left\{ \frac{w_{SZ}^{(K)}(z)}{w_Z^{(K)}(z)} > A_0 \right\} \quad \left\{ \frac{w_{SZ}^{(K)}(z)}{w_Z^{(K)}(z)} < A_0 \right\}$$

gdzie $w_{SZ}^{(K)}(z)$ i $w_Z^{(K)}(z)$ oznaczają asymptotyczne rozkłady normalne otrzy-

mane zgodnie z Twierdzeniem I dla K -elementowego łańcucha Markowa i dla macierzy przejścia równych odpowiednio $\mathbf{P}_{SZ}$ i $\mathbf{P}_Z$.

Cała trudność oszacowania jakości kwantowego układu decyzyjnego, a więc i doboru metodą kolejnych prób jego optymalnego wariantu, sprowadza się zatem do znalezienia parametrów dwóch rozkładów normalnych: a_{SZ} , $\lambda''_0(0)_{SZ}$, a_Z i $\lambda''_0(0)_Z$, charakteryzujących przyrost na jeden odebrany impuls wartości średniej i dyspersji sygnału wyjściowego w przypadku sygnału przychodzącego wraz z szumem oraz w przypadku samego szumu. Wielkości te można wyznaczyć wprost w oparciu o zależności (I.8) — (I.10). Niestety, próba rozwiązania równania charakterystycznego (I.9) już dla stosunkowo prostych macierzy przejścia napotyka na duże trudności, gdyż nieznany parametr a występuje w równaniach w postaci uwikłanej. Dlatego w celu praktycznego określenia optymalnych progów kwantowania ω_i , dla wstępnych obliczeń, które następnie można uściślić w oparciu o podane wyrażenia bardziej dokładne, wygodnie będzie przyjąć za optymalne takie wartości ω_i , którym odpowiada maksymalna bezwzględna różnica przyrostów wartości średnich na jeden odebrany impuls:

$$\Delta a = |a_{SZ} - a_Z|. \quad (18)$$

W celu oszacowania wielkości a_{SZ} i a_Z można posłużyć się Twierdzeniem II. Twierdzenie to podaje praktyczną metodę określenia tzw. prawdopodobieństw granicznych dla stacjonarnego i ergodycznego łańcucha Markowa, tj. niezależnych od wartości początkowej asymptotycznych prawdopodobieństw występowania poszczególnych wartości w dostatecznie długich odcinkach obserwowanych procesów. Jeżeli zatem każdej wartości

$$u^{(s)} = \log \frac{(p_{\alpha\beta})_s}{(q_{\alpha\beta})_s}. \quad (19)$$

zmiennej przypadkowej $U^{(s)}$ przypiszemy dwa ciągi prawdopodobieństw granicznych $\bar{p}_{SZ} = (\bar{p}_1, \bar{p}_2, \dots, \bar{p}_s)$ i $\bar{p}_Z = (\bar{q}_1, \bar{q}_2, \dots, \bar{q}_s)$ odpowiadające macierzom przejścia $\mathbf{P}_{SZ}$ i $\mathbf{P}_Z$, to bez trudu obliczymy też odpowiednie uśrednione statystycznie przyrosty zmiennych na wyjściu układu zliczającego:

$$a_{SZ} = \sum_{i=1}^s \bar{p}_i u_i, \quad (20a)$$

$$a_Z = \sum_{i=1}^s \bar{q}_i u_i, \quad (20b)$$

przy czym sumowanie obejmuje wszystkie możliwe wartości wycinkowe rozważanego procesu, tj. $s = (n+1)^{k+1}$.

Wyliczając w oparciu o powyższą uproszczoną procedurę wartości różnic Δa odpowiadające różnym poziomom kwantowania ω możemy dobrać w pierwszym przybliżeniu optymalną wartość ω , przy której różnice te są największe. Dla tak obliczonych wartości progowych można następnie przystąpić do szacowania bardziej ścisłego, w którym sygnał na wyjściu układu zliczającego traktuje się jako zmienną przypadkową podlegającą asymptotycznemu rozkładowi normalnemu (I.11). W celu oszacowania parametrów tego rozkładu (wartości średnich a_z lub a_{sz} oraz wariancji σ_z^2 lub σ_{sz}^2) należy posłużyć się ponownie zależnościami (20a), (20b), a następnie posłużyć się zależnościami (I.8) i (I.9), z których po podstawieniu odpowiednich wartości średnich a_z lub a_{sz} wyliczamy wielkości $\frac{d^2}{d\theta^2} \lambda(\theta)|_{\theta=0}$ określające szukane wartości wariancji [zob. wzór (I.11)]. Jak widać, uzyskanie drugich przybliżeń wiąże się już z dużymi trudnościami rachunkowymi i praktycznie nie może się obyć bez elektronowej maszyny liczącej. Opracowanie programu dla takich obliczeń wymaga jednak dalszego sprecyzowania własności macierzy przejścia $\mathbf{P}_z$ i $\mathbf{P}_{sz}$, z którymi możemy mieć do czynienia w praktyce.

Zbadajmy bliżej postać macierzy przejścia $\mathbf{P}_{zs}$ i $\mathbf{P}_z$ dla typowych procesów, z jakimi mamy do czynienia w technice odbioru sygnałów radiolokacyjnych na tle zakłóceń. Rozważymy w tym celu dwa przypadki szczególne.

a) Łańcuch Markowa prosty ($k = 1$) stacjonarny, układ kwantujący wielopoziomowy ($n \geq 1$). Wszystkie wartości wycinkowe łańcucha możemy uszeregować jak następuje (poszczególne wartości wycinkowe należy odczytywać pionowo):

$$\begin{array}{ll} \text{Lp.: } 1\ 2\ 3\ \dots\ n+1\ \dots & \dots\ (n+1)^2 \\ \text{I cyfra: } 0\ 0\ 0\ \dots\ 0 & 1\ 1\ 1\ \dots\ 1\ \dots\ n\ n\ n\ \dots\ n, \\ \text{II cyfra: } 0\ 1\ 2\ \dots\ n & 0\ 1\ 2\ \dots\ n\ \dots\ 0\ 1\ 2\ \dots\ n \end{array} \quad (21)$$

Macierz przejścia ma zatem budowę następującą:

$$\mathbf{P} = \left\| \begin{array}{cccccccccc} p_{11} & p_{12} & \dots & p_{1,n+1} & 0 & 0 & \dots & 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 & p_{2,n+2} & p_{2,n+3} & \dots & p_{2,2(n+1)} & 0 & \dots & 0 \\ \cdot & \cdot & & \cdot & \cdot & \cdot & & \cdot & \cdot & & \cdot \\ \cdot & \cdot & & \cdot & \cdot & \cdot & & \cdot & \cdot & & \cdot \\ \cdot & \cdot & & \cdot & \cdot & \cdot & & \cdot & \cdot & & \cdot \\ 0 & 0 & \dots & 0 & 0 & 0 & \dots & 0 & 0 & \dots & p_{(n+1)^2, (n+1)^2} \end{array} \right\| \quad (22)$$

Numery wierszy i kolumn odpowiadają w niej numerom kolejnych wartości wycinkowych (21).

Macierz (22) jest macierzą nierozkładalną indeksu $\xi = 1$; opisany przez nią łańcuch Markowa jest stacjonarnym łańcuchem acyklicznym. Acy-

kliczność macierzy $\mathbf{P}$ wynika tu z istnienia niezerowych elementów diagonalnych $p_{11}, p_{2(n+1)+2, 2(n+1)+2}, p_{3(n+1)+3, 3(n+1)+3}, \dots, p_{(n+1)^2, (n+1)^2}$, które powodują, że z prawdopodobieństwem różnym od zera układ może przebywać kilka razy z rzędu w tym samym stanie.

b) Łańcuch Markowa złożony ($k \geq 1$) stacjonarny, układ kwantujący jednopoziomowy ($n = 1$). Wartości wycinkowe łańcucha tworzą zbiór następujący:

$$\left. \begin{array}{cccccc} \text{Lp.} & 1 & \dots & k-1 & k & k+1 \\ 1 & 0 & \dots & 0 & 0 & 0 \\ 2 & 0 & \dots & 0 & 0 & 1 \\ 3 & 0 & \dots & 0 & 1 & 0 \\ 4 & 0 & \dots & 0 & 1 & 1 \\ 5 & 0 & \dots & 1 & 0 & 0 \\ 6 & 0 & \dots & 1 & 0 & 1 \\ 7 & 0 & \dots & 1 & 1 & 0 \\ 8 & 0 & \dots & 1 & 1 & 1 \\ \vdots & \vdots & \dots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \dots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \dots & \vdots & \vdots & \vdots \\ 2^{k+1} & 1 & \dots & 1 & 1 & 1 \end{array} \right\} \quad (23)$$

Odpowiadająca im macierz przejścia ma postać:

$$\mathbf{P} = \left\| \begin{array}{ccccccc} p_{11} & p_{12} & 0 & 0 & \dots & 0 & 0 \\ 0 & 0 & p_{23} & p_{24} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \dots & p_{2^k, 2^{k+1}-1} & p_{2^k, 2^{k+1}} \\ p_{2^k+1, 1} & p_{2^k+1, 2} & 0 & 0 & \dots & 0 & 0 \\ 0 & 0 & p_{2^k+2, 3} & p_{2^k+2, 4} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \dots & p_{2^{k+1}, 2^{k+1}-1} & p_{2^{k+1}, 2^{k+1}} \end{array} \right\| \quad (24)$$

Również i w tym przypadku jest to macierz nierozkładalna, acykliczna ze względu na istnienie niezerowych elementów diagonalnych.

W obu omawianych przypadkach można zatem stosować I i II twierdzenie graniczne z Dodatku I. Prawdopodobieństwa graniczne występujące w wyrażeniach (20a, b) obliczymy przy pomocy wzoru (I.12). W celu ułatwienia obliczeń w Dodatku II przytoczono wyrażenia dla występujących we wzorze (I.12) dopełnień algebraicznych $P_{aa}(\lambda)$, $a = 1, 2, \dots, (n+1)^{k+1}$ w kilku prostszych przypadkach. Pozostaje nam jeszcze określić wyrażenia dla prawdopodobieństw przejścia $p_{a\beta}$, $q_{a\beta}$, występujących jako

elementy macierzy przejścia P_{SZ} i P_Z . Jak wynika z wyrażeń (10a, b), wielkości te zależą od gęstości prawdopodobieństwa sygnału odbieranego $w_{SZ}(\bar{x})$ i $w_Z(\bar{x})$. Ponadto zależą one oczywiście od przyjętych wartości progowych $\{\omega_i\}$, $i = 0, 1, 2, \dots, n$.

Przy dość ogólnych założeniach sygnał radiolokacyjny (który nas tutaj głównie interesuje) odbity od fluktuującego celu wraz z szumem ma rozkład gaussowski; obliczenie prawdopodobieństw przejścia dla tego przypadku jest szczególnie ułatwione. Rozważmy na przykład rozkład normalny dwuwymiarowy:

$$w(x_1, x_2) = \frac{1}{2\pi\sigma^2\sqrt{1-r^2}} \exp \left\{ -\frac{1}{2\sigma^2(1-r^2)} [(x_1 - a_1)^2 - 2r(x_1 - a_1)(x_2 - a_2) + (x_2 - a_2)^2] \right\}. \quad (25)$$

Gęstość prawdopodobieństwa (25) może być zgodnie z metodą Cramera [5] przedstawiona w postaci następującego szeregu:

$$w(x_1, x_2) = \frac{1}{2\pi\sigma^2} \exp \left[-\frac{1}{2} \frac{(x_1 - a_1)^2}{\sigma^2} - \frac{1}{2} \frac{(x_2 - a_2)^2}{\sigma^2} \right] \cdot \sum_{v=0}^{\infty} H_v \left(\frac{x_1 - a_1}{\sigma} \right) H_v \left(\frac{x_2 - a_2}{\sigma} \right) \frac{r^v}{v!}, \quad (26)$$

gdzie

$$H_v(x) = e^{x^2/2} \frac{d^v}{dx^v} e^{-x^2/2}. \quad (26a)$$

Przy pomocy znanej zależności rekurencyjnej dla wielomianów Hermite'a

$$H_v(x) = x H_{v-1}(x) - H'_{v-1}(x) \quad (27)$$

łatwo otrzymać następujące zależności:

$$\int e^{-x^2/2} H_v(x) dx = e^{-x^2/2} H_{v-1}(x), \quad v = 1, 2, 3, \dots, \quad (28)$$

$$\int H_0 e^{-x^2/2} dx = \sqrt{\frac{\pi}{2}} \left[1 + \operatorname{erf} \left(\frac{x}{\sqrt{2}} \right) \right], \quad (28a)$$

znacznie upraszczające procedurę wyliczania całek funkcji $w(x_1, x_2)$ przedstawionej w postaci szeregu (26). Opisany sposób postępowania można zastosować również w przypadkach rozkładów normalnych o większej liczbie argumentów. Dla ułatwienia wyliczeń w Dodatku III podano współczynniki dla wielomianów Hermite'a $H_v(x)$ do $v = 16$ włącznie.

Metodykę obliczania prawdopodobieństw przejściowych i granicznych, stanowiącą istotny moment w proponowanej metodzie oceny układu decyzyjnego, zobrazujemy na następującym przykładzie.

Rozważmy jednopoziomowy ($n = 1$) układ decyzyjny, na którego wejście podano sygnał użyteczny

$$a(t) = a = \text{const}, \quad (29)$$

oraz zakłócenie gaussowskie stacjonarne o wartości średniej równej zeru i funkcji korelacji

$$R_{zz}(\tau) = \sigma^2 e^{-a|\tau|}, a > 0. \quad (30)$$

Jest to przypadek odpowiadający w radiolokacji obserwacji słabo fluktuującego celu o znanej prędkości radialnej na tle zakłóceń pasywnych, przy stosunkowo słabych szumach własnych (np. obserwacja statku na tle odbić od falującej powierzchni morza).

Proces przypadkowy przy tak założonej postaci funkcji korelacji jest prostym procesem Markowa ($k = 1$), co wynika ze specyficznej budowy odwróconej macierzy współczynników kowariancji (zob. na przykład [6]).

Niezerowe prawdopodobieństwa przejścia (zob. Dodatek II) są wówczas określone wzorami:

$$\begin{aligned} p_{11} = p_{31} = P\{y_n = 0 | y_{n-1} = 0\} &= \frac{\pi_{00}(\omega)}{\pi_0(\omega)} = \\ &= \int_{-\infty}^{\omega} dx_n \int_{-\infty}^{\omega} dx_{n-1} w(x_n, x_{n-1}) / \int_{-\infty}^{\omega} dx_{n-1} w(x_{n-1}); \end{aligned} \quad (31a)$$

$$\begin{aligned} p_{23} = p_{43} = P\{y_n = 0 | y_{n-1} = 1\} &= \frac{\pi_{01}(\omega)}{\pi_1(\omega)} = \\ &= \int_{-\infty}^{\omega} dx_n \int_{\omega}^{\infty} dx_{n-1} w(x_n, x_{n-1}) / \int_{\omega}^{\infty} dx_{n-1} w(x_{n-1}); \end{aligned} \quad (31b)$$

$$\begin{aligned} p_{12} = p_{32} = P\{y_n = 1 | y_{n-1} = 0\} &= \frac{\pi_{10}(\omega)}{\pi_0(\omega)} = \\ &= \int_{\omega}^{\infty} dx_n \int_{-\infty}^{\omega} dx_{n-1} w(x_n, x_{n-1}) / \int_{-\infty}^{\omega} dx_{n-1} w(x_{n-1}); \end{aligned} \quad (31c)$$

$$\begin{aligned} p_{24} = p_{44} = P\{y_n = 1 | y_{n-1} = 1\} &= \frac{\pi_{11}(\omega)}{\pi_1(\omega)} = \\ &= \int_{\omega}^{\infty} dx_n \int_{\omega}^{\infty} dx_{n-1} w(x_n, x_{n-1}) / \int_{\omega}^{\infty} dx_{n-1} w(x_{n-1}). \end{aligned} \quad (31d)$$

Występujące w nich prawdopodobieństwa π_{00} , π_{01} , π_{10} , π_{11} , π_0 i π_1 należy określić w dwóch wariantach:

a) dla $w = w_{sz}$, przy wartościach średnich równych a , dyspersjach σ^2 i współczynniku korelacji

$$r = R_{zz}(T) \quad (32)$$

b) dla $w = w_Z$, przy wartościach średnich równych zeru, dyspersjach σ^2 i współczynniku korelacji r określonych jak wyżej.

Zacznijmy od przypadku pierwszego. Podstawiając zamiast $w(x_k, x_{k-1})$ wyrażenie typu (26) i wykonując odpowiednie operacje całkowania, na podstawie wzoru (28) otrzymamy:

$$\begin{aligned} \pi_{00}(\omega) &= \frac{1}{2\pi\sigma^2} \frac{\pi}{2} \left[1 + \operatorname{erf} \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right]^2 + \frac{1}{2\pi\sigma^2} \sum_{p=1}^{\infty} \frac{r^p}{p!} \int_{-\infty}^{\omega} dx_k \exp \cdot \\ &\cdot \left[- \left(\frac{x_k - a}{\sqrt{2}\sigma} \right)^2 \right] H_p \left(\frac{x_k - a}{\sigma} \right) \cdot \int_{-\infty}^{\omega} dx_{k-1} \exp \left[- \left(\frac{x_{k-1} - a}{\sqrt{2}\sigma} \right) \right] H_p \left(\frac{x_{k-1} - a}{\sigma} \right) = \\ &= \frac{1}{2\pi\sigma^2} \frac{\pi}{2} \left[1 + \operatorname{erf} \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right]^2 + \frac{1}{2\pi\sigma^2} \exp \left[- \left(\frac{\omega - a}{\sqrt{2}\sigma} \right)^2 \right] \cdot \\ &\cdot \sum_{p=0}^{\infty} H_p^2 \left(\frac{\omega - a}{\sigma} \right) \frac{r^{p+1}}{(p+1)!}. \end{aligned} \quad (33a)$$

Analogicznie:

$$\begin{aligned} \pi_{01}(\omega) = \pi_{10}(\omega) &= \frac{1}{2\pi\sigma^2} \frac{\pi}{2} \left[1 + \operatorname{erf} \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right] \left[1 - \operatorname{erf} \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right] + \\ &- \frac{1}{2\pi\sigma^2} \exp \left[- \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right] \cdot \sum_{p=0}^{\infty} H_p^2 \left(\frac{\omega - a}{\sigma} \right) \frac{r^{p+1}}{(p+1)!}, \end{aligned} \quad (33b)$$

$$\begin{aligned} \pi_{11}(\omega) &= \frac{1}{2\pi\sigma^2} \frac{\pi}{2} \left[1 - \operatorname{erf} \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right]^2 + \frac{1}{2\pi\sigma^2} \exp \left[- \left(\frac{\omega - a}{\sqrt{2}\sigma} \right)^2 \right] \cdot \\ &\cdot \sum_{p=0}^{\infty} H_p^2 \left(\frac{\omega - a}{\sigma} \right) \frac{r^{p+1}}{(p+1)!}, \end{aligned} \quad (33c)$$

$$\pi_0(\omega) = \frac{1}{2} \left[1 + \operatorname{erf} \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right], \quad (33d)$$

$$\pi_1(\omega) = \frac{1}{2} \left[1 - \operatorname{erf} \left(\frac{\omega - a}{\sqrt{2}\sigma} \right) \right]. \quad (33e)$$

Podstawiając w powyższych wzorach $a = 0$ otrzymamy odpowiednie wyrażenia również dla rozkładu w_N . Z zależności (33) i (31) otrzymamy żądane wyrażenia dla p_{ij} (odpowiednio q_{ij}), $i, j = 1, 2, 3, 4$, a stąd na podstawie wzorów wziętych z Dodatku II dla $k = 1, n = 1$ i na podstawie wzoru (I. 12) wyliczymy prawdopodobieństwa graniczne p_i i q_i , $i = 1, 2, 3, 4$.

W celu wstępnego oszacowania jakości układu należy określić jeszcze wielkości u_i występujące we wzorach (20a, b). Wartości te w przypadku realizacji algorytmu optymalnego (13) są jednak liczbowo równe stosun-

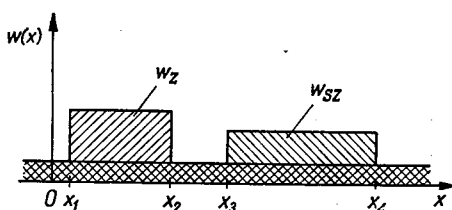
kom prawdopodobieństw π_{00} , π_{01} , π_{10} i π_{11} wyliczonym odpowiednio dla $w = w_{SZ}$ i $w = w_Z$. Inaczej mówiąc, wielkości $u^{(*)}$ [zob. wyrażenie (19)] na wyjściu urządzenia liczącego tworzą następujący zbiór wartości:

$(x-1)^{(*)}$		y_β			
		00	01	10	11
y_α	00	φ_1	φ_2	—	—
	01	—	—	φ_3	φ_4
	10	φ_1	φ_2	—	—
	11	—	—	φ_3	φ_4

gdzie:

$$\left. \begin{aligned} \varphi_1 &= \log \frac{(\pi_{00})_{SZ}}{(\pi_{00})_Z}, \\ \varphi_2 &= \log \frac{(\pi_{01})_{SZ}}{(\pi_{01})_Z}, \\ \varphi_3 &= \log \frac{(\pi_{10})_{SZ}}{(\pi_{10})_Z}, \\ \varphi_4 &= \log \frac{(\pi_{11})_{SZ}}{(\pi_{11})_Z}. \end{aligned} \right\} \quad (34)$$

Tak więc wszystkie wielkości potrzebne do wyliczenia uśrednionych przyrostów funkcji $\varphi(y_k, y_{k-1})$ na jeden odebrany impuls zostały w ten sposób określone. Przyjmując teraz kolejno różne wartości progowe ω i wykonując dla nich opisany cykl obliczeń otrzymamy ciąg wartości Δa



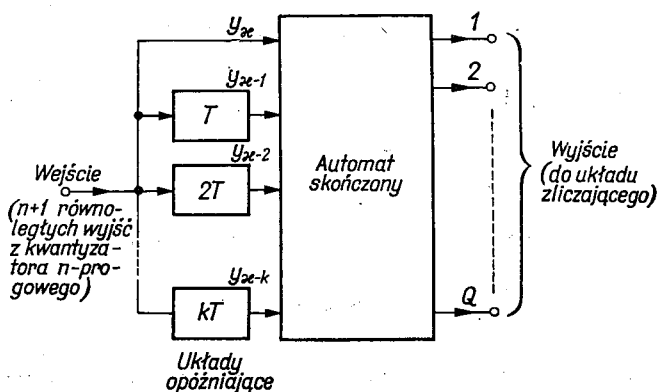
Rys. 4

[zob. wzór (18)], na których podstawie można łatwo określić optymalną (w pierwszym przybliżeniu) wartość ω . Uściślenie tej wartości wymaga, jak wspomnieliśmy, oparcia się na dokładnej postaci funkcji rozkładu prawdopodobieństwa zmiennej wyjściowej, danej przez Twierdzenie I (Dodatek I), co w praktyce jednak rzadko będzie konieczne. Może w tym miejscu powstać także wątpliwość innego rodzaju, mianowicie, czy szukane wartości optymalne ω nie są w pewien prosty sposób związane z zadaną postacią funkcji $w_{SZ}(\bar{x})$ i $w_Z(\bar{x})$. Nie ma oczywiście żadnych podstaw wykluczać możliwość istnienia takich ogólnych zależności. Następu-

jący przykład zdaje się świadczyć jednak o tym, że nie może to być żadna z sugerowanych niekiedy prostych zależności, na przykład kwantyzacja statystycznie, albo geometrycznie równomierna. Niech zadane funkcje gęstości prawdopodobieństwa w_{sz} i w_z (które w ogólnym przypadku mogą mieć postać dowolną) mają kształt przedstawiony na rys. 4. Jest intuicyjnie zrozumiałe, że dla $n = 4$ optymalne wartości progów kwantowania będą się pokrywały z punktami nieciągłości krzywych: x_1, x_2, x_3 i x_4 , gdyż przy takim wyborze punktów progowych istnieje maksymalne prawdopodobieństwo rozróżnienia sygnału i zakłóceń od samych zakłóceń na podstawie obserwacji częstości pojawiania się wartości sygnału w przedziałach (x_1, x_2) i (x_3, x_4) . Punkty te jednak mogą być rozstawione dość dowolnie.

4. UWAGI KOŃCOWE O REALIZACJI UKŁADU LICZĄCEGO

Jak wynika z przeprowadzonych rozważań, układ liczący jest czwórnikami cyfrowym, który każdej wartości wycinkowej skwantowanego sygnału $Y(t)$ na wejściu przyporządkowuje na wyjściu odpowiednią wartość funkcji $\varphi(y)$ lub $\lambda(y)$ [zob. wyrażenia (4) i (5)], wyrażonych w pewnym kodzie cyfrowym dogodnym dla przeprowadzenia operacji sumowania (odpowiednio — mnożenia). Poprzestaniemy na rozważeniu wariantu pierwszego. Ogólną postać układu liczącego przedstawia schematycznie rys. 5. Przez „wyjście” rozumie się tu zapis wartości funkcji φ wyrażonej w dowolnym kodzie q -cyfrowym (na przykład zero-jedynkowym). Ilość

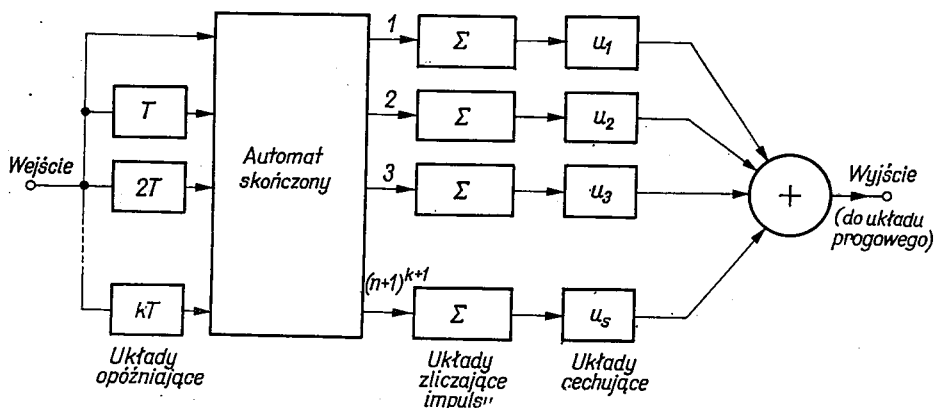


Rys. 5

cyfr Q w pojedynczym wyrażeniu kodowym określa dokładność wyliczenia wartości funkcji φ . Automat skończony stanowiący główną część układu liczącego spełnia rolę podwójną: a) rozróżnia poszczególne wartości wycinkowe sygnału wejściowego, tj. różne kombinacje zer i jedynek na zbiorze $(n+1)^{k+1}$ wejść automatu skończonego połączonych z wyjściami odpowiednich układów opóźniających, oraz b) cechuje poszczególne warto-

ści wycinkowe przyporządkowując im wyrażone w kodzie cyfrowym wartości funkcji φ . Wymienione dwie role mogą być w zasadzie rozdzielone: Q wyjść układu może odpowiadać $Q = (n + 1)^{k+1}$ wartościom wycinkowym. Układ zliczający można wtedy zastąpić przez Q równoległe pracujących liczników impulsów, przyporządkowanych poszczególnym wartościom wycinkowym procesu $Y(t)$. Operację cechowania można wówczas umieszczyć za układami zliczającymi, przed sumatorem (rys. 6).

Metodyka projektowania automatów skończonych została dość szczegółowo omówiona w literaturze (zob. na przykład monografię H. E. Kobrińskiego i B. A. Trachtenbrota [7]), dlatego nie będziemy tu omawiali jej szczegółowo. Warto jednak zwrócić uwagę na fakt, że z ograniczeniem dokładności wyliczenia funkcji φ , tj. z operacją powtórnego kwantowania sygnału, wiąże się ściśle możliwość uproszczenia struktury automatu



Rys. 6

skończonego. Wynika to stąd, że automat ten powinien rozróżniać tylko takie wartości wycinkowe procesu, dla których odpowiednie wartości funkcji φ różnią się więcej niż o kwant wielkości na wyjściu układu liczącego. Tak więc automat skończony powinien rozróżniać nie tyle wartości wycinkowe procesu, co pewne grupy tych wartości, którym odpowiadają zbliżone wartości funkcji φ .

Realizacja techniczna układów opóźniających jest również znacznie uproszczona dzięki temu, że zadaniem każdego takiego układu jest zapamiętanie pewnej informacji o charakterze dyskretnym, dotyczącej pojawienia się lub braku standartowego impulsu na jego wejściu. Rolę taką może spełniać zwykły uniwibrator o szerokości bramki przekraczającej nieco okres powtarzania T .

Na zakończenie pragnę wyrazić podziękowanie recenzentowi artykułu prof. drowi Jerzemu Seidlerowi za szereg cennych uwag krytycznych.

DODATEK I

PRZEGLĄD NIEKTÓRYCH POJĘĆ I TWIERDZEŃ TEORII DYSKRETNÝCH ŁAŃCUCHÓW MARKOWA⁸

Dyskretnym łańcuchem Markowa nazywamy ciąg zmiennych przypadkowych $Y^{(i)}$, $i = 1, 2, 3, \dots, K, \dots$ przybierających skończony lub przeliczalny zbiór wartości $y = y_1, y_2, y_3, \dots, y_n, \dots$ i powiązanych statystycznie w taki sposób, że prawdopodobieństwo warunkowe zmiennej $Y^{(i)}$ przy założeniu, że poprzedzające zmienne $Y^{(1)}, Y^{(2)}, \dots, Y^{(i-1)}$ przybierały wartości $y^{(1)}, y^{(2)}, \dots, y^{(i-1)} \in y$, zależy dla wszystkich i tylko od wartości poprzedzającej $y^{(i-1)}$. Wszystkie prawdopodobieństwa warunkowe $p_{ij}^{(\kappa)} = P\{Y^{(\kappa+1)} = y_j | Y^{(\kappa)} = y_i\}$ można przedstawić w postaci tzw. macierzy przejścia

$$P^{(\kappa)} = \begin{pmatrix} p_{11}^{(\kappa)} & p_{12}^{(\kappa)} & \dots & p_{1n}^{(\kappa)} \\ p_{21}^{(\kappa)} & p_{22}^{(\kappa)} & \dots & p_{2n}^{(\kappa)} \\ \vdots & \vdots & \ddots & \vdots \\ p_{n1}^{(\kappa)} & p_{n2}^{(\kappa)} & \dots & p_{nn}^{(\kappa)} \end{pmatrix} \quad \sum_{j=1}^n p_{ij}^{(\kappa)} = 1 \text{ dla wszystkich } i, \kappa = 1, 2, \dots \quad (I.1)$$

Jeżeli wszystkie elementy $p_{ij}^{(\kappa)}$ są niezależne od κ , to łańcuch Markowa nazywamy stacjonarnym.

Zadanie macierzy przejścia charakteryzuje w pełni związki między wartościami łańcucha obserwowanymi w różnych momentach czasu. W szczególności, prawdopodobieństwa warunkowe $p_{ij}^{(\kappa)}(k) = P\{Y^{(\kappa+k)} = y_j | Y^{(\kappa)} = y_i\}$ określa wzór:

$$P^{(\kappa)}(k) = P^{(\kappa+k-1)} \times P^{(\kappa+k-2)} \times \dots \times P^{(\kappa+1)} \times P^{(\kappa)}, \quad k = 2, 3, 4, \dots \quad (I.2)$$

W dalszym ciągu ograniczymy się do rozważania tylko łańcuchów stacjonarnych. Wzór (I.2) przybiera dla nich uproszczoną postać:

$$P(k) = P^k. \quad (I.3)$$

Jeżeli macierzy P nie można przedstawić w postaci quasi-diagonalnej

$$P = \begin{pmatrix} P_{11} & P_{12} \\ P_{21} & P_{22} \end{pmatrix}, \quad (I.4)$$

gdzie P_{ii} ($i = 1, 2$) są pewnymi kwadratowymi podmacierzami macierzy P , a P_{ij} ($i \neq j$) są takimi podmacierzami, że przynajmniej jedna z nich jest macierzą zerową, to macierz P nazywamy nierozkładalną.

Rozważmy zbiór S wartości przybieranych przez zmienne przypadkowe $Y^{(\kappa)}$ tworzące łańcuch Markowa. Wartość $y_i \in S$ nazywamy niewłaściwą, jeżeli istnieje pewna wartość $y_j \neq y_i$ oraz pewna stała ϱ taka, że $p_{ij}(\varrho) > 0$ i dla wszystkich $\mu = 1, 2, \dots$ $p_{ji}(\mu) = 0$, tzn. jeśli od wartości y_i możliwe jest przejście do innej wartości y_j , od której jednak niemożliwy jest powrót z powrotem do wartości y_i . Wartości y , które nie są niewłaściwymi, nazywamy wartościami właściwymi. Jeżeli dwie wartości y_i i y_j są wartościami właściwymi i istnieje takie $\mu_1 > 0$, że $p_{ij}(\mu_1) > 0$, to musi istnieć także $\mu_2 > 0$ takie, że $p_{ij}(\mu_2) > 0$. Taką parę wartości y_i, y_j nazywamy parą wartości wzajemnie przechodnich. Przechodność jest własnością transytywną, tzn., że jeżeli (y_i, y_j) i (y_j, y_k) tworzą pary wzajemnie przechodnie, to (y_i, y_k) tworzy również parę wzajemnie przechodnią. Wszystkie elementy zbioru S można zatem rozbić na podzbiory S_s elementów wzajemnie przechodnich w ramach każdego z tych podzbiorów i nieprzechodnich między różnymi podzbiórami. Każdy taki podzbiór nazwiemy właściwą klasą wartości. Warunek nierozkładalności macierzy przejścia oznacza, że zbiór wartości S tworzy

tylko jedną właściwą klasę wartości, w której wszystkie elementy są wzajemnie przechodnie.

Macierz przejścia $\mathbf{P}$ będziemy nazywali macierzą prawidłową, jeśli nie posiada ona liczb charakterystycznych (tj. pierwiastków równania charakterystycznego macierzy) o module równym jedności, ale różnych od 1. Jeśli ponadto $\lambda = 1$ jest pierwiastkiem prostym (pojedynczym) równania charakterystycznego

$$\Delta(\lambda) = |P - \lambda E| = 0, \quad (I.5)$$

to macierz $\mathbf{P}$ nazywamy regularną. Jeśli nierozkładalna macierz przejścia (I.5) $\mathbf{P}$ posiada tylko jedną liczbę charakterystyczną o maksymalnym module, to macierz tę nazywamy prymitywną. Macierz dodatnia $\mathbf{P} \geq 0$ jest prymitywna wtedy i tylko wtedy, gdy pewna potęga tej macierzy jest dodatnia:

$$\mathbf{P}^m \geq 0, m \geq 1. \quad (I.6)$$

Proces Markowa odpowiadający prymitywnej macierzy przejścia $\mathbf{P}$ nazywamy procesem acyklicznym. Ponieważ macierz prymitywna jest szczególnym przypadkiem macierzy prawidłowej, więc acykliczny proces Markowa jest szczególnym przypadkiem procesu prawidłowego.

Powyższe pojęcia są ściśle związane z samą budową macierzy $\mathbf{P}$. Przypuśćmy, że macierz przejścia opisująca łańcuch Markowa o jednej właściwej klasie wartości może być przedstawiona w następującej postaci:

$$\mathbf{P} = \begin{pmatrix} 0 & P_{12} & 0 & \dots & 0 \\ 0 & 0 & P_{23} & \dots & 0 \\ \vdots & \vdots & \vdots & & \vdots \\ 0 & 0 & 0 & P_{s-1, s} & \\ P_{s1} & 0 & 0 & 0 & 0 \end{pmatrix} \quad (I.7)$$

gdzie wszystkie zerowe podmacierze diagonalne są kwadratowe, $P_{12}, P_{23}, \dots, P_{s-1, s}, P_{s1}$ są macierzami przejścia odpowiednio z podklasy wartości S_1 do S_2 , z S_2 do S_3 , ..., z S_{s-1} do S_s i z S_s do S_1 , $S_1, S_2, \dots, S_s \subset S$. Macierz $\mathbf{P}$ nazywamy wtedy macierzą cykliczną indeksu s . Macierz prymitywna odpowiadająca pewnemu łańcuchowi acyklicznemu jest w sensie powyższej definicji nierozkładalną macierzą przejścia indeksu $s=1$. Acykliczność łańcucha Markowa można uważać za inaczej wyrażoną własność jego ergodyczności.

Przytoczmy teraz dwa twierdzenia graniczne, bardzo istotne dla rozważań zawartych w artykule.

Określimy następującą macierz charakterystyczną:

$$\Phi(\theta) = \| p_{\alpha\beta} e^{i\theta(x_\beta - a)} \|_1^s, \quad (I.8)$$

gdzie a jest pewną stałą. Równanie charakterystyczne macierzy $\Phi(\theta)$ ma postać

$$C[\lambda(\theta)] = |\lambda E - \Phi(\theta)| = 0. \quad (I.9)$$

Oznaczmy jego pierwiastki przez $\lambda_0(\theta), \lambda_1(\theta), \lambda_2(\theta), \dots, \lambda_g(\theta)$. Wybierzmy stałą a w wyrażeniu (I.8) tak, aby

$$a = \frac{\sum_{\beta=1}^s x_\beta \Phi_{\beta\beta}(1)}{\sum_{\beta=1}^s \Phi_{\beta\beta}(1)}, \quad (I.10)$$

gdzie $\Phi_{\beta\beta}(\lambda)$ jest dopełnieniem algebraicznym elementu c_{ij} macierzy $C[\lambda(\theta)]$ określonej wzorem (I.9). Wówczas spełnione jest następujące

Twierdzenie I

Jeśli ciąg zmiennych przypadkowych $Y^{(0)}, Y^{(1)}, Y^{(2)}, \dots$ tworzy stacjonarny i acykliczny łańcuch Markowa o wartościach należących do zbioru $y = (y_1, y_2, \dots, y_s)$, jeśli $\lambda_0''(\theta) \neq 0$ oraz stała a jest określona zależnością (I.10), to

$$\lim_{n \rightarrow \infty} P \left\{ u_1 < \frac{y_0 + y_1 + \dots + y_n - (n+1)a}{\sqrt{-(n+1)\lambda_0''(0)}} < u_2 \right\} = \frac{1}{\sqrt{2\pi}} \int_{u_1}^{u_2} e^{-u^2/2} du. \quad (\text{I.11})$$

Twierdzenie II

Jeżeli ciąg dyskretnych zmiennych przypadkowych $Y^{(0)}, Y^{(1)}, Y^{(2)}, \dots$ tworzy stacjonarny acykliczny łańcuch Markowa opisany macierzą przejścia $\mathbf{P}$, to dla dowolnych wartości y_i, y_j graniczne prawdopodobieństwa przejścia $\bar{p}_{ij}(r)$ i graniczne prawdopodobieństwa bezwarunkowe $\bar{p}_j(r)$ dążą przy $r \rightarrow \infty$ do tej samej granicy

$$\bar{p}_j = \frac{P_{jj}(1)}{\sum_{j=1}^s P_{aa}(1)} > 0, \quad j = 1, 2, \dots, s \quad (\text{I.12})$$

niezależnie od wartości wyjściowej i , przy czym

$$\sum_{j=1}^s \bar{p}_j = 1, \quad (\text{I.13})$$

a $P_{ij}(\lambda)$ jest dopełnieniem algebraicznym elementu b_{ij} macierzy charakterystycznej

$$\Psi(\lambda) = \|\lambda E - P\|. \quad (\text{I.14})$$

Łańcuch Markowa, dla którego prawdopodobieństwa graniczne nie zależą od wartości wyjściowej y_i nazywamy łańcuchem regularnym.

DODATEK II

ELEMENTY WYRAŻEŃ NA PRAWDOPODOBIENSTWA GRANICZNE
W NIEKTÓRYCH KWANTOWYCH UKŁADACH DECYZYJNYCH

Wprowadzimy następujące oznaczenia:

$$\begin{aligned} p_{ab, cd, \dots, mn} &= p'_{ab} p'_{cd} \dots p'_{mn}, \\ p_{cd}^{ab} &= p'_{ab} p'_{cd} - p'_{ad} p'_{cb}, \\ p'_{ab} &= p_{ab}, \quad a = b \\ p'_{aa} &= (p_{aa} - \lambda). \end{aligned}$$

$$k = 1, \quad n = 1$$

Wartości wycinkowe:

- | | |
|-------|-------|
| 1) 00 | 3) 10 |
| 2) 01 | 4) 11 |

Niezerowe prawdopodobieństwa przejścia:

$$p_{11}, p_{12}, p_{23}, p_{24}, p_{31}, p_{32}, p_{43}, p_{44}.$$

Podwyznaczniki macierzy charakterystycznej:

$$P_{11}(\lambda) = -\lambda^2 p'_{44} - p_{32} p_{24}^{43},$$

$$P_{22}(\lambda) = \lambda p_{11,44},$$

$$P_{33}(\lambda) = \lambda p_{11,44},$$

$$P_{44}(\lambda) = -\lambda^2 p'_{11} - p_{23} p_{51}^{12}.$$

$$k = 2, n = 1$$

Wartości wycinkowe:

- | | | | |
|--------|--------|--------|--------|
| 1) 000 | 3) 010 | 5) 100 | 7) 110 |
| 2) 001 | 4) 011 | 6) 101 | 8) 111 |

Niezerowe prawdopodobieństwa przejścia:

$$p_{11}, p_{12}, p_{23}, p_{24}, p_{35}, p_{36}, p_{47}, p_{48}, p_{51}, p_{52}, p_{63}, p_{64}, p_{75}, p_{76}, p_{87}, p_{88}.$$

Podwyznaczniki macierzy charakterystycznej:

$$P_{11}(\lambda) = -\lambda^6 p'_{88} + \lambda^4 p_{36,63,88} + \lambda^3 (p_{23,35,52,88} + p_{64,76} p_{55}^{47}) - \lambda^2 p_{24,52,75} p_{87}^{48} + p_{52} p_{64}^{23} p_{76}^{35} p_{87}^{48},$$

$$P_{22}(\lambda) = \lambda^5 p_{11,88} - \lambda^3 p_{11,36,63,88} - \lambda^2 p_{11,64,76} p_{55}^{47},$$

$$P_{33}(\lambda) = \lambda^5 p_{11,88} - \lambda^2 p_{11,64,76} p_{55}^{47} + \lambda p_{24,75} p_{51}^{12} p_{88}^{47},$$

$$P_{44}(\lambda) = \lambda^5 p_{11,88} - \lambda^3 p_{11,36,63,88} - \lambda^2 p_{23,35,52,88} p_{52}^{11},$$

$$P_{55}(\lambda) = \lambda^5 p_{11,88} - \lambda^2 p_{11,36,63,88} - \lambda^2 p_{11,64,76} p_{55}^{47},$$

$$P_{66}(\lambda) = \lambda^5 p_{11,88} - \lambda^2 p_{23,35,52,88} p_{52}^{11} + \lambda p_{24,75} p_{51}^{12} p_{88}^{47},$$

$$P_{77}(\lambda) = \lambda^5 p_{11,88} - \lambda^3 p_{11,36,63,88} - \lambda^2 p_{23,35,52,88} p_{52}^{11},$$

$$P_{88}(\lambda) = -\lambda^6 p'_{11} + \lambda^4 p_{11,36,63} + \lambda^3 (p_{11,47,64,76} + p_{23,35} p_{52}^{11}) - \lambda^2 p_{24,47,75} p_{51}^{12} + p_{47} p_{76}^{35} p_{64}^{23} p_{51}^{12}.$$

$$k = 1, n = 2$$

Wartości wycinkowe:

- | | | |
|-------|-------|-------|
| 1) 00 | 4) 10 | 7) 20 |
| 2) 01 | 5) 11 | 8) 21 |
| 3) 02 | 6) 12 | 9) 22 |

Niezerowe prawdopodobieństwa przejścia:

$$p_{11}, p_{12}, p_{13}, p_{24}, p_{25}, p_{26}, p_{37}, p_{38}, p_{39}, p_{41}, p_{42}, p_{43}, p_{54}, p_{55}, p_{56}, p_{67}, p_{68}, p_{69}, p_{71}, p_{72}, p_{73}, p_{84}, p_{85}, p_{86}, p_{97}, p_{98}, p_{99}.$$

Podwyznaczniki macierzy charakterystycznej:

$$P_{11}(\lambda) = -\lambda^6 p_{55,99} + \lambda^4 (p_{42,99} p_{55}^{24} + p_{55,73} p_{99}^{37} + p_{55}^{55} p_{55}^{68}) - \lambda^3 p_{67,72,99} p_{55}^{25} + \lambda^2 (p_{57}^{39} p_{55}^{24} p_{73}^{42} - p_{37,73} p_{55}^{55} p_{99}^{68} - p_{73,97} p_{55}^{55} p_{99}^{68} - p_{42,86} p_{99}^{68} p_{55}^{24} + p_{42,56} p_{99}^{68} p_{55}^{24} - p_{26,42} p_{99}^{68} p_{55}^{54}) + p_{37,86} \cdot p_{99}^{68} p_{73}^{42} p_{55}^{24} + p_{86,97} p_{55}^{38} p_{73}^{42} p_{55}^{24} - p_{37,56} p_{99}^{68} p_{73}^{42} p_{55}^{24} - p_{56,97} p_{99}^{38} p_{73}^{42} p_{55}^{24} + p_{26,37} p_{99}^{68} p_{73}^{42} p_{55}^{54} + p_{26,97} p_{99}^{38} p_{73}^{42} p_{55}^{68},$$

$$P_{22}(\lambda) = \lambda^5 p_{11,55,99} - \lambda^3 (p'_{11} p_{55}^{55} p_{99}^{68} + p'_{55} p_{73}^{11} p_{99}^{37}) - \lambda^2 p_{41}^{13} p_{99}^{38} p_{85}^{54} - \lambda (p_{37} p_{71}^{13} p_{55}^{55} p_{99}^{68} - p_{67} p_{71}^{13} \cdot p_{55}^{55} p_{99}^{38} + p_{97} p_{71}^{13} p_{55}^{55} p_{99}^{38}),$$

$$P_{33}(\lambda) = \lambda^5 p_{11,55,99} + \lambda^3 (p'_{11} p_{55}^{55} p_{99}^{68} + p'_{99} p_{41}^{12} p_{55}^{24}) - \lambda^2 p_{72}^{11} p_{55}^{26} p_{99}^{67} + \lambda (p_{24} p_{42}^{11} p_{55}^{55} p_{99}^{68} - p_{54} p_{42}^{11} \cdot p_{55}^{26} p_{99}^{68} + p_{84} p_{42}^{11} p_{55}^{26} p_{99}^{68}),$$

$$P_{44}(\lambda) = \lambda^5 p_{11,55,99} + \lambda^3 (p'_{11} p_{55}^{55} p_{99}^{68} + p'_{55} p_{71}^{13} p_{99}^{37}) - \lambda^2 p_{72}^{11} p_{55}^{26} p_{99}^{67} - \lambda (p_{37} p_{73}^{11} p_{55}^{55} p_{99}^{68} - p_{67} p_{73}^{11} \cdot p_{55}^{55} p_{99}^{38} + p_{97} p_{73}^{11} p_{55}^{55} p_{99}^{38}),$$

$$\begin{aligned}
P_{55}(\lambda) = & -\lambda^6 p_{11,99} - \lambda^4 (p_{71}^{13} p_{99}^{37} + p_{24,99} p_{41}^{12} + p_{11,86} p_{99}^{68}) + \lambda^3 (p_{84} p_{43}^{11} p_{99}^{38} + p_{26} p_{72}^{11} p_{99}^{67} - \\
& - \lambda^2 (p_{24,71} p_{43}^{12} p_{99}^{37} + p_{24,41} p_{72}^{13} p_{99}^{37} + p_{11,34} p_{73}^{42} p_{99}^{37} + p_{37,86} p_{73}^{11} p_{99}^{68} + p_{67,86} p_{71}^{13} p_{99}^{38} + \\
& + p_{86,97} p_{73}^{11} p_{99}^{38} + p_{42}^{11} p_{86}^{24} p_{99}^{68}) + p_{11,37,42,73} p_{86}^{24} p_{99}^{68} + p_{11,42,67,73} p_{84}^{26} p_{99}^{38} + p_{73,97} \cdot \\
& \cdot p_{42}^{11} p_{86}^{24} p_{99}^{38} + p_{24,37,72,86} p_{41}^{13} p_{99}^{68} + p_{24,67,72,86} p_{43}^{11} p_{99}^{38} + p_{24,72,86,97} p_{41}^{13} p_{99}^{38} - p_{12,37} \cdot \\
& \cdot p_{73}^{41} p_{86}^{24} p_{99}^{68} + p_{12,67} p_{73}^{41} p_{86}^{24} p_{99}^{38} + p_{13,26,37,84} p_{71}^{42} p_{99}^{68} + p_{13,26,67,84} p_{72}^{41} p_{99}^{38} + p_{13,26,84,97} \cdot \\
& \cdot p_{71}^{42} p_{99}^{68} + p_{12,43,71,97} p_{86}^{24} p_{99}^{38} + p_{13,24,37,42,71,86} p_{99}^{68} - p_{13,24,42,67,71,86} p_{99}^{38} + p_{13,24,42,71,86,97} \cdot \\
& \cdot p_{99}^{38} - p_{11,26,37,43,72,84} p_{99}^{68} + p_{11,26,43,67,72,84} p_{99}^{38} - p_{11,26,43,72,84,97} p_{99}^{68},
\end{aligned}$$

$$\begin{aligned}
P_{66}(\lambda) = & \lambda^5 p_{11,55,99} + \lambda^3 (p'_{55} p_{71}^{13} p_{99}^{37} + p'_{99} p_{41}^{12} p_{55}^{24}) - \lambda^2 p_{43}^{11} p_{84}^{55} p_{99}^{38} - \lambda (p_{13} p_{72}^{41} p_{54}^{25} p_{99}^{37} - p_{43} p_{72}^{11} \cdot \\
& \cdot p_{54}^{25} p_{99}^{37} + p_{73} p_{42}^{11} p_{54}^{25} p_{99}^{37}),
\end{aligned}$$

$$\begin{aligned}
P_{77}(\lambda) = & \lambda^5 p_{11,55,99} + \lambda^3 (p'_{11} p_{85}^{56} p_{99}^{68} + p'_{99} p_{42}^{11} p_{55}^{25}) - \lambda^2 p_{43}^{11} p_{84}^{55} p_{99}^{38} + \lambda (p_{24} p_{42}^{11} p_{86}^{55} p_{99}^{68} - p_{84} p_{42}^{11} \cdot \\
& \cdot p_{55}^{25} p_{99}^{68} + p_{54} p_{42}^{11} p_{86}^{55} p_{99}^{68}),
\end{aligned}$$

$$\begin{aligned}
P_{88}(\lambda) = & \lambda^5 p_{11,55,99} - \lambda^3 (p'_{55} p_{73}^{11} p_{99}^{37} + p'_{99} p_{42}^{11} p_{55}^{24}) - \lambda^2 p_{72}^{11} p_{55}^{26} p_{99}^{67} - \lambda (p_{41} p_{73}^{12} p_{54}^{25} p_{99}^{37} - p_{42} p_{73}^{11} \cdot \\
& \cdot p_{54}^{25} p_{99}^{37} + p_{43} p_{73}^{11} p_{54}^{25} p_{99}^{37}),
\end{aligned}$$

$$\begin{aligned}
P_{99}(\lambda) = & -\lambda^6 p_{11,55} + \lambda^4 (p_{11,68} p_{86}^{55} + p_{37,55} p_{73}^{11} + p_{42}^{11} p_{55}^{24}) - \lambda^3 p_{11,38,43} p_{86}^{54} + \lambda^2 (p_{71}^{13} p_{86}^{55} p_{68}^{37} - \\
& - p_{37,73} p_{86}^{54} p_{42}^{11} - p_{13,37} p_{86}^{54} p_{42}^{11} - p_{24,68} p_{42}^{11} p_{86}^{55} + p_{54,68} p_{42}^{11} p_{86}^{55} - p_{68,84} p_{42}^{11} p_{55}^{25}) + p_{24,73} \cdot \\
& \cdot p_{42}^{11} p_{86}^{37} p_{86}^{55} + p_{13,24} p_{72}^{41} p_{86}^{37} p_{86}^{55} - p_{54,73} p_{42}^{11} p_{86}^{37} p_{86}^{55} - p_{13,54} p_{72}^{41} p_{86}^{37} p_{86}^{55} + p_{73,84} p_{42}^{11} p_{86}^{37} p_{86}^{55} + \\
& + p_{13,84} p_{72}^{41} p_{86}^{37} p_{86}^{55}.
\end{aligned}$$

Uwaga: Wyżej przytoczone grupy podwyznaczników są symetryczne względem następującego przedstawienia indeksów:

$$\begin{array}{cccccc}
1 & 2 & \dots & s-1 & s, \\
s & s-1 & \dots & 2 & 1,
\end{array}$$

gdzie $s = (n+1)^{k+1}$.

WYKAZ LITERATURY

1. Prijom signalow pri naliczii szuma. Sbornik pieriewodow. — Moskwa 1960 (zob. w szczególności artykuły J. V. Harringtona oraz B. Blasbalga).
2. Kątecki T.: Problemy wykrywania przy wieloprogowym kwantowaniu w amplitudzie. — Biuletyn WAT, XI, nr 2 (114), 1962.
3. Wirth Wulf-Dieter: Beitrag zur Vereinfachung der Radarauswertung durch Digitalisierung. — Techn. Universität, Berlin, Dez. 1962.
4. Storz Werner: Beitrag zur digitalen Verarbeitung von Radarsignalen. — Techn. Universität, Berlin, Dez. 1962.
5. Stratonowicz R. Ł.: Izbrannyje woprosy teorii fluktuacij w radiotiehnike, Moskwa 1961.
6. Kulikowski J.: O pewnych zagadnieniach związanych z wyliczeniem elementów odwróconej macierzy kowariancji pasywnych zakłóceń radiolokacyjnych. — Rozpr. Elektrot., t. VII, zeszyt 3, 1961.
7. Kobrinskiy N. E., Trachtenbrot B. A.: Wwiedieniye w teoriiu koniecznych awtomatow. Moskwa 1962.
8. Sarymsakow T. A.: Osnovy teorii processow Markowa. Moskwa 1954.

DODATEK III

Współczynniki wielomianów Hermite'a

$$H_k(x) = \sum_l a_l^{(k)} \cdot x^l$$

λ	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
0	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
1	1	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
2	-1	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-
3	3	-3	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-
4	3	-	-6	-	1	-	-	-	-	-	-	-	-	-	-	-	-
5	-	15	-	-10	-	1	-	-	-	-	-	-	-	-	-	-	-
6	-15	-	45	-	-15	-	1	-	-	-	-	-	-	-	-	-	-
7	-	-105	-	105	-	-21	-	1	-	-	-	-	-	-	-	-	-
8	105	-	-420	-	210	-	-28	-	1	-	-	-	-	-	-	-	-
9	-	945	-	-1260	-	378	-	-36	-	1	-	-	-	-	-	-	-
0	-945	-	4725	-	-3150	-	630	-	-45	-	1	-	-	-	-	-	-
1	-	-10395	-	17325	-	-6930	-	990	-	-55	-	1	-	-	-	-	-
2	10395	-	-62370	-	51975	-	-13860	-	1485	-	-66	-	1	-	-	-	-
3	-	135135	-	-270270	-	135135	-	-25740	-	2145	-	-78	-	1	-	-	-
4	-135135	-	945945	-	-945945	-	315315	-	-45045	-	3003	-	-99	-	1	-	-
5	-	-2027025	-	4729725	-	-2837835	-	675675	-	-75075	-	4095	-	-113	-	1	-
6	14189165	-	-16216200	-	18918900	-	-7567560	-	1351350	-	-120120	5580	-	-113	-	-18	1

J. L. KULIKOWSKI

STRUCTURE ET QUALITÉS ASYMPTOTIQUES D'UNE CLASSE
DES SCHEMAS DE DECISION A QUANTIFICATION DU SIGNAL

Résumé

On a discuté le problème d'évaluation des qualités asymptotiques d'un certain schéma de décision à quantification du signal, désigné à la détection des signaux de radar sur le fond des brouillages corrélés. L'évaluation du schéma était basée sur de certaines qualités de limite des chaînes composées de Markoff. L'analyse avait suivie par des considérations supplémentaires concernant optimale structure logique du schéma de calculateur digital réalisant l'algorithme optimal du traitement du signal.

J. L. KULIKOWSKI

ÜBER DIE STRUKTUR UND ASYMPTOTISCHEN BEWERTUNGEN VON
EIGENSCHAFTEN EINER GEWISSEN KLASSE DEZEDIERENDER
QUANTENSYSTEME

Zusammenfassung

In diesem Artikel wurde die Frage der asymptotischen Bewertungen von Eigenschaften des dezедierenden Systems der Signalquantierung, welche zur Erforschung von Radiolokationssignalen in Bezug auf korelierte Störungen bestimmt ist, behandelt. Die Bewertung dieses Systems wurde auf gewisse Grenzeigenschaften gestützt. Diese Betrachtungen wurden durch Bemerkungen, welche die optimal logische Struktur des Ziffernsystems, welches den optimalen Algorhythmus der Signalbearbeitung berechnet, ergänzt.

Ю. Л. КУЛИКОВСКИ

О СТРУКТУРЕ И АСИМПТОТИЧЕСКИХ СВОЙСТВАХ ОДНОГО КЛАССА
РЕШАЮЩИХ СХЕМ С КВАНТОВАНИЕМ СИГНАЛА

Резюме

Рассматривается вопрос оценки асимптотических свойств решающего устройства с квантованием сигнала, предназначенного для обнаружения радиолокационных сигналов на фоне коррелированных помех. Метод оценки основан на некоторых предельных свойствах сложных цепей Маркова. Приводятся некоторые практические замечания, касающиеся оптимальной логической структуры цифрового вычислительного устройства, реализующего оптимальный алгоритм обработки сигнала.

621.396.619

ROMUALD NOWAK

Zagadnienie niestabilności częstotliwości w świetle zjawisk fluktuacyjnych

Rękopis dostarczono 18.7.1963

Praca omawia z teoretycznego i eksperymentalnego punktu widzenia współczesne poglądy na zagadnienie fluktuacyjnej niestabilności częstotliwości. Źródłem fluktuacyjnej niestabilności częstotliwości są szumy występujące w układach generacyjnych. Zjawiska te stanowią strukturalną granicę stabilizacji źródeł wzorcowych.

W oparciu o prace różnych autorów wyodrębniono i przedstawiono teoretyczny model drgań zwany modelem dyfuzyjnym. Model ten opisuje oddziaływanie szumu białego na Van der Polowski układ generacyjny przy określonych, dopuszczalnych ze względu na zastosowania praktyczne, założeniach upraszczających. Przytoczone zależności określają najważniejsze parametry drgań w postaci możliwie prostej i w zależności od wielkości fizycznych, łatwych do wyznaczenia w praktyce. Metodę opisu i zespół tych wielkości wybrano w taki sposób, aby wyniki można było zastosować bezpośrednio do różnego typu układów generacyjnych.

W części omawiającej stronę eksperymentalno-pomiarową zagadnienia podano w sposób usystematyzowany przegląd różnych metod i układów pomiarowych służących do pomiarów fluktuacji kąta fazowego drgań. Spośród parametrów opisujących te fluktuacje jako parametr podstawowy wybrano fluktuacyjną dewiację częstotliwości, jako wielkość umożliwiającą ekstrapolację pojęcia niestabilności częstotliwości na obszar zjawisk fluktuacyjnych.

W części końcowej pracy przeprowadzono porównanie opublikowanych wyników analitycznych i eksperymentalnych. Z porównania tego wyciągnięto wnioski o braku eksperymentalnej weryfikacji teoretycznego modelu drgań dyfuzyjnych, wskazując na przyczyny takiego stanu rzeczy. Sformułowano również założenia eksperymentu umożliwiającego doświadczalne sprawdzenie wniosków wynikających z dyfuzyjnego modelu drgań. Eksperyment taki polegałby na przeprowadzeniu odpowiednich pomiarów dla układu generacyjnego o sztucznie zwiększonym poziomie szumu białego.

1. WSTĘP

Rozwój techniki stabilizacji częstotliwości, wywołany potrzebami różnych dziedzin łączności, postawił w zupełnie nowym świetle zagadnienie niestabilności częstotliwości. Zbliżono się bowiem w tej dziedzinie do

pewnej bariery wytworzonej przez szумы i fluktuacje samego źródła drgań częstotliwości wzorcowych. Zjawiska te stanowią wewnętrzne, strukturalne ograniczenie stałości częstotliwości układów generacyjnych. Zagadnienia określenia tej granicy, a w perspektywie również jej obniżenia stanęły ostatnio na porządku dziennym. Opracowano nowe metody analityczne i eksperymentalne do badania tego typu efektów. W rezultacie okazało się, że podobnie jak optyczne źródła światła, radiotechniczne generatory drgań elektrycznych nie są źródłami monochromatycznymi. Przyczyny tego stanu rzeczy można ująć pokrótce w sposób następujący.

Drgania monochromatyczne są to drgania wyrażane następującą funkcją czasu:

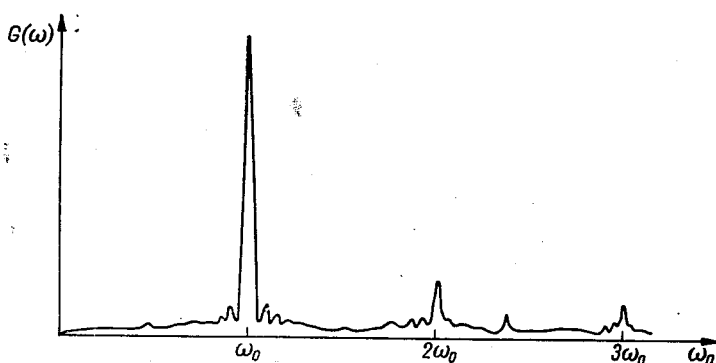
$$u(t) = A_0 \sin \Phi(t) = A_0 \sin(\omega_0 t + \varphi_0), \quad (1.1)$$

gdzie $A_0 = \text{const}$, $\omega_0 = 2\pi f_0 = \text{const}$, $\varphi_0 = \text{const}$ oznaczają odpowiednio amplitudę, pulsację i początkowy kąt fazowy drgania, a $\Phi(t)$ argument lub fazę drgań. Widmo przebiegu opisywanego powyższą zależnością jest pojedynczym prążkiem o długości proporcjonalnej do amplitudy, energii lub mocy drgania znajdującego się w punkcie $\omega = \omega_0$. Każdy rzeczywisty, tzn. generowany przez realny generator częstotliwości, sygnał wąskopasmowy, którego najprostszym przybliżeniem, a jednocześnie idealnym wzorem jest przebieg monochromatyczny, okazuje się procesem bardziej złożonym. W sygnale tym, w przypadku ogólnym można wyodrębnić zarówno składowe losowe czy aperiodyczne, jak prawie periodyczne i harmoniczne. Wynika to z oddziaływania na proces generacji różnych losowych lub zdeterminowanych czynników wewnętrznych, to jest skojarzonych ze źródłem drgań, i zewnętrznych, tzn. reprezentujących wpływy otoczenia. Niektóre z tych czynników, jak np. wpływy atmosferyczne, były znane i uwzględniane już od dawna. Inne, jak np. szумы układu generacyjnego, nabrały praktycznego znaczenia dopiero dzięki współczesnemu rozwojowi techniki generacji i stabilizacji częstotliwości. Czynniki te, zależnie od swego charakteru, oddziałują na drgania bezpośrednio lub na drodze modulacji parametrycznej. W przypadku gdy widmo procesu wywołanego przez dany czynnik zakłócający pokrywa pasmo przenoszenia obwodu drgań generatora, występuje liniowa superpozycja sygnałów monochromatycznego i zakłócającego. W przypadku przeciwnym oddziaływanie wzajemne może mieć miejsce jedynie poprzez efekty nieliniowe procesu generacji. W rezultacie dla procesów wąskopasmowych można opisać wpływ czynników zakłócających w postaci pasożytnej modulacji kąta i amplitudy drgań monochromatycznych. Występowanie losowych sygnałów zakłócających powoduje, że parametry drgań stają się wielkościami losowymi, a widmo staje się widmem ciągłym.

Na podstawie powyższych rozważań można opisać drgania rzeczywiste z uwzględnieniem wszystkich wymienionych wyżej przyczyn w postaci ogólnej jako:

$$u(t) = \sum_{k=1}^{\infty} A_k(t) \cos [\omega_k(t)t + \varphi_k(t)], \quad (1.2)$$

gdzie poszczególne parametry $A_k(t)$, $\omega_k(t)$ i $\varphi_k(t)$ są w przypadku ogólnym losowymi funkcjami czasu. Widmo takiego procesu wyrażone za pomocą funkcji gęstości mocy jednostkowej $G(\omega)$ będzie miało kształt o charakterze przedstawionym na rys. 1. Na rysunku tym wyróżniono zniekształcenia nieliniowe, pasożytniczą modulację sygnałami koherentnymi i lo-



Rys. 1. Wypadkowe widmo drgań rzeczywistych

sowymi. Jeżeli uwzględnić ponadto wpływ stanów nieustalonych, to widmo z rys. 1 należy traktować jako widmo bieżące w określonej chwili czasowej lub jako granicę, do której dąży widmo bieżące w miarę zanikania stanów nieustalonych o charakterze aperiodycznym.

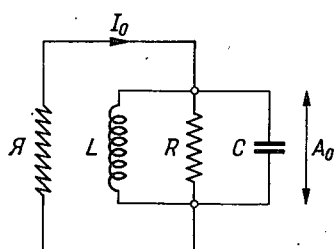
Generator częstotliwości jest z natury rzeczy układem wąskopasmowym, a więc składowe widma z rys. 1 leżące poza pasmem przenoszenia obwodu drgań są bardzo silnie tłumione. Tłumienie to można poza tym zwiększyć dodatkowo umieszczając za generatorem selektywne stopnie separujące. W związku z tym o strukturalnym ograniczeniu stałości częstotliwości źródła częstotliwości decydują składowe widma zawarte w najbliższym otoczeniu częstotliwości drgań podstawowych f_0 . Ze składowych tych można przy odpowiednim zaprojektowaniu układu generacyjnego praktycznie wyeliminować wpływ czynników zewnętrznych (jak np. modulacje od sieci zasilającej czy wpływ warunków klimatycznych). Tak więc w granicznym przypadku stałość częstotliwości źródła jest określana oddziaływaniem wewnętrznych szumów jego układu.

Niniejsza praca ma za zadanie przedstawić w sposób możliwie uporządkowany współczesny stan wiedzy w dziedzinie oddziaływania szumu białego na układy generatorów częstotliwości. W celu uproszczenia opisu

analitycznego zrezygnowano z wyprowadzenia poszczególnych zależności (wyprowadzenia te znajdują zainteresowani w pracach źródłowych), starając się jednocześnie przytoczyć wszystkie ważniejsze zależności opisujące drgania tego typu, w postaci jak najbardziej nadającej się do bezpośredniej interpretacji fizycznej i technicznej.

2. MODEL UKŁADU GENERACYJNEGO

Przy omawianiu wpływu poszczególnych zakłóceń wygodnie jest posługiwać się (w większości przypadków) Van der Polowskim modelem układu generacyjnego. Model taki przedstawiony na rys. 2 składa się



Rys. 2. Van der Polowski model generatora

z równoległego obwodu LC o wypadkowej oporności strat R odłumianego przez nieliniową oporność ujemną J o charakterystyce trzeciego stopnia:

$$i = -g_1 u + g_3 u^3. \quad (2.1)$$

W przyjętym modelu obowiązują ponadto następujące zależności podstawowe:

— pulsacja drgań monochromatycznych:

$$\omega_0 = 2\pi f_0 = \frac{1}{\sqrt{LC}}, \quad (2.2)$$

— dobroć obwodu drgań:

$$Q = \frac{R}{\omega_0 L} = R \omega_0 C, \quad (2.3)$$

— stała czasu obwodu drgań:

$$\tau = 2RC = \frac{2Q}{\omega_0} = \frac{2}{\delta_0 \omega}, \quad (2.4)$$

gdzie $\delta_0 \omega$ oznacza szerokość pasma obwodu drgań dla punktów połowy mocy.

Impedancję i rezystancję obwodu drgań można w przybliżeniu wyrazić w sposób następujący:

$$|\hat{Z}(\omega)|^2 \approx \frac{R^2}{1 + \tau_0^2(\omega - \omega_0)^2}, \quad (2.5)$$

$$R(\omega) = \operatorname{Re} \hat{Z}(\omega) \approx \frac{R}{1 + \tau_0^2(\omega - \omega_0)^2}. \quad (2.6)$$

Zależności (2.5) i (2.6) są słuszne przy niewielkich odstrojeniach od ω_0 , przy których spełniony jest warunek:

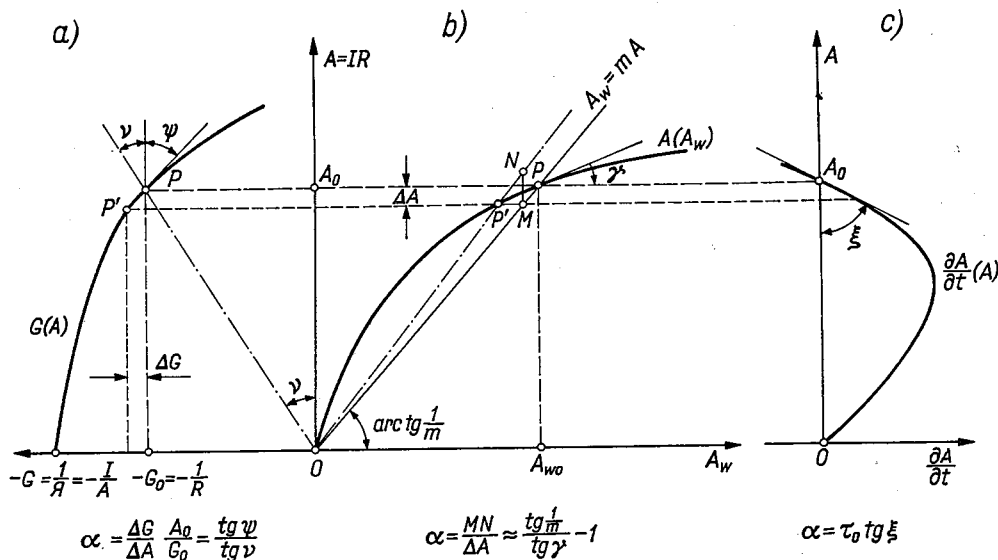
$$\frac{\omega_0 + \omega}{\omega} \approx 2. \quad (2.7)$$

Fluktuacje w oporności ujemnej $\mathcal{H}$ reprezentowane są przez napięcie szumów śrutowych, a fluktuacje w oporności dynamicznej R przez odpowiednie napięcie szumów cieplnych (zob. rozdz. 3). Zakłada się, że pozostałe elementy nie wprowadzają do układu zakłóceń szumowych.

Podstawową zaletą modelu generatora z rys. 2 jest możliwość przeprowadzenia analizy ogólnej dla różnego typu układów generacyjnych, przy czym zależności uzyskane w wyniku tej analizy wyrażają się za pomocą parametrów łatwych do wyznaczenia w praktyce. Należy jednak zdawać sobie sprawę z uproszczeń związanych z tak określonym modelem. Wśród przyjętych uproszczeń na podkreślenie zasługują dwa zagadnienia: zatarcie sposobu wprowadzenia zakłóceń do układu generacyjnego oraz wyeliminowanie wpływu obwodu automatycznego ograniczania amplitudy drgań na oddziaływanie tych zakłóceń. Przyjmując do obliczeń zdefiniowaną w rozdz. 3 wartość σ_0^2 , określającą moc jednostkową szumów wydzieloną na obwodzie drgań, można uniezależnić się od sposobu, w jakim zakłócenia wprowadzane są do układu generacyjnego. W celu uwzględnienia wpływu obwodu ograniczania amplitudy drgań należy określić reakcję układu generacyjnego z rys. 2 na niewielkie zmiany amplitudy.

Układ ograniczania amplitudy drgań generatora pracuje analogicznie jak konwencjonalny ogranicznik amplitudy. (W typowych generatorach będzie to ogranicznik siatkowy). Ogranicznik taki charakteryzuje się dwoma podstawowymi parametrami: stopniem ograniczenia i stałą czasu układu. Parametry te są od siebie w określony sposób uzależnione. Skuteczność działania ogranicznika jest największa dla bardzo wolnych, w stosunku do jego bezwładności, zmian amplitudy, tzn. dla przypadku, gdy okres obwiedni jest znacznie większy od stałej czasu ogranicznika. W miarę wzrostu szybkości zmian amplitudy (wzrost częstotliwości obwiedni) skuteczność działania ogranicznika maleje, lecz na tłumienie pasożytniczej modulacji amplitudy ma coraz większy wpływ selektywność obwodu drgań.

Stopień ograniczenia może być określony w różny sposób. Charakterystykę ogranicznika (statyczną) dla wolnych zmian amplitudy reprezentuje charakterystyka dynamiczna układu pobudzania obciążonego obwodem drgań. Na rys. 3b przedstawiono typowy kształt tej charakterystyki, wyrażonej dla generatora ze sprzężeniem zwrotnym w postaci $A = f_1(A_w)$, gdzie A i A_w oznaczają odpowiednio amplitudy napięć wyjś-



Rys. 3. Graficzna interpretacja współczynnika trwałości cyklu granicznego

ciowego i wejściowego układu pobudzania do drgań. Na rysunku tym nakreślono również prostą sprzężności zwrotnej $A_w = m A$, gdzie m oznacza moduł współczynnika sprzężenia zwrotnego. Punkt przecięcia $P(A_0, A_{w0})$ wyznacza odpowiednie amplitudy generatora monochromatycznego w stanie ustalonym. Dla modelu generatora z rys. 2 odpowiednikiem charakterystyki dynamicznej z rys. 3b jest zależność ujemnej przewodności od amplitudy. Zależność tę wyrażoną w postaci $-G = \frac{1}{R} = -\frac{I}{A} = f_2(A)$ przedstawiono na rys. 3a. W tym przypadku warunek amplitudy wyznaczony jest przez przecięcie charakterystyki przewodności z prostą $-G = -\frac{1}{R} = -\frac{I_0}{A_0} = -G_0$. Dla uzupełnienia obrazu zjawisk na rys. 3c została wykreślona funkcja $\frac{\partial A}{\partial t} = f_3(A)$ opisująca zależność szybkości zmian amplitudy w czasie od wielkości tej amplitudy [14].

W ramach przyjętych poprzednio założeń (tzn. dla małych i powolnych zmian amplitudy) skuteczność działania układu ograniczania amplitudy można wyrazić za pomocą pojedynczej liczby charakteryzującej

w sposób jednoznaczny warunki pracy danego generatora. Liczba ta, oznaczona symbolem α , będzie nie tylko wskaźnikiem stopnia ograniczenia amplitudy drgań, lecz i miarą trwałości cyklu granicznego badanego układu. W związku z tym można traktować liczbę α jako stały współczynnik, zwany współczynnikiem trwałości cyklu granicznego.

Wartość współczynnika trwałości cyklu granicznego można wyznaczyć m. in. z następującej zależności [5]:

$$\alpha = -\tau_0 \left[\frac{\partial}{\partial A} \left(\frac{\partial A}{\partial t} \right) \right]_{A=A_0} = \tau_0 \operatorname{tg} \xi, \quad (2.8)$$

gdzie kąt ξ (rys. 3c) oznacza nachylenie krzywej $\frac{\partial A}{\partial t} = f_3(A)$ w punkcie $(0; A_0)$.

Aby wyznaczyć współczynnik α z warunków amplitudy określonych przez charakterystyki dynamiczne z rysunków 3a i 3b i parametrów obwodów drgań, należy odpowiednio przekształcić znane zależności [14]. W rezultacie tych przekształceń dla generatora ze sprzężeniem zwrotnym otrzymuje się wyrażenie:

$$\alpha = \frac{MN}{\Delta A} \approx \frac{\operatorname{tg} \frac{1}{m}}{\operatorname{tg} \gamma} - 1. \quad (2.9)$$

Rys. 3b wyjaśnia w prosty sposób konstrukcję geometryczną związaną z powyższym wyrażeniem. Dla danego przyrostu amplitudy ΔA ; rzędna $A_0 - \Delta A$ wyznacza odpowiednio punkty P' i M na charakterystyce dynamicznej i prostej sprzężności zwrotnej. Przecięcie prostej OP' z prostą równoległą do osi A , przechodzącą przez punkt M , wyznacza punkt N . Stosunek długości odcinków MN i ΔA określa współczynnik α zgodnie z zależnością (2.9). Przybliżenie zastosowane we wzorze (2.9) jest słuszne przy spełnieniu warunku: $\Delta A \operatorname{tg} \gamma \ll A_0 \operatorname{tg} \frac{1}{m}$, tzn. dla dostatecznie ma-

łych wartości $\Delta A \left[\operatorname{tg} \gamma = \left(\frac{dA}{dA_w} \right)_{A_w=A_{w0}} \right]$.

W analogicznych warunkach dla generatora z opornością ujemną współczynnik α jest określony zależnością:

$$\alpha = - \left(\frac{\partial G}{\partial A} \right)_{A=A_0} \cdot \frac{A_0}{G_0} = - \left(\frac{\partial G}{\partial A} \right)_{A=A_0} \cdot A_0 R = \frac{\operatorname{tg} \psi}{\operatorname{tg} \gamma} \approx - \frac{\frac{\Delta G}{G_0}}{\frac{\Delta A}{A_0}}. \quad (2.10)$$

Odpowiednie wielkości związane z tym wyrażeniem zostały oznaczone na rys. 3a.

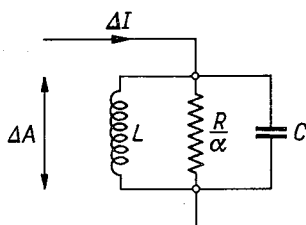
Wartość współczynnika α można wyznaczyć eksperymentalnie przez pomiar zmian amplitudy generatora $(\Delta A)_B$, wywołanej równoległym dołączeniem do obwodu drgań, o oporności dynamicznej R , opornika R_B . [5]. W tym przypadku:

$$\alpha \approx \frac{R}{R_B} \left[\frac{A}{(\Delta A)_B} - 1 \right], \quad (2.11)$$

Współczynnik trwałości cyklu granicznego można również określić na podstawie teorii synchronizacji. Jeżeli do obwodu drgań monochromatycznych opisywanych przez wyrażenie: $I_0 \sin(\omega_0 t + \varphi_0)$ doprowadzić ze źródła prądowego sygnał o postaci $\Delta I \sin(\omega_0 t + \varphi_0)$, tzn. sygnał synchroniczny i synfazowy z drganiami monochromatycznymi, to przy spełnieniu warunku: $\Delta I \ll I_0$, zmiana amplitudy napięcia na obwodzie drgań ΔA będzie proporcjonalna do amplitudy sygnału synchronizującego ΔI zgodnie z zależnością:

$$\frac{\Delta A}{\Delta I} = \frac{R}{\alpha}. \quad (2.12)$$

Oznacza to, że dla zewnętrznego sygnału zakłócającego obwód drgań przedstawia α razy mniejszą oporność dynamiczną i α razy większą szerokość pasma. W konsekwencji, rozpatrując zjawiska na płaszczyźnie fazowej drgań, łatwo zauważyć, że amplituda drgań podlegających zakłóceniom będzie dążyć do wartości A_0 (odpowiadającej promieniowi cyklu granicznego) ze stałą czasu proporcjonalną do wartości $\frac{\tau_0}{\alpha}$. Schemat zastępczy takiego obwodu przedstawiono na rys. 4.



Rys. 4. Schemat zastępczy obwodu drgań dla małych i powolnych zmian amplitudy

Jak wynika z powyższych rozważań, w celu zapewnienia w układzie generacyjnym optymalnych warunków tłumienia pasożytniczej modulacji amplitudy drgań generowanych, należy stosować duży stopień ograniczania, przy stałej czasu obwodu ograniczania τ_s , spełniającej warunek:

$$\frac{1}{\omega_0} \ll \tau_s \ll \frac{\tau_0}{\alpha}. \quad (2.13)$$

Ponieważ w typowych układach generacyjnych wartość α jest rzędu $3 \div 10$ [10], tłumienie pasożytniczej modulacji amplitudy jest stosunkowo duże. Pozwala to na znaczne uproszczenie dalszych rozważań. M. in.

można, z dostatecznym dla celów praktycznych przybliżeniem, zaniedbać wpływ tych niewielkich zmian amplitudy na częstotliwość drgań generowanych.

3. ŹRÓDŁA ZAKŁÓCEŃ LOSOWYCH

Znane są zjawiska fizyczne stanowiące źródła losowych sygnałów zakłócających dla drgań elektrycznych. Będą to różnego rodzaju szумы występujące w obwodach elektrycznych oraz potraktowane losowo oddziaływanie zmian warunków atmosferycznych i zasilania. Wśród zjawisk szumowych można wyodrębnić następujące kategorie: szумы oporników, cieplne i strukturalne, oraz szумы lampowe pochodzące od efektów: śrutowego, migotania lub rozplywu prądów pomiędzy poszczególne siatki. Spośród wyżej wymienionych źródeł sygnałów zakłócających jedynie szum cieplny w opornikach i efekt śrutowy stanowią tzw. szum biały, tzn. charakteryzują się widmem ciągłym, praktycznie w całym pasmie częstotliwości. Moc pozostałych sygnałów zakłócających jest rozłożona głównie w obszarze niskich (efekt migotania) i bardzo niskich (losowe zmiany warunków atmosferycznych i zasilania) częstotliwości.

Zgodnie z ogólnym, omówionym w rozdz. 1 obrazem fizycznym zakłócenia losowe mogą oddziaływać na źródło drgań dwoma sposobami: energia zakłóceń skupiona w pasmie niskich częstotliwości oddziałuje pośrednio poprzez element nieliniowy generatora na drodze modulacji elektrycznej lub parametrycznej drgań generowanych; energia zakłóceń zgrupowana w pasmie obwodu drgań działa bezpośrednio (nawet w przypadku pracy liniowej generatora) na źródło drgań. Jednoczesne oddziaływanie obydwoma sposobami możliwe jest tylko w przypadku szumów białych, pozostałe sygnały zakłócające wpływają jedynie poprzez nieliniowość generatora. Ponieważ w przypadku szumu białego gęstość mocy zakłóceń jest stała w całym pasmie częstotliwości, oddziaływanie bezpośrednie jest znacznie bardziej efektywne od oddziaływania pośredniego. W związku z tym dla uproszczenia opisu zjawisk pomija się w opracowaniach teoretycznych oddziaływanie pośrednie.

Opis oddziaływania zakłóceń losowych na źródło drgań zostanie przeprowadzony poniżej, w zasadzie dla przypadku szumu białego. Aby umożliwić za pomocą tych samych zależności opisanie wpływu zarówno szumów cieplnych, jak i śrutowych, wybrano jako podstawowy parametr sygnału losowego moc jednostkową zakłóceń wydzieloną na obwodzie drgań, wyrażoną w postaci dyspersji σ_0^2 lub stosunku mocy jednostkowej drgań monochromatycznych $\frac{A_0^2}{2}$ do mocy jednostkowej szumów σ_0^2 oznaczonego jako:

$$s = \frac{A_0^2}{2\sigma_0^2}. \quad (3.1)$$

Napięcie szumów cieplnych u_{nc} powstające na zaciskach opornika jest wywołane bezładnym ruchem cieplnym ładunków wewnątrz opornika. Średni kwadrat napięcia tych szumów, czyli dyspersja lub moc jednostkowa, dla częstotliwości niższych niż 10^{12} Hz wyraża się znanym wzorem Nyquista:

$$\overline{u_{nc}^2} = \sigma^2(u_{nc}) = 4k_B T_k R \Delta f, \quad (3.2)$$

gdzie:

$k_B = 1,38 \cdot 10^{-23}$ J/°K — stała Boltzmanna,

Δf — skuteczna szerokość pasma w Hz,

T_k — temperatura bezwzględna w °K,

R — oporność opornika w Ω .

W przypadku ogólnym dla dowolnego dwójnika biernego gęstość mocy jednostkowej szumów cieplnych można wyrazić jako:

$$G_n(\omega) = \frac{2}{\pi} k_B T_k R(\omega), \quad (3.3)$$

gdzie $R(\omega)$ jest rezystancją danego dwójnika.

Wykorzystując zależność (2.6) można wyznaczyć gęstość mocy jednostkowej szumów cieplnych wydzieloną w obwodzie rezonansu równoległego:

$$G_0(\omega)_c = \frac{2}{\pi} \frac{k_B T_k R}{1 + \tau_0^2(\omega - \omega_0)^2} = \frac{1}{\pi} \frac{\sigma_{0c}^2 \tau_0}{1 + \tau_0^2(\omega - \omega_0)^2}. \quad (3.4)$$

Moc jednostkowa szumów cieplnych, czyli dyspersja napięcia tych szumów na obwodzie drgań, jako całość wyrażenia (3.4) wyraża się w sposób następujący:

$$\sigma_{0c}^2 = \frac{k_B T_k}{C}, \quad (3.5)$$

gdzie C oznacza pojemność obwodu (por. rys. 2).

Traktowanie szumu cieplnego jako szumu białego jest przybliżeniem wystarczającym do dalszych rozważań. Dla częstotliwości rzędu 10^{12} Hz i wyższych należy uwzględnić tzw. efekt kwantowania, który ogranicza moc szumów cieplnych. Gęstość mocy szumu cieplnego z uwzględnieniem efektu kwantowania może być wyrażona w postaci

$$G_n(\omega) = \frac{1}{\pi^2} \frac{R h \omega}{\frac{\omega h}{e^{2\pi k_B T_k}} - 1}, \quad (3.6)$$

gdzie $h = 6,6 \cdot 10^{-34}$ J·sek — stała Plancka.

Dzięki temu ograniczeniu całkowita moc szumów cieplnych w całym pasmie częstotliwości jest stała i wynosi przy $T_k = 300^\circ\text{K}$ ok. $0,2 \mu\text{W}$.

Dla $hf \ll k_B T_k$ równanie (3.6) sprowadza się do postaci (3.3).

Przyczyną szumu śrutowego są fluktuacje prądu lampy wywołane losowym charakterem emisji i ruchu elektronów w lampie. Średni kwadrat, czyli dyspersja prądu szumów śrutowych i_{ns} dla zakresu częstotliwości, w którym można pominąć czas przelotu elektronów w lampie, wyraża wzór Schottky'ego:

$$\overline{i_{ns}^2} = \sigma^2(i_{ns}) = 2 e \bar{I}_a \lambda \Delta f, \quad (3.7)$$

gdzie:

$e = 1,6 \cdot 10^{-19} \text{C}$ — ładunek elektronu,

$\bar{I}_a$ — wartość średnia prądu anodowego lampy,

λ — współczynnik zależny od rodzaju i klasy pracy lampy. (Dla diody pracującej w warunkach nasycenia prądu anodowego $\lambda = 1$. Dla lamp typu 6J7 i 6K7 pracujących w klasie B lub C $\lambda = 0,3 \pm 0,15$ [5], [21]).

Wyznaczenie współczynnika λ dla lampy generacyjnej w typowych warunkach pracy jest połączone z dodatkową operacją uśredniania za okres drgań generowanych [5]. Intensywność efektu śrutowego w tych warunkach nie jest stała w okresie drgań. Np. przy pracy w klasie C w odstępach czasu, w których nie płynie prąd anodowy, w obciążeniu lampy nie wydzielą się moc szumów śrutowych. Z tego względu oddziaływanie szumów śrutowych jest często określane jako anizotropowe, w odróżnieniu od izotropowego charakteru efektów związanych z szumami cieplnymi.

Gęstość mocy spadku napięcia od prądu śrutowego na impedancji $Z(\omega)$ można wyrazić za pomocą zależności (2.5) w postaci:

$$G_0(\omega)_s = \frac{e \bar{I}_a \lambda}{\pi} |Z(\omega)|^2 = \frac{1}{\pi} \frac{e \bar{I}_a \lambda R^2}{1 + \tau_0^2(\omega - \omega_0)^2} = \frac{1}{\pi} \frac{\sigma_{0s}^2 \tau_0}{1 + \tau_0^2(\omega - \omega_0)^2}, \quad (3.8)$$

gdzie moc jednostkowa szumów śrutowych wydzielona w obwodzie rezonansowym:

$$\sigma_{0s}^2 = \frac{e \bar{I}_a \lambda R}{2C}. \quad (3.9)$$

Podobnie jak w przypadku szumów cieplnych, traktowanie szumu śrutowego jako szumu białego jest wystarczającym przybliżeniem dla dalszych rozważań. Należy jednak pamiętać, że skończona wartość czasu przelotu elektronów powoduje ograniczenie również mocy szumów śrutowych. Gęstość mocy $G_i(\omega)$ prądów śrutowych dla skończonego czasu przelotu elektronów Θ , może być wyrażona w postaci [17]:

$$G_i(\omega) = \frac{e \bar{I}_a \lambda}{2\pi} \frac{8}{(\omega \Theta)^2} [(\omega \Theta)^2 + 2(1 - \cos \omega \Theta - \omega \Theta \sin \omega \Theta)]; \quad (3.10)$$

dla $\omega \Theta \ll 1$ $G_i(\omega) = \frac{e \bar{I}_a \lambda}{\pi}$, skąd bezpośrednio wynika wzór Schottky'ego.

W oparciu o wzory (3.5) i (3.9) można określić dla danego obwodu stosunek mocy jednostkowych szumów śrutowego i ciepłego:

$$\frac{\sigma_{0s}^2}{\sigma_{0c}^2} = \frac{e\bar{I}_a\lambda R}{2k_B T_k} \approx 20 \bar{I}_a R \lambda \quad \text{przy } T_k = 300^\circ\text{K}. \quad (3.11)$$

Z powyższych obliczeń wynika, że w typowych układach generatorów lampowych przeważa wpływ szumów śrutowych. Dalsze rozważania będą prowadzone dla szumu białego izotropowego w postaci ogólnej, którego moc jednostkowa wydzielona na obwodzie drgań wynosi σ_0^2 . W celu obliczenia parametrów dla konkretnego rodzaju zakłóceń należy zamiast σ_0^2 podstawić odpowiednio σ_{0c}^2 lub σ_{0s}^2 , przy czym anizotropię szumu śrutowego należy wyeliminować przez odpowiednie obliczenia współczynnika λ . W związku z tym gęstość mocy $G_0(\omega)$ można na podstawie równań (3.4) i (3.8) wyrazić dla obu rodzajów szumów w postaci ogólnej:

$$G_0(\omega) = \frac{1}{\pi} \frac{\sigma_0^2 \tau_0}{1 + \tau_0^2 (\omega - \omega_0)^2}. \quad (3.12)$$

4. DYFUZYJNY MODEL DRGAŃ LOSOWYCH

W pracach teoretycznych są stosowane cztery podstawowe metody badania drgań losowych:

- Metoda oparta na bezpośredniej analizie procesów fizycznych zachodzących w obwodach generatora z zastosowaniem zasady próbkowania. Wynik ostateczny uzyskuje się z sumy odpowiedniego szeregu [10], [11]
- Metoda polegająca na rozwiązaniu równania różniczkowego układu generacyjnego z wyrazem opisującym oddziaływanie zakłóceń losowych. Rozwiązanie równania uzyskuje się za pomocą metody małego lub wolnozmiennego parametru [12], [18], [22]
- Metoda polegająca na wyznaczeniu rozkładu gęstości prawdopodobieństwa obwiedni i fazy drgań jako rozwiązanie odpowiednich równań Fokkera-Plancka [5], [25]
- Metoda traktująca układ generacyjny jako silnie odtłumiony selektywny wzmacniacz napięcia szumów [21].

Praca niniejsza jest oparta w zasadzie na wynikach osiągniętych przez Edsona [10], Bersteina [5] i Gonorowskiego [12], reprezentujących trzy pierwsze z wyżej wymienionych metod. W ten sposób prawidłowość większości przytoczonych poniżej zależności została sprawdzona za pomocą trzech niezależnych metod analitycznych.

Stworzony w taki sposób model drgań losowych został nazwany modelem dyfuzyjnym ze względu na prawo dyfuzji rządzące dyspersją odchyleń fazy drgań. Przy opracowaniu tego modelu posługiwano się następującymi założeniami upraszczającymi:

- A. Uwzględnia się oddziaływanie tylko szumu białego na układ generacyjny pomijając wpływ innych sygnałów zakłócających
- B. Ponieważ $\sigma_0 \ll A_0$, zaniedbuje się wzajemne oddziaływanie fluktuacji amplitudy i kąta fazowego drgania oraz przyjmuje się, że $\arctg \frac{\sigma_0}{A_0} = \frac{\sigma_0}{A_0}$
- C. Zakłada się symetrię krzywej rezonansowej obwodu drgań
- D. Przyjmuje się, że dobroć obwodu drgań jest wystarczająca do zapewnienia wąskopasmowości układu, tzn. $\delta_0 \omega \ll \omega_0$
- E. Rozpatruje się zmiany amplitudy i kąta fazowego drgań w odstępie czasu τ znacznie większym od stałej czasu obwodu τ_0 i okresu korekcji napięcia szumów na tym obwodzie
- F. Stosuje się metodę linearyzacji zastępczej opartą na wprowadzeniu współczynnika trwałości cyklu granicznego α (por. rozdz. 3).

W oparciu o założenie wąskopasmowości procesów badanych, oddziaływanie szumu białego na obwód rezonansowy można wyrazić bezpośrednio w postaci napięcia szumów $u_N(t)$ występującego na tym obwodzie [17], [4]:

$$u_N(t) = u_c(t) \cos \omega_0 t + u_s(t) \sin \omega_0 t \quad (4.1)$$

lub:

$$u_N(t) = U_N(t) \cos [\omega_0 t - \varphi_N(t)], \quad (4.2)$$

gdzie $U_N(t) = \sqrt{u_s(t)^2 + u_c(t)^2}$ jest amplitudą obwiedni wypadkowego napięcia szumów, a $\varphi_N(t) = \arctg \frac{u_s(t)}{u_c(t)}$ jest kątem fazowym tego napięcia; $u_s(t)$ i $u_c(t)$ są amplitudami składowych, sinusoidalnej i cosinusoidalnej napięcia szumów.

Funkcje $u_s(t)$, $u_c(t)$, $U_N(t)$ i $\varphi_N(t)$ są funkcjami losowymi wolno-zmiennymi w czasie w porównaniu z funkcjami $\sin \omega_0 t$ i $\cos \omega_0 t$. Funkcje $u_s(t)$ i $u_c(t)$ są stacjonarnymi gaussowskimi procesami losowymi o zerowej wartości średniej. W przypadku gdy krzywa przenoszenia obwodu drgań jest idealnie symetryczna, procesy $u_s(t)$ i $u_c(t)$ będą statycznie niezależne. Z tego względu, pomimo iż najbardziej prawdopodobnymi wartościami tych funkcji są wartości równe zero, prawdopodobieństwo, że obie funkcje przyjmują jednocześnie wartość zero jest bardzo małe. Proces losowy $U_N(t)$ charakteryzuje się rayleighowskim rozkładem prawdopodobieństwa, a funkcja $\varphi_N(t)$ posiada rozkład równomierny w przedziale $\pm \pi$ radianów.

Gęstość mocy średniej, czyli widmo $G_0(\omega)$ sygnału $u_N(t)$ opisywane jest dla szumu cieplnego i szumu śrutowego odpowiednio przez funkcje (3.4) i (3.8). Z uwagi na statystyczną niezależność obu składowych z wy-

rażenia (4.1), średnia moc jednostkowa napięcia szumów na obwodzie drgań będzie sumą mocy poszczególnych składowych, a więc:

$$\sigma_0^2 = \sigma^2[u_N(t)] = \frac{\overline{u_s(t)^2}}{2} + \frac{\overline{u_c(t)^2}}{2} = \frac{\overline{u_N(t)^2}}{2}, \quad (4.3)$$

przy czym:

$$\sigma_0^2 = \overline{u_s(t)^2} = \overline{u_c(t)^2}. \quad (4.4)$$

Również widmo procesów $u_s(t)$ i $u_c(t)$ będzie opisywane przez tę samą funkcję $G_0(\omega)$, określoną wzorem (3.12), co widmo procesu $u_N(t)$.

Jak wynika z powyższych rozważań, oddziaływanie szumu białego na obwód generacyjny może być wyrażone przez oddziaływanie dwóch statystycznie niezależnych, ortogonalnych składowych o losowych amplitudach $u_s(t)$ i $u_c(t)$. Przyjmując dla monochromatycznego modelu drgań generowanych w danym obwodzie zależność (1.1) przy $\varphi_0 = 0$ łatwo stwierdzić, że składowa $u_c(t)$ wywoła bezładną modulację amplitudy, a składowa $u_s(t)$ bezładną modulację kąta drgań wypadkowych.

Oba te efekty zostaną omówione kolejno, niezależnie od siebie, co jest dopuszczalne na podstawie przyjętych powyżej założeń.

Gdyby obwód generacyjny pracował w warunkach liniowych, np. w układzie wzmacniacza, powstałe wskutek liniowej superpozycji sygnałów (1.1) i (4.1) bezładne zmiany amplitudy, potraktowane jako modulacja pasożytnicza, charakteryzowałyby się średnim kwadratem głębokości modulacji $\overline{m_a^2}$ równym:

$$\overline{m_a^2} = \frac{\overline{u_c(t)^2}}{A_0^2} = \frac{\sigma_0^2}{A_0^2} = \frac{1}{2s}. \quad (4.5)$$

Ponadto obie wstęgi boczne widma takiego drgania byłyby całkowicie niezależne, a gęstość mocy jednostkowej takiego przebiegu na podstawie równań (3.4) lub (3.8) wyrażałaby się w postaci:

$$G_a(\omega)_w = \frac{A_0^2}{2} \delta(\omega + \omega_0) + \frac{1}{2\pi} \frac{\sigma_0^2 \tau_0}{1 + \tau_0^2(\omega + \omega_0)^2} {}^1). \quad (4.6)$$

Pierwszy wyraz prawej strony równania (4.6) przedstawia widmo monochromatycznej fali nośnej, a wyraz drugi widmo wstęp bocznych. Szerokość widma wstęp bocznych obliczona z powyższego równania wynosi:

$$(\delta_a \omega)_w = \frac{2}{\tau_0} = \delta_0 \omega, \quad (4.7)$$

a moc jednostkowa wstęp bocznych: $\frac{\sigma_0^2}{2}$.

Z uwagi na statystyczną niezależność wstęp bocznych widmo obwiedni $G_e(\Omega)_w$, tzn. widmo wyjścia z idealnego detektora amplitudy, będzie

¹⁾ Symbol $\delta(\cdot)$ oznacza funkcję Diraca.

energetyczną (kwadratową) sumą poszczególnych składowych ($\omega \pm \omega_0$) widma wstęp bocznych:

$$G_e(\Omega)_w = \frac{1}{\pi} \frac{\sigma_0^2 \tau_0}{1 + \tau_0^2 \Omega^2}, \quad (4.8)$$

gdzie $\Omega = |\omega - \omega_0|$.

W rzeczywistym nieliniowym generatorze należy powyższy obraz uzupełnić oddziaływaniem dwóch efektów: samoczynnego ograniczania modulacji amplitudy (zob. rozdz. 2) oraz wzmocnienia parametrycznego o częstotliwości pompowania $2\omega_0^1$.

Działanie układu ograniczania amplitudy powoduje, że dla sygnału z modulacją amplitudy schemat zastępczy obwodu drgań ma postać przedstawioną na rys. 4. W związku z tym widmo wypadkowej modulacji amplitudy będzie miało postać:

$$G_a(\omega) = \frac{A_0^2}{2} \delta(\omega - \omega_0) + \frac{1}{2\pi} \frac{\sigma_0^2 \tau_0}{\alpha^2 + \tau_0^2 (\omega - \omega_0)^2}, \quad (4.9)$$

a odpowiednia szerokość widma wstęp bocznych modulacji amplitudy:

$$\delta_a \omega = \frac{2\alpha}{\tau_0} = \alpha \delta_0 \omega, \quad (4.10)$$

przy czym moc jednostkowa wstęp bocznych wynosi: $\frac{\sigma_0^2}{2\alpha}$.

Z porównania wzorów (4.6) i (4.9) wynika, że w przypadku generatora nieliniowego widmo wstęp bocznych modulacji amplitudy jest α^2 razy mniejsze dla $\omega = \omega_0$ i α — razy szersze, niż widmo analogicznego układu liniowego, jest ono również α — razy szersze od pasma obwodu drgań generatora. Oba widma przedstawiono na rys. 5.

Wskutek wzmocnienia parametrycznego układu generacyjnego wystąpi pełna korelacja obu wstęp bocznych modulacji amplitudy. W związku z tym widmo obwiedni amplitudy będzie arytmetyczną, a nie energetyczną (kwadratową) sumą poszczególnych składowych widma wstęp bocznych o pulsacjach ($\omega \pm \omega_0$):

$$G_e(\Omega) = \frac{2}{\pi} \frac{\sigma_0^2 \tau_0}{\alpha^2 + \tau_0^2 \Omega^2}. \quad (4.11)$$

Widmo to zostało przedstawione na rys. 8.

Funkcja autokorelacji obwiedni $R_e(\tau)$ ma postać:

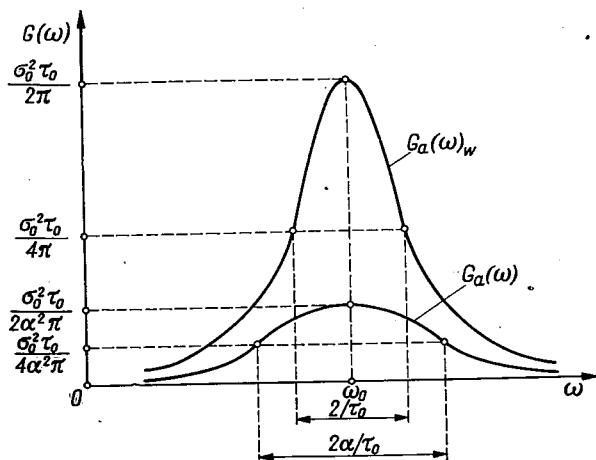
$$R_e(\tau) = \frac{\sigma_0^2}{\alpha} e^{\frac{-\alpha|\tau|}{\tau_0}}. \quad (4.12)$$

¹⁾ W generatorze z rozdz. 2 wypadkowa przewodność chwilowa zmienia się z pulsacją $2\omega_0$ [18].

Jak widać z powyższego, moc jednostkowa, czyli dyspersja napięcia obwiedni $\overline{\Delta A^2}$ jest dwukrotnie wyższa od mocy wstęg bocznych widma modulacji amplitudy:

$$\overline{\Delta A^2} = R_e(0) = \frac{\sigma_0^2}{\alpha}. \quad (4.13)$$

Reasumując wnioski wyciągnięte z przytoczonych powyżej rozważań należy podkreślić, że dzięki ograniczającym fluktuacje amplitudy właści-

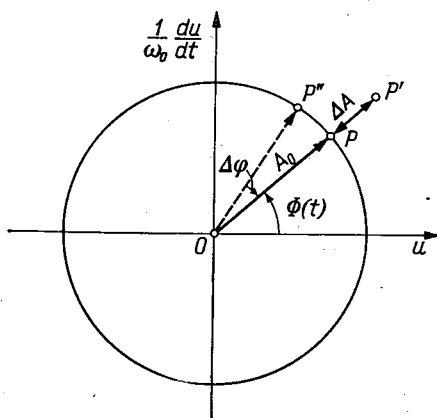


Rys. 5. Widma fluktuacyjnej modulacji amplitudy wzmacniacza $G_a(\omega)_w$ i generatora $G_a(\omega)$

wościom układu generacyjnego oddziaływanie na tej drodze sygnału zakłócającego na drganie generowane jest w poważnym stopniu osłabione. Miarą tego osłabienia jest współczynnik α .

Jak już wspomniano w rozdz. 2, współczynnik α jest również miarą trwałości cyklu granicznego drgania. Obraz drgania na płaszczyźnie fazowej przedstawiony na rys. 6 określa drgania w układzie współrzędnych $\left(u; \frac{1}{\omega_0} \frac{du}{dt}\right)$, gdzie u oznacza chwilową wartość napięcia na obwodzie generacyjnym. Dla drgania monochromatycznego obrazem cyklu granicznego będzie koło o stałym promieniu równym amplitudzie A_0 , po którym porusza się ze stałą szybkością kątową ω_0 [$\Phi(t) = \omega_0 t$] punkt P reprezentujący chwilowy stan badanego przebiegu monochromatycznego. Na płaszczyźnie fazowej można bezpośrednio zinterpretować określony model statystyczny, jaki można przypisać bezładnym fluktuacjom amplitudy. Fluktuacje te będą opisane jednoznacznie przez bezładny ruch po promieniu wodzącym punktu P' (rys. 6) określającego aktualny stan drgania losowego. Bezładny ruch punktu P' nie jest jednak ruchem swobodnym, gdyż podlega on oddziaływaniu siły sprężystości skierowanej do punktu P (leżącego na obwodzie cyklu granicznego drgań monochroma-

tycznych) o wartości $\frac{\alpha}{\tau_0} \Delta A$ dla generatora lub $\frac{1}{\tau_0} \Delta A$ dla wzmacniacza, gdzie $\Delta A = u_c(t)$ oznacza aktualne odchylenie punktu P' od P . W związku z powyższym, prawdopodobieństwo oddalenia się punktu P' od punktu P jest tym mniejsze, im większa jest wartość ΔA , a ponadto w przypadku generatora, im większa jest wartość współczynnika α .



Rys. 6. Obraz fazowy drgania

Układ generacyjny dzięki nieliniowości jest więc w stanie skutecznie korygować zmiany amplitudy wywoływane obecnością sygnałów zakłócających. Dla zakłóceń fazowych generator jest układem liniowym i nie zachodzi w nim analogiczny proces ograniczania fluktuacji fazowych. Jeżeli w generatorze monochromatycznym za pomocą oddziaływania zewnętrznego spowodować przyrost amplitudy o wartości ΔA , to po ustaleniu działania tej przyczyny układ wróci ze stałą czasową $\frac{\tau_0}{\alpha}$ do poprzedniego cyklu granicznego. Wymuszenie w analogicznych warunkach przyrostu kąta fazowego $\Delta\varphi$ spowoduje, że drganie będzie zachodzić od nowego stanu ($\omega_0 t + \Delta\varphi$) wzdłuż poprzedniego cyklu granicznego. Ujmując rzecz lapidarnie można stwierdzić, że zakłócenia fazowe powodują powstawanie luk w „życiorysie fazowym drgania”.

Do opisu fluktuacji fazowych w generatorze służy inny model statystyczny. Jest nim bezładne błądzenie po obwodzie cyklu granicznego punktu swobodnego P'' wokół punktu P (rys. 5). Prawdopodobieństwa przejścia punktu P'' w stronę punktu P lub w stronę przeciwną są w tym przypadku równe bez względu na aktualny przyrost $\Delta\varphi \approx \frac{u_s(t)}{A_0}$ kąta fazowego (względem fazy drgań monochromatycznych). Jest to model statystyczny stosowany już poprzednio do opisu procesów dyfuzji — stąd też wywodzi się nazwa opisywanego modelu drgań.

Dyspersję przyrostu kąta fazowego $\Delta\varphi_\tau$ za odstęp czasu τ , gdy $\tau \gg \tau_0$, określa tzw. prawo dyfuzji:

$$\sigma^2(\Delta\varphi_\tau) = [\overline{\Phi(t+\tau) - \Phi(t)}]^2 = 2D\tau, \quad (4.14)$$

gdzie D jest stałym parametrem, zwanym współczynnikiem dyfuzji. Równanie (4.14) jest podstawową zależnością charakteryzującą dyfuzyjny model drgań. Oznacza ono, że średniokwadratowe odchylenie fazy drgań od wartości $\omega_0 t$ w odstępie czasu τ rośnie proporcjonalnie do długości tego odstępu czasu.

Współczynnik dyfuzji D jest najbardziej charakterystycznym parametrem tego modelu, wchodzącym do wszystkich równań opisujących drgania dyfuzyjne.

Wielkość współczynnika D była wyznaczona różnymi metodami przez różnych autorów. Dla modelu generatora z rozdz. 2 można wyrazić ją w postaci:

$$D = \frac{\sigma_c^2}{A_0^2 \tau_0} = \frac{1}{2s\tau_0}. \quad (4.15)$$

Jeżeli dyspersja fazy spełnia prawo dyfuzji (4.14), to gęstość mocy drgań o takiej modulacji kąta wyraża się zależnością:

$$\begin{aligned} G_k(\omega) &= \frac{A_0^2}{2\pi} \frac{D}{D^2 + (\omega + \omega_0)^2} = \frac{1}{2\pi} \frac{\sigma_0^2 \tau_0}{\frac{\sigma_0^4}{A_0^4} + \tau_0^2 (\omega - \omega_0)^2} = \\ &= \frac{1}{2\pi} \frac{\sigma_0^2 \tau_0}{\frac{1}{4s^2} + \tau_0^2 (\omega - \omega_0)^2}, \end{aligned} \quad (4.16)$$

a szerokość widma:

$$\delta_k \omega = 2D = \frac{2\sigma_0^2}{A_0^2 \tau_0} = \frac{\delta_0 \omega}{2s}. \quad (4.17)$$

Na rys. 7 przedstawiono oba widma, modulacji amplitudy $G_a(\omega)$ i kąta $G_k(\omega)$, dla drgań dyfuzyjnych.

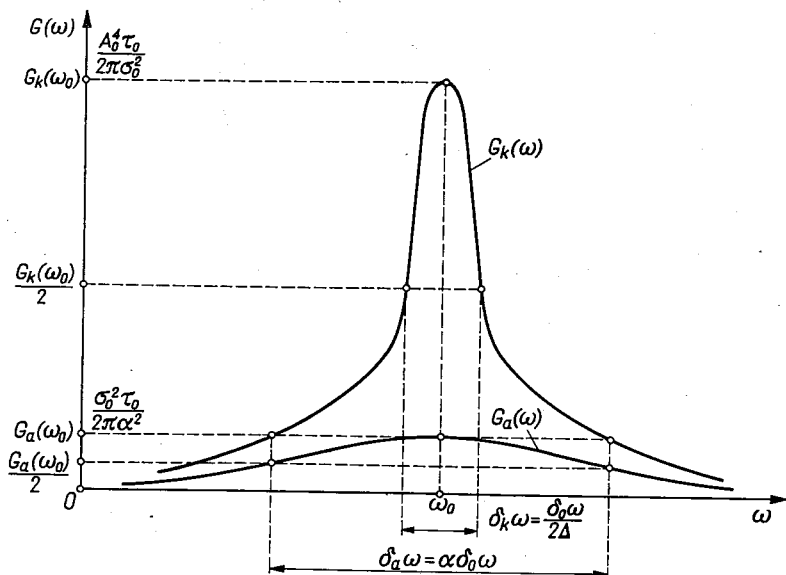
Funkcja autokorelacji obliczona z równania (4.16) ma postać:

$$R_k(\tau) = \frac{A_0^2}{2} \cos \omega_0 \tau \cdot e^{-D|\tau|}. \quad (4.18)$$

Moc jednostkowa drgań dyfuzyjnych równa się mocy drgań głównych, czyli mocy drgań monochromatycznych generatora bezszumnego, tzn.:

$$R_k(0) = \frac{A_0^2}{2}. \quad (4.19)$$

Jak wynika z powyższych rozważań, fluktuacje fazy są przyczyną rozszczepienia prążka drgań monochromatycznych na widmo ciągłe o bardzo wąskim pasmie, skupione wokół częstotliwości drgań monochromatycznych ω_0 . W związku z tym drgania dyfuzyjne nie zawierają skła-



Rys. 7. Widma modulacji amplitudy $G_a(\omega)$ i kąta $G_k(\omega)$ drgań dyfuzyjnych

dowej koherentnej. Ponadto fluktuacje fazy drgań dyfuzyjnych nie zależą zupełnie od nieliniowości generatora. Okres korelacji drgań określony jako:

$$\tau_R = \frac{1}{\delta_k \omega} \quad (4.20)$$

jest okresem czasu, w którym standardowe odchylenie fazy $\sigma(\Delta\varphi_\tau)$ osiągnie wartość jednego radiana, tzn. $\sigma(\Delta\varphi_{\tau_R}) = 1$, i jest praktyczną miarą odstępu czasowego, w którym można dokonywać pomiarów częstotliwości metodami fazowymi.

Uwzględniając wzór (4.17) widać, że dla typowych generatorów częstotliwości okres korelacji będzie dostatecznie długim odstępem czasu. W związku z tym drganie badane można traktować jako proces w wysokim stopniu zdeterminowany.

Z prawa dyfuzji można obliczyć względną wartość standardowej dewiacji częstotliwości średniej za odstęp czasu τ :

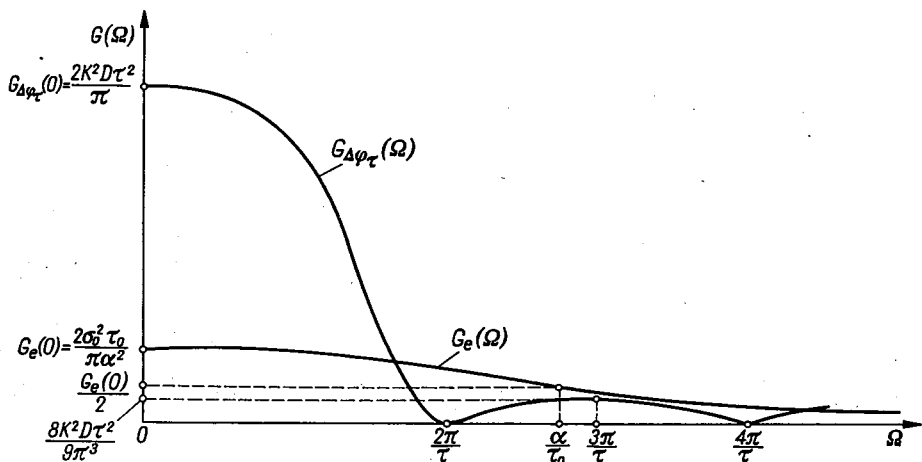
$$\frac{\sqrt{\overline{(\bar{f}_\tau - f_0)^2}}}{f_0} = \frac{\sigma(\bar{f}_\tau)}{f_0} = \frac{\sigma(\Delta\varphi_\tau)}{\Delta\varphi_\tau} = \frac{\sqrt{2D\tau}}{\omega_0 \tau} = \frac{\sigma_0 \sqrt{2}}{A_0 \omega_0 \sqrt{\tau} \tau_0} = \frac{1}{\omega_0 \sqrt{s \tau_0 \tau}} \quad (4.21)$$

Dla modelu dyfuzyjnego obliczono [13], [5] widmo energetyczne dla procesu losowego $\Delta\varphi_\tau$, czyli widmo sygnału wyjściowego z idealnego detektora fazowego, na którego wejścia przyłożono drgania dyfuzyjne o równych amplitudach, lecz przesunięte w czasie o odstęp τ :

$$G_{\Delta\varphi_\tau}(\Omega) = K^2 \frac{2D}{\pi} \left(\frac{\sin \frac{\Omega\tau}{2}}{\frac{\Omega}{2}} \right)^2, \quad (4.22)$$

gdzie K [V/rad] oznacza nachylenie charakterystyki detektora fazowego.

Na rys. 8 przedstawiono widma otrzymane przy idealnej detekcji amplitudy $G_e(\Omega)$ i kąta $G_{\Delta\varphi_\tau}(\Omega)$, określone odpowiednio zależnościami (4.11) i (4.22).



Rys. 8. Widma obwiedni amplitudy $G_e(\Omega)$ i przyrostu kąta fazowego za odstęp czasu τ $G_{\Delta\varphi_\tau}(\Omega)$ drgań dyfuzyjnych

Jak wynika z powyższych zależności, widmo $G_{\Delta\varphi_\tau}(\Omega)$ zależy od odstępu czasowego τ i zanika przy częstotliwościach $F = \frac{a}{2\pi\tau}$, gdzie a liczba naturalna. Zgodnie z przyjętymi poprzednio założeniami oba widma wykreślone zostały dla przypadku $\tau \gg \tau_0$. Szerokość widma $G_e(\Omega)$ równa jest połowie pasma obwodu z rys. 4.

W celu lepszego uwypuklenia istoty zjawisk związanych z dyfuzyjnymi modelami drgań warto jest analogicznie jak dla modulacji amplitudy, opisać w skrócie wpływ ortogonalnej składowej napięcia szumów na obwodzie wzmacniacza. Jak wynikało z poprzednich rozważań, w układzie generacyjnym nie było ujemnego fazowego sprzężenia zwrotnego, w związku z czym odchylenie fazy było opisywane przez bezładny ruch

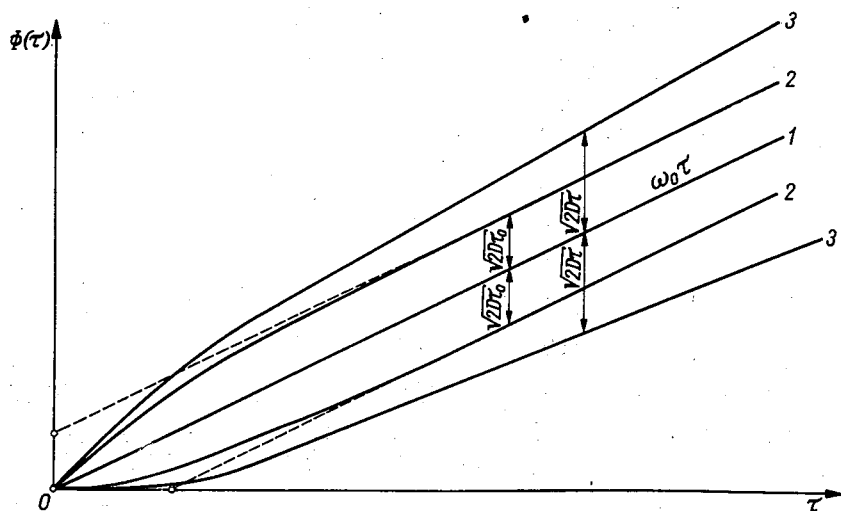
punktu swobodnego wzdłuż cyklu granicznego. We wzmacniaczu faza drgań wyjściowych jest determinowana przez sygnał wejściowy. Przyjmując, że sygnał wejściowy jest sygnałem monochromatycznym, łatwo zauważyć, że odchyleniami fazy napięcia wyjściowego wzmacniacza będzie rządzić podobny schemat statystyczny jak w przypadku modulacji amplitudy drgań generatora, tzn. będą one opisywane przez błędzenie punktu podlegającemu sile sprężystości:

$$\frac{\Delta\varphi}{2\tau_0} \approx \frac{u_s(t)}{2\tau_0 A_0}.$$

Dyspersja odchylenia fazy dla wzmacniacza przy $\tau \gg \tau_0$ wyraża się wzorem następującym [12]:

$$\sigma^2(\Delta\varphi_c)_w = 2D\tau_0 = \text{const.} \quad (4.23)$$

Na rys. 9 przedstawiono w poglądowy sposób zależności (4.14) i (4.23) opisujące dyspersję kąta fazowego w przypadku generatora i wzmacniacza. Prosta $\Phi(\tau) = \omega_0 \tau$ określa fazę drgań monochromatycznych. Krzywe 2 i 3 ograniczają odpowiednio obszary zmian kąta fazowego, określone przez standartowe odchylenie tego kąta, w przypadku wzmacniacza i generatora.



Rys. 9. Obszar zmian fazy drgań
1 — drgania monochromatyczne,
2-2 — drgania wyjściowe wzmacniacza,
3-3 — drgania dyfuzyjne generatora.

Napięcie wyjściowe wzmacniacza będzie więc zawierać składową koherentną, odpowiadającą wzmocnionemu sygnałowi wejściowemu, oraz wstęgi boczne pochodzące od fluktuacyjnej modulacji amplitudy i kąta. Ponieważ oddziaływanie zakłóceń ma miejsce w obwodzie liniowym, a za-

kłócenia amplitudy i kąta są statystycznie niezależne, wypadkowe widmo wzmacniacza $G(\omega)_w$ będzie miało postać:

$$G(\omega)_w = \frac{1}{2} A_0^2 \delta(\omega - \omega_0) + G_a(\omega)_w + G_k(\omega)_w, \quad (4.24)$$

gdzie $G_a(\omega)_w$ określona jest wzorem (4.6) a $G_k(\omega)_w$ przedstawia widmo wstęp bocznych modulacji kąta o podobnym kształcie.

Wypadkowe widmo drgań generatora $G(\omega)$ dla modelu dyfuzyjnego, w którym pomija się wzajemne oddziaływanie fluktuacji amplitudy i kąta, można opisać za pomocą wyrażenia:

$$G(\omega) = \frac{\sigma_0^2 \tau_0}{2\pi} \left[\frac{1}{\frac{\sigma_0^4}{A_0^4} + \tau_0^2 (\omega - \omega_0)^2} + \frac{1}{\alpha^2 + \tau_0^2 (\omega - \omega_0)^2} \right] \quad (4.25)$$

lub:

$$G(\omega) = \frac{A_0^2}{2\pi} D \left[\frac{1}{D^2 + (\omega - \omega_0)^2} + \frac{1}{4D^2 \alpha^2 s^2 + (\omega - \omega_0)^2} \right]. \quad (4.26)$$

Pierwszy człon reprezentuje widmo fluktuacji kąta, a drugi fluktuacji amplitudy. Wyrażenie (4.25) powstało z wyrażenia (4.6), w którym składową koherentną $\frac{A_0^2}{2} \delta(\omega - \omega_0)$ zastąpiono widmem modulacji kąta (4.16).

Aby otrzymać wykres funkcji $G(\omega)$, należy dodać obie krzywe z rys. 7. Stosunek szerokości obu widm i stosunek gęstości mocy dla $\omega = \omega_0$ na podstawie zależności (4.10) i (4.17) oraz (4.6) i (4.16) określone są odpowiednio przez następujące wyrażenia:

$$\frac{\delta_a \omega}{\delta_k \omega} = \alpha \frac{A_0^2}{\sigma_0^2} = 2\alpha s, \quad (4.27)$$

$$\frac{G_a(\omega_0)}{G_k(\omega_0)} = \frac{1}{\alpha^2} \frac{\sigma_0^4}{A_0^4} = \frac{1}{4\alpha^2 s^2}. \quad (4.28)$$

Dla punktów połowy mocy widma $G_a(\omega)$ gęstość mocy widma $G_k(\omega)$ jest dwa razy większa od odpowiedniej gęstości mocy widma $G_a(\omega)$.

Podstawowa moc procesu zawarta jest w widmie fluktuacji kąta; jest ona równa mocy ekwiwalentnego drgania koherentnego $\frac{A_0^2}{2}$. W porówna-

niu z nią moc wstęp bocznych modulacji amplitudy $\frac{\sigma_0^2}{2\alpha}$ jest znikoma. Moc wstęp bocznych modulacji amplitudy jest dostarczana przez składową kosinusoidalną źródła szumów. Straty tej mocy wynikające z czynnika α powstają w obwodach ogranicznika amplitudy. W modelu dyfuzyjnym przyjmuje się, że moc składowej sinusoidalnej napięcia szumów zostaje zużyta na „rozszerzenie” prążka drgań monochromatycznych i wobec

dużej wartości stosunku s (sygnał — szum), nie uwzględnia się tej mocy w bilansie energetycznym.

W rezultacie drgania dyfuzyjne są drganiami niekoherentnymi. Z uwagi jednak na dużą wartość okresu korelacji (4.20) stopień zdeterminowania tych drgań jest stosunkowo wysoki. Oznacza to, że dominująca część mocy drgań jest skupiona w bardzo wąskim pasmie częstotliwości.

Jak wynika z powyższego zestawienia, widmo fluktuacji amplitudy jest znacznie szersze i znacznie mniej intensywne niż widmo fluktuacji kąta drgań dyfuzyjnych. Z tego względu przy badaniach eksperymentalnych z dostatecznym dla celów praktycznych przybliżeniem można przyjąć:

$$G(\omega) \approx G_k(\omega), \quad (4.29)$$

tzn. traktować widmo fluktuacji kąta jako widmo wypadkowe drgań dyfuzyjnych.

Dyskutując wpływy poszczególnych parametrów łatwo wykazać słuszność wniosków intuicyjnie oczywistych. Im większa dobroć obwodu drgań (większa stała czasu obwodu — τ_0) i większy stosunek sygnału do szumu, tym mniejszy wpływ szumów na parametry drgań dyfuzyjnych i większy stopień zdeterminowania procesu drgań. Wnioski te wynikają bezpośrednio z przytoczonych powyżej zależności.

Ograniczenia w zastosowaniu modelu dyfuzyjnego drgań wynikają z upraszczających założeń przyjętych przy jego opracowaniu. Pominiecie wpływu modulacji skrośnej oraz wzajemnego oddziaływania modulacji amplitudy na częstotliwość drgań ogranicza zakres zastosowań do zakłóceń o bardzo małym poziomie mocy. Przyjęcie założeń o symetrii krzywej rezonansowej narzuca warunek na czas obserwacji (uśredniania), tzn. że analityczny opis jest słuszny tylko dla zmian parametrów rozpatrywanych za okres czasu znacznie większy od stałej czasu obwodu drgań ($\tau \gg \tau_0$)¹⁾. Na szczęście warunek ten nie stwarza przy obecnym stanie techniki pomiarowej żadnych trudności, ponieważ dokładne pomiary niewielkich zmian fazy lub częstotliwości wymagają stosunkowo długich odstępów czasu. Znacznie poważniejszą wadę modelu dyfuzyjnego stanowi założenie izotropowości zakłóceń. W związku z tym drgania dyfuzyjne opisują w sposób adekwatny tylko wpływ szumu cieplnego, którego oddziaływanie ma stałą wartość na całym obwodzie cyklu granicznego. Niestety, jak wynika ze wzoru (3.10), w generatorach lampowych z reguły przeważa wpływ szumów śrutowych. Operacje związane z eliminacją anizotropii szumów śrutowych za pomocą uśredniania efektów za okres drgań nie są środkiem wystarczającym, jeżeli uwzględnić wpływ szumów skojarzonych z występowaniem prądów siatki w lampie generacyjnej. W dyfuzyjnym modelu drgań nie bierze się również

¹⁾ Uproszczenie to deformuje najbardziej odległe części widma, które właśnie decydują o zjawiskach zachodzących w bardzo krótkich odstępach czasu.

pod uwagę wpływu szumów o pasmie ograniczonym do wstęgi częstotliwości akustycznych. Wszystko to sprawia, że dyfuzyjny model drgań należy traktować jako etap wstępny na drodze do poznania fizycznej strony zjawisk zachodzących poza granicą oddzielającą zdeterminowane modele drgań od modeli losowych.

W literaturze przedmiotu można spotkać modele bardziej uproszczone. W modelach tych pomija się fluktuacje amplitudy, a fluktuacje kąta charakteryzuje się możliwie najprostszymi zależnościami.

W jednym z zaproponowanych modeli [23] pomija się prawo dyfuzji przyjmując, że względna niestałość fazy $\frac{\Delta\varphi}{2\pi}$ odniesiona do jednego okresu drgań ma wartość stałą i jest parametrem charakteryzującym dane źródło drgań. W innych modelach [16], [26] przyjmuje się jako parametr podstawowy widmo wypadkowe drgań. Wyznaczając szerokość widma określa się składową modulującą odpowiadającą punktom połowy mocy. Następnie zakłada się, że oddziaływanie zakłóceń losowych jest równoważne oddziaływaniu zdeterminowanego sygnału zakłócającego o częstotliwości i amplitudzie odpowiadającej punktom połowy mocy widma ciągłego.

Istnieją również w literaturze próby bardziej wnikliwego, niż w modelu dyfuzyjnym, opisu zjawisk. Z uwagi jednak na konieczność stosowania bardzo rozbudowanego aparatu matematycznego, autorzy ograniczają się przeważnie do wyznaczenia metody postępowania, nie doprowadzając wyników do postaci umożliwiających ich praktyczne wykorzystanie.

5. METODY POMIARÓW FLUKTUACYJNEJ DEWIACJI CZĘSTOTLIWOŚCI

5.1. Fluktuacyjna dewiacja częstotliwości

Drgania losowe zostały opisane w rozdz. 4 za pomocą określonych parametrów statystycznych. Przyjęte w tym rozdziale założenia pozwalają na niezależne traktowanie fluktuacji amplitudy i fazy drgań. Z punktu widzenia stabilizacji częstotliwości fluktuacji amplitudy nie przedstawiają większego znaczenia, a ponadto mogą być ograniczone do odpowiedniego poziomu. W związku z powyższym w niniejszym rozdziale zostaną omówione zagadnienia związane wyłącznie z pomiarami efektów wywołanych przez fluktuacje fazy drgań dyfuzyjnych.

Pośród wielkości opisujących pośrednio lub bezpośrednio fluktuacje fazy przyjęto w tym rozdziale jako parametr podstawowy standardową dewiację częstotliwości. Postępowanie takie podyktowane zostało dwoma względami. Z jednej strony wybrana wielkość stanowi określoną formę ekstrapolacji na zjawiska fluktuacyjne powszechnie stosowanego w technice pojęcia niestałości częstotliwości, z drugiej strony dla modelu dyfu-

zyjnego drgań standartowa dewiacja częstotliwości, której wartość względ-
ną opisują zależności (4.21), jest określoną jednoznacznie funkcją współ-
czynnika dyfuzji D , w pełni charakteryzującego pozostałe parametry
statystyczne opisujące fluktuacje.

Przed przystąpieniem do opisu poszczególnych metod pomiarowych
należy pokrótce omówić konsekwencje związane z wprowadzeniem poję-
cia dewiacji częstotliwości do pomiarów częstotliwości.

Na podstawie znajomości sygnałów zakłócających i dróg ich oddzia-
ływania na częstotliwość drgań można przewidzieć ogólny charakter
zmian w czasie częstotliwości chwilowej $f(t)$. W funkcji tego typu można
wyodrębnić następujące składowe: składową aperiodyczną związaną
z efektem starzenia się, składowe periodyczne związane z rocznym i dobo-
wym rytmem działalności człowieka i przemian zachodzących w przy-
rodzie oraz składowe losowe.

W wyniku pomiaru częstotliwości otrzymuje się zawsze wartość śred-
nią za czas pomiaru τ_p , czyli wartość $\overline{f_{\tau_p}}$. W rezultacie zamiast funkcji
ciągłej $f(t)$ otrzymuje się zbiór dyskretnych wartości $\overline{f_{\tau_{pi}}}$, który odtwarza
jedynie przybliżony kształt funkcji $f(t)$. Przybliżenie to wynika z czte-
rech zasadniczych przyczyn:

- a) Zastąpienie funkcji ciągłej $f(t)$ zbiorem wartości dyskretnych
- b) Konieczność uśredniania za okres τ_p
- c) Błąd metody pomiarowej składający się w przypadku ogólnym
z błędu porównania oraz niestałości i niedokładności zastosowa-
nego wzorca
- d) Ograniczona czułość metody pomiarowej, powodująca, że składo-
we funkcji $f(t)$ leżące poniżej progu czułości nie zostaną odtworzone
w zbiorze wyników pomiarowych.

Powyższy skomplikowany obraz zjawisk można w określony sposób
usystematyzować w oparciu o zasadę próbkowania. Dowolna seria n po-
miarów o czasie trwania τ_p , powtarzanych regularnie w odstępach czasu
równym okresowi próbkowania $T_p \geq \tau_p$, pozwala na określenie oddzia-
ływania tych czynników periodycznych, których okres T_z spełnia wa-
runek:

$$nT_p \gg T_z \gg 2T_p, \quad (5.1)$$

oraz czynników losowych o okresie korelacji τ_R , spełniających warunek
analogiczny,

$$nT_p \gg \tau_R \gg T_p. \quad (5.2)$$

Wierność odtworzenia rzeczywistego przebiegu funkcji $f(t)$ będzie tym
lepsz, im lepszy będzie stopień realizacji warunków (5.1) i (5.2).

Składowe aperiodyczne i zdeterminowane $f_a(t)$ można wyznaczyć ze
zbioru wyników pomiarowych $\overline{f_{\tau_{pi}}}$ przy zastosowaniu tzw. metody naj-
mniejszych kwadratów [31]. Polega ona na wyznaczeniu takiej funkcji

$f_a(t)$, aby wartość średnia odchyłeń pomiędzy poszczególnymi wartościami zbioru $\overline{f_{\tau_{pi}}}$ a odpowiadającymi im wartościami funkcji $f_a(t_i)$ była równa zero, a wartość średniego kwadratu tych odchyłeń była minimalna. Warunki te można wyrazić w postaci następującej

$$\frac{1}{n} \sum_{i=1}^n [\overline{f_{\tau_{pi}}} - f_a(t_i)] = 0, \quad (5.3)$$

$$\frac{1}{n} \sum_{i=1}^n [\overline{f_{\tau_{pi}}} - f_a(t_i)]^2 = \sigma^2(\overline{f_{\tau_p}})_n T_p = \min, \quad (5.4)$$

gdzie $\sigma(\overline{f_{\tau_p}})$ oznacza standartową dewiację częstotliwości.

Dzięki wyeliminowaniu składowych aperiodycznych i zdeterminowanych można określić dewiację częstotliwości średniej w czasie τ_p w sposób następujący:

$$\Delta \overline{f_{\tau_p}} = \overline{f_{\tau_{pi}}} - f_{a\tau_p}(t_i), \quad (5.5)$$

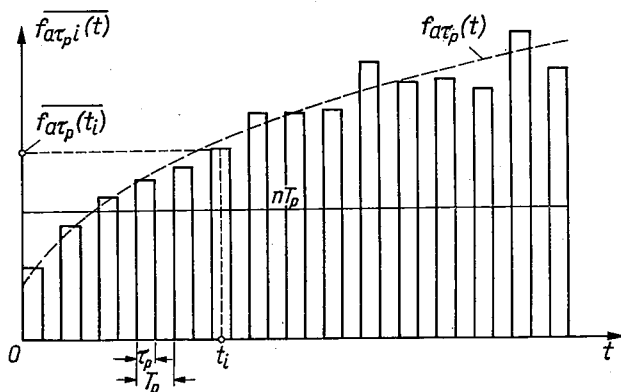
gdzie $\Delta \overline{f_{\tau_p}}$ jest dyskretną funkcją czasu.

Na podstawie funkcji dewiacji można w znany sposób [31] określić rozkład prawdopodobieństwa częstotliwości $\overline{f_{\tau_{pi}}}$. Wyznaczenie rozkładu ma istotne znaczenie, od niego bowiem zależy liczbowa interpretacja standartowej dewiacji częstotliwości $\sigma(\overline{f_{\tau_p}})$, gdzie:

$$\sigma(\overline{f_{\tau_p}}) = \sqrt{(\Delta \overline{f_{\tau_p}})^2}. \quad (5.6)$$

Zastosowanie metody najmniejszych kwadratów do zbioru wyników pomiarowych $\overline{f_{\tau_{pi}}}$ zostało w sposób poglądowy zilustrowane na rys. 10.

W oparciu o powyższe ustalenia można również sprecyzować pojęcie krótko i długookresowej oraz fluktuacyjnej niestabilności częstotliwości.



Rys. 10. Zastosowanie metody najmniejszych kwadratów do zbioru wyników pomiarowych

Pojęciem fluktuacyjnej dewiacji lub niestałości częstotliwości określa się wielkość dewiacji lub niestałości częstotliwości pochodzącej od zakłóceń o wartościach T_z lub τ_R , porównywalnych z okresem drgań podstawowych T_0 . Przy $T_z \gg T_0$ lub $\tau_R \gg T_0$ będzie to niestałość lub dewiacja krótkookresowa, natomiast oddziaływanie składowych aperiodycznych w badanym odstępie czasu nT_p opisuje się za pomocą funkcji $f_{ar_p}(t)$ jako długookresową niestabilność częstotliwości.

Wzory w rozdz. 4 zostały wyprowadzone przy założeniu idealnej długookresowej i krótkookresowej niestabilności częstotliwości, tzn. przy założeniu

$$f_0 = f_{ar}(t) = \text{const.} \quad (5.7)$$

Opisana powyżej metoda najmniejszych kwadratów pozwala na drodze analitycznej wyeliminować ze zbioru wyników pomiarowych wpływ tych niestabilności, co umożliwia porównanie wyników teoretycznych z eksperymentalnymi w przypadku, gdy warunek (5.7) nie może być spełniony.

5.2. Ogólne zasady pomiarów

Znane są trzy podstawowe metody pomiarów fluktuacyjnej dewiacji częstotliwości:

- a) Metoda amplitudowa, polegająca na określeniu eksperymentalnym szerokości widma wypadkowego $G(\omega)$. Szerokość tego widma dla modelu dyfuzyjnego jest w przybliżeniu równa szerokości widma modulacji kąta. Znając wartość $\delta_k \omega$ można na podstawie zależności (4.17) wyznaczyć współczynnik dyfuzji i pozostałe parametry drgań. Metoda ta pozwala na szybkie i bezpośrednie wyznaczenie wartości $\delta_k \omega$, wymaga jednak stosowania skomplikowanego układu pomiarowego.
- b) Metoda fazowa, polegająca na detekcji fazowej drgań losowych. Parametry drgań wyznacza się ze statystycznej analizy danej realizacji zależności kąta fazowego od czasu, względnie z analizy widma $G_{\Delta\varphi}(\Omega)$ napięcia wyjściowego detektora fazy. Zaletą tej metody jest możliwość synchronizacji źródła drgań badanych, za pomocą pętli fazowej regulacji częstotliwości, przez źródło drgań odniesienia, co pozwala na wyeliminowanie wpływów długookresowych.
- c) Metoda częstotliwościowa, polegająca na detekcji częstotliwościowej drgań badanych względnie na analizie statystycznej zbioru wyników $\overline{f_{pi}}$ pomiarów częstotliwości badanego źródła.

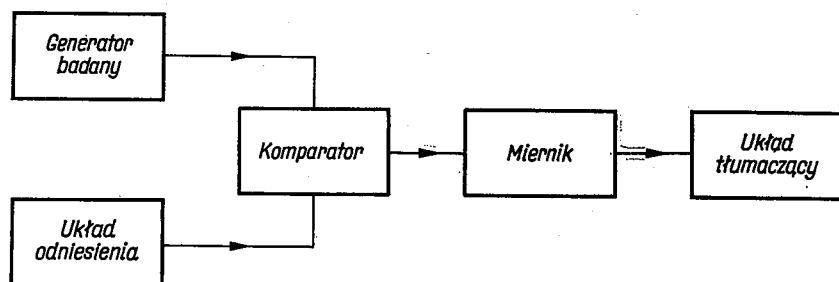
Metody fazowa i częstotliwościowa są metodami pośrednimi. W metodach tych parametry drgań wyznacza się za pomocą żmudnych i skomplikowanych operacji rachunkowych, ze względu jednak na względną pro-

stotę układu pomiarowego i możliwość zautomatyzowania tych operacji są one często stosowane.

Przy realizacji układów pomiarowych opartych na powyższych metodach występują dwie podstawowe trudności, wynikające ze specyfiki zjawisk badanych. Pierwszą trudność stanowi fakt, że w przypadku generatorów częstotliwości, wielkości wyznaczonych na drodze pomiarowej parametrów są bardzo małe. Zazwyczaj wartość tych parametrów leży poniżej progów czułości konwencjonalnych metod pomiaru modulacji amplitudy i kąta czy też analizatorów widma. Drugą trudność stanowi budowa i określenie parametrów źródła częstotliwości wzorcowych odniesienia, szczególnie gdy metoda pomiarowa wymaga zastosowania drgań odniesienia o znacznie lepszych parametrach statystycznych od drgań badanych.

5.3. Podstawowy schemat układu pomiarowego

Zasadniczy ramowy schemat logiczny, wspólny dla wszystkich metod pomiarowych, oparty zresztą na ogólnych zasadach pomiarów wielkości fizycznych, przedstawiony jest na rys. 11.



Rys. 11. Ogólny schemat układu pomiarowego

Badane drganie porównywane jest za pomocą komparatora z odpowiednio wybraną wielkością odniesienia, stanowiącą zazwyczaj wzorzec mierzonego parametru. Układ komparatora spełnia dwie zasadnicze funkcje:

- zamienia (transformuje) wielkość wybranego do pomiaru parametru lub określonej jego funkcji na sygnał pomiarowy dostosowany do charakteru użytego miernika
- wzmacnia poziom tego sygnału pomiarowego w takim stopniu, aby poziom ten znalazł się powyżej progu czułości miernika.

Miernik wyznacza wartość sygnału pomiarowego. Z uwagi na losowy charakter mierzonych wielkości, jak i na zasadę działania miernika, jest to z reguły określony parametr statystyczny lub pojedyncza realizacja wielkości mierzonej. Czas uśredniania może być określony zarówno przez zasadę działania miernika, jak i układ komparatora.

Aby ułatwić interpretację wskazań miernika przez operatora w przypadku pomiaru wielkości losowych, należy przeważnie otrzymane dane pomiarowe przedstawić w formie jak najbardziej przejrzystej. Do tego celu służy układ tłumaczący. Może on np. zarejestrować w określonym układzie współrzędnych wykres funkcji widmowej $G(\omega)$ bądź też dokonać w sposób pełno- lub półautomatyczny analizy statystycznej zarejestrowanej realizacji lub zbioru wartości danego procesu losowego.

Funkcje układu tłumaczącego mogą być w zasadzie wykonane przez operatora, zastosowanie jednak takiego układu znacznie ułatwia i przyspiesza tok pomiarów.

W rzeczywistych układach pomiarowych role poszczególnych elementów nie muszą być w tak wyraźny sposób rozgraniczone. Pewne funkcje mogą być połączone w określonych członach układu, inne mogą się okazać zbędne przy danej metodzie pomiarowej. Tym niemniej przyjęcie powyższego schematu umożliwia systematyczną analizę poszczególnych metod pomiarowych.

5.4. Układ odniesienia

Układ odniesienia w przypadku ogólnym jest źródłem lub zespołem źródeł drgań lub częstotliwości odniesienia (drgań wzorcowych dla danej metody pomiarowej). W pewnych metodach specjalnych wielkość odniesienia może być wielkością uzależnioną w określony sposób od źródła drgań badanych.

Jeżeli w układzie pomiarowym stosuje się pojedyncze źródła drgań odniesienia, to mogą zachodzić trzy przypadki:

- a) Źródło drgań odniesienia jest źródłem częstotliwości wzorcowych o parametrach odpowiednio lepszych od parametrów drgań badanych i w związku z tym może być traktowane jako źródło monochromatyczne. Dla większości generatorów częstotliwości warunki takie spełnia źródło częstotliwości wzorcowych sterowane przez atomowy lub molekularny wzorzec częstotliwości. W metodach częstotliwościowych stosuje się czasami w układzie komparatora detektor częstotliwości z rezonatorem kwarcowym w charakterze układu odniesienia.
- b) Źródło odniesienia charakteryzuje się identycznymi parametrami statystycznymi jak źródło badane. W tym przypadku wartości średnich kwadratowych określonych wielkości losowych obu generatorów są sobie równe i o $\sqrt{2}$ — razy mniejsze od wartości wypadkowej (sumowanie energetyczne).
- c) Jeżeli źródło badane w warunkach eksploatacyjnych stale współpracuje z innym, określonym źródłem drgań i wystarcza wyznaczenie parametrów wypadkowych dla obu źródeł wspólnie (np. przy badaniu łącz radiokomunikacyjnych), skojarzenie obu źródeł

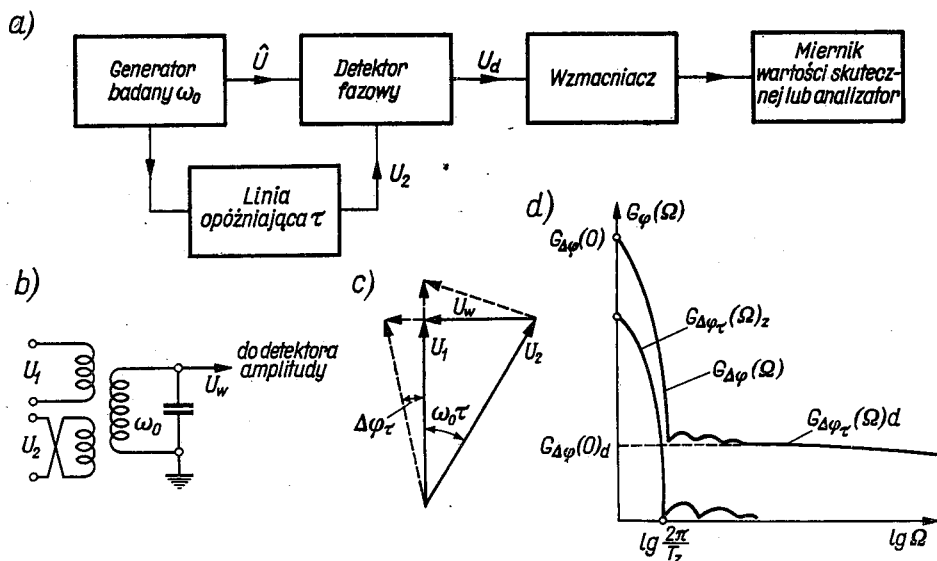
w układzie pomiarowym umożliwi realizację tak sformułowanego programu pomiarowego.

Dysponując zespołem 3 lub więcej niezależnych źródeł o dowolnych parametrach statystycznych, można metodą kolejnych wzajemnych porównań wyznaczyć parametry statystyczne każdego ze źródeł porównawczych.

Stosując w układach pomiarowych niezależne źródła odniesienia należy wyeliminować wpływ krótko i długookresowej niestałości częstotliwości źródeł porównywanych. Można tego dokonać przez odpowiedni dobór warunków pomiarowych lub przez zastosowanie specjalnych metod analitycznych, np. metody najmniejszych kwadratów do opracowania wyników pomiarowych.

Problem niestałości długo i krótkookresowej może być rozwiązany w sposób radykalny przy zastosowaniu uzależnionych od sygnału badanego źródeł sygnału odniesienia. W tym przypadku o skuteczności danej metody decyduje sposób uzależnienia obu sygnałów i wprowadzenie tą drogą zniekształcenia wielkości mierzonej.

Oryginalną koncepcję układu pomiarowego, z uzależnieniem sygnału odniesienia od sygnału badanego, zaproponował Bernstein [5]. Uproszczony schemat układu pomiarowego opartego na tej zasadzie przedstawiono na rys. 12. W układzie tym, opartym na metodzie fazowej, w specjalnym



Rys. 12. Układ pomiarowy metody fazowej z uzależnionym układem odniesienia

- a) układ pomiarowy,
 - b) schemat detektora fazowego,
 - c) wykres wskazowy,
 - d) widmo $G_{\Delta\varphi_\tau}(\Omega)$ na wyjściu detektora fazowego.
- $G_{\Delta\varphi}(\Omega)_z$ — widmo zakłóceń wolnozmiennych,
 $G_{\Delta\varphi_\tau}(\Omega)_d$ — widmo drgań dyfuzyjnych,
 $G_{\Delta\varphi}(\Omega)$ — widmo wypadkowe.

detektorze porównywane są fazy napięć doprowadzonych z generatora badanego bezpośrednio i poprzez linię opóźniającą. Czas opóźnienia wprowadzonego przez linię τ_B jest rzędu kilku μ sek. Działanie detektora fazowego (rys. 12b) polegało na przeprowadzeniu detekcji amplitudowej napięcia $\bar{U}_w$ będącego różnicą geometryczną obu napięć porównywanych $\bar{U}_w = \bar{U}_1 - \bar{U}_2$. Średnia wartość modułu napięcia $|\bar{U}_w|$ będzie funkcją średniego przesunięcia fazy $\omega_0 \tau_B$ pomiędzy napięciami $\bar{U}_1$ i $\bar{U}_2$, czyli dyspersja amplitudy U_w będzie proporcjonalna do dyspersji kąta fazowego drgań badanych za odstęp czasu τ_B . Mierzac wartość skuteczną lub kwadrat składowej zmiennej napięcia wyjściowego z detektora, U_d , można wyznaczyć współczynnik dyfuzji z zależności:

$$\sigma^2(\Delta\varphi_t) = [\varphi(t + \tau_B) - \varphi(t)]^2 = 2D\tau_B = K^2\sigma^2(U_d), \quad (5.8)$$

gdzie K — nachylenie charakterystyki detektora fazowego. W powyższym układzie, jak to poglądowo przedstawia rys. 12c, można przez odpowiedni dobór amplitud U_1 i U_2 zredukować do minimum wpływ modulacji amplitudy tych napięć na zmiany amplitudy U_w , a więc i na wynik pomiaru.

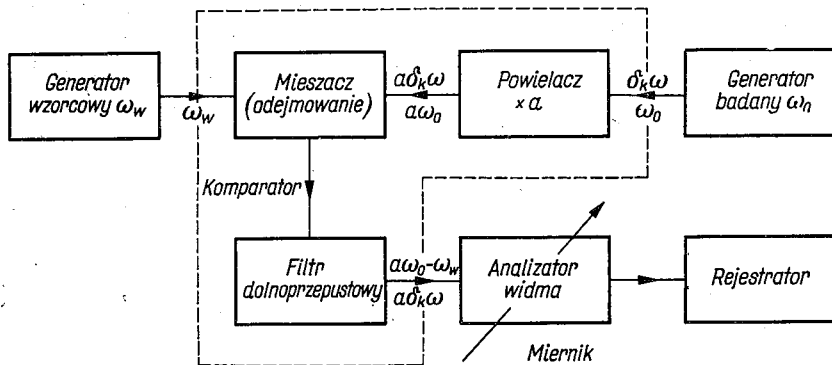
Podstawową zaletą opisanego wyżej układu jest możliwość realizowania pomiarów przy bardzo małych i regulowanych wartościach odstępu czasowego τ_B .

Jako pojedynczy niezależny układ odniesienia może być zastosowany wzorzec molekularny. Wzorce molekularne charakteryzują się pośród innych generatorów wzorcowych najmniejszą względną szerokością widma $\frac{\delta_k \omega}{\omega_0}$. Z uwagi na wysoką wartość częstotliwości nominalnej wzorce atomowe stosuje się do kontroli częstotliwości i fazy wzorców wtórnych za pomocą znanych układów syntezy częstotliwości opartych o metody fazowej regulacji częstotliwości. Częstotliwość nominalną wzorców wtórnych wybiera się przeważnie znacznie wyższą od częstotliwości nominalnej generatora badanego, aby przez powielenie tej ostatniej zwiększyć szerokość badanego widma modulacji kąta. Dysponując wzorcem tej klasy można go użyć również do heterodynowej detekcji drgań badanych lub do przesunięcia widma badanego w zakres pracy konwencjonalnego analizatora.

Schemat blokowy układu pomiarowego opartego na metodzie amplitudowej, przy pojedynczym niezależnym źródle drgań wzorcowych przedstawiono na rys. 13. Zastosowany w tym układzie analizator powinien mieć szerokość wstęgi znacznie mniejszą od powielonej a razy szerokości widma drgań badanych $a\delta_k\omega$.

Innym rodzajem układu odniesienia jest dyskryminator fazy pracujący w układzie detektora częstotliwości drgań badanych. W celu uzyskania dużej czułości i stabilności dyskryminatora wykorzystuje się do jego budowy rezonatory kwarcowe. Odnośnie niestabilności częstotli-

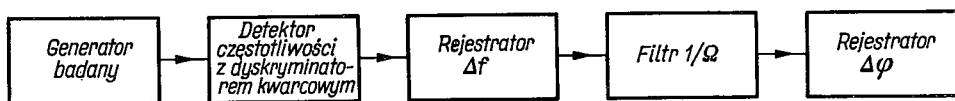
wości rezonansowej takiego dyskryminatora muszą być spełnione wszelkie wymagania dotyczące omówionej poprzednio częstotliwości wzorcowej. Z uwagi jednak na to, że dyskryminator jest biernym układem li-



Rys. 13. Zastosowanie pojedynczego, niezależnego źródła drgań wzorcowych do analizy widma

niowym, szum w jego obwodach dodaje się tylko liniowo do widma sygnału badanego. Układ pomiarowy z detektorem częstotliwości przedstawiono na rys. 14.

Jeżeli do wyjścia detektora dołączyć filtr o charakterystyce przeniesienia odwrotnie proporcjonalnej do częstotliwości, to na wyjściu takiego układu otrzymuje się napięcie proporcjonalne do dewiacji fazy. W przypadku gdy do wejścia detektora doprowadzić napięcie stanowiące rezultat przemiany częstotliwości generatora badanego i oddzielnego wzorca,



Rys. 14. Układ pomiarowy z dyskryminatorem kwarcowym

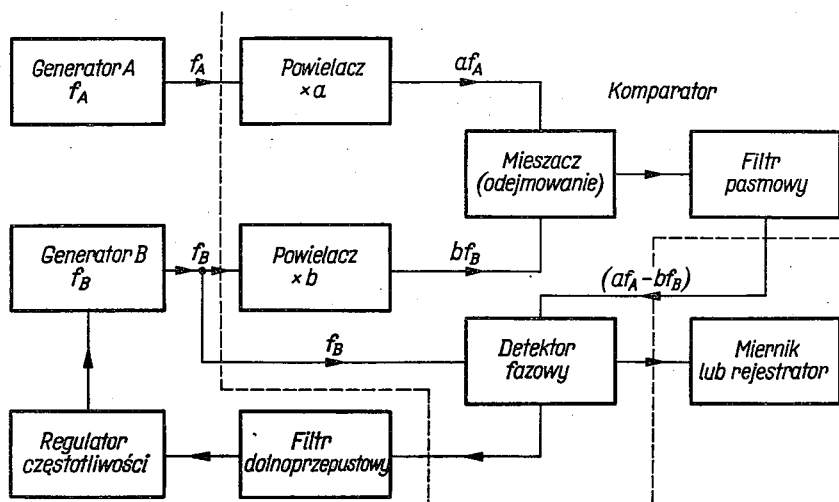
układ z rys. 14 będzie pełnił rolę fazowego miernika sygnału pomiarowego z rys. 11. W tym przypadku wymagania dotyczące stabilności częstotliwości rezonansowej dyskryminatora zostają znacznie zmniejszone.

Gdy pożądaną jest wyeliminowanie krótko i długookresowej niestabilności częstotliwości źródła badanego i wzorcowego, jak również przy pomiarach parametrów łącz radiokomunikacyjnych z detekcją koherentną, należy oba źródła częstotliwości uzależnić od siebie za pomocą odpowiedniej pętli fazowej regulacji częstotliwości. Ponieważ w układzie automatyki pracuje detektor fazowy, może być on jednocześnie wykorzystany jako element komparatora układu pomiarowego opartego na metodzie fazowej. Przykład takiego rozwiązania w zastosowaniu do łącza podano na rys. 15. Generatory A i B oznaczają generator wzbudzający nadajnika

i heterodynę odbiornika łączy. Napięcie wyjściowe detektora fazowego jest w tym układzie proporcjonalne do względnych zmian kąta fazowego obu generatorów, o ile będzie spełniony warunek synchronizmu w postaci

$$|af_A - bf_B| = \frac{m}{n} f_B, \quad (5.9)$$

gdzie f_A i f_B oznaczają odpowiednio częstotliwości generatorów A i B a m i n liczby naturalne.



Rys. 15. Układ do pomiaru względnych fluktuacji fazy generatorów współpracujących w łańcu o detekcji koherentnej

Przy interpretacji wyników w układach pomiarowych z fazową regulacją częstotliwości należy zwrócić uwagę na działanie samej pętli automatyki. Pętla ta działa jak filtr górnoprzepustowy dla składowych modulujących fazę sygnału badanego. Częstotliwość graniczna tego filtra równa się szerokości wstęgi wypadkowej całej pętli automatyki [16]. Zastosowanie układu z rys. 15 do porównania generatora badanego z generatorem wzorcowym wymaga tylko odpowiedniego doboru krotności powielenia a i b pod kątem szerokości widm tych generatorów.

Jeżeli nie dysponuje się wzorcem odpowiedniej klasy, można w charakterze układu odniesienia zastosować drugi niezależny układ generatora tego samego typu co układ badany. W tym przypadku dyspersja wypadkowa częstotliwości lub fazy równa się sumie dyspersji poszczególnych składowych:

$$\sigma_w^2 = \sigma_1^2 + \sigma_2^2. \quad (5.10)$$

Jeżeli można przyjąć, że $\sigma_1^2 = \sigma_2^2$, to:

$$\sigma_1^2 = \frac{\sigma_w^2}{2}. \quad (5.11)$$

Przy pomiarach metodami amplitudowymi można w tym przypadku przyjmować, że szerokość widma badanego równa się połowie szerokości widma wypadkowego.

W przypadku kolejnego porównywania trzech różnych niezależnych źródeł otrzymuje się 3 wartości dyspersji wypadkowych, oznaczonych jako σ_{XA}^2 , σ_{XB}^2 , σ_{AB}^2 . Dyspersję parametrów źródła badanego σ_X^2 można wyznaczyć ze wzoru [7]:

$$\sigma_X^2 = \frac{\sigma_{XA}^2 + \sigma_{XB}^2 - \sigma_{AB}^2}{2}. \quad (5.12)$$

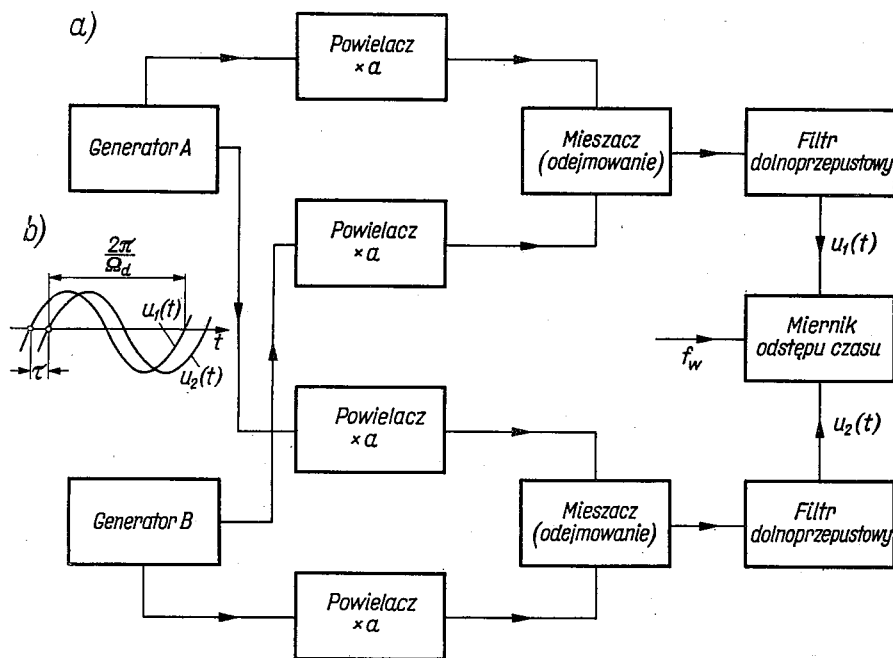
5.5. Komparator

Jak już wspomniano w rozdz. 5.3, układ komparatora spełnia dwie zasadnicze funkcje: zamienia wielkość wybranego do pomiaru parametru na sygnał pomiarowy dostosowany do charakteru zastosowanego miernika oraz odpowiednio do progu czułości tego miernika podnosi poziom sygnału pomiarowego. Ponieważ najczęściej używa się mierników przystosowanych do pomiarów napięć, komparator zawiera różnego typu detektory amplitudy fazy lub częstotliwości. Podniesienie poziomu sygnału pomiarowego przed detekcją realizuje się za pomocą powielenia i przemiany częstotliwości sygnału badanego. Do tego samego celu po detekcji stosuje się odpowiednie wzmacniacze napięcia lub mocy. Każda z powyższych operacji może być źródłem znanych zniekształceń sygnału badanego, co należy uwzględnić przy analizie dokładności i czułości metody pomiarowej. Pewne wnioski ogólne dotyczące zniekształceń amplitudy i kąta drgań powielanych, wprowadzanych przez układy powielaczy, można wyciągnąć z rozważań z rozdz. 4. Ponieważ składanie napięcia sygnału powielanego i napięcia zakłóceń szumowych skojarzonych z układem powielacza ma miejsce w liniowym obwodzie rezonansowym, to pomimo nieliniowej pracy samego stopnia powielającego, można stosować w pierwszym przybliżeniu analogię do układu wzmacniacza. Nieliniowe warunki pracy powielaczy będą przyczyną powielania widma modulacji kąta sygnału powielanego i pochodzącego z poprzednich stopni powielania oraz mogą być przyczyną wzrostu pasożytniczej głębokości modulacji amplitudy, o ile powielacze pracują bez prądu siatki. Praca w klasie C może stać się również przyczyną wprowadzenia dodatkowej modulacji kąta przy zmianach kąta przepływu. Aby zmniejszyć te efekty, stosuje się pracę przy dużych kątach przepływu, a więc przy niewielkiej krotności powielenia na jeden stopień powielacza.

Dodatkowym źródłem zakłóceń wprowadzonych przez powielacze do układu pomiarowego jest pasożytnicza modulacja napięciem sieci zasilającej, głównie przez obwody żarzenia. Z tych względów zasilanie ka-

tod lamp powielaczy odbywa się z reguły, przynajmniej w pierwszych stopniach, prądem stałym.

W celu eksperymentalnego określenia zniekształceń kąta wprowadzonych do toru pomiarowego przez powielacze układu komparatora można zastosować układ pomiarowy przedstawiony na rys. 16. W układzie tym zastosowano cztery powielacze identyczne z powielaczem ba-



Rys. 16. Układ do pomiaru własnych fluktuacji fazy powielaczy pomiarowych

a) układ pomiarowy,

b) przebiegi napięć $u_1(t)$ i $u_2(t)$

danym. Częstotliwości generatorów A i B są zbliżone do siebie, tak że na wejście miernika doprowadzone są dwa sygnały o częstotliwości dudnień Ω_d .

Wykonując serię pomiarów odstępu czasu pomiędzy chwilami odpowiadającymi takim samym fazom sygnałów różnicowych (rys. 16b), można na tej podstawie wyznaczyć dyspersję kąta przesunięcia fazowego pomiędzy tymi sygnałami, za czas odpowiadający średniemu przesunięciu fazowemu $\overline{\Delta\varphi} = \Omega_d \tau$. Dzięki podwójnej przemianie częstotliwości fluktuacje kąta drgań generowanych przez oba generatory nie wpływają na wynik pomiaru. Dyspersja fazy wprowadzona przez pojedynczy powielacz równa się 1/4 dyspersji pomierzonej.

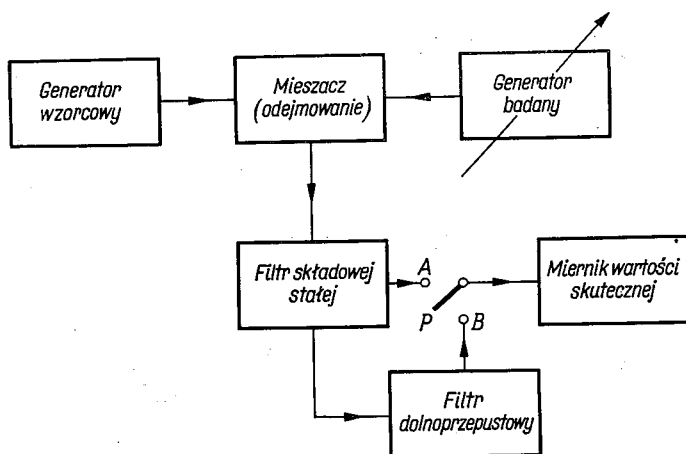
Podobne ograniczenia jak w przypadku powielaczy stosuje się do układów przemiany częstotliwości. Układy przemiany z lampami wielosiatkowymi w zasadzie nie wprowadzają zniekształceń amplitudy — każ-

dy obwód pracuje liniowo, wprowadzają jednak dodatkowe źródło szumów pochodzących od rozptywu prądów pomiędzy poszczególne siatki. Zastosowanie wysokiej krotności powielenia częstotliwości badanej utrudnia wykorzystanie przemiany wielosiatkowej z uwagi na stosunkowo wysoką częstotliwość pracy.

Do pomiarów widma modulacji amplitudy stosuje się detektory amplitudowe.

Oprócz detektorów biernych można w układach komparatora zastosować detekcję heterodynową jako czulszy rodzaj detekcji.

Przykład zastosowania detekcji heterodynowej do uproszczonej metody amplitudowej przedstawiono na rys. 17. W układzie tym mierzy się stosunek mocy sygnału do mocy wstęp bocznych oraz określa szerokość

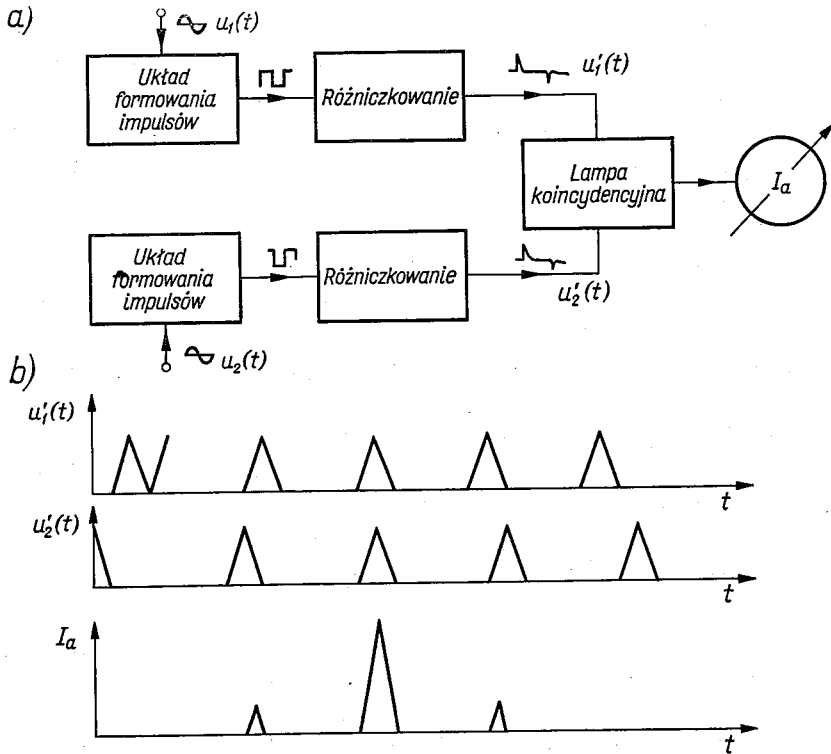


Rys. 17. Pomiar szerokości widma z detekcją heterodynową

kość widma wypadkowego [26]. Metoda powyższa wymaga, aby jeden z porównywanych generatorów, badany lub wzorcowy, był w niewielkich granicach przestrajany. W pierwszej fazie pomiarowej przełącznik P ustawiony jest w pozycji A . Na mierniku wartości skutecznej dokonuje się dwóch odczytów, przy niewielkim rozstrojeniu i przy zerze dudnień obu generatorów. Ze stosunku tych odczytów można określić stosunek mocy sygnału do mocy wstęp bocznych. W drugiej fazie pomiaru przy generatorach dostrojonych do zera dudnień ich częstotliwości średnich, w pozycji B przełącznika P , wyznacza się skuteczną szerokość widma. Szerokość tę można obliczyć ze stałej czasu filtru dolnoprzepustowego, którą reguluje się w taki sposób, aby wskazania woltomierza były o $\sqrt{2}$ razy mniejsze od wskazań odpowiadających pomiarowi pełnej mocy wstęp bocznych w fazie pierwszej.

Jeżeli chodzi o detektory kąta stosowane w układach komparatorów, to wymagania stawiane im są analogiczne jak w przypadku detektorów amplitudy. Podstawową zaletą takich układów jest duża czułość przy

prostej budowie i małych zniekształceniach. Oprócz omówionych poprzednio układów z rys. 12 i 14 warto w tym miejscu wspomnieć o bardzo czułej metodzie detekcji fazowej opartej na koincydencji impulsów [3]. Zasada działania takiego detektora jest przedstawiona na rys. 18. Oba sygnały porównywane podlegają operacji przekształcenia na bardzo wąskie impulsy. Impulsy te doprowadzone są do odpowiednich siatek lampy koincydencyjnej. Warunki pracy tej lampy dobrane są w taki sposób, że prąd anodowy płynie tylko w tych odstępach czasu, w których napięcia na obu siatkach sterujących jest dodatnie. Na rys. 18b przedsta-



Rys. 18. Koincydencyjny detektor fazowy

a) schemat układu,

b) przykłady koincydencji.

wiono poglądowo kilka różnych stanów koincydencji. Jak wynika z tego rysunku, wartość średnia prądu anodowego lampy koincydencyjnej będzie funkcją kąta różnicy faz obu sygnałów porównywanych. Dzięki dużej stromości zbocza impulsów można zapewnić dużą czułość detekcji. Jednakże operacja formowania impulsów, jak każda operacja nieliniowa w czwórniku aktywnym, jest źródłem dodatkowych zniekształceń sygnałów porównywanych. Maksymalna dewiacja fazy, możliwa do wyzna-

czenia przez wyżej opisaną metodę, zależy od czasu trwania impulsów koincydencyjnych τ_i i wyraża się zależnością:

$$\Delta\varphi \ll \omega_0 \tau_i, \quad (5.13)$$

gdzie ω_0 jest średnią pulsacją sygnałów porównywanych. Ponadto, ponieważ metoda powyższa oparta jest na dyskretnym próbkowaniu w okresie $T_p = \frac{2\pi}{\omega_0}$, w wyniku pomiaru otrzymuje się wartość różnicy kątów fazowych za odstęp czasu równy okresowi napięcia wejściowego.

5.6. Miernik

Rodzaj i czułość miernika zastosowanego w danym układzie pomiarowym zależą w ścisły sposób od całej metody pomiarowej.

Z uwagi na losowy charakter wielkości mierzonych, jako mierniki w omawianych układach pomiarowych stosuje się cztery podstawowe typy przyrządów.

- a) Analizatory widma
- b) Mierniki autokorekcji
- c) Mierniki wartości średniego kwadratu lub wartości skutecznej
- d) Rejestratory realizacji procesu losowego.

W układach pomiarowych opartych na metodzie amplitudowej do określania parametrów widma stosuje się analizatory drgań elektrycznych. Analizator jest woltomierzem selektywnym o bardzo wąskim pasmie. Ponieważ wstęga analizatora ze względów technicznych nie może być węższa od kilku Hz, bezpośrednia (tzn. bez komparatora) analiza drgań, o szerokości widma wyrażonych w ułamkach Hz, jest niemożliwa, stąd też analizator w niniejszej pracy jest traktowany wyłącznie jako miernik współpracujący z układem komparatora.

Zasadniczo rozróżnia się dwie metody analizy widma: jednoczesną i kolejną. W metodzie jednoczesnej sygnał badany oddziaływać może jednocześnie na zespół rezonatorów o częstotliwościach rezonansowych rozłożonych równomiernie w analizowanym pasmie. Przy kolejnej analizie pojedynczy rezonator analizuje kolejno poszczególne fragmenty widma badanego. Wzajemnego przesuwania pasma rezonatora względem widma badanego można dokonać za pomocą przestrajania rezonatora lub heterodynowania sygnału badanego. Analiza jednoczesna wymaga znacznie krótszego czasu niż kolejna, tym niemniej, z uwagi na większą elastyczność i ekonomię środków, stosuje się przeważnie analizatory kolejne z przemianą częstotliwości i wąskim pasmowym filtrem kwarcowym w torze częstotliwości pośredniej. W bardziej rozbudowanych rozwiązaniach analizator dostarcza jednocześnie sygnału napięciowego, liniowo zmiennego z częstotliwością analizy, do zasilania pisaka osi częstotliwości urządzenia zapisującego funkcję widmową. Aby przemiana częstotliwości

nie wprowadzała dodatkowych zniekształceń widma badanego, często stosuje się układ syntezy częstotliwości jako źródło napięcia heterodyny o bardzo małej niestabilności [1].

Szerokość pasma analizatora jest podstawowym parametrem miernika. Określa ona czas analizy poszczególnych fragmentów widma i bezwładność, czyli stałą czasu zastosowanego w analizatorze miernika napięciowego. Szerokość pasma analizatora określa także interpretację uzyskanych wyników [9]. Oddziaływanie składowych częstotliwości modulujących kąt sygnału badanego o częstotliwościach wyższych od szerokości pasma analizatora, zostanie określone przez analizator jako występowanie oddzielnych wstęg bocznych. Oddziaływanie składowych o niższych częstotliwościach zostanie określone przez analizator jako zmiana częstotliwości drgań podstawowych. Rozważania te można zresztą uogólnić na cały układ pomiarowy biorąc pod uwagę wypadkową szerokość wstęgi układu.

Analizator może być zastosowany do analizy przebiegu zmodulowanego lub też do analizy produktów detekcji tego przebiegu. Ten drugi rodzaj zastosowania wykorzystany jest w układzie z rys. 12. Dzięki analizowaniu widma sygnału wyjściowego z detektora można było wydzieleć efekty związane z szumem białym od efektów fluktuacji małej częstotliwości. Na rys. 12d przedstawiono widmo wypadkowe z wydzieleniem obu składowych: $G_{\Delta\varphi_\tau}(\Omega)_d$ oznacza widmo dyfuzyjne, a $G_{\Delta\varphi}(\Omega)_z$ oznacza widmo pochodzące od zakłóceń małej częstotliwości. Ponieważ $\frac{1}{\tau_B}$ było znacznie większe od zakresu analizowanych częstotliwości, to funkcja $G_{\Delta\varphi_\tau}(\Omega)_d \approx G_{\Delta\varphi_\tau}(0)_d = \text{const.}$ Stąd na podstawie zależności (4.22) można było określić współczynnik dyfuzji. Stosując zamiast analizatora miernik wartości skutecznej otrzymuje się w tym przypadku dyspersję wypadkowych fluktuacji fazy.

Jako mierniki współpracujące z układem komparatora stosuje się często również mierniki autokorelacji. Są to przyrządy określające funkcję autokorelacji $R_x(\tau)$ sygnału wejściowego $u_x(t)$, czyli realizujące następującą operację analityczną:

$$R_x(\tau) = \overline{u_x(t) u_x(t + \tau)}. \quad (5.14)$$

Określenie funkcji autokorelacji pozwala na wyznaczenie funkcji widmowej na podstawie przekształcenia Fouriera.

Zastosowanie bezpośrednio woltomierzy wartości skutecznej, jako mierników wartości średniego kwadratu sygnału pomiarowego z komparatora lub mierników pomocniczych, wymaga w układach pomiarowych dostosowania stałej czasu przyrządu do wypadkowej szerokości wstęgi całego toru pomiarowego. Jeżeli stała czasu miernika nie będzie znacznie większa od odwrotności szerokości wstęgi, wskazówka przyrządu będzie wykonywać bezładne wahania wokół wartości średniej, co znacznie

utrudnia proces pomiarowy. Jako mierniki mogą być stosowane z powodzeniem przyrządy ciepłne odznaczające się z natury dużą stałą czasu lub konwencjonalne mierniki mocy.

W układach pomiarowych metody częstotliwościowej stosuje się często mierniki okresu lub częstotliwości sygnału wyjściowego. Spośród znanych mierników do celów analizy statystycznej częstotliwości badanej najbardziej nadają się przyrządy zliczające. Podstawowymi zaletami zliczających mierników czasu i częstotliwości w tych zastosowaniach są: ściśle określony czas uśredniania τ_p , równy czasowi zliczania dla pojedynczego pomiaru, oraz łatwość zautomatyzowania i rejestracji wyników pomiarowych. Szczegółową analizę błędu pomiarowego czasomierzy tego typu ze szczególnym uwzględnieniem wpływu zakłóceń szumowych można znaleźć w literaturze [29].

Osobnym zagadnieniem jest cechowanie miernika w układzie pomiarowym. Zagadnienie to jest bardzo złożone wobec dużej liczby typów stosowanych mierników. W przypadku skalowania torów pomiarowych zakończonych miernikiem wartości skutecznej można posługiwać się bezpośrednio generatorami szumów. Zasadniczym elementem takiego generatora jest zwykle dioda wolframowa pracująca w warunkach nasycenia, której prąd szumów może być dokładnie wyznaczony analitycznie i którego natężenie może być w prosty sposób regulowane przez zmianę zasilania diody.

Zastosowanie w układzie pomiarowym analizatorów widma lub mierników autokorelacji dostarcza znacznie więcej informacji o procesie badanym niż określenie średniej wartości. Funkcja autokorelacji i funkcja widmowa związane są ze sobą przez przekształcenie Fouriera, a więc pod względem informacji są sobie równoważne. Znajomość tych funkcji pozwala na określenie wpływu poszczególnych składowych oddziałujących na drgania badane. Ma to doniosłe znaczenie praktyczne, gdyż teoretyczny model drgań dyfuzyjnych uwzględnia tylko oddziaływanie szumów białych. Przez odpowiednią analizę statystyczną sygnału wyjściowego można wyodrębnić wpływ szumów o ograniczonej wstędze, takich jak np. efekt migotania czy szumy strukturalne w opornikach, a nawet wyodrębnić wpływ zakłóceń periodycznych np. pochodzących od sieci zasilającej.

Rejestracja danej realizacji procesu losowego może być zastosowana we wszystkich trzech metodach pomiarowych. Rejestracja taka może być prowadzona w sposób ciągły lub dyskretny. Znane metody analityczne umożliwiają na podstawie takiego zapisu określenie parametrów statystycznych badanego procesu losowego.

W metodach częstotliwościowych do rejestracji dyskretniej szczególnie nadają się zliczające mierniki czasu i częstotliwości, a to z uwagi na ich wspomnianą wyżej dużą czułość i łatwość rejestracji wyników pomiarowych.

5.7. Układ tłumaczący

Miernik układu pomiarowego dostarcza informacji o badanym przebiegu w różnych postaciach, zależnych od konkretnego zastosowania. Jeżeli sam miernik nie rejestruje wielkości mierzonych podstawowych, zadaniem układu tłumaczącego jest rejestracja wyników pomiarowych w celu umożliwienia ich dalszego opracowania. Ta część układu tłumaczącego spełnia rolę „pamięci”. Dalsze funkcje układu tłumaczącego zależą od stopnia zautomatyzowania układu pomiarowego i związanego z tym udziału operatora w opracowaniu wyników w postaci ostatecznej. Można tu tylko pokrótce wymienić najczęściej stosowane operacje wykonywane przez układ tłumaczący. Będą to operacje określane przez przekształcenie Fouriera dla przejścia od funkcji autokorelacji do funkcji widmowej lub operacje związane z analizą statystyczną zarejestrowanej realizacji, polegające na wyznaczeniu dyspersji lub rozkładu prawdopodobieństwa. Do funkcji układu tłumaczącego może należeć również eliminacja wpływu krótko- i długookresowej niestałości częstotliwości źródeł badanych metodą najmniejszych kwadratów.

Do wymienionych wyżej celów wykorzystuje się różnego typu maszyny matematyczne analogowe i cyfrowe wraz z wyposażeniem pomocniczym. Zastosowanie tych maszyn znacznie upraszcza opracowanie wyników pomiarowych.

6. WYNIKI NIEKTÓRYCH PRAC EKSPERYMENTALNYCH

Aby umożliwić konfrontację wyników z obliczeniami analitycznymi, zamieszczono poniżej dwa przykłady obliczeniowe dla modelu dyfuzyjnego zaczerpnięte z pracy Edsona [10]. W przykładach tych uwzględniono tylko oddziaływanie szumu cieplnego na dwa charakterystyczne układy generacyjne: generator mikrofalowy z klistronem i generator kwarcowy. Parametry obu generatorów zamieszczono w tablicy 1. Wyniki obliczeń zamieszczono w tablicy 2.

Cutler [9] wykonał pomiary dla specjalnie zaprojektowanego pod kątem małej niestabilności fluktuacyjnej generatora kwarcowego. Rezonator kwarcowy 1 MHz cięcia AT o dobroci $2 \cdot 10^6$ i mocy traconej $2 \mu\text{W}$ pracował w termostacie o dokładności $\pm 0,01^\circ\text{C}$. W układzie generacyjnym zastosowano obwód automatycznej regulacji amplitudy. Częstotliwość drgań jest powielana do 5 MHz, a następnie sygnał wyjściowy z powielacza zostaje odfiltrowany w specjalnym filtrze kwarcowym o szerokości wstęgi 50 Hz. W wyniku otrzymano następujące wartości:

$$\frac{\sigma(\bar{f}_\tau)}{f_0} = 3 \cdot 10^{-11}; \quad \text{przy} \quad \tau = 2 \text{ sek}, \quad \frac{\delta_k \omega}{\omega_0} = 6 \cdot 10^{-10}.$$

Bersztejn w swoich pracach eksperymentalnych [5] badał generator LC o częstotliwości nominalnej 18 MHz. Pojemność obwodu o dobroci 80 wynosiła 30 pF. Pomierzony współczynnik dyfuzji wynosił średnio $2,5 \cdot 10^{-6}$, co odpowiada $\delta_k \omega = 5 \cdot 10^{-6}$ [rad/sek].

Tablica 1

**Parametry generatorów do przykładów
obliczeniowych wg Edsona**

Parametr	Jednostka	Generator klitronowy	Generator kwarcowy
T_k	°K	$2 \cdot 10^4$	$4 \cdot 10^2$
C	pF	2	$40 \cdot 10^6$
R	kΩ	5	5
a	—	5	5
f_0	MHz	$16 \cdot 10^3$	1,6
Q	—	10^3	$2 \cdot 10^6$
$P = \frac{A_0^2}{2R}$	mW	50	$10 \cdot 10^{-3}$

Barnes [1] w swych pracach podaje widmo generatora wzorcowego sterowanego za pomocą masera amoniakalnego. Względna szerokość tego widma jest rzędu $3 \cdot 10^{-10}$.

Tablica 2

Wyniki obliczeń wg Edsona

Parametr	Jednostka	Generator klitronowy	Generator kwarcowy
$\frac{\sigma(\bar{f}_\tau)}{f_0}; \tau=1 \text{ sek}$	—	$1,7 \cdot 10^{-12}$	$8,3 \cdot 10^{-15}$
$\frac{\delta_a \omega}{\omega_0}$	—	$5 \cdot 10^{-3}$	$2,5 \cdot 10^{-6}$
$\frac{\delta_k \omega}{\omega_0}$	—	$2,8 \cdot 10^{-13}$	$7 \cdot 10^{-22}$

Strandberg [26] za pomocą układu pomiarowego z rys. 17 wyznaczył względną szerokość widma wypadkowego generatora klitronowego w paśmie x jako 10^{-7} . Obliczona na tej podstawie maksymalna względna dewiacja częstotliwości wynosiła $3 \cdot 10^{-9}$.

Tanzman mierzył standartową dewiację częstotliwości całego zespołu generatorów kwarcowych o częstotliwościach nominalnych zawartych

w przedziale 100 kHz ÷ 5 MHz. Jako układ odniesienia stosował generator z pierścieniowym rezonatorem Essena. Uzyskane wyniki przedstawiono w tablicy 3 [29].

Przytoczone w pracy Tanzmana wykresy zależności standartowej dewiacji częstotliwości od czasu pomiaru τ nie spełniają zależności (4.21).

Tablica 3
Wyniki pomiarów standartowej dewiacji częstotliwości
średniej wg Tanzmana

τ	sek	$8 \cdot 10^{-3}$	0,08	0,8	8,0
$\frac{3\sigma(f_\tau)}{f_0}$	$\times 10^8$	1,3 ÷ 3,4	0,3 ÷ 1,0	0,04 ÷ 0,2	0,03 ÷ 0,07

Jak wynika z powyższego, zresztą bardzo pogładowego, zestawienia danych teoretycznych z eksperymentalnymi, istnieje między nimi poważna rozbieżność.

7. WNIOSKI KOŃCOWE

Pomimo ogromnego postępu zarówno w dziedzinie teoretycznej¹⁾, jak i eksperymentalnej, nie można w chwili obecnej uznać stanu znajomości zjawisk fluktuacyjnych w układach generacyjnych za zadowalający. Ogromna różnorodność zjawisk oraz niezmiernie małe rzędy wielkości, jakie występują w tej dziedzinie, są tego stanu rzeczy najważniejszą przyczyną.

Z przytoczonych powyżej rozważań wynika niedwuznacznie brak jakiegokolwiek korelacji pomiędzy teorią a eksperymentem. Brak również, poza pracami Berszteina, poważniejszej próby uzupełnienia tej luki, a i w tych pracach zgodność co do rzędu wielkości była uznawana za zadowalającą. Jako główną przyczynę takiego stanu rzeczy należy uznać trudność w odtworzeniu w układzie eksperymentalnym warunków zgodnych z założeniami danego modelu teoretycznego, a można przypuszczać, że przy obecnym stanie wiedzy nie da się stworzyć takiego modelu, który by przynajmniej w dostatecznym przybliżeniu obejmował wszystkie podstawowe zjawiska fizyczne.

Dodatkową przyczyną utrudniającą korelację wyników teoretycznych i eksperymentalnych są duże błędy pomiarowe. Błędy te spowodowane są niskim poziomem wielkości mierzonych, przeważnie tego samego rzędu co próg czułości danego układu pomiarowego.

¹⁾ Warto w tym miejscu wspomnieć o leżących odłogiem takich tematach, jak źródła szumów w rezonatorach kwarcowych, w specjalnych lampach mikrofalowych, w elementach półprzewodnikowych itp.

Przy tworzeniu określonego modelu drgań losowych przyjmuje się, że losowy sygnał zakłócający działa na źródło monochromatyczne. W takim ujęciu wszelkie założenia pomocnicze określające dany model sprowadzają się do charakterystyki tego sygnału. Przede wszystkim rozważa się oddziaływanie pojedynczego, określonego sygnału zakłócającego — eliminując wpływ innych zakłóceń. Ponadto przyjmuje się, że poziom wybranego sygnału zakłócającego jest znacznie mniejszy od poziomu drgań monochromatycznych źródła badanego.

W związku z powyższym, aby umożliwić eksperymentalną weryfikację opracowań teoretycznych, wydaje się celowe przeprowadzenie badań eksperymentalnych nad specjalną makietą układu generacyjnego, która pozwoliłaby na ograniczenie skutków wyżej wspomnianych przyczyn. Z zestawienia teoretycznych założeń upraszczających wynika bezpośrednio sposób realizacji takiego eksperymentu. Ponieważ zakłóceń losowych, wywołanych przez różnego rodzaju szumy, jako zakłóceń wewnętrznych — tzn. skojarzonych bezpośrednio ze strukturą źródła drgań — nie można wyeliminować bez naruszenia w sposób istotny warunków pracy generatora, należy w makiecie podnieść sztucznie poziom wybranego rodzaju szumów. Poziom tych szumów należy wybrać w taki sposób, aby z jednej strony wyeliminować praktycznie wpływ zakłóceń innego rodzaju, a z drugiej strony, aby spełnić podstawowe założenia teoretyczne badanego modelu drgań. Rozwiązanie takie ułatwi ponadto zrealizowanie samych pomiarów zwiększając poziom wielkości mierzonych.

Opracowując program badań eksperymentalnych należy wybrać do sprawdzenia najbardziej rozpowszechniony i najpełniej opracowany pod względem teoretycznym model drgań losowych. Modelem takim jest niewątpliwie model dyfuzyjny. Inne wspomniane poprzednio modele bądź opisują tylko pewne fragmenty zjawisk towarzyszących drganiom losowym, bądź też znajdują się dopiero w fazie opracowywania.

Dyfuzyjny model drgań losowych opisuje oddziaływanie szumu białego na monochromatyczne źródło drgań. W rzeczywistości generatorze występują dwa rodzaje szumu białego: szum cieplny, którego źródłem jest oporność strat obwodu drgań, oraz szum śrutowy wprowadzany przez układ pobudzania. Szum śrutowy ma charakter anizotropowy; intensywność jego oddziaływania zależy od chwilowej wartości prądu anodowego, a więc nie jest stała w ciągu okresu drgań generowanych. W związku z tym do makiety eksperymentalnej należy wprowadzić szum cieplny, tzn. należy zwiększyć poziom szumów w oporności strat obwodu drgań. Pociąga to za sobą ograniczenie maksymalnej częstotliwości pracy makiety generatora do wartości kilkuset kiloherców — tylko bowiem w takim zakresie można zrealizować obwód drgań o dobroci rzędu 100 przy oporności strat rzędu kilkudziesięciu omów. Zbyt mała oporność strat utrudnia ekonomiczne wprowadzenie zakłóceń do obwodu.

Osobne zagadnienie stanowi wybór poziomu wprowadzonych zakłóceń. Średniokwadratowa wartość napięcia szumów cieplnych na obwodzie drgań dla ww. zakresu częstotliwości jest rzędu $1 \mu V$, a odpowiednia wartość dla szumów śrutowych o rząd wielkości większa. Dążąc do polepszenia warunków pracy aparatury pomiarowej oraz biorąc pod uwagę, że odpowiednie wartości napięć dla innych sygnałów zakłócających, wywołanych np. przez efekt migotania czy przydźwięk sieci zasilającej, są trudne do oszacowania, należy w makiecie zrealizować możliwie duży poziom szumów cieplnych. Maksymalna wartość napięcia tych szumów ograniczona jest dysponowaną mocą źródła szumów oraz założeniami związanymi z wybranym modelem drgań. Konwencjonalne generatory szumów pozwalają na podniesienie mocy szumów o 16 dB. Ze względu na stabilność układu makiety należy unikać stosowania zbyt dużego wzmocnienia w torze sygnału zakłócającego. Założenia teoretyczne postulują, aby stosunek napięcia szumów do napięcia drgań monochromatycznych $\frac{\sigma_0}{A_0}$ pozwalał na zastąpienie funkcji $\arctg \frac{\sigma_0}{A_0}$ i $\arcsin \frac{\sigma_0}{A_0}$ przez wartość tego stosunku oraz aby nieliniowość układu pobudzenia była do pominięcia dla amplitudy napięcia szumów. Warunki powyższe będą w makiecie spełnione w sposób zadowalający, jeżeli napięcie szumów będzie rzędu stu miliwoltów.

Naszkicowana powyżej koncepcja programu badań eksperymentalnych powinna umożliwić bardziej precyzyjne porównanie wyników teoretycznych z wynikami pomiarów i na tej podstawie lepiej ocenić wartość praktyczną dyfuzyjnego modelu drgań losowych¹⁾. Osiągnięte na tej drodze rezultaty będą przedmiotem następnej publikacji.

Na zakończenie autor pragnie wyrazić serdeczne podziękowanie profesorom Januszowi Groszkowskiemu i Stanisławowi Ryżko za przejrzenie niniejszej pracy oraz szereg cennych uwag umożliwiających sprostowanie wielu nieścisłości oraz bardziej precyzyjne ujęcie przedstawionych tu zagadnień.

Politechnika Warszawska
Katedra Urządzeń Radiotechnicznych

¹⁾ Już w toku opracowania redakcyjnego niniejszej pracy ukazała się publikacja Strebela [2], w której wykorzystano zasadę sztucznego zwiększenia poziomu szumów w makiecie układu generacyjnego. Opublikowane wyniki eksperymentalne dotyczą jednak w zasadzie tylko znacznie odległych od punktu ω_0 części widma i w związku z tym są mało interesujące z punktu widzenia stabilizacji częstotliwości.

WYKAZ LITERATURY

1. Barnes J. A.: A High Resolution Amonia — Maser-Spectrum — Analyser. — IRE Trans on Instr., t. 1—10, Nr 1, 1961.
2. Barnes J., Mockler R.: The Power Spectrum and its Importance in Precise Frequency Measurments. — IRE Trans on Instr., t. 1—9, Sept. 1960 s. 149.
3. Backer G.: Ein Impulsschwebungsverfahren für genaue Frequenz — und Phasenmessungen and phasenstarre Frequenztransformation. — Frequenz, t. 12, 1958, Nr 3, s. 82.
4. Blaquiere A., Grivet P.: Nonlinear Effects of Noise in Electronics Clocks. — PIEEE, t. 51, 1963, Nr 11, s. 1606.
5. Bersztejn I. L.: Fluktuacija amplitudy i fazy łampowego gienieratora., — Izw. Ak. Nauk ZSRR, t. 14, 1950, Nr 2, s. 145.
6. Blackman R. B., Tukey J. W.: The Measurment of Power Spectra from the Point of View of Communication Engineering. — Bell Syst. Tech. Jour., t. 37, 1958, Nr 1, s. 185 i Nr 3, s. 485.
7. Chafin L. R.: A Technique for Short Term Oscillator Stability Measurments. — PIRE, t. 48, 1960, koresp. s. 1914.
8. Charkiewicz A. L.: Spiektry i analiz. — Gos. Izd. Fiz. Mat. Lit., Moskwa 1962.
9. Cutler L. S.: A Frequency Standart of Exceptional Spectral Purity and Long Term Stability. — IRE Inst. Conv., 1961, cz. 9, s. 183.
10. Edson W. A.: Noise in Oscillators. — PIRE, t. 48, 1960, s. 1454, Nr 8.
11. Glay M. J. E.: Monochromaticity and Noise in a Regenerative Electrical Oscillator. — PIRE, t. 48, 1960, Nr 8, s. 1473.
12. Gonorowski J. S.: Jeszcze o fluktuacji fazy w łampowym gienieratorie. — Rad. i El., t. 2, 1957, Nr 10, s. 1279.
13. Gorelik G. S.: Nieliniynje kolebanija, interferencija i fluktuacija — Izw. Ak. Nauk. ZSRR, Fiz. t. 14, 1950, Nr 2, s. 187.
14. Groszkowski J.: Wytwarzanie drgań elektrycznych. — Warszawa 1958, PWN.
15. Kuźniecowa R. J., Stratonowicz P. Ł., Tichonow W. I.: Wozdziejstwie eliektriczeskich fluktuacji na łampowyj gienierator. — Żur. Eksp. i Tior. Fiz., t. 28, 1955, Nr 5, s. 509.
16. Malling L. R.: Phase Stable Oscillators for Space Communications Including Relationstrip Between the Phase Noise, the Short-Term Stability and the Q of the Oscillator. — PIRE, t. 50, 1962, Nr 8, s. 1656.
17. Middleton D.: Statistical Communication Theory. 1960.
18. Mullen J. A.: Background Noise in Nonlinear Oscillators. — PIRE, t. 48, 1960, Nr 8, s. 1467.
19. Newman J., Fey L., Attkinson W. R.: Spectrum Analysis of Extremely Low Frequency Variations of Quartz Oscillators. — PIEEE, t. 51, 1963, Nr 2, s. 379.
20. Pierce J. R.: Physical Sources of Noise. — PIRE, t. 44, 1956, Nr 5, s. 601.
21. Spälti A.: Der Einfluss des Thermischen Widerstandrauschens und des Schroteffektes auf die Störmodulation von Oszillatoren. — Bull. Schweiz. El. Ver, 1948, s. 419.
22. Strebel B.: An Investigation Relating to Spälti's Theory of Noise Produced in Oscillators. — NTZ Comm. J., t. 2, 1963, Nr 6, s. 284.
23. Rytow: Fluktuacja w awtokolebiatielnich sistiemach tomsonowskiego typu. — Żur. Eks. Tior. Fiz., t. 29, 1955, Nr 3/9, s. 384.

24. Sann Klaus H.: Phase Stability of Oscillators. — PIRE, t. 49, 1961, Nr 2, Koresp. s. 527.
25. Stewart J. L.: Frequency Modulation Noise in Oscillators. — PIRE, t. 44, 1956, Nr 3, s. 372.
26. Stranberg M. W. P.: Precise Specification of Frequency. — Microw. Jour., t. 3, 1960, Nr 8, s. 45.
27. Stratonowicz R. L.: Izbranyje woprosy teorji fluktuacji w radiotekhnike. — Moskwa 1961, Sow. Radio.
28. Tang L. C.: The Behavior of Nonlinear Oscillating System in the Presence of Noise. — PIRE, t. 48, 1960, Nr 8, Koresp. s. 1493.
29. Tanzman H. D.: Short Term Stability Measurements. — IRE Trans. Conv., 1961, cz. 9, s. 191.
30. Wołodin B. G., Ganin M. P. i inni: Rukowodstwo dla inżynierow po rieszenii zadacz teorji wierojatnostiej. — Leningrad 1962, Gos. Sojuz. Izdat. Sud. Prom.
31. Zabortinskij M. E., Zilberman P. E.: O fluktuacjach w kwarcewych gienieratorach. — Dokł. Ak. Nauk ZSRR, Fiz. t. 119, 1958, s. 918.

R. NOWAK

PROBLEM OF FREQUENCY INSTABILITY IN CONNECTION WITH FLUCTUATION PHENOMENA

Summary

In the paper modern views on the question of fluctuation frequency instability from the theoretical and experimental point of view are discussed. The noises occurring in oscillator circuits are the source of this instability. These phenomena determine the structural border of stabilisation of standard sources.

A theoretical oscillation model called diffusion model of oscillation was introduced according to publications of different authors. The above model describes the effects of white noise on Van der Pol's oscillator circuit with determined assumptions admissible in practical application. The described relations determine the most important oscillation parameters in their most simple form and depending on physical quantities easily determined in practice. The description method and the unit of quantities were chosen in such a way, that the results can be applied directly to different types of oscillation circuits.

In the part describing the measuring-experimental side of the problem a systematical review of different methods and of measuring circuits serving to the measurements of the phase fluctuation were presented. From parameters describing these fluctuations a fluctuation deviation of oscillations was chosen as a fundamental parameter. The above deviation enables the extrapolation of frequency instability idea on the area of fluctuation phenomena.

In the final part of the paper a comparison of published analytical and experimental results was given. From this comparison a conclusion was drawn that the theoretical model of diffusion oscillations is not confirmed by experiments and the reasons of it were mentioned. At the same time the assumptions of an experiment enabling the confirmation of conclusions resulting from the diffusion oscillation model were formulated. Such an experiment would consist in calculation of respective measurements for the oscillator circuit with artificially augmented level of white noise.

Р. НОВАК

ВОПРОСЫ НЕСТАБИЛЬНОСТИ ЧАСТОТЫ В СВЕТЕ ФЛУКТУАЦИОННЫХ ЯВЛЕНИЙ

Резюме

В статье представляются с теоретической и экспериментальной точки зрения современные взгляды на вопрос флуктуационной неустойчивости частоты. Источником флуктуационной неустойчивости частоты являются шумы, выступающие в генерационных схемах. Эти явления составляют структурный предел стабилизации стандартных источников.

На основании работ разных авторов выделена и представлена теоретическая модель колебаний, называемая диффузионной моделью. Эта модель описывает воздействие белого шума на Ван дер Полевскую генерационную схему, при определенных допусках с точки зрения практических применений, упрощающих предпосылках. Приведение зависимости определяют самые важные параметры колебаний в возможно несложном виде и в зависимости от физических величин, легко определяемых в практике. Метод описания и состав этих величин выбран таким образом, чтобы результаты можно было применить непосредственно для разного рода генерационных схем.

В части описывающей экспериментально-измерительную сторону вопроса приведен систематизированный обзор разных методов и измерительных схем, предназначенных для измерений флуктуации фазы колебаний. Среди описывающих эти флуктуации параметров в качестве основного параметра выбрана флуктуационная девиация частоты, так как эта величина дает возможность экстраполяции понятия неустойчивости частоты на область флуктуационных явлений.

В заключительной части работы произведено сравнение опубликованных аналитических и экспериментальных результатов. С этого сравнения сделан вывод о недостатке в экспериментальной верификации теоретической модели диффузионных колебаний, показывая причины такого состояния вещей. Сформулировано также основные данные эксперимента, дающего возможность экспериментальной проверки выводов, вытекающих из диффузионной модели колебаний. Такой эксперимент состоял бы в проведении соответствующих измерений для генерационной схемы с искусственно увеличенным уровнем белого шума.

621.396.6

STANISŁAW STUCHŁY

Mikrofalowy tłumik obrotowy jako pierwotny wzorzec tłumienia

Rękopis dostarczono 17.5.1963

Podano zasadę działania falowodowego tłumika obrotowego i przeanalizowano możliwości użycia tego rozwiązania jako pierwotnego wzorca tłumienia. Opisano źródła błędów oraz podano wzory pozwalające na określenie teoretycznej dokładności przyrządu.

1. WSTĘP

Pierwotny wzorzec tłumienia jest elementem mikrofalowym nie wymagającym skalowania ani przy pomocy niezależnych pomiarów mocy, ani też przez porównanie z innym wzorcem mikrofalowym; jego skalowanie powinno opierać się na pomiarze wielkości mechanicznych.

W miernictwie mikrofalowym przyrządy tego rodzaju znajdują szerokie zastosowanie. Układy pomiarowe służące do pomiarów tłumienia i skalowania tłumików mikrofalowych na ogół wykorzystują wzorzec pierwotny. Rozszerzenie zakresu pomiarowego wzorcowych bolometrycznych mierników mocy mikrofalowej w kierunku mocy średnich wymaga również zastosowania tłumika wzorcowego. Generatory szumów wykorzystujące lampy jarzeniowe i stosowane w zakresie mikrofal do badania odbiorników posiadają stały poziom sygnału i w większości praktycznych układów powinny współpracować z wzorcowym tłumikiem regulowanym. Również generatory sygnałów wzorcowych są na wyjściu wyposażone w tłumik wzorcowy.

W ogólnym przypadku wymagania stawiane tłumikowi wzorcowemu można ująć w następujące punkty: możliwie duża dokładność, zależna od metody pomiaru tłumienia i dokładności układu mechanicznego; impedancja wejściowa i wyjściowa, możliwie zbliżone do impedancji charakterystycznych zastosowanych przewodnic falowych; możliwie duża powtarzalność, zależna od stałości elementu tłumiącego i precyzji ustawienia mechanicznego; określony zakres prawomocności skalowania, zależny od zmian krzywej skalowania wywołanych zmianami temperatury, ciśnienia, wilgotności oraz częstotliwości w określonym paśmie pracy.

Na podstawie powyżej sformułowanych wymagań nietrudno stwierdzić, że spośród znanych rozwiązań jedynie tłumik reaktancyjny oraz tłumik obrotowy spełniają postawione wymagania.

Tłumikiem reaktancyjnym przyjęto nazywać odcinek falowodu pracującego poniżej częstotliwości granicznej, wprowadzający tłumienie nieabsorbcyjne [1]. Wiadomo, że dla częstotliwości mniejszych od częstotliwości granicznej, dla ustalonego rodzaju pola w falowodzie natężenia pól maleją eksponentalnie wzdłuż osi falowodu, natomiast faza zmienia się bardzo nieznacznie w wyniku strat w ściankach. Przez zmianę długości odcinka takiego falowodu można uzyskać zmianę tłumienia w szerokich granicach, a więc rozwiązanie to można wykorzystać do budowy wzorcowego tłumika regulowanego. W budowie tłumików reaktancyjnych wykorzystuje się na ogół falowody kołowe, ze względu na uproszczenie konstrukcji i jej technologiczność. W praktycznym rozwiązaniu można zastosować zarówno rodzaj podstawowy w falowodzie kołowym, a więc H_{11} , jak również następny, wyższy, E_{01} . Dla rodzaju H_{11} elementami pobudzającymi są pętle, a więc sprzężenie posiada charakter indukcyjny, zaś dla rodzaju E_{01} dyski o charakterze pojemnościowym.

Stała tłumienia falowodu pracującego poniżej częstotliwości granicznej dla danego rodzaju pola, wyraża się wzorem [3]:

$$\alpha = \frac{2\pi}{\lambda_c} \sqrt{1 - \left(\frac{\lambda_c}{\lambda}\right)^2}, \quad (1)$$

dla falowodu kołowego

$$\lambda_c = \frac{2\pi r \sqrt{\epsilon_w}}{S_{mn}}; \quad (2)$$

gdzie:

λ -długość fali w wolnej przestrzeni, λ_c -graniczna długość fali, r -promień falowodu, S_{mn} -pierwiastki funkcji Bessela dla H_{11} $S_{11} = 1,841$, zaś dla E_{01} $S_{01} = 2,405$, ϵ_w -względna stała dielektryczna materiału wypełniającego falowód.

Zakładając, że $\frac{\lambda_c}{\lambda} \ll 1$ i $\epsilon_w \approx 1$ (powietrze) (założenie to można spełnić w praktyce),

$$\alpha \approx \frac{2\pi}{\lambda_c} \left[1 - \frac{1}{2} \left(\frac{\lambda_c}{\lambda} \right)^2 \right] \approx \frac{2\pi}{\lambda_c} = \frac{S_{mn}}{r} \frac{[\text{Nep}]}{[\text{m}]}. \quad (3)$$

Z równania (3) widać, że w pewnym przybliżeniu stała tłumienia nie zależy od częstotliwości i jest zależna jedynie od wykorzystywanego rodzaju pola oraz rozmiarów falowodu. Dla liniowego zakresu tłumienie wyraża się zatem wzorem,

$$A = \alpha l = \frac{S_{mn}}{r} l \quad [\text{Nep}]. \quad (4)$$

Zależność stałej tłumienia od rozmiarów falowodu można określić jako,

$$\frac{d\alpha}{\alpha} = - \left[\frac{1}{1 - \left(\frac{\lambda_c}{\lambda}\right)^2} \right] \frac{dr}{r}. \quad (5)$$

Wyrażenie to dla $\frac{\lambda_c}{\lambda} \ll 1$ upraszcza się do postaci

$$\frac{d\alpha}{\alpha} = - \frac{dr}{r}. \quad (6)$$

Zmiany stałej tłumienia w funkcji częstotliwości określa wzór

$$\frac{d\alpha}{\alpha} = 2 \left(\frac{\lambda_c}{\lambda} \right)^2 \frac{d\lambda}{\lambda}. \quad (7)$$

Jeżeli tłumik pracuje w zakresie długości fali $\lambda_1 \div \lambda_2$, to długość fali, dla której przeprowadza się obliczenia, jest

$$\frac{1}{\lambda_d^2} = \frac{1}{2} \left(\frac{1}{\lambda_1^2} + \frac{1}{\lambda_2^2} \right), \quad (8)$$

zaś zmiana stałej tłumienia na krańcach zakresu wynosi

$$\delta\alpha = \pm \frac{2\pi}{\lambda_c} \left[\left(\frac{\lambda_c}{\lambda_1} \right)^2 - \left(\frac{\lambda_c}{\lambda_2} \right)^2 \right] \frac{[\text{Nep}]}{[\text{m}]}. \quad (9)$$

Jeżeli spełniony jest warunek $\frac{\lambda_c}{\lambda} \ll 1$, zaś średnica falowodu jest utrzymana z dużą dokładnością, można uzyskać bardzo wysoką dokładność tłumienia, zależną głównie od dokładności pomiaru długości.

Równanie (4) postuluje liniową zależność tłumienia wyrażonego w dB od długości odcinka falowodu. Zależność ta jest spełniona z zadowalającą dokładnością dopiero przy tłumieniach rzędu $20 \div 30$ dB. Powodem tego jest powstawanie wyższych rodzajów pola w pobliżu elementów sprzęgających. Specjalne filtry umieszczone w sąsiedztwie elementów pobudzających skutecznie zmniejszają poziom tych wyższych rodzajów, jednak w praktyce początkowo tłumienie rzadko jest mniejsze od 20 dB.

Na podstawie tej krótkiej analizy można określić zalety i wady tłumika reaktancyjnego zastosowanego jako tłumik wzorcowy. Do jego zalet należy zaliczyć liniową zmianę tłumienia w funkcji długości odcinka falowodu, praktycznie niezależną od częstotliwości. Tłumik nie wymaga skalowania, a przez odpowiednią konstrukcję i dobór materiałów (współczynnik rozszerzalności termicznej) można w znacznym stopniu uniezależnić zmianę wartości tłumienia od wpływu warunków atmosferycznych. Najpoważniejszymi wadami tego rozwiązania są: duże tłumienie początkowe oraz trudności szerokopasmowego dopasowania elementów pobu-

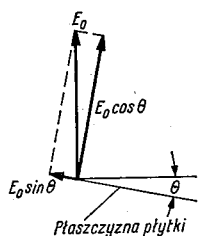
dzających (szczególnie w zakresie częstotliwości powyżej 3 GHz), jak również stosunkowo mała wytrzymałość mocowa tych ostatnich. Tłumiki reaktancyjne są stosowane z dużym powodzeniem jako wzorce pierwotne wszędzie tam, gdzie duże tłumienie początkowe nie odgrywa zasadniczej roli, zaś układ pracuje w wąskim pasmie częstotliwości.

Tłumik obrotowy, którego analiza zostanie przedstawiona w następnej części, nie posiada wad charakterystycznych dla tłumika reaktancyjnego, szczególnie w zakresie częstotliwości powyżej 3 GHz. Dla mniejszych częstotliwości rozmiary tego rodzaju tłumika wzrastają do nierozsądnych wartości.

2. ZASADA DZIAŁANIA

Tłumik obrotowy jest zbudowany na falowodzie kołowym, którego rozmiary są tak dobrane, że dla ustalonego pasma częstotliwości w falowodzie kołowym pobudza się jedynie rodzaj podstawowy H_{11} . W rodzaju tym, jak wiadomo [6], przenoszenie energii może odbywać się w dwu rodzajach posiadających tę samą długość graniczną fali, lecz o wektorach pola elektrycznego wzajemnie prostopadłym. Fale elektromagnetyczną o dowolnej polaryzacji można rozłożyć na dwie składowe wzajemnie prostopadłe, jak na rys. 1. Płaszczyzny, w których leżą te składowe nazywane są polaryzacjami podstawowymi, choć ich wybór jest całkowicie dowolny.

Jeśli w płaszczyźnie jednej z polaryzacji podstawowych umieścić bardzo cienką płytkę dielektryczną pokrytą warstwą oporową, to składowa prostopadła do tej warstwy będzie przenoszona z bardzo małym tłumie-

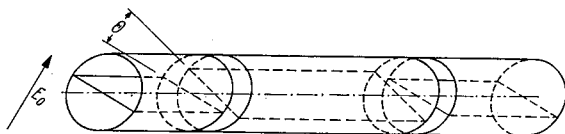


Rys. 1. Rozkład wektora E_0 na składowe

niem, natomiast składowa równoległa do płytki będzie bardzo silnie tłumiona. Amplituda natężenia pola elektrycznego składowej poprzecznej fali padającej wynosi E_0 (rys. 1), a po przejściu przez odcinek falowodu zawierający płytkę absorbcyjną zmniejszy się, do wartości $E_0 \cos \theta$, gdzie kąt θ jest zawarty między płaszczyzną płytki a płaszczyzną polaryzacji fali padającej.

Tłumik obrotowy składa się, jak pokazano na rys. 2, z trzech odcinków falowodu kołowego, w których umieszczono płytki absorbcyjne. Dwa odcinki skrajne są nieruchome, zaś odcinek środkowy jest obrotowy. W od-

cinku pierwszym wektor pola elektrycznego jest prostopadły do płytki, a więc fala przechodzi z pomijalnym tłumieniem. W części środkowej płytki jest ustawiona pod kątem Θ w stosunku do płytek nieruchomych tak, że składowa $E_0 \sin \Theta$ jest absorbowana przez płytkę, natomiast składowa $E_0 \cos \Theta$ przechodzi z pomijalnym tłumieniem. W trzecim odcin-



Rys. 2. Rysunek schematyczny wyjaśniający budowę tłumika obrotowego

ku proces ten powtarza się i tłumik opuszcza fala, której amplituda jest proporcjonalna do $E_0 \cos^2 \Theta$. Ponieważ w układzie tym obowiązuje zasada odwracalności, te same zjawiska zachodzą również dla przeciwnego kierunku poruszania się fali elektromagnetycznej.

Tłumienie wszystkich odcinków wyraża się zatem wzorem

$$A(\Theta) [\text{dB}] = -40 \lg \cos \Theta + A_0, \quad (10)$$

gdzie A_0 straty wtrącenia dla $\Theta = 0$ są wywołane przez szereg czynników, lecz ich wartość nie zależy od wartości kąta Θ .

Przez zmianę kąta Θ można — jak widać — zmieniać wartość tłumienia w szerokich granicach. Z równania (10) wynika również, że wartość tłumienia nie zależy od częstotliwości i tłumienia płytek absorbujących. Wniosek ten jest słuszny tylko w pewnych granicach i będzie zweryfikowany w dalszej części pracy. Najistotniejszym wnioskiem jest jednak to, że tłumienie dla tego tłumika zależy *bezpośrednio* od wartości kąta Θ , a więc wielkości mechanicznej. Pozwala to na uzyskanie dużej dokładności tłumienia przez nałożenie odpowiednio wysokich wymagań na układ pomiaru kąta skreślenia płytki obrotowej.

Zmiany tłumienia przebiegają wg krzywej $\lg \cos \Theta$, której nachylenie jest równe $\tg \Theta$, a więc rośnie gwałtownie przy wzroście kąta Θ . W wyniku uzyskuje się bardzo dużą dokładność dla małych wartości tłumienia, natomiast istnieją poważne ograniczenia dokładności na drugim krańcu zakresu.

3. DOKŁADNOŚĆ

Krzywa skalowania tłumika określona równaniem (10) jest zależnością teoretyczną, w której pominięto błędy nieodzownie występujące w praktyce. W niniejszym paragrafie zostaną rozważone źródła błędów wraz z przybliżoną analizą matematyczną, która pozwala na ilościową ocenę dokładności omawianego tłumika.

Błędy występujące w pracy tłumika obrotowego można w sposób całkowicie umowny podzielić na dwie zasadnicze grupy: błędy, których źródła znajdują się w układzie mechanicznym, oraz błędy, których źródła leżą w układzie elektrycznym.

Z grupy błędów mechanicznych należy wymienić: luzy i niedokładności układu mechanicznego służącego do obrotu odcinka środkowego wraz z układem napędu skali, niedoskonałość wykonania mechanicznego skali oraz niedokładność obliczenia skali z równania (10). Pierwsze dwa zależą głównie od rodzaju zastosowanej konstrukcji mechanicznej. Skala tłumika może być sprzężona bądź bezpośrednio z odcinkiem środkowym, bądź też za pośrednictwem przekładni. Oceny błędu można dokonać na podstawie równania (10) pamiętając, że

$$\frac{dA}{d\theta} \frac{[dB]}{[^\circ]} = 0,3 \operatorname{tg} \theta. \quad (11)$$

Ze względu na znaczne nachylenie krzywej skalowania, równanie (11), dla uzyskania dokładności pomiaru tłumienia rzędu ułamków procenta wymagana dokładność ustawienia odcinka środkowego wynosi kilka minut kątowych. Niedokładność obliczenia skali można sprowadzić do bardzo małych wartości stosując do obliczeń dostatecznie dokładne tablice funkcji trygonometrycznych, należy jednak podkreślić, że niedokładności te muszą być brane pod uwagę przy ekstremalnych wymaganiach postawionych dokładności pomiaru tłumienia.

Grupa błędów elektrycznych jest znacznie obszerniejsza; dla przejrzystości można ją rozbić na szereg punktów.

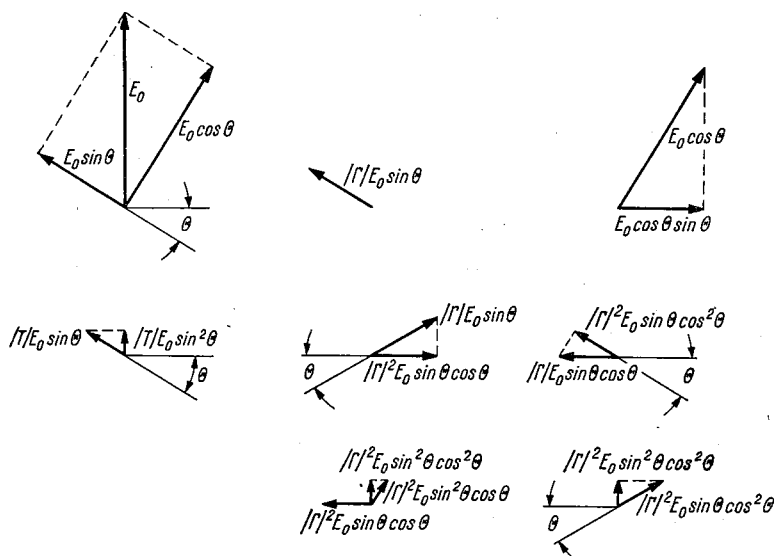
— Odbicia energii występujące w transformatorach rodzaju. Tłumik obrotowy jest, jak wiadomo, zbudowany na falowodzie kołowym, natomiast w układach pomiarowych rzadko stosuje się ten rodzaj falowodu. Występuje więc konieczność zastosowania transformatorów rodzaju. Szeroko rozwinięta teoria i technika tych ostatnich pozwala na budowę transformatorów o bardzo małym WFS w szerokim pasmie częstotliwości. Odbicia te wpływają jedynie na straty wtrącenia, które są niezależne od kąta θ , zaś ich wpływ na dokładność pomiaru tłumienia można ocenić na podstawie np. [2] str. 401.

— Odbicie energii w punktach połączenia odcinka obrotowego i odcinków stałych. Połączenie elektryczne tych odcinków wykonuje się zwykle przy pomocy dławików radialnych bądź współosiowych. Przy odpowiedniej budowie można tu uzyskać bardzo mały WFS w szerokim pasmie częstotliwości, ponadto odbicia te nie wpływają podobnie jak poprzednie na dokładność krzywej skalowania, a jedynie powiększają straty wtrącenia.

— Nieprostokątne ustawienie nieruchomych płytek absorbcyjnych w stosunku do płaszczyzny polaryzacji fali w falowodzie kołowym powo-

duże powstanie dodatkowego tłumienia, którego wartość nie zależy od kąta θ , a przyczynia się jedynie do wzrostu strat wtrącenia. Straty te dla każdej z płytek zależą od $\sin \alpha$, gdzie α jest kątem pomiędzy płytką nieruchomą a płaszczyzną polaryzacji fali.

— Odbicia energii od zakończeń płytek absorbcyjnych. Wprowadzenie do falowodu kołowego płytki absorpcyjnej powoduje powstanie odbić, których wartość zależy od oporności powierzchniowej płytki, kąta ustawienia płytki w stosunku do polaryzacji rodzaju H_{11} . Odbicia są tym większe, im większy jest ten kąt i im mniejsza jest oporność powierzchniowa płytki. Autorowi nie są znane równania pozwalające określić te zależności, jak również teoretyczne metody dopasowania płytek absorbcyjnych w falowodzie kołowym. W praktyce zakończenia płytek są wykonane w kształcie rybiego ogona i przy dostatecznej długości odcinka do-



Rys. 3. Graficzne wyprowadzenie wzoru określającego wpływ wielokrotnych odbić wewnętrznych na dokładność tłumika

pasowującego metoda ta daje zadowalające rezultaty. Odbicia od płytek mają dwójaki wpływ na dokładność tłumika. Po pierwsze w znacznym stopniu decydują o wartości wejściowego WFS całego tłumika (szczególnie płytka obrotowa), przy czym WFS ten rośnie przy wzroście tłumienia (wzrost kąta θ). Po drugie wywołują wielokrotne odbicia wewnętrzne, zależne od wartości współczynnika odbicia zakończenia płytki $|I|$ oraz kąta θ . Graficzne wyprowadzenie związku, który pozwala na ocenę powstałego błędu podano na rys. 3, przy czym dla uproszczenia założono, że $|I|$ jest dla wszystkich zakończeń jednakowa, zaś rozkład fazowy

najbardziej niekorzystny. W wyniku wielokrotnych odbić wewnętrznych na wyjściu pojawia się fala o amplitudzie

$$E'_{wyj} = 2 E_0 |\Gamma|^2 \sin^2 \Theta \cos^2 \Theta, \quad (12)$$

co daje błąd wyrażony w decybelach,

$$\Delta A[\text{dB}] \approx 8,7 (2 |\Gamma|^2 \sin^2 \Theta - |\Gamma|^4 \sin^4 \Theta) \quad (13)$$

dla $|\Gamma| \ll 1$.

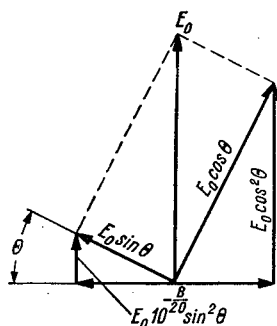
— Skończona wartość tłumienia płytki obrotowej dla składowej równoległej do płytki jest źródłem dalszych błędów w stosunku do krzywej teoretycznej z równania (10). Jeżeli tłumienie składowej równoległej do płytki wynosi B [dB], to na wyjściu tłumika pojawia się składowa określona równaniem wyprowadzonym graficznie na rys. 4

$$E''_{wyj} = E_0 10^{-\frac{B}{20}} \sin^2 \Theta, \quad (14)$$

a błąd w krzywej skalowania wynosi

$$\Delta A[\text{dB}] \approx 8,7 \left(10^{-\frac{B}{20}} \text{tg}^2 \Theta - \frac{1}{2} 10^{-\frac{B}{10}} \text{tg}^4 \Theta \right). \quad (15)$$

Z równania (15) wynika, że do uzyskania możliwie małego błędu konieczne jest zapewnienie dużej wartości B . Autorowi nie są znane równania pozwalające ustalić zależność tłumienia od oporności powierzchniowej,



Rys. 4. Rysunek wyjaśniający powstawanie składowej $E_0 10^{-\frac{B}{20}} \sin^2 \Theta$

jednak wiadomo, że rośnie ono przy zmniejszeniu oporności powierzchniowej. Wybór oporności musi być kompromisem pomiędzy wymaganym tłumieniem a dopuszczalnymi odbiciami od zakończeń płytki. W praktyce przyjmuje się na ogół oporność powierzchniową w granicach $100 \div 300 \Omega/\square$.

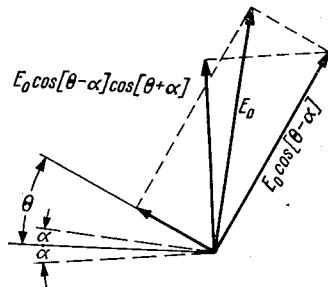
— Wzajemne skrzyżowanie płytek nieruchomych również powoduje odchyłki od teoretycznej krzywej skalowania. Jeśli płytki są skrzyżowane względem siebie o kąt $\pm \alpha$, to na wyjściu pojawia się składowa wyprowadzona na rys. 5

$$E'''_{wyj} = E_0 \cos^2 \Theta \cos^2 \alpha - E_0 \sin^2 \Theta \sin^2 \alpha \quad (16)$$

dając błąd w decybelach

$$\Delta A [\text{dB}] \approx 8,7 \left[-\sin^2 \alpha (1 + \operatorname{tg}^2 \Theta) - \frac{\sin^4 \alpha}{2} (1 + \operatorname{tg}^2 \Theta)^2 \right]. \quad (17)$$

Błąd ten można zmniejszyć przez uważne ustawienie płytek, przy czym wymagana dokładność jest rzędu pojedynczych minut kątowych.



Rys. 5. Graficzne wyprowadzenie wzoru określającego wpływ skręcenia płytek nieruchomych na dokładność tłumika

— Skończona wartość tłumienia płytek nieruchomych posiada drugorzędny wpływ na dokładność tłumika. Tłumienie to musi być jednak dostatecznie duże w celu usunięcia składowych nieprostopadłych do płytek. W praktyce wartość $C > 40$ dB w większości wypadków uważa się za wystarczającą.

— Powstawanie wyższych rodzajów pola może być źródłem pewnych błędów, jednak zagadnienie to nie zostało dostatecznie zbadane. Ogólnie wiadomo, że wprowadzenie do falowodu płytki absorpcyjnej powoduje lokalne pola wyższych rodzajów, jednak ich wpływ na przebieg krzywej skalowania tłumika nie jest znany. W dostępnych danych doświadczalnych błąd ten nie został uchwyciony.

4. STABILNOŚĆ

Głównym źródłem niestabilności wszystkich tłumików absorpcyjnych są zmiany tłumienia wynikłe ze zmian oporności jednostkowej płytki absorpcyjnej. Dla tłumika obrotowego zależność tłumienia od oporności powierzchniowej można wyrazić wzorem [8]

$$\Delta A [\text{dB}] = \frac{10 \operatorname{tg}^2 \Theta}{10^{\frac{B}{20}} + \operatorname{tg}^2 \Theta} \lg \frac{1}{1 - 2 \frac{dr}{r}}, \quad (18)$$

gdzie $\frac{dr}{r}$ jest zmianą oporności na jednostkę powierzchni wynikłą ze zmian temperatury, wilgotności itp. Dla płytek pokrytych chromonikielem z powłoką fluorku manganu zmiany oporności powierzchniowej w warunkach laboratoryjnych nie przekraczają kilku procent. Z równa-

nia (18) wynika, że dla dostatecznie dużych wartości B , co jest w praktyce możliwe, tłumik obrotowy jest bardzo stabilny.

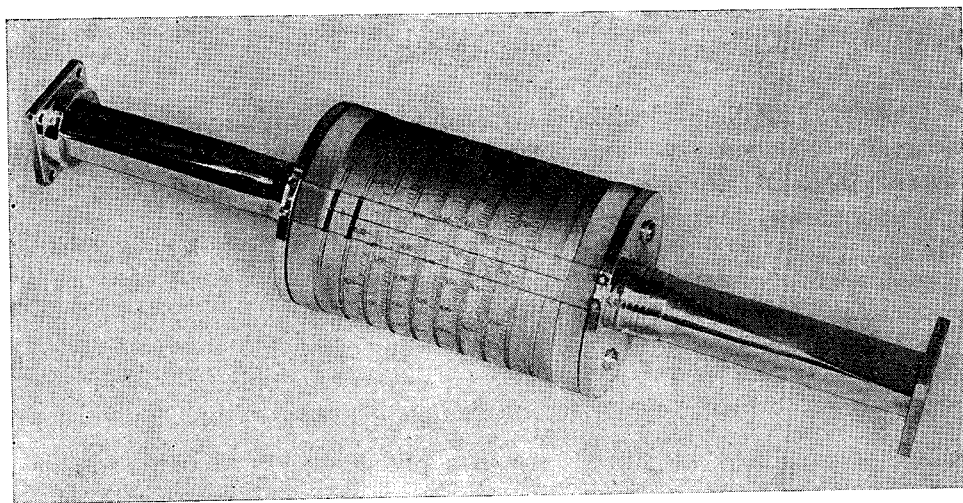
Zmiany rozmiarów elementów konstrukcji mechanicznej tłumika w funkcji temperatury mogą być również źródłem błędów, są one jednak bardzo małe i muszą być brane pod uwagę jedynie przy ekstremalnych wymaganiach dotyczących dokładności.

5. WYNIKI DOŚWIADCZALNE

W oparciu o przedstawioną powyżej teorię i analizę działania tłumika obrotowego autor opracował w Katedrze Techniki Fal Ultrakrótkich Politechniki Warszawskiej tłumik wzorcowy przystosowany do zakresu częstotliwości pasma X ($8600 \div 9600$ MHz).

Prace nad konstrukcją samego tłumika wzorcowego poprzedziły badania zachowania się płytek absorbcyjnych w falowodzie kołowym pracującym w rodzaju H_{11} . Celem tych prac było dobranie optymalnej z punktu widzenia tłumienia i dopasowania oporności powierzchniowej. Na podstawie szeregu pomiarów wiadomo, że dla wybranego pasma częstotliwości i średnicy falowodu kołowego 23,8 mm optymalną wartością oporności powierzchniowej jest około $150 \Omega/\square$. W trakcie tych pomiarów dobrano również doświadczalnie rozmiary klina dopasowującego płytkę absorbcyjną.

Do napędu i odczytu położenia obrotowej sekcji falowodu kołowego użyto bardzo precyzyjnej, bezluzowej przekładni śrubowej o przełożeniu 100/3. Taka konstrukcja umożliwiła zastosowanie spiralnej skali o długości około 2 m, na której dokładność odczytu równa była dokładności tłumika. Wygląd zewnętrzny opisywanego tłumika przedstawia rys. 6.



Rys. 6. Wygląd zewnętrzny tłumika obrotowego na pasmo X

Dokładność tłumika w zakresie $0 \div 40$ dB wynosi $0,5\%$ lub $0,01$ dB, w zależności od tego, która wartość jest większa. W określeniu dokładności zawarte są wszystkie przedstawione w poprzedniej części źródła błędów.

W okresie przeprowadzania prac doświadczalnych nie dysponowano niestety żadnym innym tłumikiem wzorcowym w tym pasmie częstotliwości, dlatego też sprawdzenia skalowania dokonano przez porównanie dwóch takich samych tłumików. Wybranim przedziałem (np. $0 \div 3$ dB) tłumienia jednego z tłumików sprawdzono cały zakres tłumienia drugiego egzemplarza. W celu zwiększenia czułości układu pomiarowego stosowano detektor różnicowy, pozwalający na odczyt zmian tłumienia poniżej $0,01$ dB.

Pomiary przeprowadzono dla kilku egzemplarzy tłumików w wielu punktach skali i przy różnych częstotliwościach w zakresie $8600 \div 9600$ MHz. Opisana metoda była jedyną możliwą do zastosowania. Potwierdzenie przyjętej metody pomiaru można znaleźć w pracy [5] opublikowanej po zakończeniu prac nad opisywanym tłumikiem.

Wykryte błędy skalowania były zawarte w granicach określonych na wstępie. W trakcie prac doświadczalnych stwierdzono znaczny wpływ nierównoległości płytek nieruchomych na dokładność tłumika, jednakże po dokładnym ustawieniu tych płytek błędy te udało się zmniejszyć do wartości niewykrywalnej w posiadanym układzie pomiarowym.

Po wykonaniu próbnej serii opisywanych tłumików opracowana w Katedrze dokumentacja została przekazana do przemysłu, który podjął produkcję tych tłumików.

6. ZAKOŃCZENIE

W pracy przedyskutowano błędy tłumika obrotowego, podając ich źródła oraz zależności analityczne pozwalające na obliczenie ich wartości. Na podstawie przeprowadzonej analizy można stwierdzić, że tłumik obrotowy może być stosowany jako wzorzec pierwotny. Przy użyciu precyzyjnej konstrukcji mechanicznej, właściwym dopasowaniu płytek absorbcyjnych i przy zapewnieniu dostatecznego tłumienia płytki obrotowej, można uzyskać dokładności rzędu ułamków procentu.

Dalszej analizy wymaga zagadnienie szerokości pasma i związanej z tym sprawy powstawania wyższych rodzajów pola, jak również ocena przesunięcia fazy, jakie wprowadza tłumik, a którego wartość niewątpliwie zależy od wartości tłumienia (położenie płytki obrotowej).

Pewne dodatkowe zwiększenie dokładności w stosunku do znanych rozwiązań można by uzyskać przez nałożenie warstwy oporowej na płaszczynie rozciągniętego walca dielektrycznego wypełniającego falowód. Taka

konstrukcja pozwala ponadto na zwiększenie równoległości płytek nieruchomych (równania [16] i [17]) i zapewnienie precyzyjnego ustawienia punktu zerowego przy równoczesnym znacznym zmniejszeniu gabarytów tłumika. Ze względu na dość znaczne rozmiary geometryczne zasada tłumika obrotowego jest wykorzystywana zwykle na falach krótszych od 10 cm.

Autor pragnie wyrazić podziękowanie Kierownikowi Katedry Radiolokacji Politechniki Warszawskiej prof. drowi inż. Stanisławowi Sławińskiemu za umożliwienie wykonania tej pracy oraz opiekę i pomoc okazaną w trakcie jej wykonywania.

*Zakład Doświadczalny
Budowy Aparatury Naukowej PAN*

WYKAZ LITERATURY

1. Gawron T., Panecki M., Stuchły S., Wiśniewski F.: Technika mikrofalowa — Prowadnice falowe i ich elementy — Terminy i określenia — Projekt normy. Prace PIT, rok XI, Nr 32/33, str. 99 ÷ 106.
2. Ginzton E. L.: Miernictwo mikrofalowe. PWT, Warszawa 1961.
3. Griesheimer R. N.: rozdz. 11 w książce pod red. Montgomery C. C.: Technique of Microwave Measurements, — Rad. Lab. Ser., t. 11, N. York 1947.
4. Hand B. P.: Electronics, 30 (1954), zeszyt 1, str. 184.
5. Mariner P. F.: An Absolute Microwave Attenuator. — Proc. IEE, t. 109, cz. B, Nr 49, wrzesień 1962.
6. Montgomery C. G., Dicke R. H., Purcell E. M.: Principle of Microwave Circuits. — Rad. Lab. Ser., t. 8, N. York 1948.
7. Southworth G. C.: Principles and Applications of Waveguide Transmission, N. York 1950.
8. Steinhart R.: Ein frequenzunabhängiges und weitgehend phasenreines Dämpfungsglied der Mikrowellentechnik. — NTZ, 1957, zeszyt 6, str. 294—297.

S. STUCHŁY

ATTENUATEUR VARIABLE POUR MICROONDES COMME ÉTALON PRIMAIRE D'AFFAIBLISSEMENT

Résumé

On a présenté le principe d'action d'un atténuateur variable guide d'onde et on a discuté les possibilités d'utilisation de cet appareil comme un étalon primaire d'affaiblissement. On a décrit les sources d'erreur et on a donné les formules permettant la détermination théorique de la précision de l'appareil.

S. STUCHŁY

DER MIKROWELLEN- DREHDÄMPFER ALS PRIMÄRES DÄMPFUNGSMUSTER

Zusammenfassung

Hier wurde die wellenführende Funktionierungsgrundlage des Drehdämpfers angegeben sowie die Anwendungsmöglichkeit dieser Lösung als primäres Dämpfungsmuster analysiert (erwogen). Es wurden auch die Fehlerquellen beschrieben, sowie Formeln, welche die theoretische Genauigkeit des Instrumentes zu ermitteln erlauben, angegeben.

С. СТУХЛЫ

МИКРОВОЛНОВЫЙ ВРАЩАТЕЛЬ В КАЧЕСТВЕ ПЕРВИЧНОГО СТАНДАРТА ЗАТУХАНИЯ

Резюме

Приведен принцип действия волноводного вращательного аттенюатора и проанализированы возможности применения такого прибора в качестве первичного стандарта затухания. Описаны источники погрешностей и приведены формулы, позволяющие определить теоретическую точность прибора.

SPIS TREŚCI

W. Fuliński: Nowe układy kompensacyjno-różnicowe do sprawdzania przekładników napięciowych	259
A. Staniszewski: Metody pomiaru kąta obciążenia maszyny synchronicznej	307
S. Piątek: Siła magnetomotoryczna wirnika turbogenerators potrzebna dla danych warunków pracy turbogenerators przy różnych systemach chłodzenia	331
W. Lidmanowski: Niektóre zagadnienia wyładowań ślizgowych	343
J. S. Zieliński: O obliczaniu przepięć wewnętrznych przy pomocy modeli matematycznych	361
Z. Szczepański: Wyładowania niezupełne jako przyczyna starzenia dielektryku elektrycznych kondensatorów wysokonapięciowych	383
J. L. Kulikowski: O strukturze i własnościach asymptotycznych pewnej klasy kwantowych układów decyzyjnych	427
R. Nowak: Zagadnienie niestabilności częstotliwości w świetle zjawisk fluktuacyjnych	453
S. Stuchły: Mikrofalowy tłumik obrotowy jako pierwotny wzorzec tłumienia	501

CONTENTS

W. Fuliński: New Circuits for Differential Testing of Potential Transformers Using Potentiometer Principle	302
S. Piątek: The Rotor-MMF Needed for Given Working Conditions of Turbogenerators with Various Cooling Systems	341
W. Lidmanowski: Some Aspects of Surface Discharges	358
J. S. Zieliński: Calculation of Internal Overvoltages with Analog Computer	381
R. Nowak: Problem of Frequency Instability in Connection with Fluctuation Phenomena	499

TABLE DE MATIÈRES

W. Fuliński: Nouveaux systèmes de mesure des transformateurs de tension utilisant les méthodes différentielle et d'opposition	303
A. Staniszewski: Methodes de mesure de l'angle de charge de la machine synchrone	329
S. Piątek: Force magnétomotrice de rotor de turbogénérateur nécessaire aux conditions données de fonctionnement de turbogénérateur à différents systèmes de refroidissement	341

W. Lidmanowski: Quelques problèmes concernant les décharges superficielles	358
Z. Szczepański: Décharges incomplètes comme la cause du vieillissement des diélectriques des condensateurs énergétiques haute tension	424
J. L. Kulikowski: Structure et qualités asymptotiques d'une classe des schémas de décision à quantification du signal	452
S. Stuchły: Atténuateur variable pour microondes comme étalon primaire d'affaiblissement	512

INHALT

W. Fuliński: Neue Spannungswandler — Prüfschaltungen nach dem Kompensations-Differentialprinzip	303
A. Staniszewski: Messmethoden für Lastwinkelmessungen an Synchronmaschinen	330
S. Piątek: Die magnetomotorische Kraft des Turbogenerators, die für gegebene Arbeitsbedingungen des Turbogenerators bei verschiedenen Kühlsystemen erforderlich ist	342
W. Lidmanowski: Einige Probleme von Gleitentladungen	359
J. S. Zieliński: Die Berechnung innerer Überspannungen mit Hilfe mathematischer Modelle	381
Z. Szczepański: Unvollständige Entladungen als Ursache von Alterung des Dielektrikums in energetischen Hochspannungskondensatoren	425
J. L. Kulikowski: Über die Struktur und asymptotischen Bewertungen von Eigenschaften einer gewissen Klasse dezedierender Quantensysteme	452
S. Stuchły: Der Mikrowellen-Drehdämpfer als primäres Dämpfungsmuster	513

СОДЕРЖАНИЕ

В. Фулинский: Новые компенсационно-дифференциальные схемы для проверки измерительных трансформаторов напряжения	605
А. Станишевский: Методы измерения угла нагрузки синхронной машины	330
С. Пионтэк: Магнитодвижущая сила ротора турбогенератора необходимая для данных условий работы турбогенератора при разных системах охлаждения	342
В. Лидмановски: Некоторые вопросы скользящих разрядов	359
Е. С. Зелиньски: К расчёту внутренних перенапряжений с применением математического моделирования	382
З. Щепаньски: Неполные разряды в качестве причины старения диэлектрика энергетических конденсаторов высокого напряжения	426
Ю. Л. Куликовски: О структуре и асимптотических свойствах одного класса решающих схем с квантованием сигнала	452
Р. Новак: Вопросы неустойчивости частоты в свете флуктуационных явлений	500
С. Стухлы: Микроволновый вращатель в качестве первичного стандарта затухания	513

POLSKA AKADEMIA NAUK
INSTYTUT PODSTAWOWYCH PROBLEMÓW TECHNIKI

ROZPRAWY ELEKTROTECHNICZNE

TOM X • ZESZYT 4

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1964

RADA REDAKCYJNA

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

KOMITET REDAKCYJNY

Redaktor Naczelny

PROF. WITOLD NOWICKI

Sekretarz

JULIUSZ MIERZEJEWSKI

ADRES REDAKCJI

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

PAŃSTWOWE WYDAWNICTWO NAUKOWE
Warszawa, Młodowa 10

Nakład 700 (571 + 129)	Oddano do składania 3.VIII.1964 r.
Ark. wyd. 12,50, ark. druk. 11,0 + 2 wkl.	Podpisano do druku w grudniu 1964 r.
Papier ilustr. kł. III, 80 g, 70×100	Druk ukończono w grudniu 1964 r.
Zamówienie nr 1229/64	Cena zł 25,— Z-32

Drukarnia im. Rewolucji Październikowej, Warszawa

JERZY OWCZAREK

Zarys współczesnej teorii selsynów nadawczo-odbiorczych wskaźnikowych

Rękopis dostarczono 17.6.1963

Praca stanowi krytyczne zestawienie aktualnego dorobku w dziedzinie teorii nadawczo-odbiorczych selsynów ze stykami ślizgowymi, opublikowanego w poszczególnych pozycjach literatury powołanych w załączonym wykazie bibliografii, oraz własnych prac autora. Uwzględniono prace opublikowane do 1963 r. włącznie. To zestawienie obejmuje całokształt rozpatrywanego zagadnienia, co ułatwia zorientowanie się w nim i zapobiega opieraniu się na przypadkowo dobranych fragmentarycznych wiadomościach, podawanych w oddzielnych artykułach. Korzystanie z wrywkowych wiadomości opublikowanych ponadto w różnych językach jest nie tylko kłopotliwe, lecz czasem niebezpieczne z tego względu, że niektóre z nich wskutek niedostatecznie wszechstronnej analizy zjawisk podają błędne wnioski, których niesłuszność została później udowodniona przez innych autorów. Stwierdzenia podane w niniejszej pracy są bądź udowodnione teoretycznie na miejscu lub w powołanej literaturze, bądź sprawdzone doświadczalnie. Dla zapobieżenia zbytniemu zwiększeniu objętości pracy większość wyprowadzeń matematycznych została z konieczności pominięta, jednakże w każdym z takich przypadków podano źródło, w którym można znaleźć szczegółowe uzasadnienie przytoczonego stwierdzenia.

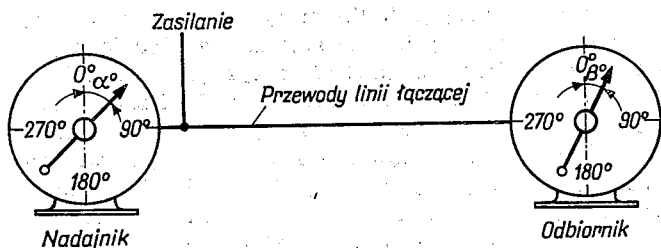
1. WSTĘP

Do sprzężenia na drodze elektrycznej dwóch lub kilku wałów, które nie mogą być sprzężone mechanicznie, stosuje się układy zestawione z różnego rodzaju maszyn elektrycznych specjalnej konstrukcji. Takie układy nazywane są łączami. Przeznaczeniem łącza jest natychmiastowe (synchroniczne) przeniesienie położenia kąтового wału jednej maszyny (nadajnika) względem stojana tejże maszyny do drugiej maszyny (odbiornika), znajdującej się w pewnej odległości i połączonej z nadajnikiem przewodami elektrycznymi. Położenie kątowe nadajnika zadawane jest urządzeniem sterującym. Dla obrócenia wału odbiornika o odpowiedni kąt powinien być w odbiorniku wytworzony pewien moment obrotowy, zwany momentem synchronizującym lub momentem przenoszonym przez łącze. W przypadku idealnego łącza kąt obrotu wału odbiornika spowodowany działaniem tego momentu odpowiadałby ściśle kątowi obrotu wału nadajnika. W łączach rzeczywistych przeciwstawiają się momen-

towi synchronizującemu momenty pasożytnicze, jak na przykład moment tarcia lub inne momenty szkodliwe, o których będzie mowa dalej, oraz moment obciążenia; powodują one powstanie kąta niezgodności, stanowiącego dla tego obciążenia uchyb kątowy. Jeżeli na przykład wirnik nadajnika (rys. 1) obracamy w prawo o kąt α° , to wirnik odbiornika obróci się o kąt β° , wówczas różnica tych kątów:

$$\delta^\circ = \alpha^\circ - \beta^\circ \quad (1)$$

nazywana kątem niezgodności będzie miarą uchybu przy danym obciążeniu wału odbiornika. Przy nawrotnej pracy łącza opóźnianie się wału odbiornika o uchyb kątowy przy obrocie w obydwie strony powoduje powstanie martwej strefy łącza, której wartość kątowa wynosi podwójną wartość uchybu kątowego.



Rys. 1. Kąt niezgodności łącza

Przy omawianiu teorii łącza dla uproszczenia zależności kąt odbiornika jest rozpatrywany w odniesieniu do kąta nadajnika, czyli uwzględnia się tylko kąt niezgodności (przy $\alpha^\circ = 0$, $|\delta^\circ| = |\beta^\circ|$). Takie postępowanie można stosować, jeżeli przenoszony moment i kąt niezgodności nie zależą od położenia wału nadajnika (kąta α°), co ma miejsce w większości łączy.

Najważniejszą charakterystyką łącza jest charakterystyka wytwarzanego w nim momentu synchronizującego w funkcji kąta niezgodności, zaś najważniejszym parametrem, charakteryzującym jakość łącza, jest stromość charakterystyki momentu w zakresie małych kątów niezgodności, określona pojęciem sztywności S lub współczynnikiem sztywności S' . Współczynnikiem sztywności S' nazywamy pochodną momentu względem kąta niezgodności, przy wartości tego kąta dążącej do zera:

$$S' = \left(\frac{dM}{d\delta^\circ} \right)_{\delta \rightarrow 0} \quad (2)$$

Graficzną interpretacją współczynnika sztywności jest styczna do charakterystyki momentu w punkcie $\delta^\circ = 0$. W praktyce posługujemy się częściej pojęciem sztywności S interpretowanym jako wartość momentu przy kącie niezgodności $\delta^\circ = 1^\circ$.

$$S = (M)_{\delta^\circ = 1^\circ} \quad (3)$$

Ze względu na to, że — jak to wyniknie z dalszych rozważań — charakterystyki momentu łączy mają postać sinusoid odkształconych przy większych kątach niezgodności, ich początkowy odcinek jest prawie prostoliniowy, wobec czego sztywność S dla dostatecznie małych kątów niezgodności może być obliczona praktycznie jako:

$$S = \frac{M}{\delta^0} \quad [\text{cmG}^\circ] \text{ lub } [\text{cmG/radian}],$$

a z drugiej strony, ze względu na to, że dla małych kątów niezgodności słuszna jest zależność

$$M = M_{\max} \cdot \sin \delta^\circ,$$

więc:

$$S = M_{\delta=1^\circ} = M_{\max} \sin \delta^\circ$$

a

$$S' = \left(\frac{dM}{d\delta^\circ} \right)_{\delta^\circ \rightarrow 0} = M_{\max} \left(\frac{d \sin \delta^\circ}{d\delta^\circ} \right)_{\delta^\circ \rightarrow 0} = M_{\max}.$$

Stąd

$$S = S' \sin 1^\circ,$$

czyli

$$S = 0,0175 S'. \quad (4)$$

Jest oczywiste, że im większa jest sztywność łączy, tym mniejszy będzie jego uchyb przy danym momencie obciążenia i momentach pasywnych, a więc tym lepsza jego jakość.

Zakres pracy statycznej maszyn łączy mieści się jak i w przypadku innych maszyn elektrycznych na rosnącej części charakterystyki momentu — od zera do momentu maksymalnego. Praktyczny zakres pracy łączy mieści się w granicach małych kątów niezgodności, do około $5 \div 10^\circ$. Wynika to z przeznaczenia łączy, którym jest synchroniczne przekazywanie położenia wałów z możliwie dużą dokładnością, a więc przy minimalnym kącie niezgodności, co jest szczególnie ważne przy nawrotnej pracy łączy. Z tego względu maszyny łączy oblicza się w założeniu wykorzystania cieplnego przy małych kątach niezgodności. Również i przebiegi charakterystyk są praktycznie ważne w zakresie tych samych małych kątów niezgodności. Należy zaznaczyć, że w tym zakresie charakterystyka momentu może być uważana za prostoliniową (z uchybem mniejszym od 0,75% [36]), a moment na wale obliczany przez pomnożenie wartości sztywności przez wartość kąta niezgodności.

Maszynowe łączy synchroniczne mogą być zasilane prądem stałym i przemiennym o rozmaitych częstotliwościach i liczbach faz. Najszersze rozpowszechnienie znalazły łączy zasilane prądem przemiennym. Przy

zasilaniu wielofazowym, a właściwie prawie wyłącznie trójfazowym, pracują indukcyjne silniki pierścieniowe większych mocy, tworząc łąca zwane wałami elektrycznymi, przeznaczone przede wszystkim do użytkowania ściśle synchronicznego obracania się kilku maszyn. Wały elektryczne ze względu na wielkość stosowanych przy ich zestawianiu maszyn nie wchodzi w zakres niniejszego opracowania i nie będą rozpatrywane. Zauważymy jedynie, że charakterystyki momentów tych łączy zależą od kierunku wirowania maszyn względem kierunku wirowania pola magnetycznego, wytwarzanego przez ich trójfazowe uzwojenia zasilające. Łąca zasilane jednofazowym prądem przemiennym bywają zestawiane z maszyn miniaturowych, a ich cechą charakterystyczną jest niezależność charakterystyk od kierunku obrotów. Spowodowane jest to tym, że zasilające jednofazowe uzwojenie nie wytwarza pola wirującego w określonym kierunku, lecz oscylujące wzdłuż określonej osi.

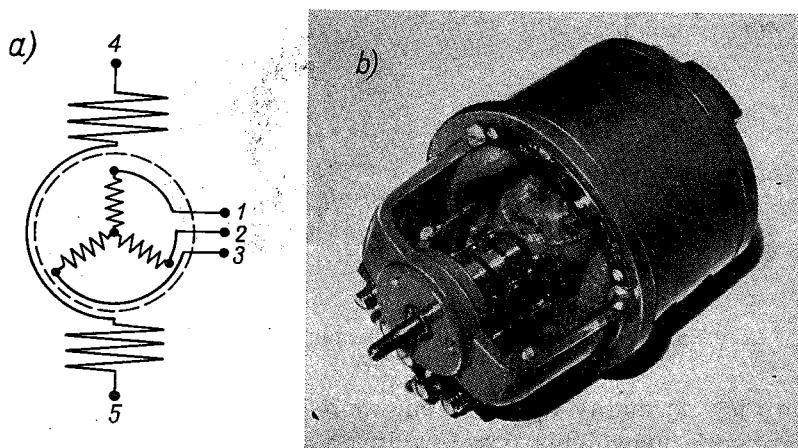
Łąca synchroniczne mogą być używane do zdalnego przekazywania położenia kąowego lub momentu obrotowego. W pierwszym przypadku wał odbiornika niesie na sobie tylko starannie wyważoną wskazówkę, wobec czego możemy uważać, że pożyteczny moment obciążenia nie istnieje, a uchyb powodowany jest tylko tarciem i momentami pasożytniczymi. W drugim przypadku — przekazywania momentu — poza momentami pasożytniczymi momentowi synchronizującemu przeciwstawia się moment pożyteczny spowodowany bezpośrednim lub pośrednim obciążeniem wału odbiornika. Oba rodzaje pracy łąca mogą odbywać się w stanie quasistatycznym — przy powolnych zmianach położenia wałów — i kinematycznym — przy wirowaniu wałów z określoną prędkością obrotową. W praktyce częściej spotyka się pracę łączy, która może być uważana za quasistatyczną. Zależności teoretyczne występujące przy rozważaniu tego rodzaju pracy są prostsze, a przejście do określenia stanu kinematycznego umożliwia odpowiednie, również stosunkowo nieskomplikowane wzory i rozważania. Z tego względu wzory określające zachowanie się maszyn łąca wyprowadzane są z reguły dla pracy quasistatycznej.

2. ZASADA BUDOWY I DZIAŁANIA

Selsyn nadawczo-odbiorczy jest indukcyjną maszyną elektryczną o obwodzie pierwotnym zasilanym jednofazowo i nie połączonym z nim elektrycznie obwodzie wtórnym, nawiniętym trójfazowo. Jednofazowe zasilanie obwodu pierwotnego zapewnia niezależność charakterystyk selsyna od kierunku obrotów. Trójfazowe uzwojenie wtórne daje najbardziej statyczną pracę łąca i uniezależnia wartości momentów od położenia uzwojeń pierwotnych względem wtórnych [46]. Momenty przenoszone przez jednofazowo zasilane łące selsynowe o trójfazowych uzwojeniach wtórnych zależą więc tylko od kąta niezgodności δ° .

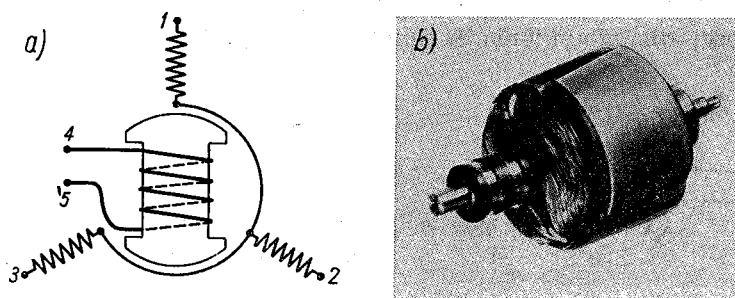
Zasilane uzwojenie pierwotne może być umieszczone na wirniku lub na stojanie, a wtórne odpowiednio na stojanie lub wirniku. Zasada działania i przebieg zjawisk elektromagnetycznych zachodzących w selsynie nie zależą od tych różnic konstrukcyjnych.

Uzwojenie pierwotne selsyna ma wytworzyć pole magnetyczne oscylujące wzdłuż osi o określonym w przestrzeni kierunku. Dla spełnienia tego zadania w przypadku umieszczenia obwodu pierwotnego na stojanie



Rys. 2. Selsyn z wydatnymi biegunami na stojanie

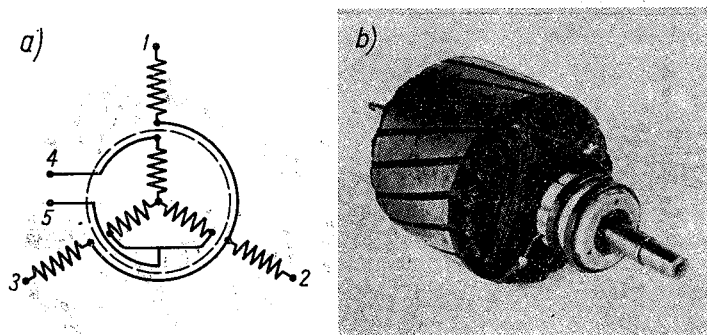
uzwojenie to wykonuje się w postaci cewek nawiniętych na wydatnych biegunami — rys. 2. Widzimy, że poza wytoczeniem stojana przewidzianym na umieszczenie wirnika, wewnątrz selsyna pozostaje przy tym



Rys. 3. Selsyn z wydatnymi biegunami na wirniku

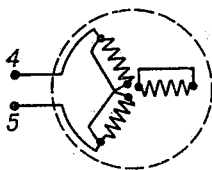
rozwiązaniu stosunkowo dużo niewykorzystanego miejsca. Nieruchome uzwojenie pierwotne jest zasilane poprzez zaciski umieszczone na obudowie selsyna. Trzy końce uzwojenia wtórnego umieszczonego w tym przypadku na wirniku wyprowadzane są na zewnątrz poprzez trzy pierścienie i ślizgające się po nich szczotki. W przypadku umieszczenia ob-

wodu pierwotnego na wirniku jego zasilanie odbywa się poprzez szczotki ślizgające się po dwóch pierścieniach. Dla wytworzenia pola magnetycznego oscylującego wzdłuż określonego kierunku — osi magnetycznej obwodu wirnika — wirnik wykonuje się albo z biegunami wydajnymi i skupionym uzwojeniem, jak na rys. 3, albo z biegunami utajonymi i uzwojeniem rozłożonym, rys. 4. W tym ostatnim przypadku jest on



Rys. 4. Selsyn z utajonymi biegunami na wirniku

wykonany jak normalny wirnik bębnowy maszyny elektrycznej nawinięty trójfazowo — uzwojenia poszczególnych faz są przesunięte względem siebie w przestrzeni o 120° . Fazy połączone są tak, aby stworzyć strumień o określonym w przestrzeni kierunku, a prostopadły do niego kierunek objąć zwartym obwodem. Należy podkreślić, że przy takim połączeniu prądy płynące w uzwojeniach wirnika mają zawsze jedną i tę samą fazę w czasie. Działanie wirnika z utajonymi biegunami jest równoważne działaniu wirnika z wydajnymi biegunami: nierównomierność szczeliny powietrznej lub działanie zwartego obwodu powodują różnicę przewodności dla strumienia magnetycznego w kierunkach podłużnym



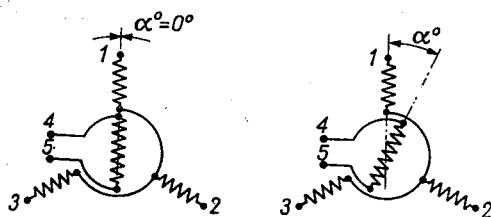
Rys. 5. Odmiana połączenia faz wirnika z utajonymi biegunami

i poprzecznym wirnika. Założenie uzwojenia skupionego czy rozłożonego nie wpływa na tok rozważań. Najchętniej stosowany układ połączeń faz wirnika, spełniający omówione wymagania, jest podany na rys. 4. Inny układ połączeń uzwojenia wirnika z biegunami utajonymi jest przedstawiony na rys. 5. Łatwo udowodnić, że obydwa układy są identyczne pod względem elektromagnetycznym pod warunkiem zachowania stosunku zwojów w fazie 2 : 1,5. Możliwe są również inne połączenia

uzwojeń, jak na przykład — dwie fazy połączone szeregowo, a trzecia zwarta itp.

Przy umieszczeniu obwodu pierwotnego na wirniku otrzymuje się mniejsze objętości i ciężar selsyna przy tej samej sztywności, a ponadto zmniejsza się liczba styków ślizgowych stanowiących główne źródło szkodliwych momentów tarcia. Dlatego to nowocześniejsze rozwiązanie jest korzystniejsze i właśnie ono będzie wyłącznie rozpatrywane w dalszym ciągu, tym bardziej że pod względem elektromagnetycznym wszystkie spotykane i omówione rozwiązania konstrukcyjne są równoważne, dzięki czemu postać rozważań jest identyczna i dla dowolnego rozwiązania. Dlatego parametry dotyczące obwodu pierwotnego będą oznaczane indeksem, którego pierwszą cyfrą będzie 1, niezależnie od miejsca umieszczenia tego obwodu, a obwodu wtórnego — 2.

Jeżeli uzwojenia pierwotne zasilic napięciem przemiennym, to w obwodzie magnetycznym selsyna powstanie oscylujący strumień magnetyczny i we wszystkich uzwojeniach obwodu wtórnego zostaną indukowane siły elektromotoryczne, których wielkość będzie zależała od położenia tych uzwojeń względem kierunku pola magnetycznego wytworzo-



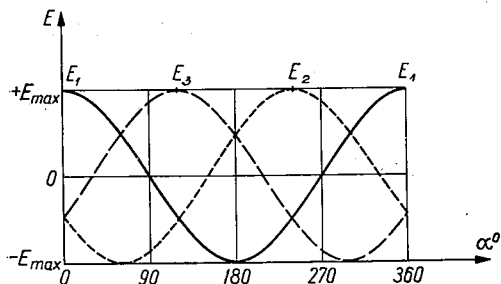
Rys. 6. Wzbudzenie się selsyna

nego przez obwód pierwotny, a fazy będą zgodne lub przeciwne (przesunięte o 180°). Na przykład w przypadku zgodności osi uzwojenia 1 z kierunkiem pola magnetycznego (rys. 6) indukowana w nim siła elektromotoryczna będzie maksymalna E_m , natomiast semne w uzwojeniach 2 i 3 będą wynosiły po $\frac{1}{2} E_m$ i będą miały przeciwny zwrot, czyli będą przesunięte w fazie o 180° . Przy obracaniu wirnika selsyna niosącego obwód pierwotny siły elektromotoryczne w poszczególnych uzwojeniach wtórnych umieszczonych na stojanie będą się zmieniały w miarę zmiany kąta α° między osią magnetyczną wirnika i uzwojeniem 1 stojana. Ich wartości skuteczne będą określone wzorami:

$$\begin{aligned} E_{21} &= E_m \cos \alpha^\circ, \\ E_{22} &= E_m \cos (\alpha^\circ - 120^\circ), \\ E_{23} &= E_m \cos (\alpha^\circ + 120^\circ), \end{aligned} \quad (5)$$

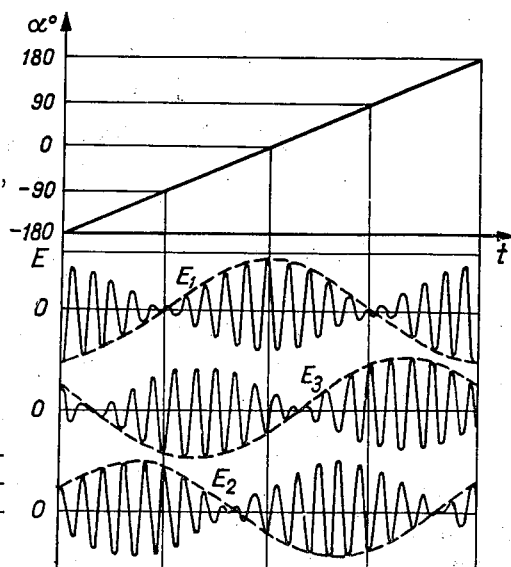
których interpretację graficzną podaje rys. 7.

Podczas obracania selsyna z pewną stałą prędkością obrotową wartości chwilowe semnych obwodu wtórnego odpowiednio zgodne lub przeciwne w fazie zmieniają się w czasie jak na rys. 8. Siły elektromotoryczne są modulowane sinusoidalnymi obwiedniami, których częstotliwość jest określona prędkością obrotową, a wzajemne przesunięcia fazowe pozostają równe 120° .



Rys. 7. Siły elektromotoryczne wzbudzane w uzwojeniach wtórnych selsyna

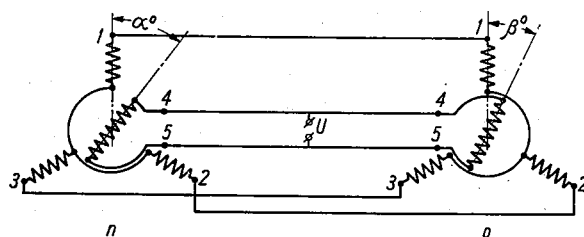
Nadawczo-odbiorcze łącze selsynowe można zestawić z dwóch lub kilku jednakowych albo rozmaitych selsynów. Rozpatrywanie zjawisk zachodzących w łączy przeprowadza się zazwyczaj zakładając, że łącze jest zestawione z dwóch identycznych jednostek. Wartości wielkości



Rys. 8. Zmiany w czasie wartości chwilowych sił elektromotorycznych wzbudzanych w uzwojeniach wtórnych selsynów podczas obrotów wirnika

otrzymanych dla takiego łącza są nazywane wartościami charakterystycznymi i są miarodajne dla każdego z identycznych selsynów łącza. Przejście do odpowiednich wartości łącza zestawionego z różnych jednostek przy znajomości wartości charakterystycznych każdej jednostki umożliwia odpowiednie wzory, które zostaną przytoczone dalej.

Jeżeli wirniki dwóch identycznych selsynów zasilić z jednego źródła napięcia przemiennego, a uzwojenia stojanów połączyć przeciwko sobie, jak na rys. 9, to przy jednakowym położeniu uzwojeń wirników względem uzwojeń stojanów (na rys. 9 odpowiadałoby to przypadkowi $\alpha^\circ = \beta^\circ$) suma sił elektromotorycznych, działających w każdym z trzech obwodów złożonych z uwojeń stojanów i przewodów łączących, będzie równała się zeru. Jeżeli natomiast wirnik jednego z selsynów obrócić w dowolnym kierunku o pewien kąt, położenie uzwojeń faz sto-



Rys. 9. Łączy selsynowe nadawczo-odbiorcze

jana względem pola wirnika w tym selsynie zmieni się, wielkości sił elektromotorycznych indukowanych w tych uzwojeniach zmienią się także, siły elektromotoryczne w obwodach stojanów przestaną się równoważyć i w wyniku działania wypadkowych sił elektromotorycznych, proporcjonalnych do cosinusa kąta niezgodności osi uzwojeń wirników, popłyną w obwodach stojanów prądy. Te prądy we współdziałaniu z polem magnetycznym wirnika wytworzą moment, usiłujący obrócić wirniki w kierunku zanikania momentu, czyli do zgodnego położenia, podobnie jak to ma miejsce przy pracy równoległej maszyn synchronicznych, a w szczególności przy współpracy prądnicy synchronicznej z silnikiem synchronicznym.

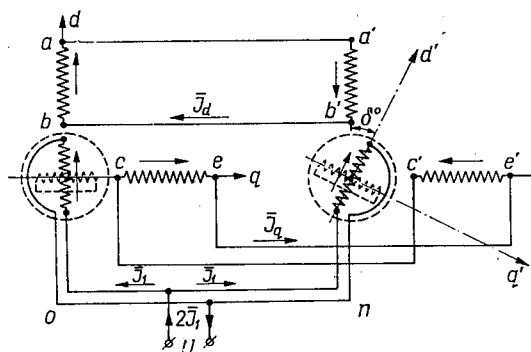
Jeżeli obracany wirnik zostanie odchylony od położenia równowagi i utrzymany w tym nowym położeniu, to wirnik drugiego selsyna przyjmie natychmiast odpowiednio zgodne położenie. Selsyn, którego wirnik jest ustawiany w zadane położenie kątowe, nazywa się selsynem nadawczym lub nadajnikiem (*n*), selsyn powtarzający to położenie kątowe nazywa się selsynem odbiorczym lub odbiornikiem (*o*).

Konstrukcja odbiornika jest w zasadzie identyczna z konstrukcją nadajnika; jedynie w przypadku selsynów zasilanych napięciami o niskich częstotliwościach, rzędu 50–60 Hz, stosuje się czasem mechaniczne tłumiki drgań. Przy tych częstotliwościach zasilania podczas występowania większych prędkości obrotowych lub chwilowych przyspieszeń kątowych wirnika w selsynie odbiorczym mogą powstać drgania. Jeżeli nie zostaną one dostatecznie szybko i skutecznie stłumione powstaje moment asynchroniczny, który powoduje obracanie się wirnika odbior-

nika z pewną stałą prędkością, niezależną od zachowywania się nadajnika. Tendencja do występowania drgań niskoczęstotliwościowych selsynów zależy od ich konstrukcji i parametrów elektromagnetycznych oraz sztywności i występuje nie zawsze. Przy większych częstotliwościach zasilania — rzędu 300÷500 Hz — zjawisko drgań praktycznie nie występuje i tłumiki mechaniczne nie są potrzebne. Dokładniejszemu omówieniu sprawy tłumienia drgań selsynów będzie poświęcony punkt 3.1.5.

3. ELEMENTY TEORII

Podczas rozwoju teorii selsynów wyodrębniły się dwie metody ich badania i opisywania zachodzących zjawisk. Jedna traktująca selsyny jako maszyny synchroniczne w oparciu o teorię dwóch reakcji, [30], [33]; druga — traktująca selsyny jako maszyny indukcyjne w oparciu o teorię pól wirujących [46], [47]. Pierwsze z podanych publikacji są pracami inauguracyjnymi daną metodę analizy selsynów, drugie — najbardziej pełnymi jej ujęciami. Każda z obydwu metod pozwala na odpowiednio łatwiejsze otrzymanie jednych z rozpatrywanych zależności przy skomplikowaniu postaci innych, tak że obydwie uzupełniają się wzajemnie. Wyrażenia matematyczne ważniejszych wyników końcowych otrzymanych obydwoma metodami dają się stosunkowo łatwo sprowadzić do siebie, co stanowi wzajemny sprawdzian słuszności przeprowadzanych rozważań.



Rys. 10. Zastępczy dwufazowy układ uzwojeń łąca

Pełna analiza zjawisk zachodzących w selsynie, a więc dotycząca dowolnie dużych kątów niezgodności, prowadzi do bardzo skomplikowanych zależności matematycznych. Dla uproszczenia tych zależności zamiast rzeczywistego trójfazowego układu wtórnych uzwojeń rozpatruje się równoważny dwufazowy układ uzwojeń sprowadzonych do wzajemnie prostopadłych osi skierowanych zgodnie z podłużną (d) i poprzeczną (q) osią uzwojenia pierwotnego (rys. 10). Można udowodnić, że

przy pomocy odpowiednich wzorów redukcyjnych dla dowolnego stanu pracy selsyna parametry uzwojeń wszystkich jego obwodów dadzą się sprowadzić do tych nieruchomych osi. Otrzymane parametry zastępczych obwodów można wyrazić przez parametry rzeczywistego selsyna dające się dla niego obliczyć i pomierzyć.

Wypisując układ równań dla poszczególnych obwodów układu zastępczego i rozwiązując go względem badanego parametru można określić jego zachowanie się w dowolnych założonych warunkach, a więc i przy dowolnych kątach niezgodności. Ponieważ jednak przytaczanie wspomnianych zależności spowodowałoby znaczne przekroczenie przewidzianej objętości artykułu, pominiemy je, ograniczając się do ważnych w zastosowaniach praktycznych, a jednocześnie najprostszych, wyprowadzeń i podania bez dowodów wniosków z ogólnej teorii w miarę ich przydatności do dalszych rozważań. Szczegółowe wyprowadzenia tej teorii można znaleźć w powołanej literaturze, a przede wszystkim w [47] i [33].

W dalszym ciągu rozważań podstawowe występujące w rzeczywistości oporności selsyna dające się dlań obliczyć i pomierzyć zostaną oznaczone jak następuje:

A. Po stronie pierwotnej

- A.1. Dla jednofazowego uzwojenia wzbudzającego oporności — pozorna, czynna i bierna odpowiednio:

$$Z_1 = R_1 + jX_1, \quad (6a)$$

z tym, że oporność bierna składa się z oporności samoindukcji i oporności rozproszenia:

$$X_1 = X_{1d} + X_{1s}; \quad (6b)$$

oporność bierna indukcji wzajemnej tego uzwojenia i fazy uzwojenia wtórnego przy zgodnych kierunkach ich osi:

$$X_{a1}. \quad (6c)$$

- A.2. Dla ewentualnego zwartego obwodu tłumiącego odpowiednio:

$$\begin{aligned} Z_{t1} &= R_{t1} + jX_{t1}, \\ X_{t1} &= X_{t1d} + X_{t1s}, \\ X_{at1}. \end{aligned} \quad (7)$$

B. Po stronie wtórnej

- B.1. Dla trójfazowego uzwojenia synchronizującego oporności fazowe:

$$Z_2 = R_2 + jX_2. \quad (8)$$

- B.2. Oporności synchroniczne przy braku zwartego obwodu tłumiącego:

$$\begin{aligned} \text{podłużna} \quad Z_d &= R_d + jX_d, \\ \text{poprzeczna} \quad Z_q &= R_q + jX_q; \end{aligned} \quad (9)$$

oporności synchroniczne przy istnieniu zwartego obwodu tłumiącego:

$$\begin{aligned} \text{podłużna} \quad Z'_d &= R'_d + jX'_d, \\ \text{poprzeczna} \quad Z'_q &= R'_q + jX'_q, \end{aligned} \quad (10)$$

z tym że:

$$\left. \begin{aligned} R'_d &= R_2 + \xi_1 R_1 \\ X'_d &= X_d - \xi_1 X_1 \end{aligned} \right\}, \quad \text{gdzie } \xi_1 = \frac{3}{2} \cdot \frac{X_{d1}^2}{Z_1^2}; \quad (11)$$

$$\left. \begin{aligned} R'_q &= R_2 + \xi_{11} R_{11} \\ X'_q &= X_q - \xi_{11} X_{11} \end{aligned} \right\}, \quad \text{gdzie } \xi_{11} = \frac{3}{2} \cdot \frac{X_{d11}^2}{Z_{11}^2}. \quad (12)$$

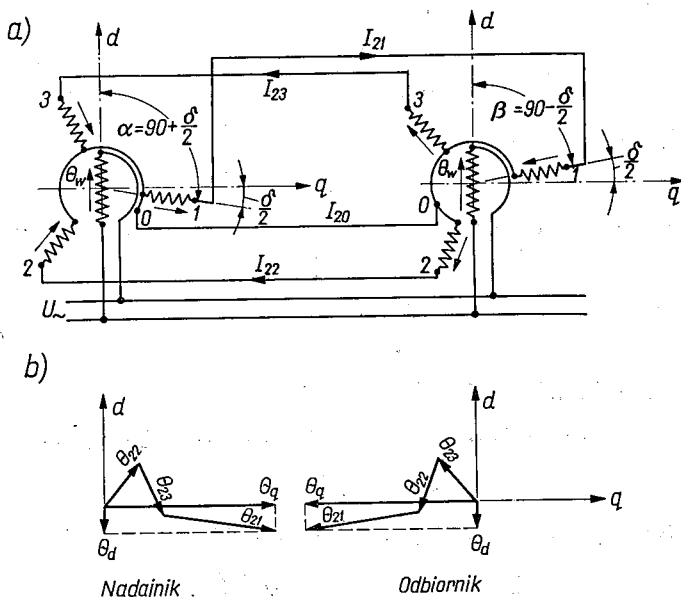
3.1. Łącze charakterystyczne (symetryczne)

3.1.1. Praca statyczna

Rozpatrzmy łącze charakterystyczne w warunkach pracy quasistatycznej i przy małym kącie niezgodności, jako przypadek najczęściej spotykany w praktyce i mogący posłużyć jako podstawa do wyprowadzenia zależności dla innych warunków pracy łącza.

Dla uproszczenia zależności założymy wzajemne położenie uzwojeń selsynów jak na rys. 11. Kąty przesunięć uzwojeń selsynów wyniosą

$$\alpha^\circ = 90^\circ + \frac{\delta^\circ}{2}; \quad \beta^\circ = 90^\circ - \frac{\delta^\circ}{2};$$



Rys. 11. Prądy i przepływy w uzwojeniach selsynów przy kącie niezgodności δ

różnice napięć w fazach uzwojeń wtórnych [zob. wzory (5)] wyniosą:

$$\begin{aligned} \hat{E}_{21} &= \hat{E}_{21o} - \hat{E}_{21n} = \hat{E}_{2f} [\cos \beta^\circ - \cos \alpha^\circ] = 2\hat{E}_{2f} \cdot \sin \frac{\delta^\circ}{2} = \Delta \hat{E}_2, \\ \hat{E}_{22} &= \hat{E}_{22o} - \hat{E}_{22n} = \hat{E}_{2f} [\cos (\beta^\circ - 120^\circ) - \cos (\alpha^\circ - 120^\circ)] = \\ &= -\hat{E}_{2f} \sin \frac{\delta^\circ}{2} = -\frac{\Delta \hat{E}_2}{2}, \\ \hat{E}_{23} &= \hat{E}_{23o} - \hat{E}_{23n} = \hat{E}_{2f} [\cos (\beta^\circ + 120^\circ) - \cos (\alpha^\circ + 120^\circ)] = \\ &= -\hat{E}_{2f} \sin \frac{\delta^\circ}{2} = -\frac{\Delta \hat{E}_2}{2}. \end{aligned} \quad (13)$$

Największa różnica napięć wystąpi przy założonym kącie niezgodności w fazie 1, w fazach 2 i 3 różnica napięć jest o połowę mniejsza i przeciwnie skierowana. Prądy wyrównawcze płynące w uzwojeniach faz wtórnego obwodu wyniosą:

$$\begin{aligned} I_{21} &= \frac{\Delta \hat{E}_{21}}{2Z_2} = \frac{\Delta \hat{E}_{2f}}{Z_2} \cdot \sin \frac{\delta^\circ}{2}, \\ I_{22} &= \frac{\Delta \hat{E}_{22}}{2Z_2} = -\frac{\Delta \hat{E}_{2f}}{2Z_2} \cdot \sin \frac{\delta^\circ}{2}, \\ I_{23} &= \frac{\Delta \hat{E}_{23}}{2Z_2} = -\frac{\Delta \hat{E}_{2f}}{2Z_2} \cdot \sin \frac{\delta^\circ}{2}. \end{aligned} \quad (14)$$

Prąd w przewodzie zerowym $I_{20} = I_{21} + I_{22} + I_{23} = 0$, z czego wynika, że ten przewód zerowy nie jest potrzebny. W rzeczywistości nie stosuje się przewodu zerowego w łączach selsynowych, a na rysunku podano go tylko dla ułatwienia rozumowania.

Na rys. 11 podane są również wykresy wektorowe przepływów Θ_{21} , Θ_{22} i Θ_{23} w uzwojeniach wtórnego obwodu synchronizującego selsynów nadawczego i odbiorczego. Z rysunku wynika, że wypadkowy przepływ jest skierowany wzdłuż osi uzwojenia fazy 1, w której płynie maksymalny prąd, i jest półtorakrotnie większy od przepływu fazy 1. Wartość średnia wypadkowego przepływu Θ_{2w} przeliczona z wartości skutecznej wyniesie:

$$\Theta_{2w} = 1 \frac{1}{2} \Theta_{21} = \frac{3}{2} \cdot \frac{2\sqrt{2}}{\pi} \cdot 2z_2 k_{z2} I_{21} = 2,7 \frac{E_{2f} z_2 k_{z2}}{Z_2} \sin \frac{\delta}{2},$$

gdzie z_2 i k_{z2} odpowiednio liczba zwojów i współczynnik uzwojenia dla uzwojeń obwodu wtórnego.

Jeżeli kierunek zgodny z osią magnetyczną uzwojenia wzbudzenia nazwiemy kierunkiem podłużnym (d), a prostopadły do niego — poprzecznym (q), to rozkładając wypadkowy przepływ na te kierunki otrzymamy:

Przepływ podłużny:

$$\Theta_d = \Theta_{2w} \sin \frac{\delta}{2} = 2,7 \frac{E_{2f} z_2 k_{z2}}{Z_2} \sin^2 \frac{\delta}{2} = 1,35 \frac{E_{2f} \cdot z_2 k_{z2}}{Z_2} (1 - \cos \delta). \quad (15)$$

Przepływ poprzeczny:

$$\Theta_q = \Theta_{2w} \cos \frac{\delta}{2} = 2,7 \frac{E_{2f} z_2 k_{z2}}{Z_2} \sin \frac{\delta}{2} \cos \frac{\delta}{2} = 1,35 \frac{E_{2f} z_2 k_{z2}}{Z_2} \sin \delta. \quad (16)$$

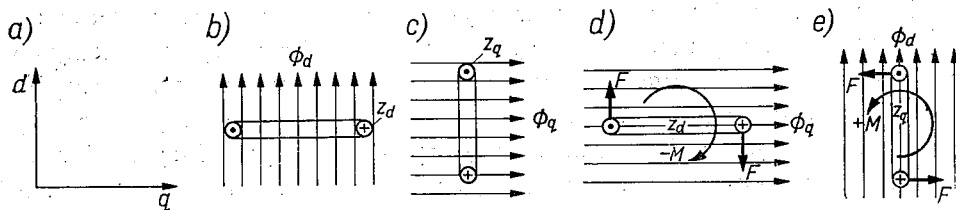
Przepływ podłużny wytwarza strumień magnetyczny skierowany zarówno w nadajniku, jak i w odbiorniku przeciwnie do strumienia wzbudzenia. Przepływ poprzeczny wytwarza strumień magnetyczny skierowany prostopadłe do strumienia wzbudzenia.

Przy małych kątach niezgodności przepływy podłużne są znikomo małe w porównaniu do poprzecznych. Na przykład przy $\delta^\circ = 5^\circ$

$$\frac{\Theta_d}{\Theta_q} = \frac{1 - \cos 5^\circ}{\sin 5^\circ} = 0,026. \quad (17)$$

Z przytoczonego rozumowania wynika, że przy małych kątach niezgodności maksymalna wartość prądu wyrównawczego występuje w fazie 1, której oś magnetyczna skierowana jest prawie zgodnie z osią poprzeczną selsyna. Uzwojenie tej fazy wytwarza więc prawie wyłącznie przepływ poprzeczny. Dlatego oporność pozorna Z_2 występująca we wzorach odpowiada oporności synchronicznej poprzecznej Z_q znanej z teorii maszyn synchronicznych.

Obecnie możemy rozpatrzyć powstawanie momentu synchronizującego w łączu. Przyjmujemy zwrot dodatni osi d i q jak na rys. 12a.



Rys. 12. Współdziałanie podłużnych i poprzecznych przepływów z podłużnym i poprzecznym strumieniem magnetycznym

Kierunki przepływów będziemy uważali za dodatnie, jeżeli wytwarzane przez nie strumienie będą skierowane zgodnie z tymi osiami. Dodatni kierunek obrotu i momentu obrotowego — jak zazwyczaj przeciwny do ruchu wskazówek zegara. Na rys. 12b i c pokazane są kombinacje wzajemnego usytuowania podłużnych i poprzecznych strumieni z zezwojami o podłużnych i poprzecznych kierunkach osi magnetycznych. Uzwojenie o podłużnie usytuowanej osi magnetycznej z podłużnym strumieniem

(rys. 12b) i o poprzecznie usytuowanej osi z poprzecznym strumieniem (rys. 12c) nie wytwarzają momentu obrotowego. Natomiast uzwojenie o podłużnie usytuowanej osi z poprzecznym strumieniem (rys. 12d) i o poprzecznie usytuowanej osi z podłużnym strumieniem (rys. 12c) wytwarzają momenty — pierwszy ujemny, a drugi — dodatni. Wypadkowy moment synchronizujący wyniesie więc:

$$M = k_1(\Phi_d \Theta_q - \Phi_q \Theta_d). \quad (18)$$

Jak wynika z (17), dla małych δ można uważać $\Theta_d \approx 0$ i wtenczas

$$M = k_1 \Phi_d \Theta_q. \quad (19)$$

Dla rozłożonych sinusoidalnie i sinusoidalnie oscylujących w czasie strumienia i przepływu wielkość momentu synchronizującego może być określona ze wzoru

$$M = \frac{\pi}{8} \cdot \frac{10^{-3}}{9,81} \Phi_d \Theta_q \cos \psi \quad [\text{cm G}], \quad (20)$$

gdzie:

$$\Phi_d = \frac{E_{2f} \cdot 10^8}{4,44 z_1 k_{1z} f} \text{ — strumień magnetyczny podłużny,}$$

$$\Theta_q = 1,35 \frac{E_{2f} z_2 k_{z2}}{Z_q} \sin \delta \text{ — przepływ poprzeczny,} \quad (21)$$

$$\cos \psi = \frac{X_q}{Z_q} \text{ — kąt przesunięcia pomiędzy } \Phi_d \text{ i } \Theta_q.$$

Ponadto:

$$Z_q = \sqrt{R_q^2 + X_q^2}, \quad (22)$$

gdzie R_q i X_q czynne i bierne składowe synchronicznej oporności poprzecznej Z_q .

Podstawiając (21) i (22) do (20) otrzymamy po uporządkowaniu:

$$M = 1217 \frac{E_{2f}^2}{f} \cdot \frac{X_q}{R_q^2 + X_q^2} \sin \delta \quad [\text{cm G}]. \quad (23)$$

Dzięki założeniu małych kątów niezgodności możemy $\sin \delta$ zastąpić wartością δ w radianach otrzymując:

$$M = 1217 \frac{E_{2f}^2}{f} \cdot \frac{X_q}{R_q^2 + X_q^2} \cdot \delta \quad [\text{cm G}]. \quad (24)$$

Sztywność łączy wyniesie:

$$S = 1217 \frac{E_{2f}^2}{f} \cdot \frac{X_q}{R_q^2 + X_q^2} \left[\frac{\text{cm G}}{\text{radian}} \right]. \quad (25)$$

Po zastąpieniu radianów stopniami kątowymi:

$$S = 1217 \cdot \frac{2\pi}{360} \cdot \frac{E_{2f}^2}{f} \cdot \frac{X_q}{R_q^2 + X_q^2} = 21,2 \frac{E_{2f}^2}{f} \cdot \frac{X_q}{R_q^2 + X_q^2} \quad [\text{cm G}/^\circ]. \quad (26)$$

Z rys. 11 wynika, że przepływ poprzeczny nadajnika jest dodatni, a odbiornika — ujemny; odpowiednie znaki będą miały również wytwarzane przez nie momenty synchronizujące starające się sprowadzić ich wirniki do zgodnego położenia. Można udowodnić również, że momenty synchronizujące nadajnika i odbiornika będą zawsze równe sobie i dlatego łącze selsynowe w podanym podstawowym układzie połączeń może być tylko przenośnikiem momentu, a nie może samo wytworzyć jakiegokolwiek momentu.

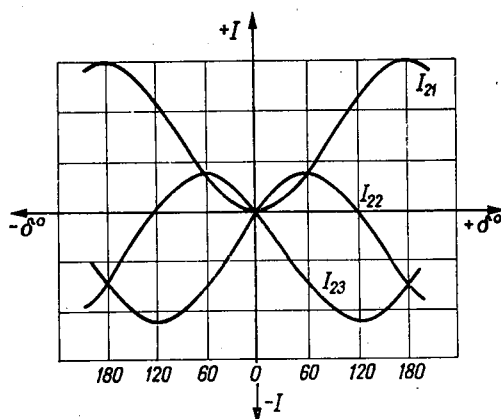
Po stronie pierwotnej selsyna może znajdować się zwarty obwód tłumiący, którego rola zostanie omówiona później, o oporności pozornej $Z_{t1} = R_{t1} + jX_{t1}$ oraz oporności biernej indukcji wzajemnej z fazą uzwojenia wtórnego przy zgodności ich osi — X_{at1} .

Jeżeli selsyn zawiera taki obwód tłumiący, to wzór (26) przekształci się do postaci:

$$S = 21,2 \frac{E_{2f}^2}{f} \cdot \frac{X'_q}{R_q'^2 + X_q'^2} \quad [\text{cm G}/^\circ], \quad (27)$$

gdzie $R_q' = R_2 + \xi_{t1}R_{t1}$ oraz $X_q' = X_q - \xi_{t1}X_{t1}$, a z kolei

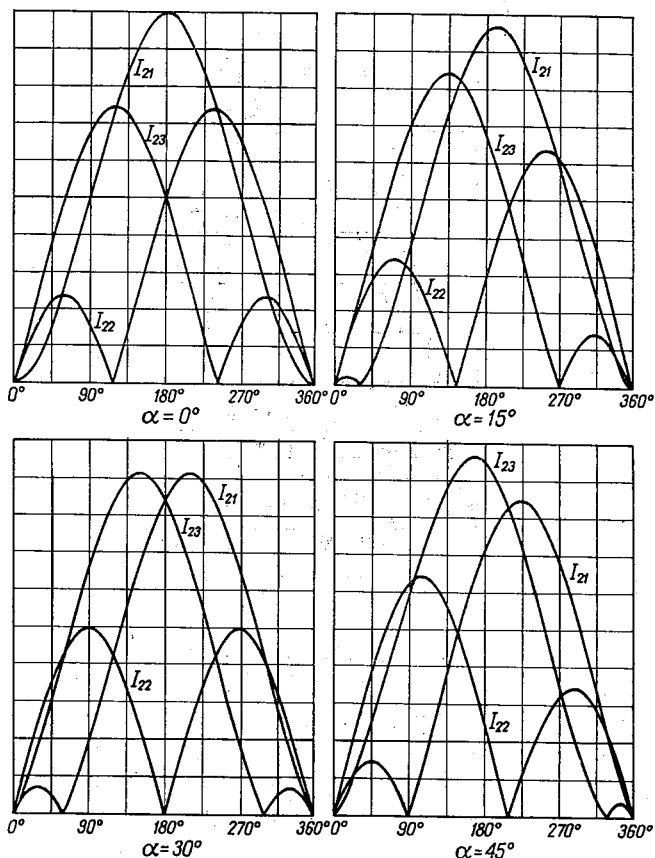
$$\xi_{t1} = \frac{3}{2} \cdot \frac{X_{at1}^2}{R_{t1}^2 + X_{t1}^2} = \frac{3}{2} \frac{X_{at1}}{Z_{t1}^2}. \quad (28)$$



Rys. 13. Kształt przebiegów prądów w trójfazowych uzwojeniach wtórnych

Kształt przebiegów prądów w trójfazowym uzwojeniu wtórnym selsyna w zależności od kąta niezgodności przy założeniu $\alpha^\circ = 0^\circ$ (zob. rys. 9) podaje rys. 13.

Na podstawie rys. 13 można wykreślić przebiegi bezwzględnych wartości prądów w uzwojeniach wtórnych selsyna przy $\alpha^\circ = 0^\circ$. Można udowodnić, że dla kątów $\alpha^\circ \neq 0^\circ$ wartości ekstremalne krzywych przebiegów poszczególnych prądów przesuwają się względem wartości podanych na rys. 13 wzdłuż osi odciętych o kąt równy kątowi α° , w kierunku zgodnym z jego zmianą. W kierunku osi rzędnych natomiast, krzywe przesuwają się w sposób zapewniający ich przejście przez zero dla



Rys. 14. Przesunięcia przebiegów prądów w trójfazowych uzwojeniach wtórnych przy $\alpha \neq 0$

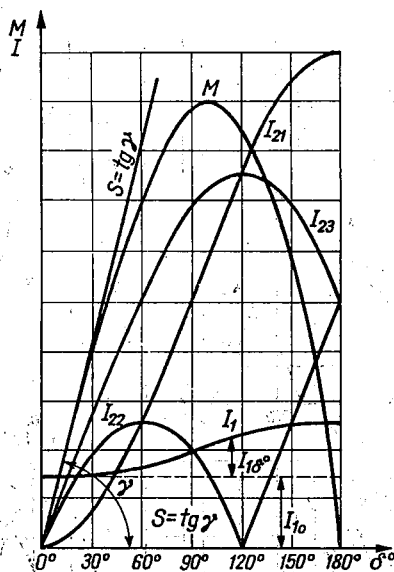
$\delta^\circ = 0^\circ$ przy nowym miejscu rozmieszczenia ekstremów. Odpowiednie wykresy zmian prądów wtórnych dla kilku wybranych wartości kąta α° podaje rys. 14.

Przebiegi wielkości charakteryzujących selsyn, w zależności od kąta niezgodności δ° , podane są na rys. 15. Przebiegi prądów wtórnych odpowiadają przypadkowi $\alpha^\circ = 0^\circ$.

Poza omówionymi poprzednio oznaczeniami: I_{10} — składowa prądu magnesującego biegu jałowego, $I_{1\delta}$ — składowa kompensująca oddziaływanie obwodu wtórnego.

Na podstawie przytoczonych uproszczonych rozważań można stwierdzić:

1. Przy zgodnym położeniu wirników maszyn łączy zarówno w odbiorniku, jak i w nadajniku występuje tylko strumień podłużny.



Rys. 15. Przebiegi momentów i prądów selsyna w funkcji kąta niezgodności δ

2. Przy istnieniu kąta niezgodności w uzwojeniach obwodu wtórnego łączy powstają przepływy wywołujące strumienie poprzeczny i podłużny. Strumień podłużny przeciwdziałający strumieniowi wzbudzenia reprezentuje oddziaływanie obwodu wtórnego i jest kompensowany zwiększonym przepływem uzwojenia pierwotnego (wzrost prądu wzbudzenia). Strumień poprzeczny we współdziałaniu z podłużnym przepływem (strumieniem wzbudzenia) wytwarzają momenty synchronizujące, dążące do sprowadzenia wirników selsynów do zgodnego położenia katowego względem ich stojanów.

3. Przy małych kątach niezgodności prądy wyrównawcze płynące w uzwojeniach wtórnych wytwarzają prawie wyłącznie strumień poprzeczny, który we współdziałaniu z podłużnym strumieniem wzbudzenia wytwarza moment synchronizujący.

Ze wzoru (19) wynika, że warunkiem koniecznym powstania momentu synchronizującego jest istnienie przepływu poprzecznego. Jeżeli przez uzwojenia wtórne selsynów będzie płynął prąd wyrównawczy, nie wywołujący przepływu poprzecznego, to nie powstanie żaden moment

obrotowy. Stan taki może wystąpić przy różnicy napięć znamionowych obwodów wtórnych nadajnika i odbiornika przy kącie niezgodności $\delta^\circ = 0^\circ$. Mimo że pod wpływem tej różnicy napięć we wszystkich fazach będą płynęły prądy wyrównawcze nawet przy braku kąta niezgodności, to wytworzą one wyłącznie przepływ podłużny i moment synchronizujący przy $\delta^\circ = 0$ będzie również wynosił zero, tak samo, jak w rozpatrzonym wyżej przypadku równości wtórnych napięć znamionowych selsynów. W niektórych rozwiązaniach łączy nadawczo-odbiorczych stosowane jest łączenie selsynów o różnych napięciach wtórnych. W tym przypadku strumień poprzeczny, powstający przy danym małym kącie niezgodności, jest większy niż przy równości napięć wtórnych, gdyż większe są różnice napięć. Prowadzi to do zwiększenia momentu przy danym kącie niezgodności, a więc i sztywność łącza. Ta zaleta jest równoważona wadą, jaką jest stały przepływ prądu nawet przy zgodnym położeniu selsynów nagrzewającego je niepotrzebnie. W przypadku więc prawidłowego zaprojektowania selsynów przy uwzględnieniu należytego wykorzystania materiałów, a więc również minimum ciężaru i kosztu, zestawianie łączy selsynowych z jednostek o różnych napięciach wtórnych nie powinno być stosowane.

Sztywność selsyna (26) jest proporcjonalna do kwadratu siły elektromotorycznej, a tym samym kwadratu podłużnego strumienia wzbudzenia, gdyż $\Phi_d = cE_f$ (zob. [21]). Ponadto jest proporcjonalna do wyrażenia, które oznaczmy:

$$Y = \frac{X_q}{R_q^2 + X_q^2}. \quad (29)$$

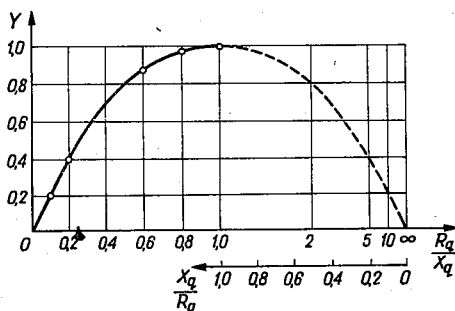
Wartość tego wyrażenia będzie zmieniała się przy zmianach stosunku wchodzących weń parametrów $R_q = X_q$. Przy zmianie X_q wartość Y osiągnie maksimum dla $X_q = R_q$, gdyż:

$$\begin{aligned} \frac{dY}{dX_q} &= \frac{R_q^2 - X_q^2}{R_q^2 + X_q^2} = 0, \\ R_q^2 - X_q^2 &= 0, \\ X_q &= R_q. \end{aligned} \quad (30)$$

Zmianę wartości członu Y przy zmianie stosunku $\frac{R_q}{X_q}$ ilustruje rys. 16, na którym maksymalną wartość Y przyjęto za 1.

Przy obranej na tym rysunku skali osi odciętych krzywa przybiera najbardziej przejrzysty kształt. Do wartości $R_q = X_q$ skala jest równomierna. Przy tej wartości krzywa osiąga maksimum. Dalszą część skali dla $\frac{R_q}{X_q} > 1$ otrzymuje się przez przekształcenie równomiernej skali

stosunków $\frac{R_q}{X_q}$ na skalę odwrotności tego stosunku $\frac{X_q}{R_q} = 1 : \frac{R_q}{X_q}$. Przy przyjęciu takiej skali krzywa ma symetryczną budowę względem wartości maksymalnej. Dzięki temu widać wyraźnie, że zmniejszenie się wartości Y jest takie samo przy n -krotnym stosunku $\frac{R_q}{X_q}$ co i przy $\frac{X_q}{R_q}$.



Rys. 16. Zmiana członu Y w zależności od zmian stosunku $\frac{R_q}{X_q}$

W praktyce ważna jest część krzywej dla zakresu zmiany wartości stosunku $\frac{R_q}{X_q}$ od 0 do 1, gdyż wszystkie spotykane rozwiązania selsynów mają ten stosunek zawarty w podanych granicach. Dalsza część krzywej ma znaczenie raczej orientacyjne — obrazując jedynie kształt funkcji $Y = f(X_q)$ dla całego zakresu zmienności zmiennej niezależnej.

Wpływ zmiany R_q wynika natychmiast z budowy wzoru (29) (R_q^2 w mianowniku): wartość Y przy każdej wartości X_q będzie rosła ze zmniejszaniem się R_q .

Jako wniosek z powyższych rozważań można sformułować podstawowe prawo, którego uwzględnienie ma decydujący wpływ przy projektowaniu selsyna: dla otrzymania maksymalnej sztywności należy zapewnić występowanie największego dopuszczalnego strumienia wzbudzenia (podłużnego) najmniejszej oporności poprzecznej R_q i równej jej oporności synchronicznej poprzecznej $X_q = R_q$, a jeżeli to ostatnie nie jest możliwe — należy dążyć do możliwego zbliżenia wartości X_q do R_q .

Rozpatrując zależności występujące w pierwotnym obwodzie selsyna, a więc w obwodzie wzbudzenia, stwierdzamy, że składowa bierna prądu wzbudzenia pobieranego z sieci idzie na wytworzenie pola magnetycznego selsyna, zaś składowa czynna — na straty w miedzi uzwojeń i stali pakietu. Wynika z tego, że selsyn jest tym lepszy, im mniejszy jest współczynnik mocy $\cos \varphi_1$ pobieranej przez niego, gdyż pożyteczną jest tylko moc bierna potrzebna do wytworzenia strumienia magnetycznego. Wniosek ten dyktuje postępowanie w przypadku potrzeby

odciążenia źródła zasilającego łącznie selsynowe. Mianowicie przyłączenie równoległe do uzwojenia pierwotnego kondensatora o odpowiednio dobranej pojemności może bardzo wydatnie zmniejszyć pobór prądu (oraz mocy) ze źródła zasilającego. Na przykład w przypadku miniaturowego selsyna 500 Hz o sztywności 2 [cmG/°] (Askania 127—205), pobierającego normalnie prąd wzbudzenia 0,55A, równoległe włączenie do obwodu zasilania kondensatora 5 μ F powoduje zmniejszenie tego prądu do 0,1A.

3.1.2. Praca kinematyczna

Poza rozpatrzonym rodzajem pracy quasistatycznej łącznie selsynowe może pracować kinematycznie, przenosząc synchronicznie obrót wału selsyna nadawczego. Ten rodzaj pracy jest często nazywany w literaturze „dynamicznym”. Ponieważ jednak chodzi w tym przypadku o ruch ustalony, właściwsze jest przyjęte tu określenie — „praca kinematyczna”. Podczas pracy kinematycznej prędkości obrotowe zawierają się przeważnie w granicach do około 500÷600 obr/min, rzadziej są wyższe — aż do synchronicznych.

Podczas rozważań zmian sztywności przy pracy kinematycznej prędkość ich wirowania rozpatruje się w postaci prędkości względnej v stanowiącej stosunek rozpatrywanej prędkości do prędkości synchronicznej (3000 obr/min przy 50 Hz; 24000 obr/min przy 400 Hz i 30000 obr/min przy 500 Hz). Ze wzrostem względnej prędkości obrotowej sztywność łączy maleje. Dokładne obliczenia oparte na ogólnej teorii selsynów wykazują, że sztywność kinematyczna z uwzględnieniem spotykanych różnic w stosunkach parametrów selsynów zmniejsza się dla zakresu prędkości wirowania od $v = 0$ do $v = 0,2$ (600 obr/min dla 50 Hz, 4800 obr/min dla 400 Hz i 6000 obr/min dla 500 Hz) co najwyżej do 95%₀ swej wartości przy pracy statycznej. Wynika z tego, że prace selsynów przy prędkościach wirowania nie przekraczających podanych wartości można z dostateczną dla praktyki dokładnością (z uchybem $\leq 5\%$) uważać za quasistatyczną, wyczerpując w ten sposób znaczną większość przypadków pracy łączy przy przenoszeniu ruchu obrotowego.

Dla prędkości względnych w zakresie od $v = 0,2$ do $v = 0,65$ (1950 obr/min przy 50 Hz, 15500 obr/min przy 400 Hz i 19500 obr/min przy 500 Hz) i kątów niezgodności δ° do 40° dostatecznie dokładnie można określić wielkości momentu, a zatem i sztywności z prostego empirycznego wzoru Ellera:

$$M_{kin} = M_{stat, max} \cdot \sin \delta \cdot \cos \frac{\pi}{2} \cdot v. \quad (31)$$

Dla pozostałych przypadków pracy łączy przenoszących ruch obrotowy o prędkościach bliskich synchronicznym zmuszeni jesteśmy korzystać ze znacznie bardziej skomplikowanych wzorów wyprowadzonych na podstawie ogólnej teorii selsynów. Można mianowicie udowodnić, że moment na wale odbiornika przy synchronicznym wirowaniu nadajnika wyrazi się wzorem:

$$M_{osyn} = 21,2 \frac{E_f^2}{f R_2} \cdot \frac{1}{\varrho^2 + \mu^2} [\sqrt{\mu^2 + 1} \sin(\delta + \varphi) - 1] \quad [\text{cm G}], \quad (32)$$

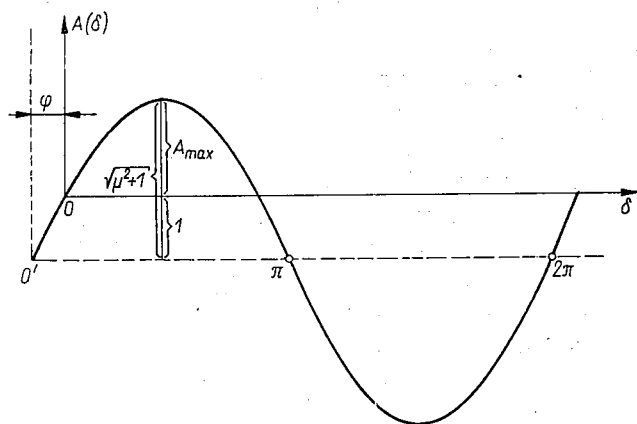
gdzie:

$$\mu = \frac{X'_d + X'_q}{R_2}; \quad \varrho = \frac{R'_d + R'_q - R_2}{R_2}; \quad \varphi = \arctg \frac{1}{\mu},$$

natomiast R'_d , X'_d oraz R'_q i X'_q są określone wzorami (11) i (12). Moment zaś na wale nadajnika wyrazi się wzorem:

$$M_{nsyn} = -21,2 \frac{E_f^2}{f \cdot R_2} \cdot \frac{1}{\varrho^2 + \mu^2} [\sqrt{\mu^2 + 1} \sin(\delta - \varphi) + 1] \quad [\text{cm G}]. \quad (33)$$

Ze wzorów (31) i (33) wynika, że momenty nadajnika i odbiornika przy synchronicznej prędkości obrotowej są sinusoidalnie zależne od kąta niezgodności przy dowolnych wartościach parametrów selsyna i mo-



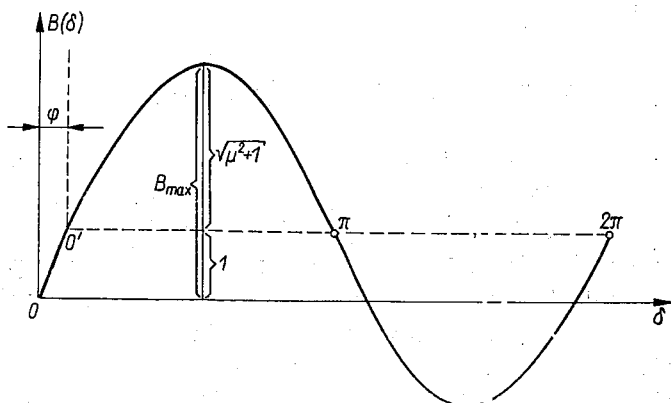
Rys. 17. Zależność momentu synchronizującego odbiornika od kąta niezgodności przy synchronicznym wirowaniu selsynów łączy

gą być łatwo przedstawione graficznie, jak na rysunkach 17 i 18, gdzie przez $A(\delta)$ i $B(\delta)$ oznaczono składniki wzorów na momenty M_{osyn} i M_{nsyn} zależne od kąta niezgodności:

$$\begin{aligned} A(\delta) &= \sqrt{\mu^2 + 1} \sin(\delta + \varphi) - 1, \\ B(\delta) &= \sqrt{\mu^2 + 1} \sin(\delta - \varphi) + 1. \end{aligned} \quad (34)$$

skąd wyrażenia na momenty synchronizujące przyjmują odpowiednio postać:

$$\begin{aligned} M_{osyn} &= 21,2 \frac{E_f^2}{f \cdot R_2} \cdot \frac{1}{\varrho^2 + \mu^2} \cdot A(\delta) \quad [\text{cm G}], \\ M_{nsyn} &= 21,2 \frac{E_f^2}{f \cdot R_2} \cdot \frac{1}{\varrho^2 + \mu^2} \cdot B(\delta) \quad [\text{cm G}]. \end{aligned} \quad (35)$$



Rys. 18. Zależność momentu synchronizującego nadajnika od kąta niezgodności przy synchronicznym wirowaniu selsynów łączy

Sztywność określająca dokładność nadążania odbiornika przy synchronicznych obrotach nadajnika wyrazi się wzorem:

$$S_{syn} = \left(\frac{dM_{odb}}{d\delta} \right)_{\delta=0} = 21,2 \frac{E_f^2}{f \cdot R_2} \cdot \frac{\sqrt{\mu^2 + 1}}{\varrho^2 + \mu^2} \cos \varphi \quad [\text{cm G}/^\circ],$$

a ponieważ $\cos \varphi = \frac{1}{\sqrt{1 + \operatorname{tg}^2 \varphi}} = \frac{\mu}{\sqrt{\mu^2 + 1}}$, to

$$S_{syn} = 21,2 \frac{E_f^2}{f \cdot R_2} \cdot \frac{\mu}{\varrho^2 + \mu^2} \quad [\text{cm G}/^\circ]. \quad (36)$$

Porównując wartości otrzymane z tego wzoru z otrzymanymi ze wzoru (26) można określić zmniejszenie się sztywności danego selsyna przy jego synchronicznym wirowaniu.

Na przykład dla selsyna o utajonych biegunach ze zwartym obwodem tłumiącym w osi poprzecznej o parametrach spełniających równanie:

$$X'_d = X'_q; \quad R'_d = R'_q = 2R_2 \quad \text{ i } \quad X'_q = R'_q$$

otrzymamy $\varrho = 3$; $\mu = 4$

oraz zmniejszenie się sztywności synchronicznej w stosunku do statycznej:

$$\frac{S_{syn}}{S} = \frac{1}{R_2} \cdot \frac{\mu}{\varrho^2 + \mu^2} \cdot \frac{R_q'^2 + X_q'^2}{X_q'} = \frac{4}{3^2 + 4^2} \cdot \frac{8 R_2^2}{2 R_2^2} = 0,64.$$

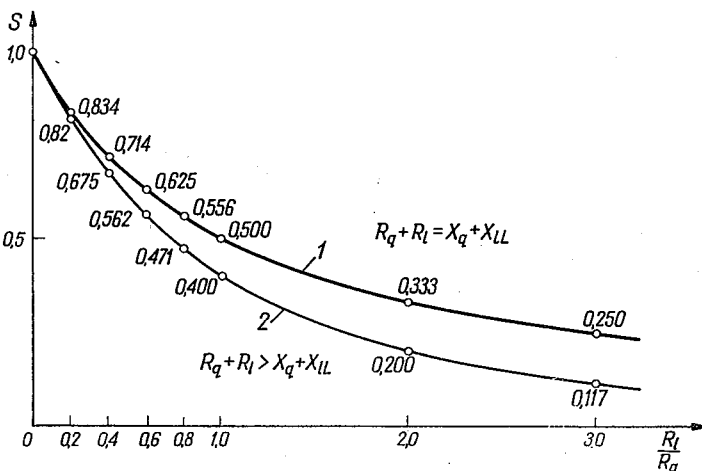
Analizując przebiegi prądów można udowodnić, że przy synchronicznej liczbie obrotów prądy we wtórnym uzwojeniu synchronizującym wytwarzają pole magnetyczne wirujące z prędkością równą podwójnej synchronicznej.

3.1.3. Wpływ linii łączącej selsyny

Wszystkie wzory charakteryzujące selsyny są wyprowadzane przy założeniu, że oporność łączących je przewodów równa się zeru, co jest równoznaczne z tym, że jest ona pomijalna w stosunku do oporności uzwojeń selsyna. Założenie to jest spełnione w wielu przypadkach łączy. Jeżeli natomiast oporności przewodów łączących są porównywalne z opornościami uzwojeń selsynów, to zaczynają wpływać na zmniejszenie sztywności łączy. Wielkość tego wpływu można określić na podstawie wzorów (25) lub (27), kładąc zamiast R_q i X_q odpowiednio wartości $R_q + R_l$ oraz $X_q + X_{lL}$, gdzie R_l i X_{lL} — oporności czynna i bierna indukcyjna linii łączącej selsyny łączy. Oczywiście jest, że w tym przypadku warunek (30) na uzyskanie maksimum sztywności łączy przekształci się do postaci:

$$R_q + R_l = X_q + X_{lL} \quad (37)$$

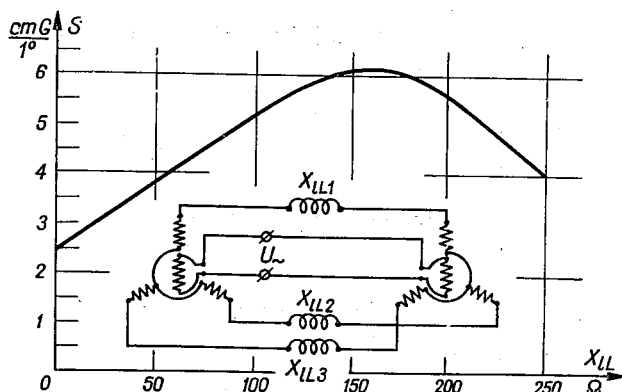
z tym, że wniosek o zmniejszaniu się sztywności łączy ze wzrostem obydwu tych równych sobie oporności, podany w omówieniu (30), pozostaje



Rys. 19. Spadek sztywności łączy przy zmianach oporności czynnej linii łączącej obwody wtórne

w mocy. Przebieg zmniejszania sztywności w tym przypadku ilustruje krzywa 1 na rys. 19.

Jednakże zazwyczaj linie łączące selsyny o niewielkich przekrojach przewodów mają oporność czynną znacznie większą od biernej indukcyjnej ($R_l \gg X_{lL}$). Zmniejszenie się sztywności ze wzrostem oporności jest w takim przypadku jeszcze szybsze; przykładowy przebieg $S = f\left(\frac{R_l}{R_q}\right)$ dla $R_l > X_{lL}$ podaje krzywa 2 na rys. 19. Przywrócenie sztywności maksymalnie możliwej do uzyskania przy danej oporności linii, czyli powrót do linii 1 z rys. 19, jest możliwe przy szeregowym włączeniu do obwodu wtórnego (w sposób uwidoczniiony na schemacie rys. 20) dławików o oporności biernej przewyższającej znacznie czynną.



Rys. 20. Wpływ dodatkowej indukcyjności w linii obwodu wtórnego na sztywność łącza

Przykładową ilustrację wpływu zastosowania takiego układu połączeń podaje wykres na rys. 20 zaczerpnięty z wyników badań [31], a dotyczy charakterystycznego łącza z selsynów krajowej produkcji SO-2 i linii łączącej ich obwody wtórne o oporności $R_l = 150 \Omega$. Z wykresu widać wyraźnie, że zastosowanie dławików o odpowiedniej oporności indukcyjnej pozwala uzyskać dwu i pół krotny wzrost sztywności łącza bez jakichkolwiek zmian parametrów samych selsynów, z których zestawiono łącze.

Znaczny wpływ na sztywność łącza ma również symetria oporności przewodów linii łączącej. Niejednakowe wartości czynnych i biernych oporności poszczególnych przewodów linii powodują znaczne obniżenie sztywności łącza.

Nieco inaczej przedstawiają się zależności w łączu zawierającym linię długą, gdzie znaczny wpływ ma oporność bierna pojemnościowa. Ten przypadek występuje przy wzajemnych odległościach selsynów łącza rzędu kilku do kilkudziesięciu kilometrów. Jako linie długie po-

winny być traktowane w tym przypadku linie łączące uzwojenia selsynów zarówno pierwotne, jak i wtórne.

Równania obciążonej linii długiej mają postać:

$$\begin{aligned}\bar{U}_1 &= \bar{U}_2 \cosh \gamma l + \bar{I}_2 \bar{Z}_l \sinh \gamma l, \\ \bar{I}_1 &= \bar{U}_2 \frac{\sinh \gamma l}{\bar{Z}_l} + \bar{I}_2 \cosh \gamma l,\end{aligned}\quad (38)$$

gdzie:

- $\bar{U}_1$ i $\bar{I}_1$ — napięcie i prąd na początku linii,
- $\bar{U}_2$ i $\bar{I}_2$ — „ „ „ na końcu linii,
- $\bar{Z}_l$ — oporność falowa linii,
- γ — współczynnik rozchodzenia się fali,
- l — długość linii.

W przypadku łączy selsynowych jako linie długie stosowane są przeważnie linie kablowe, dla których

$$\begin{aligned}\bar{Z}_l &= \sqrt{\frac{1}{2} \frac{R_l}{j\omega C_l}} = \sqrt{\frac{R_l}{\omega C_l}} \cdot \frac{1-j}{2}, \\ \gamma &= \sqrt{\frac{1}{2} jR_l \omega C_l} = \sqrt{R_l \omega C_l} \cdot (1+j),\end{aligned}\quad (39)$$

gdzie:

- R_l — jednostkowa oporność jednego przewodu linii,
- C_l — jednostkowa pojemność linii.

Przy znanym obciążeniu na końcu linii $\bar{Z}_{obc} = \frac{\bar{U}_2}{\bar{I}_2}$ z (38) można otrzymać zależności napięć i prądów na początku i na końcu linii

$$\begin{aligned}\bar{U}_1 &= \bar{U}_2 \left(\cosh \gamma l + \frac{\bar{Z}_l}{\bar{Z}_{obc}} \cdot \sinh \gamma l \right), \\ \bar{I}_1 &= \bar{I}_2 \left(\cosh \gamma l + \frac{\bar{Z}_{obc}}{\bar{Z}_l} \cdot \sinh \gamma l \right),\end{aligned}$$

a z nich — oporność wejściową linii

$$\bar{Z}_{wejs\acute{c}} = \bar{Z}_l \frac{\bar{Z}_{obr} \cosh \gamma l + \bar{Z}_l \sinh \gamma l}{\bar{Z}_l \cosh \gamma l + \bar{Z}_{obc} \sinh \gamma l}.$$

Schemat zastępczy dwuprzewodowej linii długiej określonej wzorami (38) może mieć postać czwórnika jak na rys. 21 o pozornych oporności i przewodności, wyrażonych wzorami:

$$\begin{aligned}\hat{z}_1 &= \bar{Z}_l \frac{\cosh \gamma l - 1}{\sinh \gamma l}, \\ \hat{y}_0 &= \frac{\sinh \gamma l}{\bar{Z}_l}.\end{aligned}\quad (40)$$

Schemat zastępczy trójprzewodowej linii długiej jest podany na rys. 22.

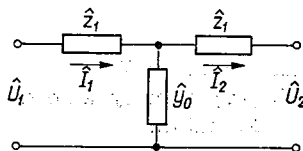
Dla linii kablowych o długości do 10—20 km i częstotliwości zasilania 50 Hz $\gamma l \ll 1$, więc wzory (40) mają w tym przypadku znacznie prostszą postać:

$$\begin{aligned}\hat{z}_1 &= R_l l, \\ \hat{y}_0 &= j\omega C_l l,\end{aligned}\quad (41)$$

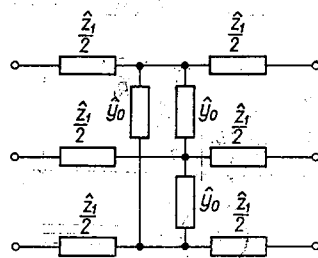
a schemat zastępczy przedstawi się jak na rys. 23, z tym, że:

$$r_1 = \frac{R_l l}{2}; \quad c_1 = C_l l;$$

widać, że przy niskich częstotliwościach i stosunkowo niedługich kablowych liniach łączących przeważający wpływ na pracę łącza ma oporność czynna linii.

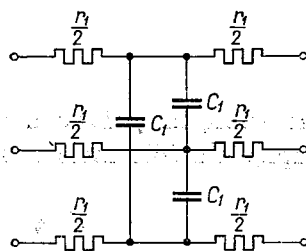


Rys. 21. Schemat zastępczy dwuprzewodowej kablowej linii długiej



Rys. 22. Schemat zastępczy trójprzewodowej linii długiej

W oparciu o szczegółową teorię selsynów można udowodnić, że w obecności linii długich przy opornościach poprzecznych nadajnika i odbiornika $\hat{Z}_{qn}$ i $\hat{Z}_{qo}$ znacznie mniejszych od oporności falowej linii



Rys. 23. Schemat zastępczy trójprzewodowej kablowej linii długiej

$\hat{Z}_l$ sztywność łącza zmniejsza się szybko ze wzrostem oporności czynnej linii R_l . Jednakże jeżeli wspomniane oporności spełniają zależność

$$\frac{\hat{Z}_{qn} \cdot \hat{Z}_{qo}}{\hat{Z}_l} = -1, \quad (42)$$

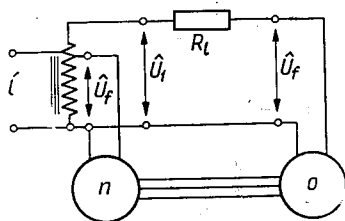
to sztywność łącza jest niezależna od długości linii. To zjawisko tłumaczy się wystąpieniem rezonansu indukcyjności uzwojeń wtórnych selsyna z pojemnością linii długiej.

Spełnienie zależności (42) jest realne dla selsynów o równych czynnych i biernych opornościach poprzecznych $R'_q = X'_q$, a więc jak zostanie stwierdzone dalej — zasilanych niskimi częstotliwościami. Ze względu na to, że oporności falowe linii kablowych niskich częstotliwości są rzędu kilkuset lub nawet tysięcy omów, to dla spełnienia zależności (42) można stosować specjalne transformatory dopasowujące.

Linie kablowe we wtórnym obwodzie łącza powinny być oczywiście symetryczne pod względem pojemności między żyłami.

Linia długa w obwodzie pierwotnym łącza powoduje występowanie spadku napięcia i przesunięcie jego fazy w miarę oddalania się od źródła zasilania.

Dla kompensacji spadku napięcia przy źródle zasilania można instalować autotransformator w układzie z rys. 24. Napięcie wtórne auto-



Rys. 24. Układ kompensacji spadku napięcia w linii przy pomocy autotransformatora

transformatora U_1 przy znajomości oporności linii R_l , mocy pobieranej przez odległy selsyn P , jego współczynnika mocy $\cos \varphi$ i napięcia znamionowego obwodu pierwotnego selsynów U_f można określić ze wzoru:

$$U_1 = U_f \sqrt{1 + a(a + 2 \cos \varphi)}, \quad (43)$$

gdzie

$$a = \frac{R_l \cdot P}{U_f^2}.$$

Kąt przesunięcia ψ fazy napięcia zasilającego oddalony selsyn względem napięcia źródła można określić ze wzoru:

$$\sin \psi = \frac{a \sin \varphi}{\sqrt{1 + a(a + 2 \cos \varphi)}}, \quad (44)$$

przy czym to napięcie wyprzedza napięcie zasilania.

Można udowodnić, że w przypadku gdy oddalonym od źródła zasilania selsynem jest odbiornik, przesunięcie fazy nie wpływa na zmianę sztywności łącza. Natomiast jeżeli od źródła zasilania oddalone są nadajniki, sztywność łącza gwałtownie spada ze zwiększaniem się kąta ψ .

Ponieważ spadek napięcia zasilającego zmniejsza również sztywność łącza, co wynika bezpośrednio z zasady jego działania, to oba przejawy działania linii powinny być skompensowane. Spadek napięcia można kompensować przy pomocy transformatorów lub włączonych równolegle kondensatorów zmniejszających prąd magnesujący, a więc i spadek napięcia w linii. Szeregowe włączenie w obwód zasilający odbiornika znajdującego się w pobliżu źródła zasilania oporności czynnej równej oporności falowej linii może skompensować kąt przesunięcia fazowego ψ .

3.1.4. Wzmocnienie momentu przy pomocy przesunięcia fazowego napięć wzbudzenia selsynów łącza

W rozdz. 3.1 ułożono równania określające łącze selsynowe w oparciu o teorię maszyn synchronicznych (sprowadzanie parametrów do osi d i q). Jak wspomniano jednakże, można również układać podobne równania w oparciu o teorię maszyn indukacyjnych. Wynikiem takiego postępowania i to przy założeniu ogólnego przypadku pracy łącza, a mianowicie dowolnych (niekoniecznie jednakowych) parametrów uzwojeń nadajnika i odbiornika, ich napięć wzbudzenia E_{1n} i E_{1o} i przesunięć fazowych (w czasie) tych napięć — odpowiednio φ_n i φ_o , można otrzymać wzory na momenty powstające w nadajniku i odbiorniku przy danych kątach wychyleń nadajnika i odbiornika α° i β° w następującej postaci [44]:

dla nadajnika —

$$M_n = 1217 \frac{(z_2 k_{z2})_n}{z_{1n}} \cdot \frac{(z_2 k_{z2})_o}{z_{1o}} \cdot \frac{E_{1n} \cdot E_{1o}}{f \cdot Z} \left[\cos(\varphi_o - \varphi_n) - \frac{X}{Z} \sin(\varphi_o - \varphi_n) \frac{R}{Z} \right] \cdot \sin(\alpha - \beta) \quad [\text{cm G}], \quad (43)$$

dla odbiornika —

$$M_o = 1217 \frac{(z_2 k_{z2})_n}{z_{1n}} \cdot \frac{(z_2 k_{z2})_o}{z_{1o}} \cdot \frac{E_{1n} \cdot E_{1o}}{f \cdot Z} \left[\cos(\varphi_o - \varphi_n) \frac{X}{Z} + \sin(\varphi_o - \varphi_n) \frac{R}{Z} \right] \cdot \sin(\beta - \alpha) \quad [\text{cm G}]. \quad (44)$$

We wzorach (43) i (44) poza stosowanymi uprzednio, ewentualnie omówionymi oznaczeniami, X , R i Z są sumarycznymi wartościami fazowymi zawierającymi oporności nadajnika, odbiornika i linii łączącej.

Łatwo zauważyć, że wzory (43) i (44) przy założeniu zasilania jednakowych nadajnika i odbiornika z tego samego źródła, a więc po podstawieniu

$$\frac{E_{1n} \cdot (z_2 k_{z2})_n}{z_{1n}} = \frac{E_{1o} \cdot (z_2 k_{z2})_o}{z_{1o}} = E_{2f} \quad \text{i} \quad \varphi_o = \varphi_n$$

przekształca się do postaci (23), która służyła nam dotychczas do rozpatrywania łącza charakterystycznego. Przeciwnie skierowanie M_n i M_o wynika z różnicy ich znaków przy ujednoczeniu postaci kąta niezgodności

$$\begin{aligned}\sin(\alpha - \beta) &= \sin \delta, \\ \sin(\beta - \alpha) &= -\sin \delta.\end{aligned}$$

Postać wzorów (43) i (44) umożliwia analizę wpływu przesunięcia fazowego napięć zasilających selsyny na momenty nadajnika i odbiornika. Jeżeli na przykład łącze spełnia warunek równości oporności czynnej i biernej $R = X$ równoważny warunkom (30) i (37), to:

$$\frac{X}{Z} = \frac{R}{Z},$$

ale jednocześnie:

$$\frac{X}{Z} = \sin \varphi, \quad \frac{R}{Z} = \cos \varphi,$$

więc $\sin \varphi = \cos \varphi$, co jest możliwe tylko w przypadku $\varphi = 45^\circ$, kiedy to $\sin \varphi = \cos \varphi = \frac{\sqrt{2}}{2}$.

W tym przypadku wyrażenie $\left[\cos(\varphi_o - \varphi_n) \frac{X}{Z} \mp \sin(\varphi_o - \varphi_n) \frac{R}{Z} \right]$ występujące we wzorach (43) i (44) możemy przepisać w postaci:

$$[\cos(\varphi_o - \varphi_n) \cos 45^\circ \mp \sin(\varphi_o - \varphi_n) \sin 45^\circ] = \cos \left[\frac{\pi}{2} \pm (\varphi_o - \varphi_n) \right]; \quad (45)$$

natomiast same wzory (43) i (44) w postaci:

$$\begin{aligned}M_n &= C \frac{E_{1n} \cdot E_{1o}}{f \cdot Z} \cdot \sin(\alpha - \beta) \cdot \cos \left[\frac{\pi}{4} + (\varphi_o - \varphi_n) \right], \\ M_o &= C \frac{E_{1n} \cdot E_{1o}}{f \cdot Z} \cdot \sin(\alpha - \beta) \cdot \cos \left[\frac{\pi}{4} - (\varphi_o - \varphi_n) \right].\end{aligned} \quad (46)$$

Z postaci wzorów (46) wynika, że istnieje możliwość otrzymywania nierównych sobie wartości momentów nadajnika i odbiornika przez zmiany wzajemnego przesunięcia fazy napięć zasilających ich pierwotne obwody.

Jeżeli łącze będzie spełniało warunek $|M_n| < |M_o|$, a więc gdy moment otrzymywany na wale odbiornika będzie większy od momentu przykładanego do wału nadajnika, to w odróżnieniu od dotychczas rozpatrywanych przypadków będzie ono pracowało nie tylko jako prze-

nośnik, lecz również jako wzmacniacz momentu. Ze wzorów (46) wynika, że jeżeli

$$\varphi_o - \varphi_n = \frac{\pi}{4}, \quad (47)$$

to moment na wale odbiornika będzie maksymalny dla danego kąta niezgodności $(\alpha - \beta)$, gdyż $\cos\left(\frac{\pi}{4} - \frac{\pi}{4}\right) = 1$, natomiast moment na wale nadajnika będzie równy zero — $\cos\left(\frac{\pi}{4} + \frac{\pi}{4}\right) = 0$. W rzeczywistości do walu nadajnika trzeba będzie przyłożyć moment równy jego własnemu momentowi tarcia.

Omówione zjawisko można wytłumaczyć fizycznie rozpatrując wzajemne przesunięcia przepływów i strumieni wytwarzających momenty obrotowe w nadajniku i odbiorniku w przypadku, gdy faza φ_o składowej biernej napięcia wzbudzenia E_{1o} będzie o $\frac{\pi}{4}$ wyprzedzała fazę φ_n składowej biernej napięcia wzbudzenia E_{1n} . Przypominając sobie zasadę działania selsyna stwierdzamy, że przy $R = X$ przepływ spowodowany obrotem wirnika nadajnika będzie opóźniał się w czasie o $\frac{\pi}{4}$ względem φ_n .

Natomiast strumień wzbudzenia w odbiorniku będzie opóźniał się o $\frac{\pi}{4}$ względem φ_o . W ten sposób przy konwencjonalnym zasilaniu łącza, gdy $\varphi_n = \varphi_o$, strumień odbiornika opóźnia się w czasie względem przepływu nadajnika. W ten sposób przy konwencjonalnym zasilaniu łącza, gdy $\varphi_n = \varphi_o$, strumień odbiornika opóźnia się w czasie względem przepływu nadajnika o $\frac{\pi}{4}$. Aby moment obrotowy wytwarzany wskutek współdziałania strumienia i przepływu był maksymalny, strumień i przepływ powinny być w fazie ze sobą, a dla spełnienia tego powinna mieć miejsce zależność $\varphi_o - \varphi_n = \frac{\pi}{4}$, a więc ta sama, którą otrzymaliśmy z rozważań w postaci (47). Wynika z tego również, że przy przesunięciu (47) moment na wale odbiornika będzie $\sqrt{2}$ razy większy niż przy braku przesunięcia, a to dlatego, że $\cos 0 : \cos \frac{\pi}{4} = 1 : \frac{\sqrt{2}}{2} = \sqrt{2}$. W nadajniku natomiast moment jest wytwarzany przez współdziałanie przepływu wywołanego przez przesunięte napięcie odbiornika ze strumieniem wzbudzenia nadajnika, a one przy $\varphi_o - \varphi_n = \frac{\pi}{4}$ i $R = X$ będą przesunięte względem siebie o $\frac{\pi}{2}$; a ponieważ $\cos \frac{\pi}{2} = 0$, więc moment w nadajniku nie zostanie wytworzony.

Przy rozpatrywanych dotychczas zależnościach zakładaliśmy spełnienie warunku równości czynnych i biernych oporności łącza; stwierdziliśmy ponadto, że można na nie wpływać przez zmiany parametrów linii łączących, ewentualnie przez odpowiednie włączanie skupionych indukcyjności lub pojemności. W praktyce zestawiając łącze z istniejących selsynów możemy przez włączanie dodatkowych cewek lub kondensatorów w bardzo szerokich granicach zmieniać oporności bierne łącza, natomiast możliwości zmian oporności czynnych są znikome. Rozpatrzymy zatem możliwości, które daje kompensacja oporności biernych łącza, wprowadzając w odpowiedniej chwili nie poddającą się praktycznie regulacji oporność czynną łącza do wzorów na momenty.

Jeżeli wprowadzimy dodatkowe pojemności całkowicie kompensujące indukcyjności łącza, tak że przy uwzględnieniu oporności selsynów i linii łączących $X = 0$, a $Z = R$, to wzory (43) i (44) przekształcą się do postaci:

$$M_n = M_o = 1217 \frac{(z_2 k_{z2})_n}{z_{1n}} \cdot \frac{(z_2 k_{z2})_o}{z_{1o}} \cdot \frac{E_{1n} \cdot E_{1o}}{f \cdot R} \cdot \sin(\alpha - \beta) \cdot \sin(\varphi_o - \varphi_n) [\text{cm G}]. \quad (48)$$

Z postaci (48) wzoru wynika, że dla danego kąta niezgodności $(\alpha - \beta)$ momenty osiągają maksimum przy $\varphi_o - \varphi_n = \frac{\pi}{2}$, oraz, że znaki momentów nadajnika i odbiornika są jednakowe, a więc momenty obrotowe działają w zgodnym kierunku. Znaczy to, że podczas gdy moment odbiornika M_o dąży do zmniejszenia kąta niezgodności, moment działający na sterowany z zewnątrz wał nadajnika usiłuje obrócić go w kierunku zwiększenia tegoż kąta niezgodności, wywołując efekt podobny do otrzymywanego w przypadku zastosowania dodatniego sprzężenia zwrotnego. W tym przypadku łącze działa więc znów jako wzmocniacz przenoszonego momentu.

Na podstawie obydwu rozpatrzonych teoretycznych przypadków pracy łącza selsynowego jako wzmacniacza momentu można wyprowadzić najważniejsze praktycznie zależności optymalne dla konkretnego łącza, a mianowicie określić wzajemne przesunięcie fazowe napięć zasilających selsyny — $\varphi_p = \varphi_o - \varphi_n$ i najkorzystniejsze stosunki oporności czynnej i biernej, przy których jednocześnie:

- a) moment nadajnika M_n teoretycznie (bez uwzględnienia jego własnego momentu tarcia) równałby się zeru,
- b) sztywność odbiornika byłaby maksymalna.

Wprowadzając oznaczenie:

$$K = 1217 \frac{(z_2 k_{z2})_n}{z_{1n}} \cdot \frac{(z_2 k_{z2})_o}{z_{1o}} \cdot \frac{E_{1n} \cdot E_{1o}}{f}$$

i pamiętając, że

$$\frac{X}{Z} = \sin \varphi; \quad \frac{R}{Z} = \cos \varphi; \quad \frac{1}{Z} = \frac{\cos \varphi}{R}; \quad \varphi_o - \varphi_n = \varphi_p;$$

$$\sin(\alpha - \beta) = \sin \delta,$$

wzory (43) i (44) przepisujemy w postaci:

$$\begin{aligned} \frac{M_n}{\sin \delta} = M'_n &= K \frac{\cos \varphi}{R} (\cos \varphi_p \cdot \sin \varphi - \sin \varphi_p \cos \varphi) = \\ &= \frac{K}{R} \cos \varphi \cdot \sin(\varphi_p - \varphi) \end{aligned}$$

lub:

$$M'_n = \frac{K}{2R} [\sin(\varphi_p - 2\varphi) + \sin \varphi_p];$$

analogicznie

$$M'_o = \frac{K}{2R} [\sin(\varphi_p + 2\varphi) + \sin \varphi_p]. \quad (49)$$

Dla spełnienia pierwszego z postawionych warunków — osiągnięcia zerowej wartości momentu na wale nadajnika — musi być:

$$\sin(\varphi_p - 2\varphi) + \sin \varphi_p = 0. \quad (50)$$

Równanie (50) jest spełnione przy $\varphi = \frac{\pi}{2}$ lub przy $\varphi_p = \varphi$. Pierwszy pierwiastek nie ma sensu fizycznego, gdyż przy $R \neq 0$, $\cos \varphi \neq 0$ i $\varphi \neq \frac{\pi}{2}$. Warunkiem otrzymania zerowej wartości momentu na wale nadajnika pozostaje więc spełnienie zależności

$$\varphi_p = \varphi.$$

Dla znalezienia warunku spełnienia drugiego wymagania — maksymalnej sztywności odbiornika — musimy znaleźć maksimum funkcji

$$f(\varphi) = \sin(\varphi_p + 2\varphi) + \sin \varphi_p = \sin 3\varphi + \sin \varphi. \quad (51)$$

Różniczkując $\frac{df(\varphi)}{d\varphi}$ i przyrównując wynik do zera, po odrzuceniu rozwiązań odpowiadających minimum funkcji oraz rozwiązań nie mających sensu fizycznego, stwierdzamy, że maksimum $f(\varphi)$ wystąpi przy:

$$3 \cos^2 \varphi - 2 = 0, \quad \text{czyli} \quad \cos \varphi = \pm \sqrt{\frac{2}{3}},$$

skąd optymalne wartości przesunięć:

$$\varphi_{opt} = \varphi = \varphi_p = 35,5^\circ. \quad (52)$$

Łatwo sprawdzić, że przy wartościach (52):

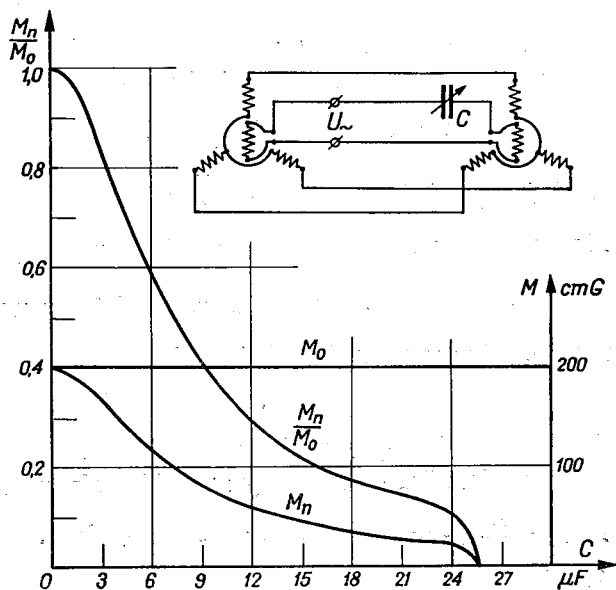
$$M'_n = 0; \quad M'_o = \frac{K}{2R} [\sin 106,5^\circ + \sin 35,5^\circ] = 1,54 \frac{K}{2R},$$

podczas gdy przy spełnieniu tylko warunku (47):

$$M'_n = 0; \quad M'_o = \frac{K}{2R} [\sin 135^\circ + \sin 45^\circ] = 1,41 \frac{K}{2R} < 1,54 \frac{K}{2R}.$$

Dla osiągnięcia optymalnych przesunięć φ i φ_p w realnym łączy trzeba doprowadzić do spełnienia zależności $\frac{X}{R} = \arctg 35,5^\circ = 0,717$, co wobec przeważnie występujących w praktyce $R > X$ nie nastręcza zbyt dużych trudności i jest łatwiejsze od osiągnięcia zależności $R = X$ wymaganej przy $\varphi = 45^\circ$. Łatwiej jest również osiągnąć wyprzedzający kąt φ_p o wartości $35 \div 36^\circ$ niż o wartości 45° .

Przytoczone rozważania teoretyczne zostały w całej rozciągłości potwierdzone wynikami przeprowadzonych doświadczeń. Na rys. 25 poda-



Rys. 25. Wpływ przesunięcia fazowego napięć wzbudzenia selsynów łączy

ny jest przykładowo wynik prób doboru pojemności przesuwającej napięcie wzbudzenia odbiornika łączy zestawionego z seryjnych selsynów krajowej produkcji, które miały na celu redukcję momentu potrzebnego do napędu wału nadajnika teoretycznie do wartości zerowej, a praktycznie do wartości własnego momentu tarcia nadajnika. Odbiornik jest obciążony stałym momentem oporowym [31]. Z rysunku widać, że posta-

wione zadanie może być rozwiązane przez zastosowanie nadzwyczaj prostego układu (jeden dodatkowy kondensator), mimo że parametry łącza nie spełniają dokładnie ani warunku (47), ani (52).

Doświadczalnie stwierdzono również, że niewielkie wzajemne przesunięcia faz zasilania nadajnika i odbiornika rzędu $\varphi_p = \pm 10^\circ$ wywołują zmiany momentu synchronizującego, a więc i sztywności odbiornika o $\pm 15 \div 20\%$.

Rozpatrzone układy łącza wskaźnikowego zapewniające możliwość wzmocnienia momentu sterującego są szczególnie przydatne w teledystrybucji, w przypadkach gdy moment przyrządu pierwotnego, którego wskazania trzeba przekazać na odległość, wystarcza jedynie dla przezwyciężenia momentu tarcia selsyna nadajnika, nie wystarczyłby natomiast do napędu wskaźników wtórnych — selsynów odbiorczych — w klasycznym układzie przenoszenia wskazań.

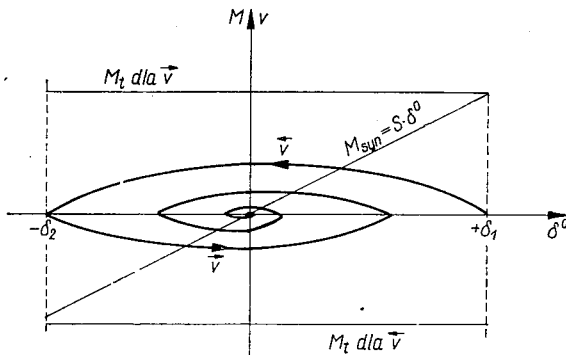
3.1.5. Drgania selsynów i ich tłumienie

Z omówionej w rozdz. 2 zasady działania łącza wynika, że przy obracaniu się wirnika nadajnika wirnik odbiornika podąża za nim i po pewnym czasie powinien zająć odpowiednio zgodne położenie. Im czas uzgodnienia położenia jest krótszy, tym szybciej możemy otrzymać prawidłowy odczyt, a więc tym lepsza jest praca łącza. Dla ustalenia, od czego zależy czas uzgodnienia położenia wirników selsynów, rozpatrzmy zjawiska zachodzące podczas przebiegu stanu nieustalonego nadawania.

Założmy, że przy zahamowanym wirniku nadajnika obrócimy wirnik odbiornika o pewien niewielki kąt $+\delta_1^\circ$ względem położenia równowagi, a następnie puścimy. W pierwszej chwili obrócony wirnik znajduje się w spoczynku, więc jego prędkość obrotowa równa się zeru, działa nań natomiast synchronizujący moment obrotowy $+S\delta^\circ$. Pod wpływem tego momentu wirnik zacznie się obracać, a jego prędkość obrotowa zacznie wzrastać. W chwili osiągnięcia zgodnego położenia wirnik będzie obracał się dalej z rozpędu, mimo że moment synchronizujący w tym położeniu będzie równał się zeru. Kąt niezgodności zmieni znak, a wraz z nim i moment synchronizujący, który od tej chwili będzie przeciwdziałał obrotowi wirnika. Po pewnym czasie wirnik obróci się o kąt $-\delta_2$ względem położenia równowagi i zatrzyma się, a następnie cały cykl zjawisk powtórzy się w przeciwnym kierunku. Przez cały czas ruchu wirnika od $+\delta_1$ do $-\delta_2$ będzie nań działał moment tarcia w łożyskach o zwrocie przeciwnym do kierunku obrotów. Zmianę omówionych wielkości ilustruje rys. 26 dla przypadku $|\delta_1| > |\delta_2|$, czyli drgań tłumionych.

Podczas drgań wirnik obraca się okresowo z pewną prędkością, więc w jego uzwojeniach wtórnych poza normalnie występującą *sem* transformacji indukuje się również *sem* rotacji, powodująca powstawanie

asynchronicznego momentu obrotowego na zasadzie analogicznej jak w przypadku silnika jednofazowego. Rozpatrując łącznie z nadajnikiem o nieruchomym wirniku i odbiornikiem o wirniku obracającym się z pewną prędkością, możemy traktować ten ostatni jako odwrócony jednofazowy silnik asynchroniczny, w którego zwarte uzwojenie wtórne wtrącono jako dodatkową oporność uzwojenie wtórne nadajnika.



Rys. 26. Wykres drgań tłumionych selsyna odbiorczego

Zgodnie z teorią silnika jednofazowego (a w szczególności jego odmiany wykonawczej) wytwarzany przezeń moment zależy od stosunku wypadkowych parametrów obwodu wtórnego, a mianowicie oporności czynnej R_{2w} i biernej — X_{2w} . Jeżeli $R_{2w} < X_{2w}$, to moment asynchroniczny przy prędkości obrotowej $n \neq 0$ jest dodatni (o zwrocie zgodnym z kierunkiem obrotów wirnika; maszyna nie ma momentu rozruchowego, ale po wprowadzeniu wirnika w ruch moment asynchroniczny doprowadzi maszynę do określonej prędkości obrotowej i będzie ją utrzymywał. Jeżeli natomiast $R_{2w} \geq X_{2w}$, to moment asynchroniczny przy prędkości $n \neq 0$ jest ujemny (o zwrocie przeciwnym do kierunku obrotów wirnika); po wprowadzeniu wirnika w ruch moment asynchroniczny będzie zawsze hamował jego obrót przybierając tym większe wartości, im większa będzie wymuszająca prędkość obrotowa.

Na wirnik selsyna odbiornika działają: moment tarcia M_t — stały i skierowany zawsze przeciwnie do kierunku obrotów, moment synchronizujący, proporcjonalny do kąta niezgodności $M_s = -T_s \delta$, moment asynchroniczny (tłumiący), proporcjonalny do prędkości obrotowej wirnika selsyna $M_a = \mp T_a \frac{d\delta}{dt}$ oraz moment bezwładności, proporcjonalny do przyspieszenia $M_J = J \frac{d^2\delta}{dt^2}$. Współczynniki T_s i T_a nazywane są współczynnikami momentu synchronizującego¹⁾ i momentu tłumiącego, z tym,

¹⁾ Współczynnik momentu synchronizującego jest oczywiście (z definicji) równy sztywności selsyna S . Postać zapisu — T_s przyjęto dla uzyskania przejrzystego kształtu dalszych wzorów.

że $T_s < 0$ (moment jest zawsze skierowany w kierunku zmniejszenia δ), natomiast $T_a \geq 0$ (moment może być skierowany przeciwnie lub zgodnie do kierunku obrotów, w zależności od stosunku R_{2w} i X_{2w}). Równanie momentów działających na wirnik można zapisać w postaci:

$$J \frac{d^2 \delta}{dt^2} + T_a \frac{d\delta}{dt} + T_s \delta = \pm M_t. \quad (53)$$

Jest to równanie różniczkowe układu wykonującego drgania skrętne.

Jeżeli moment tarcia możemy pominąć, wobec jego małości, zakładając $M_t = 0$, co odpowiada większości praktycznych wykonawstw selsynów odbiorników bez tłumików, to równanie (53) staje się liniowym i jego rozwiązanie przy $T_s > \frac{T_a}{4J}$ będzie:

$$\delta = \delta_m e^{-\lambda t} \cos(\nu t + \varphi), \quad (54)$$

gdzie:

$$\lambda = \frac{T_a}{2J} \text{ — współczynnik tłumienia,}$$

$$\nu = \sqrt{\frac{T_s}{J} - \lambda^2} \text{ — pulsacja drgań tłumionych,}$$

δ_m i φ — amplituda i faza drgań zależne od warunków początkowych.

Z otrzymanych zależności wynika, że drgania odbiornika mogą być trzech rodzajów:

przy $\lambda > 0$; $T_a > 0$ — drgania tłumione;

$\lambda = 0$; $T_a = 0$ — drgania ustalone (nie tłumione); (55)

$\lambda < 0$; $T_a < 0$ — drgania o wzrastającej amplitudzie (rozbieganie się selsyna).

Tak więc dla selsynów z pomijalnym momentem tarcia o tłumieniu drgań decyduje moment asynchroniczny

$$M_a = -T_a \frac{d\delta}{dt}.$$

Jeżeli będziemy uwzględniali moment tarcia, co jest konieczne szczególnie przy rozpatrywaniu selsynów z tłumikami mechanicznymi, to dzieląc stronami równanie (53) przez J możemy je zapisać w postaci:

$$\frac{d^2 \delta}{dt^2} + 2\lambda \frac{d\delta}{dt} + \nu_0 (\delta \pm B) = 0,$$

gdzie:

$$B = \frac{M_t}{T_s} \text{ — martwa strefa określona stosunkiem momentu}$$

i sztywności ($|T_s| = |S|$),

$$\nu_0 = \sqrt{\frac{T_s}{J}} \text{ — pulsacja drgań swobodnych.}$$

Rozwiązanie tego nieliniowego równania znane z teorii drgań ma postać

$$\delta_n = (\delta_0 - B)A^n - 2B \frac{A(1 - A^{n-1})}{1 - A} - B, \quad (56)$$

gdzie:

δ_0 — amplituda początkowa,

δ_n — amplituda n -tego jednostronnego wychylenia,

$A = e^{-\pi \frac{\lambda}{\nu}} < 1$ — dekrement tłumienia

przy czasie jednostronnego wychylenia

$$t_0 = \frac{\pi}{\nu} = \frac{\pi}{\sqrt{\nu_0^2 - \lambda^2}}.$$

Przy założeniu $0 < \delta_n < B$ oznaczającym, że n -te wychylenie było ostatnim (wirnik selsyna znalazł się w strefie martwej), czas całkowitego stłumienia drgań wyrazi się wzorem

$$t_{tłum} = nt_0,$$

który po przekształceniach przyjmie postać

$$t_{tłum} = \frac{\lg \left[1 + \frac{\delta_0}{2B} (1 - A) \right]}{0,434 \lambda}. \quad (57)$$

Z otrzymanych zależności wynika ciekawa właściwość łączy będąca skutkiem nieliniowości drgań spowodowanej wpływem tarcia. Mianowicie z (57) wynika, że przy $\lambda < 0$, $A > 1$ i w przeciwieństwie do poprzednich wniosków (55) czas tłumienia ma wartość skończoną pod warunkiem, że $\delta_0 < \frac{2B}{A-1}$ (otrzymujemy drgania tłumione), natomiast

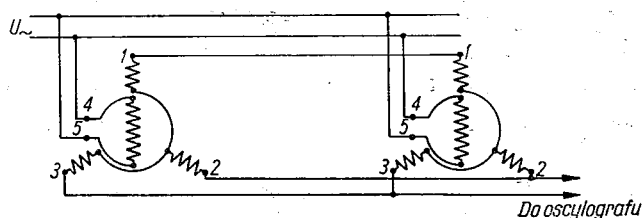
wartość nieskończoną przy $\delta_0 \geq \frac{2B}{A-1}$ (rozbieganie się selsyna). Tak więc nawet przy $\lambda < 0$ drgania są tłumione do pewnej wartości początkowego wychylenia wirnika, a dopiero po jej przekroczeniu następuje rozbieganie się selsyna.

Tak więc w przypadku selsynów z nie pomijalnym momentem tarcia o tłumieniu drgań decyduje poza momentem asynchronicznym również moment tarcia wpływający na wielkość martwej strefy.

Wszystkie współczynniki równania (53) można określić doświadczalnie dla danego selsyna. Mianowicie: moment zamachowy J — jedną z metod stosowanych w mechanice, a więc metodą rozbiegu lub mierzenia czasu trwania okresu sprężystych drgań skrętnych; T_a — oscylografując napięcia liniowe faz uzwojeń wtórnych selsynów, które w położe-

niu równowagi równe są zeru (układ z rys. 27) i przeliczając dekrement tłumienia z kolejnych amplitud wg wzoru (56); T_s — z definicji jest równe sztywności łącza; M_t — jedną z metod pomiaru małych momentów stosowanych przy badaniach maszyn miniaturowych.

Z przytoczonych rozważań wynika, że tłumienie drgań selsynów może odbywać się na drodze elektrycznej pod działaniem momentu asynchronicznego M_a i na drodze mechanicznej z wykorzystaniem momentu tarcia M_t .



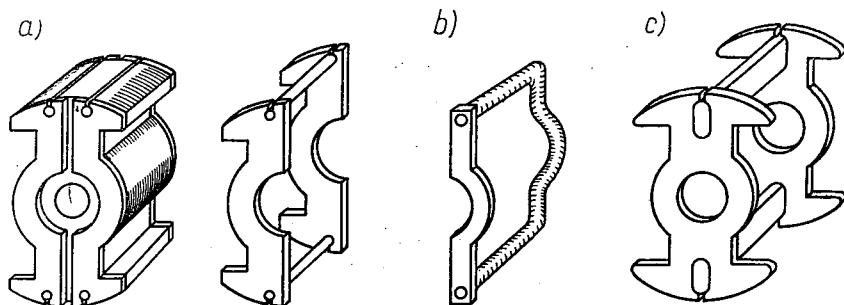
Rys. 27. Układ do rejestracji drgań selsyna

Tłumienie na drodze elektrycznej polega na przekształceniu energii mechanicznej drgań wirnika w energię elektromagnetyczną, która z kolei wywiązuje się w postaci ciepła w uzwojeniach tłumiących i zostaje rozproszona. Jako uzwojenia tłumiące mogą być traktowane zwarte uzwojenia obwodu pierwotnego o osi prostopadłej do kierunku strumienia wzbudzenia, czyli uzwojenia poprzeczne. Podczas wyprowadzania wzorów na moment i sztywność w rozdz. 3.1.1 stwierdzono, że przy małych kątach niezgodności, takich jakie występują podczas drgań selsynów, przepływy w osi podłużnej obwodu wtórnego nie odgrywają praktycznie żadnej roli, natomiast momenty powstają w wyniku współdziałania strumienia wzbudzającego z przepływem poprzecznym. Siły elektromotoryczne rotacji indukowane są podczas drgań selsyna również tylko w elementach uzwojenia, których oś jest zgodna z osią poprzeczną selsyna. Jeżeli te uzwojenia będą zwarte, to pod wpływem indukowanych sił elektromotorycznych popłyną w nich prądy nagrzewające je, a więc odbierające energię drgań, czyli tłumiące te drgania. Jeżeli nadajnik jest również wyposażony w zwarte uzwojenie poprzeczne, to ono jest transformatorowo połączone z poprzecznym uzwojeniem odbiornika. W ten sposób obydwie poprzeczne uzwojenia przyjmują wspólnie udział w tłumieniu drgań stanowiąc obwody umożliwiające przepływ poprzecznych prądów tłumiących. Z tego względu korzystne jest wyposażenie również i nadajnika w zwarte poprzeczne uzwojenia tłumiące.

W przypadku wirnika z biegunami utajonymi konstrukcja uzwojenia wzbudzającego (zob. rozdz. 2) zapewnia istnienie uzwojeń tworzących zwarte obwody w poprzecznej osi selsyna. W tym przypadku spełniają one podwójne zadanie. Z jednej strony zapewniają odpowiednie

parametry w osi poprzecznej (X_q), a z drugiej działają jako uzwojenie tłumiące.

W przypadku wirnika z biegunami wydawnymi uzwojenie wzbudzące nie ma tej właściwości. Dlatego wirniki z biegunami wydawnymi są w miarę potrzeby wyposażane w specjalne zwarte zwoje tłumiące. Na rys. 28 uwidocznione są trzy rozwiązania konstrukcyjne takich uzwojeń:



Rys. 28. Uzwojenia tłumiące dla wirników z wydawnymi biegunami

a. uzwojenie z dwóch izolowanych od siebie zwojów, składających się z dwóch prętów i dwóch płytek o kształtach odpowiadających blasze wirnika; wadą tego rozwiązania są cztery połączenia lutowane każdego zwoju, zwiększające opory przejścia;

b. uzwojenie różniące się od poprzedniego wyeliminowaniem dwóch połączeń w każdym ze zwojów; można je stosować do określonej grubości przewodu pozwalającej na uwidocznione wyginanie, przy grubszych przewodach okrągłych lub profilowych nieunikniona jest konstrukcja a;

c. uzwojenie podobne do rozwiązania a, jednakże składające się z jednego zwoju swartego.

Pręty i płytki uzwojenia swartego powinny być izolowane od żelaza izolacją żłobkową ewentualnie odpowiednimi przekładkami.

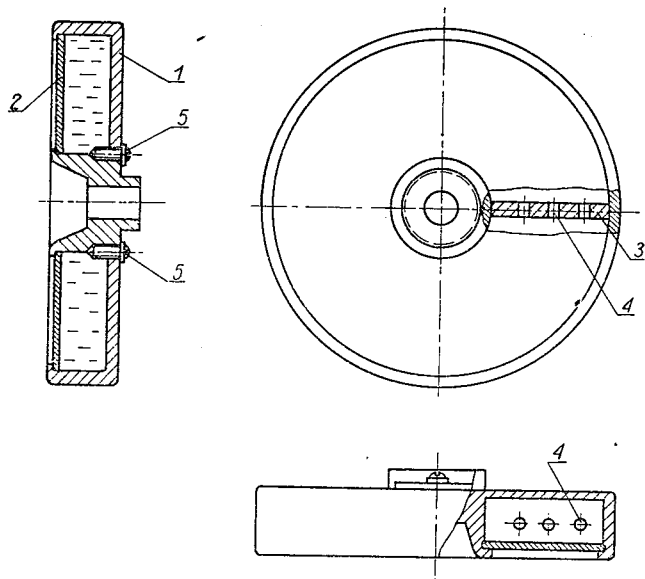
Selsyny zaopatrzone w zwarte uzwojenie tłumiące w osi poprzecznej obwodu wzbudzającego mają znacznie lepsze właściwości niż przy braku takiego uzwojenia. Wykazują się one większą sztywnością przy pracy we wskaźnikowych łączach nadawczo-odbiorczych oraz mogą być stosowane jako nadajniki w łączach transformatorowych o odbiornikach z wydawnymi biegunami.

Selsyny bez uzwojeń tłumiących na wirniku mają często dodatni moment asynchroniczny wywiązujący energię podobnego rzędu co i moment tarcia. Takie selsyny łatwo wpadają w drgania nietłumione, ewentualnie tłumione bardzo wolno. Dla uniknięcia tej wady można je wyposażać w tłumiki mechaniczne. Również niektóre mniej udane konstrukcje selsynów z uzwojeniami tłumiącymi mają zbyt wolno tłumione drgania. I w tym przypadku tłumik mechaniczny może polepszyć ich właści-

wości. Należy jednak zaznaczyć, że tłumik nie jest rozwiązaniem wygodnym, gdyż poza stosunkowo znacznym skomplikowaniem budowy mechanicznej może wprowadzać dodatkowe uchyby szczególnie przy pracy dynamicznej.

Podstawowymi elementami każdego tłumika mechanicznego są bezwładna masa zamachowa i urządzenie cierne. Masa zamachowa jest osadzona na wale wirnika w sposób pozwalający na jej obrót względem wirnika z pewnym momentem tarcia, zapewnionym przez urządzenie cierne. Zasada działania tłumika mechanicznego polega na tym, że przy powstaniu drgań masa zamachowa wskutek dużej bezwładności pozostaje prawie nieruchoma i na wale selsyna pojawia się dodatkowy moment tarcia na skutek jego obrotu względem tej masy. Na przewyciężenie dodatkowego momentu tarcia wywołanego przez tłumik zużywana jest część energii kinetycznej drgającego wirnika, co prowadzi do szybszego zanikania tych drgań.

Najprostszym, lecz również najmniej skutecznym, tłumikiem mechanicznym jest tłumik rtęciowy. Budowa takiego tłumika jest uwidoczniona na rys. 29. Składa się on z korpusu 1, pokrywki 2, dwóch prze-

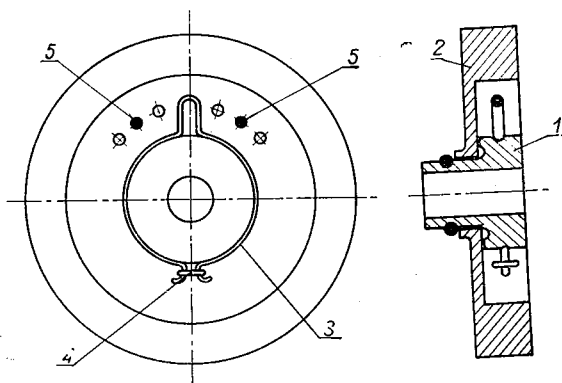


Rys. 29. Budowa rtęciowego tłumika drgań

gródek 3 z otworami 4, umocowanych wewnątrz korpusu, i rtęci wypełniającej korpus. Pokrywka zostaje zaprasowana w korpusie przez rozwałcowanie widocznych na rysunku zewnętrznej i wewnętrznej krawędzi korpusu. Korpus jest wypełniany rtęcią przez otwory zaslepiane wkrętami 5. Przy drganiach selsyna rtęć przelewa się z jednej komory tłumika do drugiej przez otwory 4 w przegródkach 3. Moment tłumiący na wale selsyna powstaje na skutek tarcia rtęci o ścianki przegródek,

korpusu i pokrywki. Tak więc i w tym tłumiku bezwładną masą zamachową jest rtęć, a urządzeniem ciernym przegródki i ścianki powierzchni wewnętrznej jej zbiornika. Podstawowymi wadami tłumika rtęciowego są znaczny ciężar, wymaganie hermetyczności i mała pewność ruchu przy stosunkowo niewielkiej skuteczności tłumienia. Z tego względu tłumiki rtęciowe nie są obecnie stosowane i nie mogą być zalecane. Można je jednak spotkać w starszych rozwiązaniach selsynów.

Znacznie lepsze właściwości wykazuje dwustopniowy tłumik mechaniczny uwidoczniiony na rys. 30. Składa się on ze sztywno osadzonej na wale selsyna tulei 1, na której obraca się z niewielkim tarcie koło zamachowe 2 umocowane w sposób uwidoczniiony na rysunku, oraz z urządzenia ciernego w postaci sprężyny 3 ze spinaczem 4 obracającej się ze znacznie większym tarcie w wycięciu tulejki 1.

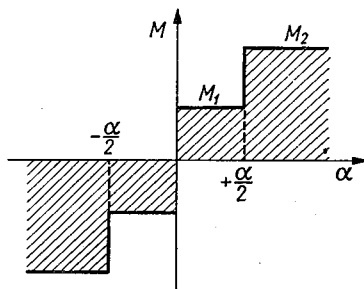


Rys. 30. Dwustopniowy tłumik cierny

Widoczny na rysunku występ sprężyny opiera się przy obracaniu się wału selsyna o kołki 5. Przy nagłych obrotach wału selsyna, kiedy tarcza zamachowa pozostaje jeszcze w spoczynku z powodu swej bezwładności, w granicach kąta $\pm \frac{\alpha}{2}$, gdzie α kąt między kołkami, na wał działa moment tarcia M_1 spowodowany przesuwaniem się pasowanej obrotowo tulejki wewnątrz koła zamachowego (rys. 31). Przy obrocie o kąt większy od $\pm \frac{\alpha}{2}$ występ sprężyny opiera się o kołek i zaczyna działać dodatkowy moment tarcia wywoływany w urządzeniu ciernym, a spowodowany przesuwaniem się zaciśniętej na tulei sprężyny 3. Całkowity moment tarcia wyniesie M_2 (rys. 31). Z rysunku widać dwustopniowość działania opisywanego tłumika. Moment tarcia M_1 jest stały i zależny od materiałów, z których wykonano tarczę i tuleję, pasowania między nimi i ciężaru tarczy. Moment M_2 można łatwo regulować przez zmianę wygięcia i docisku sprężyny 3. Zmianę kąta α można otrzymać odpowiednio przedstawiając kołki. Przy odpowiednim doborze masy zamachowej ką-

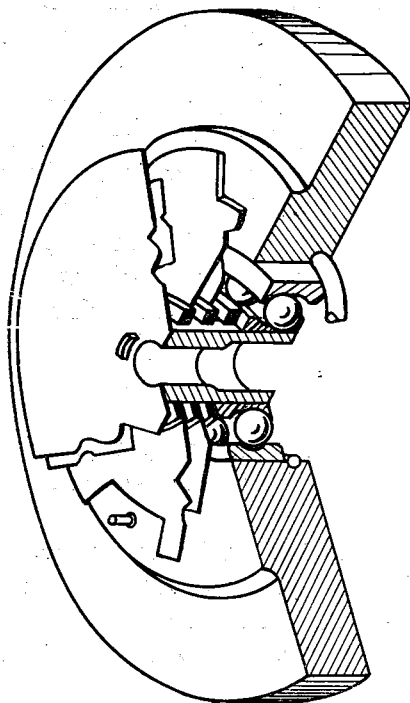
ta α i intensywności tarcia w elemencie ciernym można osiągnąć bardzo skuteczne tłumienie w czasie kilkakrotnie krótszym niż przy stosowaniu tłumika rtęciowego — rzędu ułamka sekundy. Zaletami tego tłumika są

Rys. 31. Wykres momentów tarcia dwustopniowego tłumika ciernego



skuteczność działania, stosunkowo nieznaczny ciężar, prostota konstrukcji i możliwość łatwej regulacji i pomiaru momentu tarcia drugiego stopnia.

Na rys. 32 pokazany jest tłumik pracujący na podobnej zasadzie, różniący się osadzeniem tarczy zamachowej na łożysku kulkowym, co



Rys. 32. Czterostopniowy tłumik cierny

zmniejsza tarcie pierwszego stopnia tłumienia, oraz liczbą stopni tłumienia, wynoszącą — jak widać z rysunku — cztery. Ograniczniki swobodnego ruchu tarczy mogą być przesuwane umożliwiając regulację

kąta α w granicach od kilku do 45° . Ten tłumik umożliwia bardziej precyzyjną regulację czasu i charakterystyki tłumienia i wprowadza mniejsze uchyby dzięki mniejszej wartości momentu M_1 , który jest właściwie pomijalny. Jednakże jego konstrukcja jest znacznie bardziej skomplikowana niż tłumika omówionego poprzednio, co stanowi w wielu przypadkach dużą wadę.

Stosowanie tłumików mechanicznych może być zalecane tylko przy pracy quasistatycznej łącza z niewielkimi przyspieszeniami kątowymi. Wtenczas uchyb wprowadzany przez nie jest niewielki. Natomiast przy nierównomiernym obracaniu się wału tłumiki mechaniczne mogą znacznie pogarszać pracę łącza. Na przykład przy obracaniu się nadajnika z prędkością zmienną sinusoidalnie, na wale odbiornika powstaną przyspieszenia zmienne cosinusoidalnie. Pod wpływem tych przyspieszeń tłumik mechaniczny wprowadzając okresowo dodatkowy udarowy moment tarcia na wale odbiornika będzie znacznie zniekształcał charakter jego ruchu.

3.1.6. Uchyby łącza

Moment synchronizujący selsyn odbiornik zależy od kąta niezgodności, a przy zgodnym położeniu kątowym wirników nadajnika i odbiornika równa się zeru. Dlatego istnienie nawet najmniejszych momentów hamujących powoduje powstawanie pewnego uchybu określającego dokładność przekazywania wskazań przez łącze. Momenty oporowe, a więc i uchyby, zależą od rodzaju pracy łącza. Rozróżniamy uchyby statyczne występujące przy niewielkich prędkościach kątowych obrotu wirników selsynów i uchyby kinematyczne — przy większych.

Uchyby statyczne charakterystyczne dla wskaźnikowej pracy łącza selsynowego, które dla małych kątów niezgodności określa zależność:

$$\Delta\delta = \frac{M_h(\delta)}{S}, \quad (58)$$

są powodowane w głównej mierze następującymi przyczynami:

1. tarciami w łożyskach i styku szczotkowym (dla selsynów ze szczotkami),
2. niedokładnością wyważenia wirnika i wskazówki,
3. momentami reluktancyjnymi spowodowanymi niejednakową przewodnością magnetyczną stali obwodu wtórnego,
4. obecnością zwartych obwodów w uzwojeniu wtórnym oraz w stali obwodu wtórnego,
5. momentami reluktancyjnymi wywołanymi obecnością zębów na wirniku i stojanie,
6. niedokładnym zachowaniem cylindryczności wirnika, nierównomiernościami szczeliny powietrznej i innymi usterkami technologicznymi.

Uchyby kinematyczne występujące podczas pracy łącza, przy której wały selsynów wirują ze stałą lub zmienną prędkością obrotową, zależą od sztywności łącza, bezwładności wirnika, tłumienia drgań, prędkości obracania się nadajnika i innych czynników.

Na wielkość uchybu wpływają również poprzez sztywność wahania napięcia zasilającego i częstotliwości.

Uchyby przy pracy statycznej

Wartości uchybów statycznych łącza określamy obracając wirnik nadajnika w zakresie jednego pełnego obrotu co pewien określony kąt równomiernej skali, początkowo w jednym kierunku — zgodnie ze wskazówkami zegara — a następnie po przejściu pełnych 360° , w przeciwnym kierunku — przeciwnie do wskazówek zegara. Wskazania odbiornika odczytujemy na takiej samej skali. Różnica wskazań nadajnika i odbiornika określa uchyb statyczny. Typowe krzywe uchybów statycznych otrzymane z pomiarów są podane na rysunkach 33a i 33b. Połowa sumy wartości bezwzględnych maksymalnych uchybów obydwu znaków nazywa się uchybem maksymalnym selsyna:

$$\Delta\delta_{\max} = \frac{|+\Delta\delta_m| + |-\Delta\delta_m|}{2}.$$

W zależności od wielkości uchybu maksymalnego dzielimy selsyny na 3 klasy dokładności:

klasa 1	— $\Delta\delta_{\max}$ dla nadajników $\leq 0,25^\circ$;	dla odbiorników $\leq 0,75^\circ$
„ 2	— $\Delta\delta_{\max}$ „ „ $\leq 0,5^\circ$;	„ „ $\leq 1,5^\circ$
„ 3	— $\Delta\delta_{\max}$ „ „ $\leq 1,0^\circ$;	„ „ $\leq 2,5^\circ$

Rozpatrzmy obecnie wpływ poszczególnych wymienionych poprzednio przyczyn powodujących powstawanie uchybu statycznego.

Wpływ tarcia

Wielkość wpływu tarcia w łożyskach i styku szczotkowym M_t na uchyb zależy od stosunku sztywności odbiornika do jego momentu tarcia, zwanego współczynnikiem statycznej dobroci selsyna D_s :

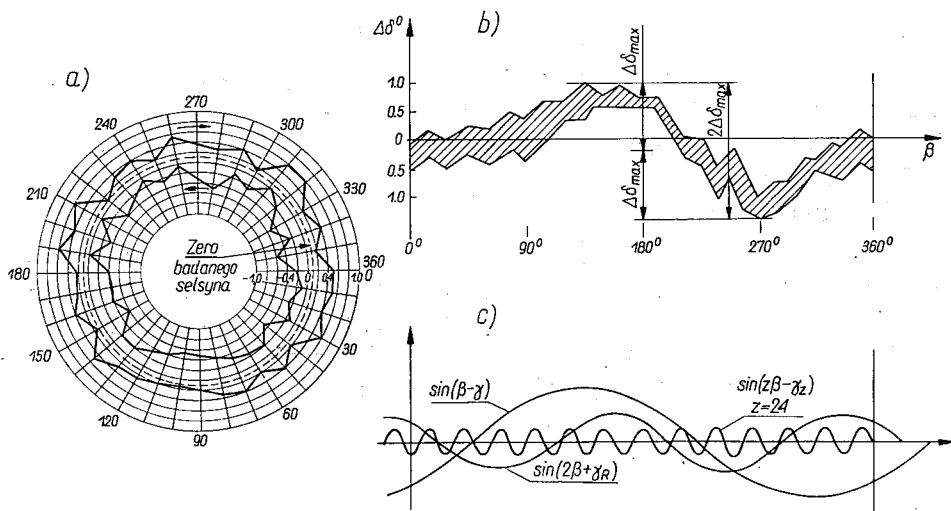
$$D_s = \frac{S}{M_t(\delta)_{\max}}. \quad (59)$$

Martwa strefa selsyna B — kąt, w granicach którego wirnik odbiornika może zająć dowolne położenie przy danym położeniu wirnika nadajnika — jest określona zależnością odwrotną:

$$B(\delta) = \pm \frac{M_t(\delta)}{S}. \quad (60)$$

Ze względu na to, że moment tarcia jest skierowany zawsze przeciwnie do kierunku ruchu, to przy doprowadzaniu wirnika nadajnika do danego

położenia kąтового z jednej strony, a następnie z drugiej, otrzymamy niedochodzenie odbiornika do odpowiednio zgodnego położenia również kolejno z obydwu stron. Różnica między tymi dwoma położeniami wirnika odbiornika widoczna na rys. 33 jako obszar między krzywymi uchybów zdejmowanych przy obrotach w obydwu kierunkach reprezentuje



Rys. 33. Wykresy uchybów statycznych. a) i b) — dwie formy rejestracji c) wyznaczenie składowych sinusoidalnych krzywej z pozycji b)

wielkość martwej strefy selsyna. Martwa strefa maleje oczywiście ze zmniejszaniem się liczby styków szczotkowych (przewaga układu z jednofazowym obwodem pierwotnym na wirniku) i jest najmniejsza dla selsynów bezstykowych (które nie są rozpatrywane w tym opracowaniu).

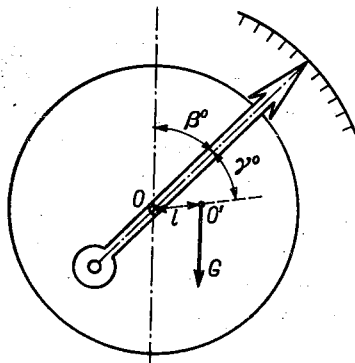
Wykonywane obecnie selsyny mają współczynnik statycznej dobroci D_s w granicach $0,25 \div 1,5 \frac{1}{\text{stopień}}$. Odbiorniki o $D_s \leq 0,5$ należy uważać za niedokładne, gdyż uchyb statyczny spowodowany tylko tarciem przekracza w nich 1° , a ponadto obserwuje się powolne pełzanie wskazówki do ustalonego położenia po zatrzymaniu się wirnika nadajnika, co zwiększa stałą czasu łącza i utrudnia odczyt.

Odbiorniki o $D_s \geq 1$ można uważać za dokładne; ich uchyb spowodowany tarciem jest mniejszy od 1° .

Moment tarcia nie jest związany jednoznacznie zależnością z kątem obrotu wirnika selsyna, wobec czego uchyby spowodowane przezeń nie wykazują okresowej zmienności. Należy zaznaczyć, że momenty tarcia wzbudzonego selsyna są nieco mniejsze niż nie wzbudzonego, z powodu wibracji wywołanych przepływem przemiennych prądów i strumieni.

Wpływ niedokładności wyważania wirnika i wskazówki

Jeżeli środek ciężkości elementów ruchomych selsyna O' nie leży na osi obrotu O (rys. 34), to na wirnik selsyna, poza momentem synchronizującym, działa siła ciężkości G , powodująca powstawanie uchybu



Rys. 34. Powstawanie momentu od niewyważenia wirnika

statycznego. Moment M_G spowodowany działaniem tej siły przy oznaczeniach jak na rys. 34 wyniesie:

$$M_G = l \cdot G \cdot \sin(\beta \pm \gamma). \quad (61)$$

Moment M_G jest równoważony przez moment synchronizujący odpowiadający pewnemu kątowi niezgodności $\Delta\delta_G$, stanowiącemu uchyb spowodowany niewyważeniem

$$\Delta\delta_G \cdot S = lG \sin(\beta \pm \gamma),$$

stąd

$$\Delta\delta_G = \frac{lG}{S} \cdot \sin(\beta \pm \gamma). \quad (62)$$

Uchyb statyczny spowodowany niewyważeniem elementów ruchomych selsyna jest więc sinusoidalnie zmienny z obrotem wirnika odbiornika, przy czym na jeden obrót przypada jeden pełny okres zmienności. Zmniejszenie uchybu może być osiągnięte dokładnym wyważeniem samego wirnika wraz ze wszystkimi sztywno połączonymi z nim ruchomymi elementami układu przekazywania wskazań lub momentu. Maksymalny moment od niewyważenia nie może w żadnym wypadku przekraczać wartości 10% sztywności:

$$M_{G_{max}} = lG \leq 0,1 S. \quad (63)$$

Szczególnie duży wpływ na dokładność pracy łącza ma niewyważenie w przypadku występowania wibracji i wstrząsów, co jest częstym

zjawiskiem w warunkach pracy selsynów. Uchyb wzrasta bowiem proporcjonalnie do przyspieszenia wibracji lub wstrząsu. Przy wstrząsach z przyspieszeniem 4g na przykład uchyb $\Delta\delta_G$ wzrasta 4-krotnie w stosunku do wartości spoczynkowej.

Wpływ niejednorodności właściwości magnetycznych stali

Przewodność magnetyczna blach elektrotechnicznych λ , z których pakietuje się stojany i wirniki selsynów, wykazuje różnicę przewodności wzdłuż i w poprzek kierunku walcowania dochodzące przy indukcjach stosowanych w selsynach do 20÷25%. Jeżeli pakietowanie wykonać przy zachowaniu zgodnego kierunku walcowania blach, to przewodność wzdłuż tego kierunku λ_{max} będzie znacznie większa niż w poprzek λ_{min} , co spowoduje wystąpienie uprzywilejowanych magnetycznie kierunków, do których będzie przyciągany wzbudzony wirnik. Moment M_λ spowodowany działaniem tego przyciągania wyrazi się w zależności od kąta obrotu wirnika odbiornika β wzorem

$$M_\lambda = C(\lambda_{max} - \lambda_{min}) \sin(2\beta \pm \gamma_1), \quad (64)$$

gdzie γ_1 — kąt między zerowym położeniem odbiornika a kierunkiem maksymalnej przewodności magnetycznej λ_{maks} , zaś uchyb spowodowany działaniem tego momentu, wzorem:

$$\Delta\delta_\lambda = \frac{C(\lambda_{max} - \lambda_{min})}{S} \cdot \sin(2\beta \pm \gamma_1). \quad (65)$$

Uchyb statyczny spowodowany niejednorodnością przewodności magnetycznej wtórnego obwodu selsyna jest więc sinusoidalnie zmienny z obrotem wirnika odbiornika, przy czym na jeden obrót przypadają dwa pełne okresy zmienności. Zmniejszenie uchybu jest możliwe przez ujednolicenie przewodności magnetycznej wzdłuż obwodu szczeliny powietrznej. Osiąga się to przez wachlarzowe pakietowanie blach, przy którym każda blacha jest obracana względem poprzedniej o jedną podziałkę żłobkową w jednym i tym samym kierunku.

Dla całkowitego zlikwidowania momentu M_λ należy, aby liczba blach pakietu N była:

$$\text{przy parzystej liczbie żłobków } z — N_p = \frac{n}{2} z, \quad (66)$$

$$\text{przy nieparzystej liczbie żłobków } z — N_n = n \cdot z,$$

gdzie n dowolna liczba całkowita.

Jeżeli zachowanie tych warunków nie jest możliwe, należy tak pakietować blachy, aby przy nieparzystej liczbie żłobków wykonać całkowitą liczbę obejść obwodu ($n \cdot 360^\circ$), a przy parzystej — całkowitą liczbę obejść połowy obwodu ($n \cdot 180^\circ$).

Wpływ zwartych obwodów po wtórnej stronie selsyna

Zwarte obwody, jak stwierdzono przy rozpatrywaniu teorii działania selsyna, zmniejszają przewodność magnetyczną w kierunku zgodnym z ich osią. Istnienie więc takich zwartych obwodów jest równoznaczne z istnieniem omówionej poprzednio niejednorodności przewodności magnetycznej wzdłuż obwodu szczeliny powietrznej selsyna, niezależnie od tego, czy taki zwarty obwód powstał przez zwarcie jednego lub kilku zwojów uzwojenia wtórnego, czy przez zwarcie blach pakietu, na którym to uzwojenie jest nawinięte. Moment M_{zw} spowodowany działaniem takiego zwartego obwodu będzie więc podobnie jak moment M_λ wyrażony wzorem:

$$M_{zw} = M_{zw \max} \sin(2\beta \pm \gamma_2) \quad (67)$$

i będzie wykazywał te same właściwości.

Ze względu na to, że zwarcie nawet jednego zwoju uzwojenia wtórnego znacznie zwiększa uchyb selsyna, należy szczególnie starannie zwoić selsyn i kontrolować wykonane uzwojenia na brak obecności zwartych zwojów. Zwarte obwody mogą powstać również w pakiecie stali wskutek niedokładności izolacji poszczególnych blach lub nieodpowiedniej obróbki gotowego pakietu. Dla uniknięcia tego blachy przed pakietowaniem powinny być dokładnie izolowane, a po spakietowaniu nie należy obrabiać pakietu, na który będą nawinięte wtórne uzwojenia selsyna. Dostateczną gładkość powierzchni szczeliny powietrznej od strony tego pakietu może zapewnić pakietowanie na odpowiednio stolerowanej oprawce (wewnętrznej dla selsyna z obwodem wtórnym na stojanie), oczywiście pod warunkiem odpowiednio ostrego wykrojnika lub po stosowaniu gradowania blach po wykrojeniu. Ewentualne zwarte obwody, powstające przy obróbce pakietu, na które będzie nawinięte wzbudzające uzwojenie pierwotne, aczkolwiek nie są pożądane ze względu na zwiększenie strat, to jednak nie wpływają na zwiększenie uchybów selsyna.

Wpływ momentów reluktancyjnych spowodowanych uzębieniem stojana i wirnika

W przypadku spakietowania blach stojana i wirnika bez odpowiednich skosów żłobków powstają momenty reluktancyjne, zależne od wzajemnego położenia zębów obydwu pakietów. Największe wartości przybiera ten moment w przypadku selsynów z wydatnymi biegunami, szczególnie w chwilach, gdy kraniec bieguna ustawia się naprzeciwko poszczególnych zębów pakietu obwodu wtórnego. W przypadku selsynów z utajonymi biegunami reluktancyjne momenty spowodowane uzębieniem pakietów są z reguły nieco mniejsze. Z uwagi jednak na stosunkowo znaczne wielkości tych momentów należy zarówno w selsy-

nach z utajonymi, jak i wydatnymi biegunami bezwzględnie stosować skos żłobków stojana o jedną podziałkę żłobkową wirnika lub skos żłobków wirnika o jedną podziałkę żłobkową stojana. Przykładowe doświadczenie wykazało, że zastosowanie skosu żłobków spowodowało około 40-krotne zmniejszenie omawianego szkodliwego momentu.

Dalsze obniżenie momentów reluktancyjnych od uzębienia w selsynach z wydatnymi biegunami może być osiągnięte przez takie dobieranie długości łuku nabiegunnika, aby stanowiła ona całkowitą wielokrotność podziałki żłobkowej obwodu wtórnego przy stosowaniu nieparzystej liczby żłobków tego obwodu.

Rozpatrując przyczyny powstawania reluktancyjnego momentu od uzębienia pakietów można wywnioskować, że wykres jego zmienności w funkcji kąta obrotu selsyna wyrazi się w przybliżeniu wzorem

$$M_z = M_{z \max} \cdot \sin(z\beta + \gamma_z), \quad (68)$$

czyli na jeden obrót przypadnie z pełnych okresów zmienności tego momentu.

Moment reluktancyjny spowodowany uzębieniem pakietów bądź pozostałości tego momentu nakładają się na omówione poprzednio krzywe uchybów o pojedynczym i podwójnym okresie zmienności na jeden obrót selsyna. W wyniku tego otrzymuje się kształt wykresu uchybów przedstawiony na rysunkach 33a i 33b.

Wpływ momentów spowodowanych nierównomiernością szczeliny powietrznej i innymi usterkami technologicznymi

Nierównomierność szczeliny spowodowana niedokładnym zachowaniem cylindryczności wirnika lub jego ekscentrycznym umieszczeniem względem powierzchni wewnętrznej stojana powoduje moment reluktancyjny o podwójnym okresie zmienności na obrót selsyna. Taki sam okres zmienności jest powodowany ewentualną nierównością liczby zwojów w zezwojach poszczególnych faz uzwojeń wtórnych lub nierównością ich czynnych oporności. Tak więc:

$$M_{\text{techn}} = M_{\text{techn max}} \sin(2\beta + \gamma_4). \quad (69)$$

Ogólnie można stwierdzić, że sumaryczny szkodliwy moment działający na wirnik selsyna i powodujący powstawanie jego uchybu statycznego składa się z czterech podstawowych składowych różniących się okresem zmienności:

$$M_h(\beta) = M_t(\beta) + M_{G \max} \sin(\beta \pm \gamma) + M_{R \max} \sin(2\beta \pm \gamma_R) + M_{z \max} \sin(z\beta \pm \gamma_z), \quad (70)$$

a mianowicie: a) momentów tarcia w łożyskach i stykach ślizgowych M_t o przypadkowej zmienności w funkcji kąta obrotu β , b) momentu od

niewyważenia ruchomych elementów selsyna M_G o pojedynczym okresie zmienności na jeden obrót selsyna, c) momentu M_R o podwójnym okresie zmienności na jeden obrót selsyna, będącego wypadkową momentów reluktancyjnych od nierównomierności przewodności magnetycznej pakietu M_z , od zwartych obwodów M_{zw} i momentów od usterek technologicznych M_{techn} , oraz d) momentu M_z o z -krotnym okresie zmienności na jeden obrót selsyna — reluktancyjnego od uzębienia pakietów stali selsyna.

W wyniku analizy zdjętej doświadczalnie krzywej uchybów statycznych można otrzymać odpowiednie składowe sinusoidalne wzoru (70), jak to pokazano na rys. 33c dla przypadku krzywej z rys. 33b.

Uchyby przy pracy kinematycznej

Ze wzrostem prędkości wirowania selsynów zmniejsza się moment synchronizujący i sztywność łącza, jak stwierdzono w punkcie 3.1.2. Wielkość spadku sztywności jest proporcjonalna do prędkości względnej wirowania selsynów łącza. Ponieważ przy danej prędkości wirowania selsynów łącza ich względna prędkość maleje ze wzrostem częstotliwości zasilania, oczywiste jest, że dla łączy przeznaczonych do pracy kinematycznej lepiej jest stosować selsyny wysokoczęstotliwościowe.

Wartości uchybów występujące przy kinematycznej pracy łącza nadawczo-odbiorczego można znaleźć rozwiązując nieliniowe równanie różniczkowe ruchu wirnika selsyna przy narzuconym rodzaju ruchu i przyjęciu niezbędnych założeń upraszczających.

Najczęściej spotykanymi w praktyce rodzajami pracy kinematycznej łącza jest wirowanie z prędkością jednostajną i z prędkością sinusoidalnie zmienną.

Przy obracaniu się wirników selsynów łącza z jednostajną prędkością obrotową n — podobnie jak w przypadku pracy statycznej — na uchyb wpływają momenty tarcia i sinusoidalnie zmienne momenty reluktancyjne. Poprzednio stwierdzono, że te ostatnie mają pojedynczy, podwójny lub krotny liczbę zębów z okres zmienności na jeden obrót selsyna, a ich wartości przy odpowiednim zaprojektowaniu selsyna są niewielkie. Dlatego uchyby łącza przy pracy kinematycznej na ogół niewiele różnią się od statycznych; jedynie przy prędkościach wirowania określonych wzorem

$$n_r = \frac{\nu}{k} \quad (k = 1, 2, z),$$

gdzie ν częstotliwość drgań własnych ruchomej części selsyna, zależna od sztywności selsyna i bezwładności jego wirnika, mogą wystąpić niebezpieczne zjawiska rezonansowych drgań wirnika, powodujące bardzo duże uchyby pracy łącza. Również przy prędkościach wirowania n , niewiele różniących się od rezonansowych n_r , obserwuje

się tym większy wzrost uchybu, im bliższe do siebie są wartości n i n_r . W tych przypadkach należy albo zmienić prędkość wirowania selsynów łącza, albo przekonstruować ruchomą część selsynów tak, aby rozstroić rezonans.

Przy obracaniu się wirników selsynów łącza z sinusoidalnie zmienną prędkością obrotową kąt obrotu wirników jest określony zależnością

$$\alpha = \alpha_m \sin \Omega t.$$

Jeżeli wprowadzimy pojęcie współczynnika kinematycznej dobroci selsyna

$$D_k = \frac{S}{J} \quad (71)$$

analogicznie do przyjętego poprzednio pojęcia współczynnika statycznej dobroci selsyna D_s , to uchyb łącza przy pracy kinematycznej w założeniu nieuwzględniania reluktancyjnych momentów można oszacować według wzoru

$$\Delta \delta'_{kin} = \sqrt{\left(\frac{\alpha_m \Omega}{D_k}\right)^2 + \left(\frac{\pi}{2D_s}\right)^2}. \quad (72)$$

W szczególnych przypadkach, gdy współczynnik dobroci kinematycznej jest znacznie większy niż statycznej — $D_k \gg D_s$, uchyby przy rozpatrywanym rodzaju pracy kinematycznej są proporcjonalne do statycznych i mogą być określane na podstawie znajomości tych ostatnich.

Momenty reluktancyjne podobnie jak w przypadku pracy łącza przy obracaniu się wirników z jednostajną prędkością mogą powodować występowanie dodatkowych uchybów przy przechodzeniu przez prędkości obrotowe odpowiadające rezonansowym. Najbardziej niebezpieczny jest przypadek, przy którym prędkość rezonansowa odpowiada maksymalnej prędkości kątowej obrotu wirników, gdyż w tym przypadku przedłuża się czas występowania zjawiska rezonansu.

Na zwiększenie uchybów łącza przy pracy kinematycznej o sinusoidalnie zmiennej prędkości obrotowej mogą wpływać również tłumiki mechaniczne. Mianowicie działanie ograniczników swobodnego ruchu masy zamachowej tłumika, występujące przed każdą zmianą kierunku przyspieszenia, spowoduje występowanie dodatkowych obciążeń udarowych wirnika, zwiększających uchyb pracy łącza. Obliczenie tych dodatkowych uchybów można wykonać rozpatrując równania ruchu wirnika i tłumika przy uwzględnieniu wielkości momentów tarcia w tłumiku i kąta rozstawienia ograniczników swobodnego ruchu jego masy zamachowej. Przy odpowiednim doborze tych wielkości można osiągnąć zmniejszenie wpływu działania tłumika na uchyb selsyna do dopuszczalnych granic.

Uchyby spowodowane zmianami napięcia i częstotliwości zasilania

Zmiany napięcia i częstotliwości zasilania wpływają na uchyb selsyna poprzez wpływ na jego sztywność.

Wpływ zmian napięcia zasilającego na sztywność selsyna wynika ze wzoru (25):

$$S = c \frac{E_{2f}^2}{f} \cdot \frac{X_q}{R_q^2 + X_q^2}.$$

Można od razu stwierdzić, że zmiana sztywności danego selsyna jest proporcjonalna do kwadratu zmiany napięcia zasilającego. A ponieważ, jak stwierdzono poprzednio, uchyb jest odwrotnie proporcjonalny do sztywności, wnioskujemy, że spadek napięcia zasilającego prowadzi do wzrostu uchybów selsyna. Z tego względu należy przestrzegać, aby wahania napięcia zasilającego nie przekraczały $\pm 10\%$ (dopuszczalnych w niektórych układach zasilających łącza selsynowe), a lepiej, aby były możliwie jak najmniejsze.

Spadek częstotliwości zasilania powoduje proporcjonalną zmianę oporności biernej obwodu wzbudzenia selsyna, a więc pewien wzrost prądu wzbudzającego i strumienia magnetycznego. Siła elektromotoryczna E_{2f} proporcjonalna do iloczynu tego strumienia i częstotliwości nieco maleje. Oporność bierna poprzeczna obwodu wtórnego X_q maleje proporcjonalnie do częstotliwości. Jak widać ze wzoru (25), wypadkowy wpływ spadku częstotliwości na sztywność selsyna nie jest oczywisty i zależy od parametrów decydujących o tym, czy przewagę będzie miało zmniejszenie się siły elektromotorycznej E_{2f} (występującej we wzorze w kwadracie), czy bezpośrednie działanie czynnika f w mianowniku i zmniejszenie się oporności X_q . W praktyce zmniejszenie się częstotliwości zasilania prowadzi przeważnie do pewnego wzrostu sztywności, a więc zmniejszenia uchybu. Dopuszczalne wahania częstotliwości wynoszą często w praktycznych układach $\pm 5\%$.

Ponieważ w większości przypadków wahania napięcia i częstotliwości zasilania są jednakowokierunkowe, to wpływy obydwu na sztywność selsyna i uchyb w pewnym stopniu wzajemnie się kompensują, jednakże przy przewadze wpływu napięcia.

Należy podkreślić, że przytoczone rozumowanie może dotyczyć jedynie niewielkich wahań omówionych parametrów zasilania wokół wartości znamionowych i w żadnym przypadku nie może być uogólniane na zmianę wartości znamionowych tych parametrów.

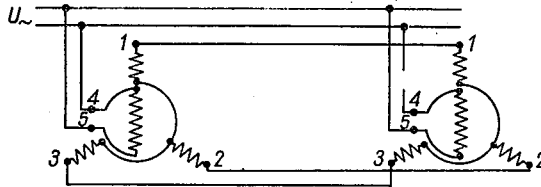
3.1.7. Uszkodzenia i wadliwe połączenia łącza

W praktyce może wystąpić kilka rodzajów uszkodzeń łącza nie przerywających jego pracy, lecz powodujących określone zmiany w tej pra-

cy. Powstanie takich uszkodzeń może prowadzić do mylnych odczytów otrzymanych z pozornie sprawnego łącza. Dlatego należy je rozpatrzeć i nauczyć się rozpoznawać po odpowiednich objawach zewnętrznych.

1. Przerwa zasilania obwodu wzbudzenia jednego z selsynów (rys. 35).

Jeżeli przerwany jest obwód wzbudzenia odbiornika, to przy zasilaniu nadajnika tylko w nim powstaje strumień magnetyczny indukujący odpowiednie siły elektromotoryczne w uzwojeniach wtórnych. W uzwo-



Rys. 35. Łącze z przerwą w obwodzie pierwotnym jednego z selsynów

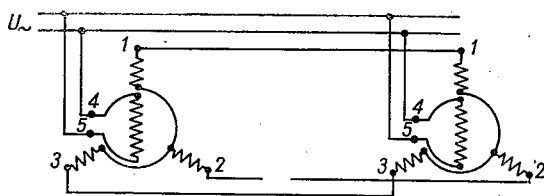
jeniach wtórnych odbiornika z braku strumienia wzbudzającego nie indukują się żadne siły elektromotoryczne. Z tego powodu pod wpływem sił elektromotorycznych wzbudzonych w nadajniku w obwodach wtórnych popłyną prądy niezależne od kąta niezgodności. Te prądy przepływając przez wtórne uzwojenia odbiornika wytworzą składowe strumienie, których suma będzie skierowana tak samo jak strumień wzbudzający w nadajniku. W ten sposób nastąpi przeniesienie kierunku oscylującego pola magnetycznego z nadajnika do odbiornika. Wirnik odbiornika będzie ustawiał się tak, aby kierunek jego największej przewodności magnetycznej, a więc kierunek podłużnej osi d , pokrył się z kierunkiem przeniesionego oscylującego pola magnetycznego.

Łącze będzie pracowało w sposób podobny do normalnego, jednakże występujący w tym przypadku reluktancyjny moment synchronizujący będzie znacznie mniejszy niż przy pracy normalnej. Będzie on mianowicie proporcjonalny do różnicy przewodności w osiach podłużnej i poprzecznej, a więc tym samym do różnicy $X_d - X_q$, występującej zarówno w selsynach z wydatnymi biegunami, jak i z utajonymi biegunami przy istnieniu zwartego obwodu tłumiącego w osi poprzecznej. Zmniejszy się oczywiście sztywność i zwiększy uchyb łącza. Ponadto wirnik odbiornika obrócony o 180° pozostanie w ustalonym położeniu równowagi, gdyż i w tym przypadku kierunek największej przewodności magnetycznej i przeniesionego strumienia będą zgodne, a moment reluktancyjny — równy zeru.

Tak więc w przypadku przerwy w obwodzie wzbudzenia selsyna-odbiornika w zakresie jednego obrotu jego wirnika wystąpią dwa położenia trwałej równowagi przesunięte względem siebie o 180° .

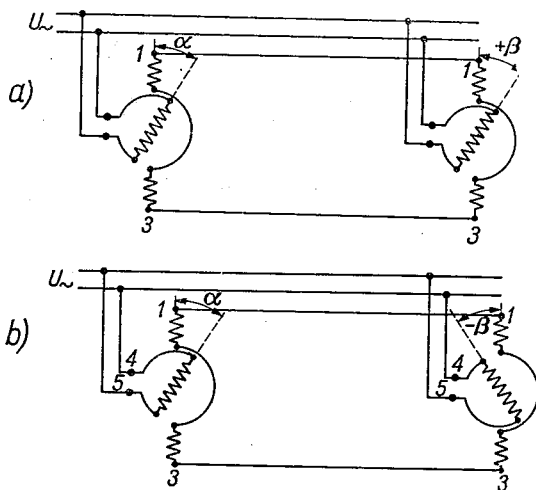
Jeżeli przerwa wystąpi w obwodzie wzbudzenia nadajnika przy wzbudzonym odbiorniku, to zewnętrzne objawy uszkodzenia pozostaną takie same.

2. Przerwa jednego z przewodów w obwodzie wtórnym łączy rys. 36



Rys. 36. Łączy z przerwą w obwodzie wtórnym

W tym przypadku wtórny obwód łączy można uważać za jednofazowy zastępując dwie szeregowo połączone fazy jedną równoważną im, jak na rys. 37. Z rysunku wynika, że wypadkowa siła elektromotoryczna E_2 działająca w obwodzie wtórnym łączy, a stanowiąca różnicę sił



Rys. 37. Obwód równoważny do podanego na rys. 36. Dwa położenia równowagi trwałej

elektromotorycznych wzbudzonych w uzwojeniach wtórnych nadajnika (E_{2n}) i odbiornika (E_{2o}) wynosi:

$$E_2 = E_{2n} - E_{2o} = E_{2m} (\cos \alpha - \cos \beta),$$

jeżeli $\alpha = \beta$, to $E_2 = 0$, we wtórnym uzwojeniu nie płynie prąd i układ pozostaje w stanie trwałej równowagi. Jeżeli $\alpha \neq \beta$, to w obwodzie wtórnym pojawia się prąd wywołujący moment synchronizujący usiłujący sprowadzić różnicę $(\cos \alpha - \cos \beta)$ do zera. Równanie

$$\cos \alpha - \cos \beta = 0$$

ma dwa rozwiązania dla kąta β , a mianowicie

$$\begin{aligned}\beta_1 &= \alpha, \\ \beta_2 &= -\alpha,\end{aligned}$$

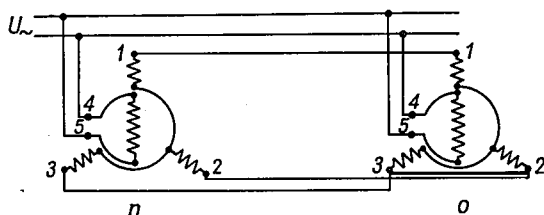
gdyż $\cos \alpha = \cos (-\alpha)$.

Każdemu położeniu wirnika nadajnika α odpowiadają więc dwa położenia trwałej równowagi wirnika odbiornika przesunięte względem siebie o kąt 2α .

Wirnik odbiornika może obracać się w lewo lub w prawo względem położenia równowagi, niezależnie od kierunku obrotu nadajnika, a w zależności od przypadkowych przyczyn.

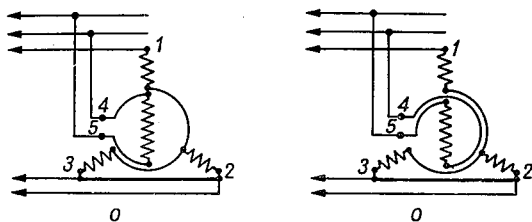
Wielkości momentu synchronizującego, a więc i sztywności, będą w tym przypadku nieco mniejsze niż normalnie wskutek braku działania jednej z trzech faz uzwojenia wtórnego.

3. Zwarcie dwóch przewodów w obwodzie wtórnym łączy (rys. 38)



Rys. 38. Łączy ze zwarcie dwóch faz obwodu wtórnego

Również w tym przypadku łączy staje się jednofazowym po stronie wtórnej. Poprzednio stwierdzono, że objawy zewnętrzne, charakteryzujące takie łączy, to nieokreślony kierunek obrotu odbiornika i nieco zmniejszona sztywność. Jednakże w rozpatrywanym przypadku istnieje zwarty obwód utworzony z uzwojeń faz 2 i 3, a moment wywołany przez prąd płynący w tym obwodzie we współdziałaniu ze strumieniem wzbudzenia jest większy od momentu łączy z jednofazowym uzwoje-



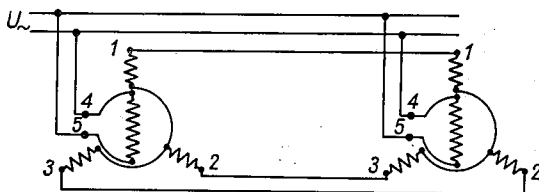
Rys. 39. Położenia trwałej równowagi odbiorników łączy z rys. 38

niem wtórnym i określa zachowanie się układu. Ten decydujący moment zależy od położenia wirnika odbiornika, a nie zależy od położenia wirnika nadajnika. Mianowicie przy $\beta = 0^\circ$ i $\beta = 180^\circ$ moment od zwartych uzwojeń 2 i 3 równa się zero (rys. 39). W obydwu przypadkach

uzwojenia faz 2 i 3 zajmują symetryczne położenie względem strumienia wzbudzenia i dlatego indukowane w nich siły elektromotoryczne będą równe sobie i przeciwnie skierowane, w zwartym obwodzie faz nie popłynie więc prąd i selsyn pozostanie w położeniu trwałej równowagi. Jeżeli obrócić wirnik odbiornika o pewien kąt od tego położenia, równowaga sił elektromotorycznych wzbudzonych w uzwojeniach faz 2 i 3 zostanie zachwiana i w ich zwartym obwodzie popłynie prąd stwarzający moment przywracający poprzednie położenie równowagi.

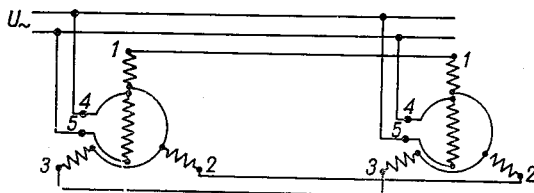
Zewnętrznym objawem zwarcia dwóch faz selsyna łącza będzie więc to, że w zakresie jednego obrotu wirnika odbiornika wystąpią dwa położenia równowagi trwałej przesunięte względem siebie o 180° i niezależne od położenia wirnika nadajnika.

4. Skrzyżowanie połączenia faz wtórnego obwodu łącza (rys. 40)



Rys. 40. Łącze ze skrzyżowanymi przewodami obwodu wtórnego

Z teorii maszyn indukcyjnych wiadomo, że kierunek obracania się pola magnetycznego przenoszonego przez trójfazowy układ uzwojeń przy skrzyżowaniu dwóch przewodów doprowadzających zostaje odwrócony. Zjawisko to występuje niezależnie od przyczyny wywołującej powstawanie obracającego się pola; dlatego również w przypadku łącza selsynowego ze skrzyżowanym połączeniem faz wtórnego obwodu kierunek obracania się wirnika odbiornika będzie przeciwny do kierunku obracania się wirnika nadajnika, przy zachowaniu bez zmian pozostałych właściwości prawidłowo połączonych układu.



Rys. 41. Łącze ze skrzyżowanymi przewodami obwodu pierwotnego

5. Odwrócone przyłączenie uzwojenia wzbudzenia jednego z selsynów (rys. 41)

Jeżeli skrzyżować przewody doprowadzające napięcie wzbudzenia do pierwotnego obwodu jednego z selsynów, to kierunek strumienia wzbú-

dzenia zmieni się w nim o 180° względem tegoż kierunku w drugim selsynie. Wskutek tego wirnik odbiornika obróci się o 180° i pozostanie przy dalszej pracy łączy przesunięty o ten stały kąt względem wirnika nadajnika. Po wyzerowaniu selsynów łączy będzie pracowało poprawnie, a omówione odwrócone połączenie uzwojeń wzbudzenia selsyna może być praktycznie nie uważane za wadliwe.

3.2. Łączy niesymetryczne

Niesymetrycznym będziemy nazywali łączy zestawione więcej niż z dwóch identycznych jednostek lub z dowolnej liczby różnych jednostek selsynowych o różnych sztywnościach charakterystycznych, a więc i różnych wymiarach. Jeżeli łączy będzie zestawione z dwóch niejednakowych selsynów, z których jeden będzie pełnił rolę nadajnika, a drugi odbiornika, nazwiemy je łączy pojedynczym. Jeżeli więcej niż z dwóch selsynów — nadajnika i kilku odbiorników — nazwiemy je łączy wielokrotnym. Zależności występujące przy rozpatrywaniu łączy wielokrotnego są ogólniejsze i łatwo wyprowadzić z nich przypadek niesymetrycznego łączy pojedynczego. Dlatego najpierw rozpatrzmy łączy wielokrotne.

3.2.1. Łączy wielokrotne

Jeden selsyn nadawczy może napędzać większą liczbę odbiorczych. Otrzymywane w tym przypadku wielojednostkowe łączy nazywa się łączy wielokrotnym. Sztywność takiego łączy różni się od sztywności charakterystycznej poszczególnych jego jednostek. Tę wypadkową sztywność łączy S_l można znaleźć po przeprowadzeniu szczegółowej analizy zastępczego dwufazowego łączy, której wyniki podamy bez dowodu.

Jeżeli zestawimy łączy wielokrotne z identycznych jednostek, to wypadkowa sztywność łączy będzie zmieniała się względem sztywności charakterystycznych nadajnika i odbiorników $S_{no} = S_n = S_o$ zgodnie ze wzorem:

$$S_l = S_{no} \frac{2}{1+n}, \quad (73)$$

gdzie n — liczba odbiorników. Zmniejszanie się sztywności łączy w miarę wzrostu liczby odbiorników ilustruje rys. 42.

Jeżeli zestawimy łączy wielokrotne z nadajnikiem różniącym się od odbiorników, przy czym sztywność charakterystyczna nadajnika będzie S_n , a odbiorników S_o , to wypadkowa sztywność łączy wyrazi się wzorem:

$$S_l = \frac{2S_n S_o}{S_n + nS_o}. \quad (74)$$

Ze wzoru (74) wynika, że jeżeli chcemy, aby łączy wielokrotne zachowało sztywność równą sztywności charakterystycznej odbiorników, to należy zastosować nadajnik o sztywności $S_n = n \cdot S_o$ — n -krotnie większej od sztywności odbiorników. Następny wniosek co do parametrów łączy wielokrotnego w tym przypadku można wyciągnąć, jeżeli rozpatrzmy fakt, że przy narzuconym znamionowym napięciu wtórnym dla zachowania sztywności odbiorników trzeba, aby przez ich uzwojenia płynął ten sam prąd, co podczas pracy w łączy charakterystycznym; będzie to możliwe, jeżeli oporność uzwojenia wtórnego nadajnika będzie n -krotnie mniejsza od oporności odpowiednich uzwojeń odbiorników. Również i podczas dalszych rozważań zgodnie z przytoczonym rozumowaniem zmianom sztywności nadajnika będzie towarzyszyła odwrotnie proporcjonalna zmiana oporności pozornej ich wtórnych obwodów.

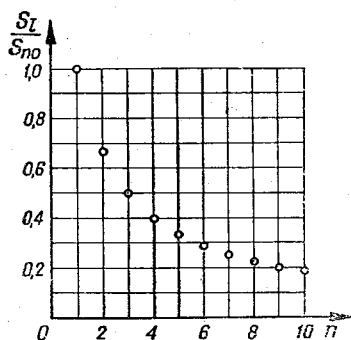
Oznaczając

$$\frac{S_o}{S_n} = k \quad (75)$$

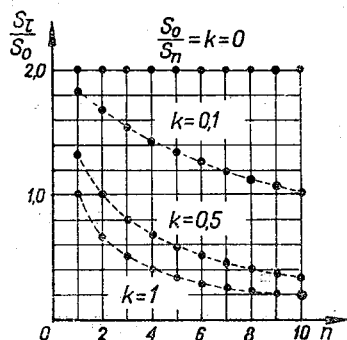
wzór (74) można przepisać w postaci:

$$S_l = S_o \frac{2}{1 + nk} \quad (76)$$

Na rys. 43 podano zależność sztywności łączy względem sztywności charakterystycznej odbiorników od ich liczby, przy rozmaitych wartościach k . Widać, że im mniejsza jest wartość k , czyli im większy i sztywniejszy jest nadajnik w stosunku do odbiorników, tym mniejszy wpływ



Rys. 42. Zmniejszanie się sztywności łączy zestawionego z jednakowych selsynów ze wzrostem liczby odbiorników



Rys. 43. Zmniejszanie się sztywności łączy wielokrotnego w miarę jego rozbudowy

ma zwiększanie liczby odbiorników i tym większą do określonej granicy otrzymuje się wypadkową sztywności łączy. W szczególności, dla $k = 1$ — czyli dla identycznych odbiorników i nadajnika — zależność przekształca się w podaną na rys. 42; natomiast dla wartości granicznej

$k = 0$ — czyli sztywności nadajnika nieskończenie większej od sztywności odbiorników — wypadkowa sztywność łącza staje się niezależna od liczby odbiorników i wynosi podwójną wartość sztywności charakterystycznej odbiorników. Tak więc łącze wielokrotne nie może mieć wypadkowej sztywności większej od podwójnej sztywności odbiorników. Widać również, że gdyby k zaczęło przybierać wartości większe od jedności, czyli zastosowalibyśmy nadajnik o sztywności mniejszej niż sztywność odbiorników, sztywność wypadkowa zmniejszyłaby się jeszcze bardziej, co byłoby oczywiście szkodliwe. Z tego właśnie powodu zaleca się, aby sztywność charakterystyczna nadajnika łącza wielokrotnego była zawsze większa od sztywności odbiorników.

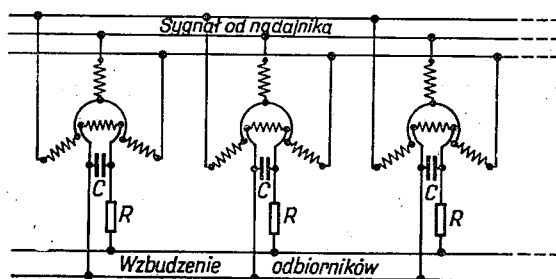
Ze szczegółowej analizy pracy łącza wielokrotnego wynika również, że wskutek połączenia obwodów wtórnych odbiorników w jeden układ występuje wzajemne ich oddziaływanie na siebie, co może wpływać na zwiększenie uchybów poszczególnych jednostek. Jeżeli na przykład na wał jednego z odbiorników będzie działał moment hamujący pożyteczny (obciążenie) lub pasożytniczy (zwiększone tarcie, zacięcie się wskazówki itp.), który zwiększy jego kąt niezgodności względem położenia nadajnika, to zmieniony skutek tego rozptyw prądu we wspólnym wtórnym obwodzie synchronizującym spowoduje powstanie większych kątów niezgodności (uchybów) również i na innych odbiornikach, których wały nie są poddane żadnym wpływom zakłócającym prawidłowość wskazań.

Jeżeli odbiorniki łącza wielokrotnego nie są identyczne, to największą sztywność ma ten z nich, w którym indukowana wtórna siła elektromotoryczna jest większa, a oporności pozorne faz uzwojenia pierwotnego są mniejsze. Przypadkowe zwiększenie kąta niezgodności tego najsztwniejszego odbiornika będzie miało większy zniekształcający wpływ na wskazania pozostałych odbiorników, niż miałyby takie samo zwiększenie kąta niezgodności innego, mniej sztywnego odbiornika.

Ograniczenie wzajemnego wpływu odbiorników jest możliwe przez zastosowanie nadajnika o większej sztywności lub ograniczenie momentów oddawanych przez odbiorniki. Doświadczenia wykazały, że włączenie do obwodów wzbudzenia odbiorników szeregowo oporności i boczniowo pojemności, jak na rys. 44, może spowodować znaczne obniżenie wzajemnego oddziaływania odbiorników bez obniżenia ich sztywności, przy jednoczesnym znacznym zmniejszeniu czasów przebiegów nieustalonych przy nagłych zmianach położenia (na przykład przy uwolnieniu osi odbiornika odchylonej o maksymalny kąt niezgodności). Przy stosowaniu takiego połączenia dodatkowych oporności i pojemności trzeba oczywiście odpowiednio dopasować napięcie sieci zasilającej.

Przesuwając fazę napięcia zasilającego poszczególne odbiorniki można na zasadzie omówionej w rozdz. 3.1.4 regulować wzajemne wpływy na siebie poszczególnych odbiorników. Jednakże zmniejszeniu wpływu

pierwszego odbiornika na drugi towarzyszy zawsze zwiększenie wpływu drugiego na pierwszy. Na przykład przy całkowitym wyeliminowaniu wpływu pierwszego odbiornika na drugi przez przesunięcie jego napięcia zasilającego o 45° , wpływ drugiego odbiornika na pierwszy wzrośnie



Rys. 44. Układ zmniejszający wzajemne oddziaływanie odbiorników łączy wielokrotnego

około $\sqrt{2}$ -krotnie w stosunku do przypadku zgodnej fazy napięć zasilających uzwojenia wzbudzenia tych selsynów. Wykorzystując te zjawiska można zestawić łączy wielokrotne z odpowiednio dobranymi sztywnościami poszczególnych odbiorników i wzajemnymi ich oddziaływaniami, stosując przesuwники fazowe lub odpowiednio dobrane kondensatory. Jednakże w tym przypadku należy dokładnie przeanalizować współdziałanie poprzesuowanych w fazie prądów i strumieni, gdyż znaczne ich przesunięcia fazowe wprowadzają do układu zniekształcenia, które powodują, że realne przebiegi mogą różnić się znacznie od wyidealizowanych teoretycznych przewidywań.

3.2.2. Łączy pojedyncze

W niesymetrycznym łączy pojedynczym, a więc zestawionym z jednego selsyna nadawczego o sztywności charakterystycznej S_n i jednego odbiorczego o sztywności S_o , sztywność wypadkową można łatwo obliczyć ze wzoru (74) kładąc $n = 1$. Wzór przyjmuje postać:

$$S_{11} = \frac{2S_n \cdot S_o}{S_n + S_o}; \quad (77)$$

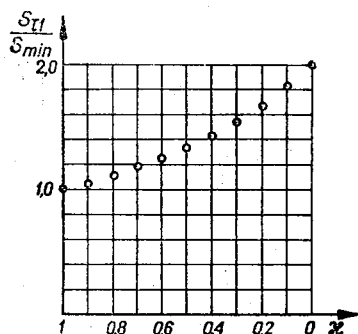
widzimy, że wzór jest całkowicie symetryczny względem obydwu sztywności charakterystycznych S_n i S_o . Wynika z tego, że przy łączy zestawionym z dwóch selsynów o różnych sztywnościach charakterystycznych sztywność wypadkowa będzie taka sama, niezależnie od tego, który z selsynów będzie użyty jako nadajnik, a który jako odbiornik.

Jeżeli teraz przez κ oznaczymy stosunek mniejszej ze sztywności charakterystycznych selsynów łącza — S_{min} — do większej, to wzór (77) można będzie przekształcić do postaci:

$$S_{l1} = S_{min} \frac{2}{1 + \kappa}, \quad (78)$$

z tym, że zawsze $\kappa \leq 1$.

Zależność S_{l1} od κ przedstawia rys. 45. Przy zmianie κ od 1 (łącze charakterystyczne z selsynów o jednakowych sztywnościach) do 0 (jeden



Rys. 45. Zmiany sztywności niesymetrycznego łącza pojedynczego ze zmianą stosunku sztywności jednostek

z selsynów o sztywności nieskończenie większej niż drugi) sztywność łącza rośnie od wartości sztywności słabszego selsyna do wartości skończonej — wynoszącej dwukrotną wartość sztywności słabszej jednostki.

W ten sposób sztywność wypadkowa niesymetrycznego łącza pojedynczego zawiera się zawsze w granicach:

$$S_{min} \leq S_{l1} \leq 2 S_{min}. \quad (79)$$

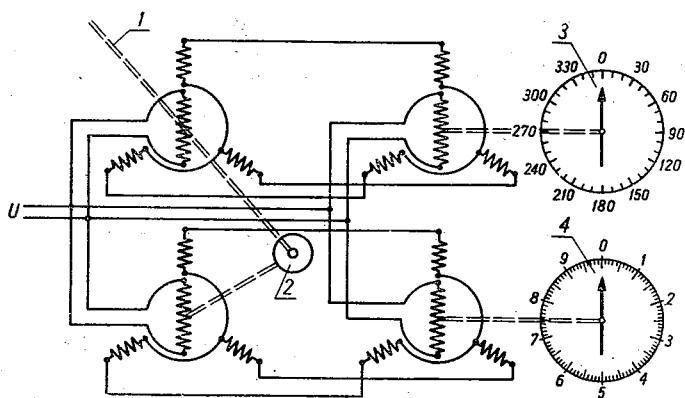
Wartość górnej granicy — $2 S_{min}$ — jest oczywiście w praktyce nieosiągalna.

3.2.3. Łącze dwutorowe

Dokładność przekazywania położenia kąowego obserwowanego wału jest ograniczona uchybem łącza i odczytu wskazania na skali odbiornika. Jak wspomniano, selsyny zaliczane do najdokładniejszej klasy I dają uchyb łącza dochodzący do $0,75^\circ$. Jeżeli do tego dodamy uchyb odczytu podziałki skali, która w wielu przypadkach praktycznego zastosowania łączy nie może być wykonana z dokładnością właściwą podziałkom precyzyjnych przyrządów pomiarowych, to może się zdarzyć, że otrzymamy wypadkową dokładność mniejszą od wymaganej w danym urządzeniu.

Dokładność zdalnego odczytania położenia wału przy pomocy selsynowego łącza nadawczo-odbiorczego można zwiększyć stosując przekładnię przyspieszającą obrót selsyna nadawczego. Przy zastosowaniu odpowiednio precyzyjnie wykonanej przekładni (szczególnie z automatycznym kasowaniem luzów), której uchyb własny jest pomijalny, błąd określenia położenia wału zmniejsza się praktycznie tylokrotnie, ilekrotna jest przekładnia. Wskazania łącza przestają być jednakże jednoznaczne, ponieważ jednemu obrotowi obserwowanego wału odpowiada krotność obrotów selsyna równa przekładni. Utrudnia to odczytanie i może powodować błędy. Ponadto w przypadku nagłego obrotu obserwowanego wału wirnik selsyna odbiorczego, wskutek swej bezwładności, może nie nadążyć za otrzymywanym od nadajnika sygnałem elektrycznym. Jeżeli kąt niezgodności w pewnej chwili, wskutek nagłej zmiany położenia wału nadajnika, przekroczy wartość odpowiadającą maksimum momentu charakterystyki łącza, to wał odbiornika osiągnie położenie zgodności z wałem nadajnika, opuszczając jeden lub kilka obrotów tego ostatniego. W celu uniknięcia powstania takiego zjawiska stosuje się tak zwane łącza dwutorowe.

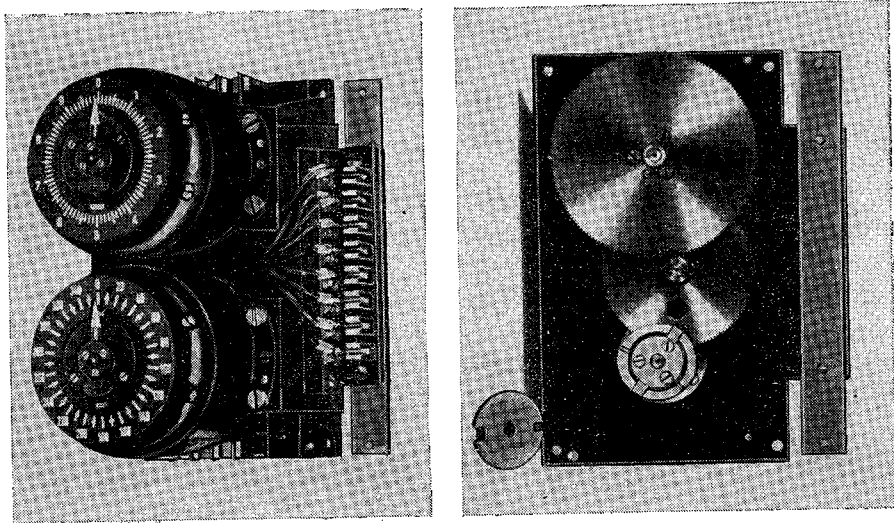
Łącze dwutorowe składa się z dwóch par selsynów niezależnie połączonych ze sobą elektrycznie. Jeden selsyn nadawczy jest sprzężony bezpośrednio z obserwowanym wałem, drugi selsyn nadawczy — przez prze-



Rys. 46. Układ połączeń łącza dwutorowego. 1 — wał obserwowany, 2 — przekładnia mechaniczna 1:36, 3 — podziałka odbiornika toru wskazań przybliżonych, 4 — podziałka odbiornika toru wskazań dokładnych

kładnię mechaniczną. Selsyny odbiorcze połączone odpowiednio ze swoimi nadawczymi są od siebie niezależne. Selsyn odbiorczy połączony z nadawczym sprzężonym bezpośrednio daje przybliżone odczytanie położenia obserwowanego wału. Drugi selsyn odbiorczy połączony z nadaw-

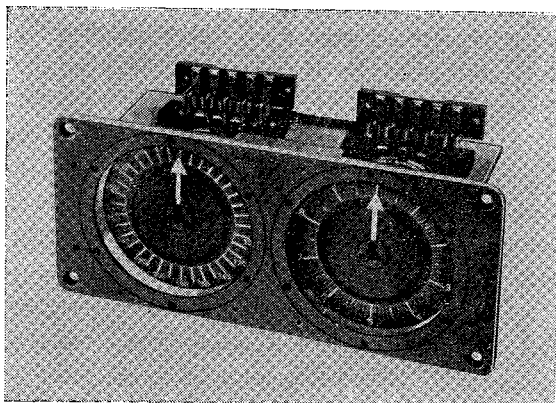
czym sprzężonym przez przekładnię daje odczytanie dokładne określające położenie obserwowanego wału w częściach podziałki skali przybliżonej. Dla dokładnego toru stosowane są przekładnie 1:10, 1:18, 1:36, 1:20 itp. pozwalające na odczytywanie położenia w wartościach stopni kątowych i ich części dziesiętnych, stopni i minut kątowych, „tysięcznych” oraz innych miar kątowych. Takie rozwiązanie przy umożliwieniu dużej dokładności określenia położenia obserwowanego wału eliminuje trudności odczytania i obawy powstania błędu przez „przepuszczenie” obrotu selsyna toru dokładnego, co w tym przypadku nie wpływa na odczytaną wartość. Układ połączeń i fotografie łączy dwutorowego podają rys. 46, 47 i 48.



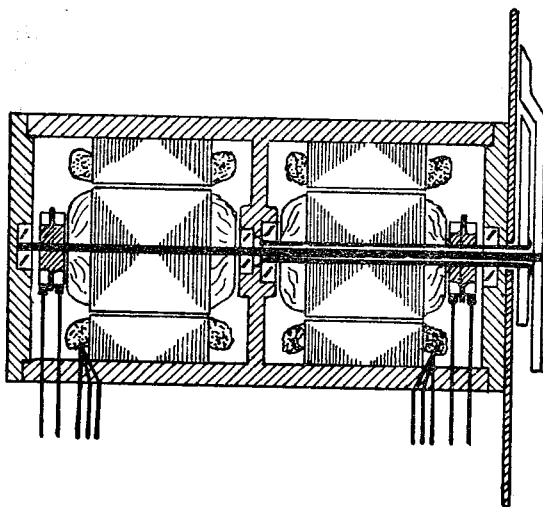
Rys. 47. Nadajniki łączy dwutorowego

Przy pomocy przytoczonego przykładowo na rysunkach rozwiązania łączy dwutorowego można bez trudu jednoznacznie określać położenie wału z uchybem mniejszym niż $0,5^\circ$, mimo zastosowania skal i wskazówek wygodnych do obserwacji, lecz nie odznaczających się szczególną precyzją wykonania.

Schemat innego wykonania konstrukcyjnego selsynów odbiorczych łączy dwutorowego przedstawiono na rys. 49. W tym przypadku dzięki zastosowaniu drażnionego wału jednego z selsynów otrzymano przybliżone i dokładne wskazania położenia obserwowanego wału na przyrządzie dwuwskazówkowym o jednej tarczy i ewentualnie podwójnej koncentrycznej podziałce.



Rys. 48. Odbiorniki łącza dwutorowego



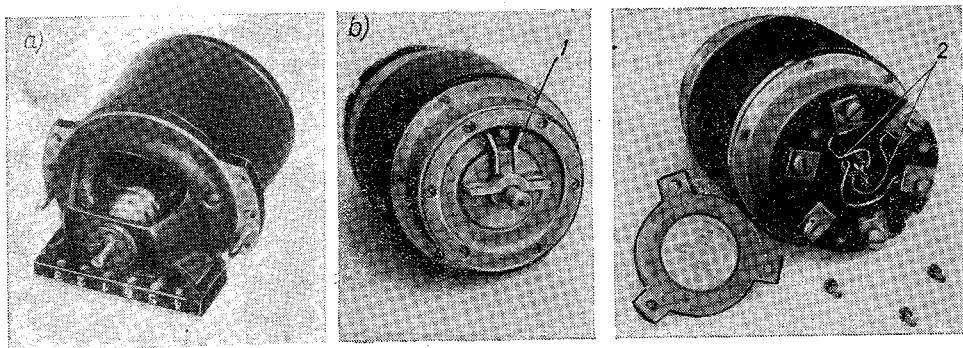
Rys. 49. Szkic konstrukcji odbiorników łącza dwutorowego ze współśrodkowymi wskaźówkami

3.2.4. Łącze różnicowe

3.2.4.1. ŁĄCZE RÓŻNICOWE ELEKTRYCZNE

Łącze różnicowe elektryczne jest modyfikacją łącza nadawczo-odbiorczego, umożliwiającą przekazywanie na drodze elektrycznej sumy lub różnicy obrotów dwóch wałów nadawczych. Zestawienie takiego łącza wymaga zastosowania dodatkowo selsyna o specjalnej konstrukcji, zwanego selsynem różnicowym. Selsyn różnicowy różni się od nadawczo-

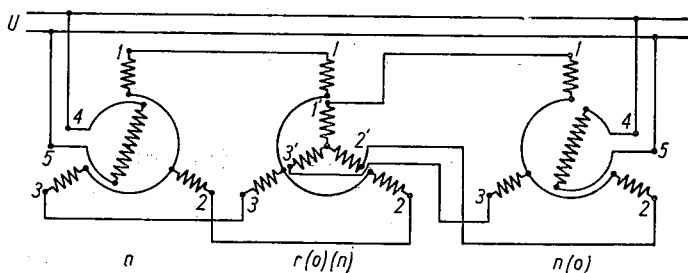
—odbiorczego z uzwojeniem pierwotnym umieszczonym na wirniku z utajonymi biegunami (rys. 4) tym, że wszystkie trzy fazy uzwojenia wirnika połączone w gwiazdę są wyprowadzone do zacisków zewnętrznych. Wyprowadzenie to wykonuje się przeważnie przy pomocy trzech pierścieni ślizgowych osadzonych na wale selsyna i odpowiedniego układu szczotek. W niektórych przypadkach, gdy przy pracy łączy przewiduje się jedynie niewielkie wychylenia wirnika selsyna różnicowego od określonego średniego położenia, zamiast pierścieni stosuje się połączenie uzwojeń faz z zaciskami na obudowie selsyna przy pomocy trzech giętkich przewodów. Dzięki takiemu rozwiązaniu uzyskuje się znaczne zmniejszenie tarcia przy ruchu wirnika selsyna oraz eliminuje się całkowicie możliwość iskrzenia, które może być niebezpieczne ze względu na warunki zewnętrzne w miejscu pracy selsyna. W przypadku połączenia wirnika z zaciskami stojana przy pomocy giętkich przewodów, na jego wale lub pakiecie wirnika umieszcza się zawsze pewnie zamocowany ogranicznik obrotu. Obydwie omówione konstrukcje selsyna różnicowego ilustruje rys. 50. Ogranicznik obrotu umocowany jest w pokazanym przypadku na wale selsyna.



Rys. 50 Selsyny różnicowe. a) z pierścieniami ślizgowymi i szczotkami, b) z ogranicznikiem obrotu na wale — 1 i z przewodami giętkimi — 2

Zasada działania łączy różnicowego zostanie rozpatrzona na przykładzie układu podanego na rys. 51. Siły elektromotoryczne indukowane w uzwojeniach wtórnych wzbudzonego selsyna nadawczego n wywołają prądy w obwodzie połączonych ze sobą faz stojanów tego selsyna i selsyna różnicowego r . Przepływy fazowe wytworzone przez te prądy sumując się w selsynie różnicowym przeniosą doń kierunek pola magnetycznego odpowiadający położeniu wirnika nadajnika n . W przypadku zgodności kierunków uzwojeń stojana i wirnika selsyna różnicowego (jak na rys. 51) przepływy fazowe zostaną przetransformowane i przekazane dalej z zachowaniem wzajemnego stosunku wielkości, wypadkowy kie-

runek pola magnetycznego zostanie zachowany bez zmiany, a odbiornik nie odczuje obecności selsyna różnicowego. Jeżeli natomiast wirnik selsyna różnicowego zostanie obrócony o pewien kąt względem stojana, to wskutek zmiany wzajemnych sprzężeń poszczególnych uzwojeń przetransformowane w selsynie różnicowym przepływy fazowe zmieniają wzajemny stosunek wielkości tak, że kierunek pola magnetycznego przeka-



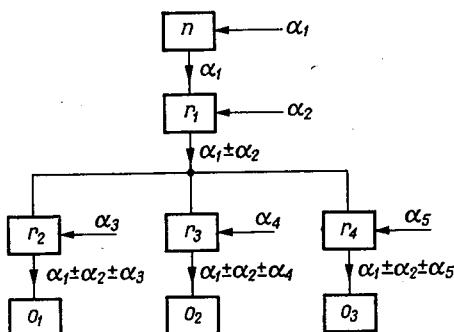
Rys. 51. Łącze różnicowe

zywany dalej do selsyna odbiorczego zostanie obrócony względem kierunku określonego przez położenie wirnika nadajnika n o kąt równy kątowni obrotu wirnika selsyna różnicowego r . Pod działaniem tego nowego kierunku pola magnetycznego selsyn odbiorczy obróci się o kąt odpowiadający sumie kątów obrotu wałów nadajnika i selsyna różnicowego. Z zasady działania łącza selsynowego i omówionych poprzednio przypadków niewłaściwego połączenia selsynów łącza wynika, że w przypadku skrzyżowania dowolnych dwóch przewodów selsyna różnicowego, selsyn odbiorczy będzie się obracał odpowiednio do różnicy kątów obrotu selsyna nadawczego i różnicowego.

W tym samym łączu obydwie selsyny nadawczo-odbiorcze mogą być używane jako nadawcze. W tym przypadku odbiorczym będzie selsyn różnicowy. Wskazania jego będą odpowiadały sumie lub różnicy kątów obydwu selsynów nadawczych w zależności od połączenia przewodów, jak to zostało omówione poprzednio. Możliwe jest również zestawienie rozgałęzionych układów z selsynami różnicowymi. Przykład jednego z takich rozwiązań pokazany jest na schemacie blokowym rys. 52.

Ścisła teoretyczna analiza zależności zachodzących w łączu różnicowym stanowi jeden z najbardziej skomplikowanych działów teorii selsynów. Fragmenty jej opracowania można znaleźć w [23], [44] i [46]. Dalej zostaną przytoczone bez dowodu podstawowe zależności występujące w tym łączu, które mogą być przydatne w praktyce obliczeniowej i eksploatacyjnej. Dotyczą one najprostszego, z teoretycznego punktu widzenia, przypadku łącza zestawionego z dwóch identycznych selsynów nadawczo-odbiorczych i selsyna różnicowego o jednakowych oporno-

ściach czynnych i biernych obwodu pierwotnego i wtórnego (stojana i wirnika). Większość spotykanych w praktyce łączy może być uważana z dostatecznym przybliżeniem za tak właśnie zbudowane.



Rys. 52. Przykładowy schemat blokowy rozgałęzionego łączy różnicowego. n — nadawnik, r — selsyny różnicowe, o — odbiorniki

Zachowując poprzednio stosowane oznaczenia parametrów selsynów nadawczo-odbiorczych i oznaczając parametry selsyna różnicowego:

- $R_{(r)}$ — oporność czynna fazy uzwojenia stojana lub wirnika,
- $X_{(r)}$ — oporność indukcyjna fazy uzwojenia stojana lub wirnika,
- $X_{(m)}$ — oporność indukcji wzajemnej uzwojeń stojana i wirnika przy zgodności ich osi,
- $X_{(s)}$ — oporność rozproszenia,

przy czym oczywiście $X_{(r)} = X_{(m)} + X_{(s)}$,

otrzymamy wzór na sztywność łączy różnicowego, w którym selsyn różnicowy pełni rolę odbiornika o postaci:

$$S_{r_o} = 2 \cdot 21,2 \frac{E_f^2}{f} \cdot X_{(m)} \times \\ \times \frac{(R'_q + R_{(r)})^2 + (X'_q + X_{(s)} + 2X_{(m)}) \cdot (X'_q + X_{(s)})}{[(R'_d + R_{(r)})^2 + (X'_d + X_{(s)} + 2X_{(m)})^2][(R'_q + R_{(r)})^2 + (X'_q + X_{(s)})^2]} \quad [\text{cm G/}^\circ]. \quad (80)$$

Dla zbadania wpływu poszczególnych parametrów jednostek selsynowych na sztywność rozpatrywanego łączy oznaczmy:

$$\frac{X'_q + X_{(s)}}{R'_q + R_{(r)}} = y; \quad \frac{X_{(m)}}{R'_q + R_{(r)}} = k. \quad (81)$$

Ponadto z dostateczną dla praktycznych przeliczeń dokładnością można przyjąć:

$$\frac{R'_d + R_{(r)}}{R'_q + R_{(r)}} \approx 1; \quad \frac{X'_d + X_{(s)}}{R'_q + R_{(r)}} \approx y. \quad (82)$$

Wówczas wzór na sztywność łącza różnicowego (80) przekształci się do postaci:

$$S_{r_o} = \frac{21,2}{2} \cdot \frac{E_f^2}{f} \cdot \frac{1}{R'_q + R_{(r)}} \cdot F(k, y), \quad (83)$$

gdzie:

$$F(k, y) = 4k \frac{1 + y(y + 2k)}{[1 + (y + 2k)^2](1 + y^2)} \quad (84)$$

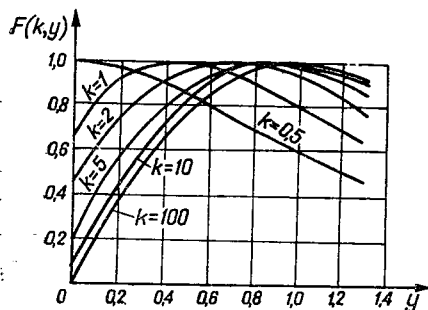
jest bezwymiarową funkcją, której wykres dla różnych wartości k jest podany na rys. 53.

Ostatnia postać wzoru na sztywność (83) wraz z wykresem na rys. 53 pozwalają na szybkie określenie sztywności łącza różnicowego, przy znajomości parametrów selsynów łącza oraz na wyciągnięcie niektórych wniosków, które należy uwzględnić przy projektowaniu łącza różnicowego, a mianowicie:

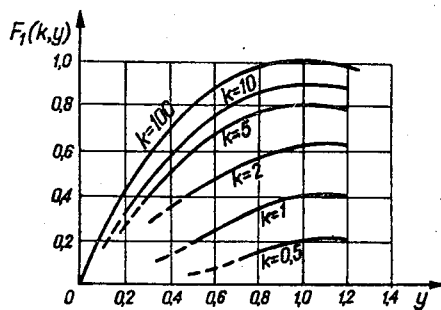
1. Należy możliwie zmniejszyć oporność czynną uzwojeń faz selsyna różnicowego.

2. Celowe jest wykonywanie selsyna różnicowego z możliwie małą szczeliną powietrzną.

3. Dla danej wartości k krzywe $F(k, y)$ na rys. 53 określają optymalną wartość y , a tym samym i rozproszenie uzwojeń faz selsyna różnicowego oraz parametry nadajników.



Rys. 53. Wykres funkcji $F(k, y)$ przy różnych wartościach k



Rys. 54. Wykres funkcji $F_1(k, y)$ przy różnych wartościach k

W przypadku gdy selsyn różnicowy jest drugim nadajnikiem, a rolę odbiornika pełni selsyn nadawczo-odbiorczy, postać wzoru na sztywność łącza analogiczna przytoczonej dla poprzedniego przypadku (83) przedstawia się jak następuje:

$$S_{r_n} = \frac{21,2}{2} \cdot \frac{E_f^2}{f} \cdot \frac{1}{R'_q + R_{(r)}} \cdot F_1(k, y), \quad (85)$$

gdzie:

$$F_1(k, y) = 4k \frac{y(y + 2k) - 1 + 2\alpha'_k}{[1 + (y + 2k)^2](1 + y^2)^2} \quad (86)$$

z tym, że:

$$\alpha'_k = \frac{R'_q - R_2}{R'_q + R_2} \quad (87)$$

Wykres funkcji $F_1(k, y)$ podany jest na rys. 54 w założeniu małych wartości α'_k .

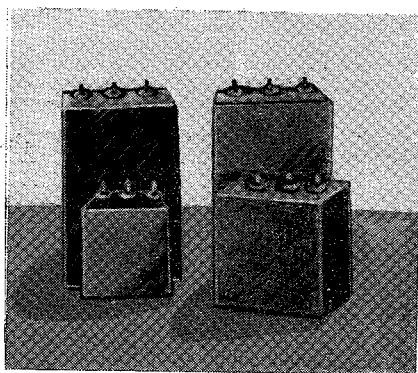
Z porównania wzorów (83) i (85) widać, że sztywność łącza z selsynem różnicowym pełniącym rolę nadajnika jest mniejsza niż w przypadku wykorzystania go jako odbiornika. Dla zmniejszenia tej różnicy należy i w tym przypadku dążyć do zwiększenia wartości k , na przykład poprzez zmniejszenie szczeliny powietrznej.

Ze szczegółowej analizy zależności występujących w łączy różnicowym wynika ponadto, że najkorzystniejsze jest zasilanie obydwu selsynów wzbudzanych bezpośrednio z sieci zasilającej napięciami pierwotnymi zgodnymi ze sobą w fazie oraz przypadek spełnienia zależności (30), a mianowicie $R_q = X_q$ przez obydwa trójfazowe odcinki łącza, z tym, że wartości oporności czynnych i biernych dla obwodu zawartego między stojanem pierwszego nadajnika a stojanem selsyna różnicowego, niekoniecznie muszą być równe odpowiednim wartościom dla obwodu zawartego między wirnikiem selsyna różnicowego a stojanem drugiego selsyna wzbudzanego bezpośrednio z sieci zasilającej. Przy spełnieniu tych warunków otrzymuje się minimum obciążenia wałów nadajników przy maksymalnej sztywności na wale odbiornika.

Selsyn różnicowy czerpie prąd magnesujący z sieci zasilającej poprzez uzwojenia stojanów i wirników współpracujących z nim selsynów nadawczo-odbiorczych. Ten dodatkowy prąd płynący przez uzwojenia współpracujących selsynów wywołuje niepotrzebne nagrzewanie się tych ostatnich. Ponieważ uzwojenia selsyna są obliczane w założeniu określonego wykorzystania cieplnego, dodatkowy prąd płynący przez nie do selsyna różnicowego może zmusić do obniżenia napięcia zasilającego łącze, a przez to do zmniejszenia czynnych strumieni magnetycznych selsynów, co spowoduje w dalszej konsekwencji spadek sztywności łącza. Aby tego uniknąć, należy skompensować składową indukcyjną dodatkowego prądu przepływającego przez uzwojenia selsynów połączonych z siecią. W tym celu stosuje się układy kondensatorowe, które mogą być wykonane jako zespół trzech kondensatorów połączonych w trójkąt lub gwiazdę i umieszczonych we wspólnej obudowie, jak pokazano na rys. 55.

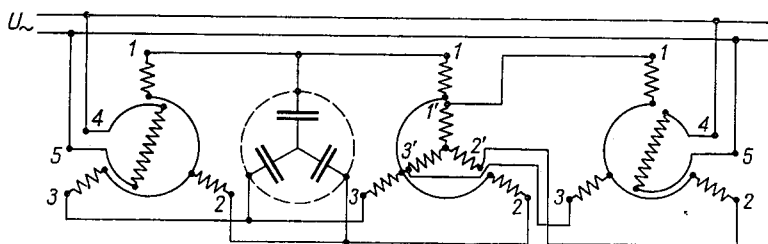
Odpowiedni dobór pojemności kondensatorów do indukcyjności uzwojeń selsyna różnicowego pozwala na znaczne zmniejszenie dodatkowego prądu indukcyjnego płynącego przez uzwojenia współpracujących sel-

synów, utrzymanie odpowiedniej wartości napięcia zasilającego łącze i osiągnięcie znacznie większej sztywności roboczej łącza różnicowego, niż w przypadku, gdy kondensatory nie są zastosowane. Działanie kondensatorów jest poza tym w pewnym stopniu równoważne zmniejsze-



Rys. 55. Kondensatory kompensujące

niu szczeliny powietrznej selsyna różnicowego, którego ważność była podkreślana we wnioskach z rozpatrywania wzorów (83), (84) i rys. 53. Rozważania teoretyczne sprawdzone doświadczeniami wykazują, że dla ograniczenia składowej indukcyjnej dodatkowego prądu płynącego w oby-



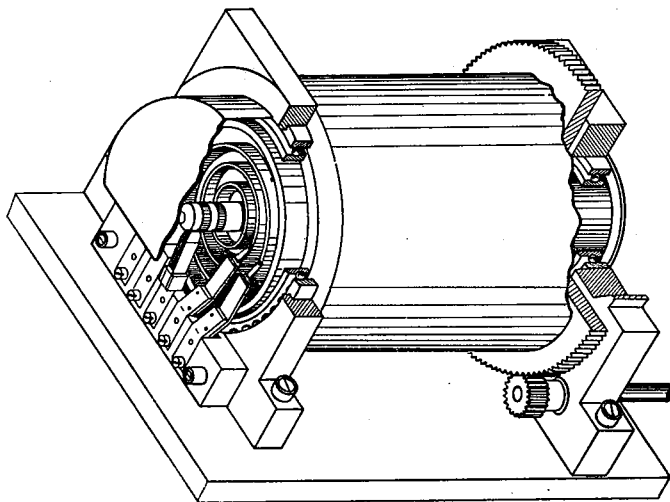
Rys. 56. Schemat łącza różnicowego z kondensatorami kompensującymi

dwu selsynach współpracujących z różnicowym, kondensatory można przyłączyć w dowolnym miejscu trójfazowej linii przewodowej, a więc zbocznikować kondensatorami jeden tylko zespół uzwojeń (na stojanie lub na wirniku) selsyna różnicowego. Układ połączeń łącza różnicowego skompensowanego kondensatorami podany jest na rys. 56.

3.2.4.2. ŁĄCZE RÓŻNICOWE MECHANICZNE

Łącze różnicowe mechaniczne stosowane jest do przekazania sumy lub różnicy obrotu dwóch wałów w tym przypadku, gdy możliwe jest ich sprzężenie mechaniczne poprzez odpowiednio wykonaną obudowę

selsyna nadawczego łączy nadawczo-odbiorczego. Takie łączy różnicowe różni się od nadawczo-odbiorczego tylko mechanicznym wykonaniem obudowy stojana nadajnika i jej obrotowym umocowaniem w miejscu zainstalowania. Dodawany lub odejmowany od położenia wirnika nadajnika kąt jest wprowadzany do łączy przez obrót w odpowiednim kierunku



Rys. 57. Przykład konstrukcji obudowy selsynów łączy różnicowych mechanicznych

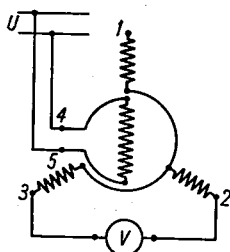
ku stojana nadajnika. Obrót ten jest urzeczywistniany przeważnie przy pomocy przekładni zębatej lub ślimakowej między zadającym wałem dodawczym a stojanem selsyna. Przykład wykonania obudowy nadajnika takich łączy podaje rys. 57.

3.2.5. Zerowanie selsynów

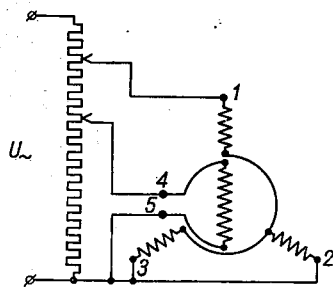
Po wmontowaniu selsynów łączy nadawczo-odbiorczego w miejscu zainstalowania należy przeprowadzić ich zerowanie. Polega ono na dokładnym ustaleniu położenia wirników selsynów względem ich stojanów, w którym łączy znajduje się w równowadze i zamocowaniu na wałach wskazówek odpowiednio do punktu „0” wygodnie ustawionych skal. W celu ułatwienia zerowania i umożliwienia najwygodniejszego ustawienia skal (przyjęło się ustawianie skal z punktem „0” skierowanym pionowo ku górze) selsyny powinny być wykonane w sposób pozwalający bądź na obracanie i zamocowywanie skali i wskazówki w dowolnym położeniu względem nieruchomego kadłuba i wału selsyna, bądź na obracanie i umocowywanie całego kadłuba selsyna w gnieździe urządzenia, niezależnie od skali przy ustalonym położeniu wskazówki na wale. Gniaz-

do urządzenia wykonane jest zazwyczaj tak, aby umożliwić łatwą wymianę jednostki selsynowej.

Dla ustalenia wspomnianego zgodnego położenia uzwojeń stojana i wirnika jednostkę nadawczo-odbiorczą lub różnicową o wirniku nie mogącym obracać się swobodnie, jak na przykład nadajnik sprzężony z obserwowanym wałem lub odbiornik obciążony pomocniczym urządzeniem, którego moment oporowy mógłby wprowadzić uchyb ustawienia, należy połączyć odpowiednio wg rys. 58 i obracając wirnik doprowadzić napięcie wskazywane przez woltomierz do zera. Woltomierz po-



Rys. 58. Układ do zerowania selsyna przy pomocy woltomierza



Rys. 59. Układ do samoczynnego zerowania się selsyna

winien mieć możliwie duży opór wewnętrzny i możliwie dużą dokładność. Samoczynne ustawienie się pojedynczego selsyna w położeniu zerowym można uzyskać łącząc go odpowiednio wg rys. 59. Potencjometr służy do doboru odpowiednich wartości napięcia pierwotnego i wtórnego. Odbiorniki o wirniku mogącym obracać się swobodnie, jak na przykład odbiorniki w łączy wskaźnikowym, ustawiają się samoczynnie w położeniu równowagi trwałej po zasileniu łącza. Ustawienie to można sprawdzić i ewentualnie skorygować przy pomocy woltomierza połączonego jak na rys. 58 do tych końcówek uzwojenia wtórnego, na których występuje najmniejsze indukowane napięcie.

4. WYTYCZNE PROJEKTOWANIA

4.1. Uwagi o konstrukcji

4.1.1. Rodzaje budowy

Z rozpatrywania zasady działania łącza nadawczo-odbiorczego wynika, że uzwojenie wzbudzające selsyna powinno być dwubiegunowe, gdyż tylko w tym przypadku na jeden obrót wału jego wirnika przypada jedno ustalone położenie równowagi trwałej łącza. Przy wielobieguno-

wym rozwiązaniu liczba ustalonych położeń równowagi trwałej łączy byłaby równa liczbie par biegunów, co w ogólnym przypadku byłoby niekorzystne.

Dwubiegunowe uzwojenie wzbudzające może być, jak wspomniano w rozdz. 2, umieszczone zarówno na stojanie, jak i na wirniku selsyna nadawczo-odbiorczego. Rozpatrzmy budowę, zalety i wady poszczególnych rozwiązań konstrukcji, zdefiniowanych i ponumerowanych jak następuje:

1. Pierwsze rozwiązanie konstrukcyjne — uzwojenie wzbudzenia na wydalnych biegunach stojana; uzwojenie synchronizujące w żłobkach cylindrycznego wirnika wyprowadzone do trzech pierścieni ślizgowych. Schemat elektryczny tego rozwiązania konstrukcyjnego ilustruje rys. 2.

2. Drugie rozwiązanie konstrukcyjne — uzwojenie wzbudzające na wirniku z wydalnymi biegunami, zasilane przez 2 pierścienie i szczotki; uzwojenie synchronizujące — w żłobkach stojana. Schemat elektryczny drugiego rozwiązania konstrukcyjnego ilustruje rys. 3.

3. Trzecie rozwiązanie konstrukcyjne — uzwojenie wzbudzające na wirniku z utajonymi biegunami, zwarta pętla w osi poprzecznej wykonana wewnątrz maszyny, zasilanie uzwojenia wzbudzenia przez dwa pierścienie i szczotki; uzwojenie synchronizujące w żłobkach stojana. Schemat elektryczny trzeciego rozwiązania konstrukcyjnego ilustruje rys. 4.

4. Czwarte rozwiązanie konstrukcyjne różni się od trzeciego tym, że wszystkie trzy fazy wirnika są wyprowadzone do trzech pierścieni ślizgowych i poprzez szczotki na tabliczkę zaciskową selsyna.

Przy pierwszym rozwiązaniu konstrukcyjnym stojan selsyna składa się z pierścieniowego jarzma i dwóch wydalnych biegunów. Przeważnie blacha stojana wykonywana jest w całości jako jeden wykrój. Przy takim rozwiązaniu utrudnione jest zakładanie cewek wzbudzenia na bieguny, więc czasem stosuje się konstrukcje z odejmowanymi biegunami. Zastosowanie odejmowanych biegunów komplikuje jednakże konstrukcję, technologię i montaż selsyna i wprowadza dodatkową szczelinę powietrzną, co z kolei powoduje zwiększenie prądu stanu jałowego selsyna i zmniejszenie pożytecznego wykorzystania jego materiałów czynnych. Z tych względów rozwiązanie z odejmowanymi biegunami należy uznać za gorsze. Niezależnie jednak od konstrukcji obwodu magnetycznego wykorzystanie objętości wewnętrznej selsyna jest niepełne. Między obydwooma cewkami wzbudzenia oraz między nimi i obudową musi pozostać wiele powietrznej przestrzeni.

Dla zwiększenia różnicy przewodności drogi strumienia w osiach podłużnej i poprzecznej selsyna stosuje się nierównomierną szczelinę powietrzną pod biegunem, a to w ten sposób, że promień łuku wytoczenia nabiegunnika wydłuża się względem promienia powierzchni zewnętrz-

nej wirnika o wielkość przekraczającą grubość szczeliny powietrznej pod środkiem bieguna. To wydłużenie dobiera się tak, aby szczelina powietrzna przy krańcu bieguna była około 10-krotnie większa niż w osi bieguna. Dzięki takiemu ukształtowaniu bieguna można osiągnąć kilkukrotne zmniejszenie strumienia poprzecznego przy nieznacznym tylko zmniejszeniu strumienia podłużnego. Ponadto wskutek zmniejszenia wartości strumienia pod krańcem bieguna zmniejsza się szkodliwe reluktancyjne oddziaływanie krawędzi bieguna. Należy jednakże zaznaczyć, że zastosowanie nierównomiernej szczeliny nie zawsze prowadzi do wymienionych korzystnych zjawisk, gdyż może przeważać wpływ niewielkiego zmniejszenia strumienia podłużnego.

Wirnik selsyna o pierwszym rozwiązaniu konstrukcyjnym, stanowiący obwód uzwojenia wtórnego, powinien być izotropowy pod względem przewodności magnetycznej. Właściwość tę osiąga się przez wachlarzowe pakietowanie blach i skos zębów omówione szczegółowo w rozdz. 3.1.6.

Styki ślizgowe umieszczone we wtórnym obwodzie (synchronizacji) są obciążone prądowo tylko podczas występowania niezgodności położeń selsynów łączy. Z tego względu niektóre źródła [44] podają, że można stosować w nich mniejsze dociski szczotek, co z kolei zmniejsza moment tarcia przypadający na jeden styk blisko o połowę względem pozostałych rozwiązań konstrukcyjnych selsynów. Należy jednak pamiętać, że napięcie na styku jest zależne od kąta niezgodności i przy najważniejszych praktycznie — małych kątach niezgodności występuje niebezpieczeństwo utraty styku, które nawet będąc chwilowym może wywołać całkowite zniekształcenie wskazań. Wobec tego zmniejszenie docisku staje się niebezpieczne, a ze względu na większą liczbę zestyków wypadkowy moment tarcia selsynów omawianego pierwszego rozwiązania konstrukcyjnego jest w praktyce większy niż przy pozostałych rozwiązaniach konstrukcyjnych.

W selsynach z wydajnymi biegunami na stojanie nie ma możliwości zastosowania tłumienia elektrycznego przy pomocy zwartego obwodu w osi poprzecznego strumienia. Niezbędne staje się zatem stosowanie w miarę potrzeby wyłącznie tłumików mechanicznych, które, jak wiadomo (rozdz. 3.1.5), nie zawsze są pewne w działaniu, komplikują konstrukcję i mogą zwiększać uchyby przy pracy kinematycznej łączy.

Wpływ wszystkich omówionych właściwości selsynów pierwszego rozwiązania konstrukcyjnego powoduje również i to, że przy danej wartości sztywności ich objętość i ciężar są większe niż selsynów o pozostałych rodzajach rozwiązań konstrukcyjnych, których wspólną cechą charakterystyczną jest umieszczenie uzwojenia wzbudzenia na wirniku.

Przy drugim rozwiązaniu konstrukcyjnym cewki wzbudzenia są nawijane na wydajnych biegunach umieszczonych na wir-

niku. Ta konstrukcja pozwala na stosowanie uzwojeń tłumiących umieszczonych w nabiegunkach (zob. rozdz. 3.1.5) i wyeliminowanie tłumików mechanicznych. Jak wiadomo, uzwojenie tłumiące może wyraźnie polepszyć właściwości selsyna.

Zmniejszenie przewodności drogi dla strumienia poprzecznego otrzymuje się w drugim rozwiązaniu konstrukcyjnym poza działaniem ewentualnego zwartego uzwojenia tłumiącego, które może składać się z 1 do 4 zezwojów również podobnie jak przy pierwszym rozwiązaniu przez ukształtowanie nierównomiernej szczeliny powietrznej. Przykład odpowiedniego wykroju blachy wirnika podany jest na rys. 62. Ponadto długość łuku nabiegunka jest zazwyczaj skracana do rozwarcia kąтового maksimum $80^\circ \div 90^\circ$. Przy obieraniu długości nabiegunka należy pamiętać, że reluktancyjny moment współdziałania zębów z krawędzią nabiegunka powodujący harmoniczną uchybu od uzębienia zależy od stosunku długości nabiegunka do podziałki żłobkowej stojana.

Twierdzenie spotykane w niektórych publikacjach omawiających właściwości selsynów z wydatnymi biegunami na wirniku, że ich sztywność jest większa niż sztywność selsynów z utajonymi biegunami (trzecie i czwarte rozwiązanie konstrukcyjne) dzięki dodatkowemu działaniu momentu reluktancyjnego wywołanego kształtem wirnika, jest, jak udowodniono w [2], niesłuszne. Występujący rzeczywiście moment reluktancyjny może tylko zmniejszać, a nigdy zwiększać, sztywność selsyna.

Zalety i wady drugiego rozwiązania konstrukcyjnego spowodowane umieszczeniem uzwojenia wzbudzenia na wirniku zostaną omówione łącznie dla drugiego, trzeciego i czwartego rozwiązania konstrukcyjnego.

Blachy stojana dla uzyskania magnetycznej izotropowości nawiniętego w ich żłobkach wtórnego obwodu selsyna powinny być oczywiście pakietowane wachlarzowo, zgodnie ze wskazówkami rozdz. 3.1.6.

Trzecie rozwiązanie konstrukcyjne, przy stojanie identycznym jak w rozwiązaniu drugim, różni się nawinięciem cewek wzbudzenia w żłobkach cylindrycznego wirnika i połączeniem ich w sposoby podany na rys. 4, zapewniający powstanie dwubiegunowego uzwojenia z utajonymi biegunami i zwartą pętlą uzwojenia w osi poprzecznej selsyna.

Przy takiej konstrukcji zwarte uzwojenie tłumiące stanowi element organiczny pierwotnego obwodu, zapewniając jednocześnie automatycznie różnice przewodności strumienia magnetycznego w podłużnej i poprzecznej osi selsyna. Unika się również automatycznie szkodliwego oddziaływania krawędzi nabiegunków wydatnych biegunów. Odpowiedni dobór stosunku liczby zębów i skosów żłobków zgodnie ze wskazówkami w rozdz. 1.1.6, pozwala na zmniejszenie szkodliwych momentów reluktancyjnych do praktycznie pomijalnych wartości.

Czwarte rozwiązanie konstrukcyjne różni się od trzeciego wyłącznie tym, że wszystkie końce faz uzwojenia wirnika połączonego w gwiazdę są poprzez pierścienie i szczotki wyprowadzone do zacisków selsyna. Różnica sprowadza się więc do dodatkowego pierścienia ze szczotkami. To rozwiązanie konstrukcyjne przewidziane jest w zasadzie dla selsyna różnicowego, gdzie istnienie trzech pierścieni wirnika jest niezbędne ze względu na zasadę działania łącza różnicowego. W szczególnych przypadkach selsyn o takiej konstrukcji może być użyty jako nadawczo-odbiorczy, jeżeli odpowiednie połączenie faz jego wirnika (odpowiadające pokazanemu na rys. 4) zostanie zrealizowane na zewnętrznych zaciskach maszyny. Takie postępowanie nie może być jednak zalecane, gdyż nie dając żadnych korzyści jest gorsze w porównaniu z trzecim rozwiązaniem konstrukcyjnym z powodu zwiększenia liczby styków, a tym samym szkodliwego momentu tarcia.

Dla drugiego, trzeciego i sztucznie połączonego czwartego rozwiązania konstrukcyjnego, których wspólną cechą jest umieszczenie uzwojeń wzbudzenia na wirniku, oporności przejścia na stykach szczotkowych są włączone w obwód pierwotny i ich zmiany nie wpływają praktycznie na dokładność przekazywania wskazań. Nawet przerwa styku zwiększa jedynie uchyb, jak omówiono w rozdz. 3.1.7, a nie zmienia w sposób zasadniczy wskazań odbiornika, jak to ma miejsce w pierwszym rozwiązaniu konstrukcyjnym. W obwodzie styków panuje ciągle stałe i stosunkowo duże napięcie, dzięki czemu chwilowe przerwy przewodzenia są mniej prawdopodobne. Mniejsza liczba zestyków prowadzi z reguły w praktyce do mniejszej wartości szkodliwego momentu tarcia.

Jedyną właściwością tych rozwiązań, która może być uważana za ich wadę, jest stały przypływ prądu wzbudzenia przez zestyki szczotkowe, jednakże uwzględniając z reguły niewielkie wartości tych prądów tę wadę można uważać za nieistotną.

Rozwiązania konstrukcyjne ze wzbudzeniem na wirniku należy więc uważać za lepsze od pierwszego rozwiązania konstrukcyjnego z uzwojeniami wzbudzania na stojanie. Natomiast decyzja wyboru pomiędzy drugim a trzecim rozwiązaniem konstrukcyjnym, a więc z wydatnymi czy z utajonymi biegunami na wirniku, zależy od wymiarów i częstotliwości zasilania selsyna.

Najkorzystniejsze ogólnie jest trzecie rozwiązanie konstrukcyjne selsyna — z utajonymi biegunami na wirniku. Eliminuje ono szkodliwe od działywanie krawędzi nabiegunnika i pozwala na łatwe sprowadzenie do pomijalnych wartości pozostałych momentów reluktancyjnych. To rozwiązanie powinno być w miarę możliwości stosowane przede wszystkim. W przypadku jednakże małych dla danej częstotliwości zasilania jednostek selsynowych zwarte uzwojenie w osi poprzecznej wirnika z utajonymi biegunami nie zapewnia dostatecznego zmniejszenia prze-

wodności magnetycznej w tej osi. Ponadto w przypadku najmniejszych jednostek selsynowych o średnicach zewnętrznych wirnika rzędu 20 mm wykonanie większej od dwu liczby żłobków staje się technicznie niewykonalne. W tych więc przypadkach jesteśmy zmuszeni zrezygnować z najkorzystniejszego trzeciego rozwiązania konstrukcyjnego i zastosować drugie — a więc o wydatnych biegunach na wirniku, starając się możliwie zmniejszyć wpływ oddziaływania krawędzi nabiegunków.

Streszczając powyższe rozważania można przyjąć, co następuje.

Pierwszego rozwiązania konstrukcyjnego selsynów z wydatnymi biegunami na stojanie, jako najniekorzystniejszego, należy unikać.

Drugie rozwiązanie konstrukcyjne selsynów z wydatnymi biegunami na wirniku należy stosować tam, gdzie nie jest możliwe stosowanie trzeciego rozwiązania, a więc w selsynach o małych dla danej częstotliwości wymiarach, w selsynach miniaturowych.

Trzecie rozwiązanie konstrukcyjne selsynów z utajonymi biegunami na wirniku jako najlepsze należy stosować we wszystkich przypadkach, gdy to jest możliwe.

Czwarte rozwiązanie konstrukcyjne selsynów z wyprowadzeniem wszystkich trzech faz uzwojenia wirnika do pierścieni należy stosować przede wszystkim do jednostek różnicowych, a tylko wyjątkowo jako rozwiązanie zastępcze do nadawczo-odbiorczych.

4.1.2. O materiałach, technologii i dokładności wykonania

Podobnie jak we wszystkich maszynach miniaturowych jednym z podstawowych wymagań stawianych materiałom czynnym jest dokładna znajomość i powtarzalność właściwości magnetycznych blach, co w praktyce nastrocza często duże trudności, gdy poszczególne partie materiału otrzymywanego z hut wykazują niepomijalne różnice pod tym względem. Do produkcji selsynów stosuje się przeważnie zwykłą blachę transformatorową o właściwościach omówionych w rozdz. 4. 3. Czasami szczególnie w miniaturowych selsynach bywa stosowany permalój, który zapewnia pewne zmniejszenie zawartości harmonicznych w przebiegu indukcji spowodowanych nieliniowością charakterystyki magnesowania. Ta nieliniowość w pobliżu zerowych wartości indukcji jest w przypadku permaloju znacznie mniejsza niż w przypadku innych materiałów magnetycznych, zaś dzięki wysokim indukcjom nasycenia punkt pracy na charakterystyce magnesowania selsyna udaje się przeważnie, a w przypadku selsynów zasilanych wysokimi częstotliwościami z reguły, utrzymać przed wystąpieniem nieliniowości spowodowanej tymże nasyceniem.

Od miedzi nawojowej nie wymaga się żadnych specjalnych właściwości poza wysoką jakością izolacji przewodów.

Specjalne wymagania stawiane są natomiast łożyskom, szczególnie dla miniaturowych selsynów, o niewielkim z konieczności momencie synchronizującym. Łożyska te muszą być selekcyjonowane pod względem momentu tarcia i luzów.

Właściwości mechaniczne materiałów, z których wykonane są elementy wchodzące w skład zestyków szczotkowych, oraz obróbka mechaniczna tych elementów, powinny również spełniać wysokie wymagania co do jakości.

Przy produkowaniu selsyna należy bezwzględnie przestrzegać najwyższej staranności wykonania, której możliwość powinna być zapewniona odpowiednim opracowaniem technologicznym. Staranność wykonania elementów i montażu jest czynnikiem wpływającym decydująco na pracę selsyna, a jej brak może zepsuć wyniki najlepiej nawet obliczonej i zaprojektowanej jednostki. Poza wysokimi wymaganiami co do właściwości używanych materiałów, wykonania uzwojeń (jednałkowa liczba przewodów w zezwoju, brak zwartych zwojów), części mechanicznej selsyna (równomierne i minimalne tarcie, równomierna szczelina, wyważenie wirnika, brak zwartych obwodów nie związanych z uzwojeniem), należy zwrócić baczną uwagę na wykonanie wykrojników do blach. Wykrojnik powinien być nie tylko dokładny i ostry, lecz musi ponadto jednoznacznie określać kierunek walcowania blachy na wyciętej blasze selsyna w celu umożliwienia wachlarzowego pakietowania tych blaszek. Zaniedbanie tego rodzaju pakietowania, tak samo jak i odpowiedniego skosu złołków, prowadzi do znacznego uwielokrotnienia uchybu selsyna.

4.2. Wpływ parametrów selsyna i zasilania na sztywność selsyna¹⁾

Dla określenia wpływu parametrów selsyna i zasilania na uzyskiwaną sztywność selsyna rozpatrzmy wzór (27) po dokonaniu pewnych przekształceń. Przede wszystkim zastąpimy siłę elektromotoryczną występującą w tym wzorze przez strumień magnetyczny, gdyż wartość strumienia jest w określony sposób powiązana z wymiarami maszyny elektrycznej, natomiast zmieniając uzwojenie można przy tym samym strumieniu uzyskać różne wartości siły elektromotorycznej E_{2f} . Pamiętając, że:

$$E = 4,44 \Phi f z k_q \cdot 10^{-8} \text{ V} \quad (88)$$

¹⁾ Rozważania przytoczone w niniejszym rozdziale są oparte na szczegółowej analizie zagadnienia przeprowadzonej przez autora i częściowo opublikowanej w [36] i [37].

możemy napisać:

$$S = 21,2 \frac{E_f^2}{f} \cdot \frac{X'_q}{R_q'^2 + X_q'^2} = C \Phi^2 f z^2 \frac{X'_q}{R_q'^2 + X_q'^2}. \quad (89)$$

Z otrzymanej postaci wzoru na sztywność wynika, że zależy ona od strumienia wzbudzanego w selsynie, czynnych i biernych oporności jego uzwojeń, liczby zwojów tych uzwojeń i częstotliwości zasilania. Wpływ częstotliwości zasilania występuje zarówno bezpośrednio, jak i w sposób ukryty poprzez pozostałe czynniki wzoru (89). Przy zmianie częstotliwości zasilania dla otrzymania określonej dopuszczalnej dla tej częstotliwości wartości strumienia magnetycznego w danym selsynie należy zmienić liczby zwojów uzwojeń. Wiadomo, że zarówno oporności czynne, jak i bierne uzwojenia zmieniają się proporcjonalnie do kwadratu zwojów. Wartość ułamka $\frac{X'_q}{R_q'^2 + X_q'^2}$ jest więc proporcjonalna do $\frac{1}{z^2}$ oraz, jak wykazano w rozdz. 3.1.1, zależna od stosunku $\frac{R_q'}{X_q'}$ w sposób ilustrowany rys. 16. Uwzględniając to możemy napisać:

$$\frac{X'_q}{R_q'^2 + X_q'^2} = C' \cdot \frac{1}{z^2} \cdot Y. \quad (90)$$

Natomiast wzór (89) przepisać w postaci:

$$S = C_1 \Phi^2 f Y. \quad (91)$$

Tę postać wzoru przyjmiemy za podstawę dyskusji i rozważymy kolejno wchodzące doń czynniki, pamiętając, że dążymy do uzyskania jak największej wartości sztywności.

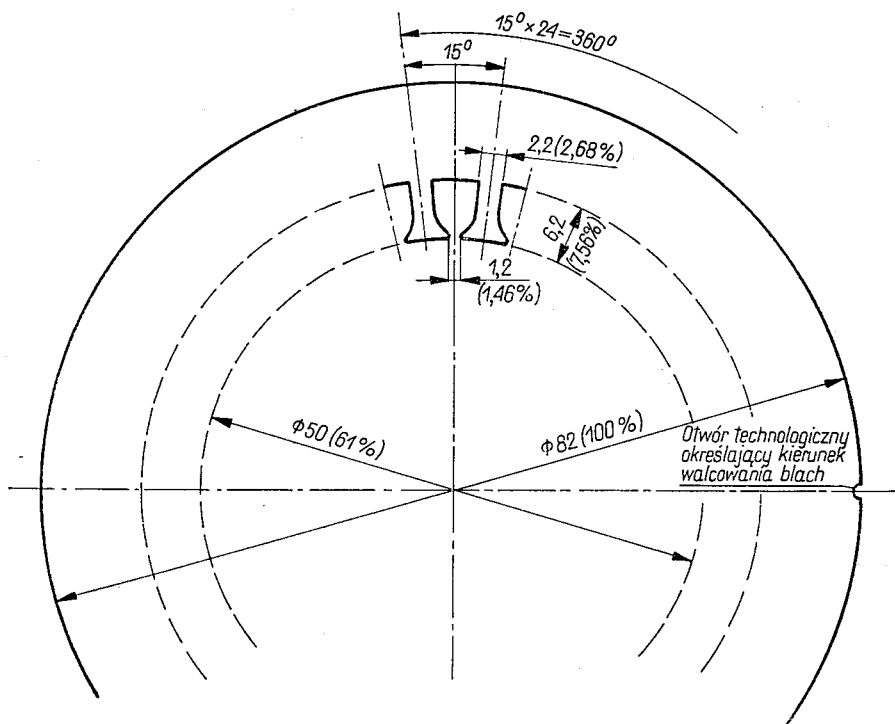
Dopuszczalną wartość strumienia określają wymiary geometryczne pakietu blach selsyna i częstotliwość zasilania. Wymiary geometryczne określone są objętością pakietu blach, stosunkiem długości pakietu l do jego średnicy zewnętrznej D_1 i wykrojem blachy. Od tych czynników zależy charakterystyka magnesowania selsyna. Natomiast od częstotliwości zasilania zależy położenie punktu pracy selsyna na tej charakterystyce, a więc wartość strumienia wytwarzanego podczas pracy selsyna. Za punkt odniesienia dla dalszych rozważań i metody obliczania przyjęto objętość pakietu blach selsyna brutto określoną wzorem:

$$V = \frac{\pi D_1^2 l}{4}. \quad (92)$$

Stosunek $\frac{l}{D_1}$ może być obierany przez konstruktora w granicach $0,5 \div 1$. Można udowodnić [36], [37], że najkorzystniejszymi wartościami tego stosunku są wartości bliskie 1. Jeżeli konstruktor nie jest związany

wymiarami istniejącego już wykrojnika blach, należy tak projektować wykroj blachy przy danej objętości pakietu, aby stosunek $\frac{l}{D_1}$ był równy lub możliwie bliski jedności. Przyjęcie innej wartości tego stosunku pociąga za sobą zmianę dopuszczalnej wartości strumienia dla danej objętości blach selsyna. Ze względów technologicznych, konstrukcyjnych i wytrzymałościowych, szczególnie przy pracy w warunkach wibracji i wstrząsów, wartość $\frac{l}{D_1} > 1$ spotyka się niezmiernie rzadko.

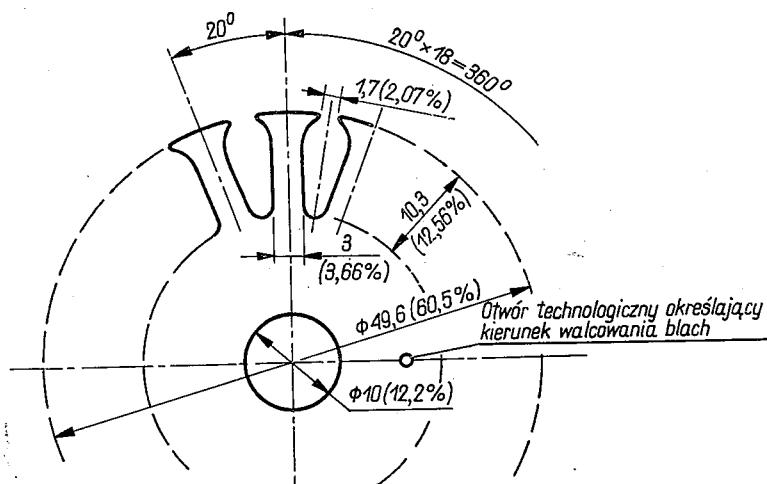
Wykroj blachy wpływa nieznacznie na wartość strumienia, jednakże najkorzystniejszy jest taki wykroj, który zapewnia możliwie równomierne nasycenie obwodu magnetycznego wzdłuż drogi przebiegu strumie-



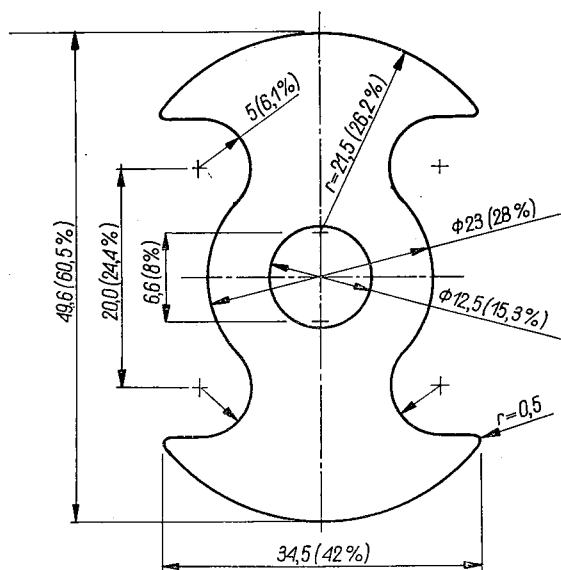
Rys. 60. Optymalny wykroj blachy stojana selsyna

nia. Wykroje blach zapewniające praktyczną równomierność nasycień podane są na rys. 60, 61 i 62. Wymiary podane są w procentach od obrotowego podstawowego wymiaru średnicy zewnętrznej pakietu. Kształt nabiegunków wirnika z biegunami wydawnymi podany na rys. 62 zapewnia wystarczająco optymalny rozkład strumieni w szczelinie powietrznej.

Na wartość dopuszczalnego strumienia wpływa w selsynie, jak w każdej maszynie elektrycznej, szczelina. Normalnie stosowana jest szczelina $\delta = 0,2$ mm, jedynie przy najmniejszych jednostkach o objętościach



Rys. 61. Optymalny wykrój blachy wirnika selsyna z biegunami utajonymi



Rys. 62. Optymalny wykrój blachy wirnika selsyna z biegunami wydatnymi

pakietu V do około 50 cm^3 , których obróbka ze względu na wymiary elementów powinna być szczególnie precyzyjna, stosuje się szczelinę $\delta = 0,1$ mm. Należy zaznaczyć, że stosowanie szczeliny $0,1$ mm dla większych objętości jest niecelowe, gdyż mimo znacznego utrudnienia obróbki nie daje to dla takich objętości dostrzegalnego wzrostu sztywności.