

5. RELATIONS BETWEEN THE OPTIMAL WEIGHTS OF NODES AND THEIR ORIGINATING AND TERMINATING TRAFFIC CAPABILITIES

The next step was to determine relations between the optimal weights found for nodes and other available values characterizing them aiming at establishing a method of their determination. Different parameters were considered here, e.g. demographical data, telecommunication data and some geographic data concerning the situation of nodes in the network.

As a result, in particular for asymmetric models, a strict correlation was found between the originating and terminating weights and outgoing and incoming traffic respectively as well as some measure of relative distance of a node in relation to other nodes in the network. This is illustrated in Figure 3 which shows the optimal weights of nodes as a function

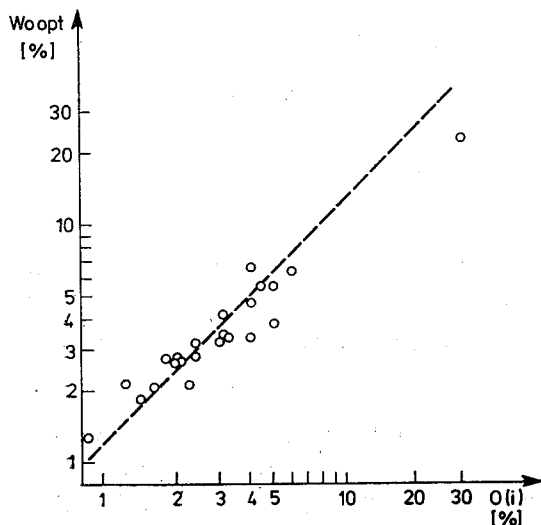


Fig. 3. Optimal originating weights of nodes as a function of the relative traffic originated at these nodes for the asymmetric exponential model

ction of the corresponding traffic for the asymmetric exponential model. A method of determining the optimal weights is evident based on condition of outgoing and incoming traffic equal to row and column totals respectively. This leads inevitably to indicating Kruithof's double factor transformation as the method of iterative calculating the optimal weights. Here row „scaling” defines a subsequent better approximation of originating weights $W_o(i)$ and column „scaling” gives better terminating weights values $W_t(i)$. The starting values for W_o and W_t may be accepted in proportion to the outgoing and incoming traffic respectively and the calculations are stopped when the above mentioned equality of corresponding totals is satisfied with required accuracy.

A similar procedure can be applied to symmetric models; in this case the intention is to obtain equality of row and column totals at each node (outgoing and incoming traffic are assumed to be equal).

The case of the conditionally asymmetric model is different. Here originating and terminating weights are as a rule equal to one another $W_o(i) = W_t(i) = W(i)$, but inco-

ming traffic is usually different from outgoing traffic and is derived from the model itself. In this case, an iterative procedure of defining the optimal weights may be indicated that has some common features with Kruithof's method. The procedure implies knowledge of the outgoing traffic at each exchange $O(i)$. Starting values for weights $W(i)$ may be accepted proportionately to this traffic $O(i)$. Then the traffic matrix can be calculated

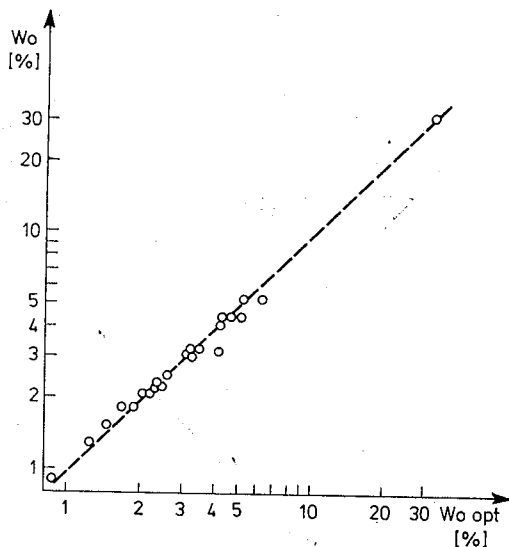


Fig. 4. Originating weights derived from Kruithof's procedure as a function of the optimal weights (for the same model)

according to the basic formula (11) for this distribution with parameters as , r defined previously. Subsequent calculations are repeated iteratively: the row total $T_o(i)$ is defined and the former weights are multiplied by $O(i)/(T_o(i))$. For the new weights the traffic matrix is calculated and the process is continued until the row totals $T_o(i)$ agree with the assumed $O(i)$ as closely as required.

The difference in comparison with Kruithof's transformation consists in the fact that the „scaling” is applied only to weights of nodes and the traffic matrix must be calculated from the basic formula in every iteration.

The results of applying methods described here to the optimal weights determination for different distribution models are presented in Table 2. In this table weights of Warsaw (for illustration) and the deviation from the reference model are given for the types of distribution corresponding to these specified in Table 1. Comparing the corresponding values of weights and deviation for asymmetric models some relatively small differences can be noticed due to the fact that in the first case they are calculated from route traffic and in the other, from the outgoing and incoming traffic values. Figure 4 illustrates this comparison in relation to the weights of nodes.

For symmetric models the difference in the deviation is significant, which proves that the assumption about equality of the outgoing and incoming traffic at each exchange would lead to considerable errors. The result for the conditionally asymmetric model lies between the results obtained for asymmetric and symmetric models. In this case,

Table 2

Results derived from originating and/or terminating traffic

Lp.	Type of model	Weight of Warsaw	Deviation
1	Symmetric uniform	$W = 21,24$	41,18
2	Asymmetric uniform	$W_0 = 24,88$ $W_t = 35,47$	23,73
3	Asymmetric gravity	$W_0 = 25,96$ $W_t = 37,0$	12,86
4	Symmetric exponential	$W = 20,58$	33,07
5	Asymmetric exponential	$W_0 = 24,37$ $W_t = 35,21$	12,11
6	Conditionally asymmetric exponential	$W = 24,0$	18,48

however, the deviation in route traffic, defined according to the formula (5) should not be the only criterion to decide whether a model is good or not.

It is to be noted that the models described are applicable to the forecast of the traffic matrix only under the assumption that the shape and parameters of f -function as well as the mechanism of differentiating $r(i, j)$ and $r(j, i)$ values in the conditionally asymmetric model remain constant during the prediction period.

6. CONCLUSIONS

The following conclusions can be formulated with regard to practical applications of the study:

- the presented method enables the comparison between different theoretical traffic distribution models, subject to determined traffic flows on routes, and the choice of the optimal model under given circumstances, taking into account the deviation from empirical data and the number of parameters describing a model,
- an iterative procedure, parallel to Kruithof's transformation is indicated for determining optimal weights in a given model, assuming the outgoing and/or incoming traffic being known,
- the new model was presented viz. the conditionally asymmetric model which requires forecasting only one value for each node and is comparable with other asymmetric models with regard to accuracy.

Conclusions concerning the specific case of the Polish telex network are as follows:

- distribution of telex traffic is well described by an exponential model (asymmetric or conditionally asymmetric),
- outgoing and incoming traffic is as a rule different,
- the distribution is in a great measure centralized.

There exists one capital centre in Warsaw, collecting in a great measure (about 30%) traffic originated at different exchanges. Furthermore, there are also some local centres, in particular the centres of great industrial areas.

REFERENCES

1. K. Palmowska, *Badania komputerowe modeli ruchu teleksowego*, (Computer investigations of telex traffic distribution models), Paper for the Polish-Czechoslovakish conference „Zastosowanie ETO w łączności” (Use of computer methods in telecommunications) held in Warsaw in 1978 published in the Proceedings of the conference.
2. D. Bear, *Some theories of telephone traffic distribution: A critical survey*, Proceedings of the 7th ITC Congress, Stockholm 1973.
3. D. Bear, *Principles of telecommunication traffic engineering*, Peter Peregrinus Ltd., London 1976
4. W. Findeisen, J. Szymanowski, A. Wierzbicki, *Teoria i metody obliczeniowe optymalizacji*, PWN, Warszawa 1977.
5. Ortega, Rheinhold, *Iterative solutions of nonlinear equations in several variables*, New York Academic Press, 1970.

K. PALMOWSKA

BADANIE MODELI ROZPŁYWU RUCHU W ODNIESIENIU DO POLSKIEJ SIECI TELEKSOWEJ

Streszczenie

Przedstawiono metodę badań porównawczych teoretycznych modeli rozptyłu ruchu opartą na wynikach pomiarów ruchu w wiązkach łącz. W metodzie tej wyznaczane są optymalne wartości parametrów opisujących model, w szczególności optymalne wagi węzłów (parametrów odnoszących się do central), przy których odchylenie modelu teoretycznego względem danych empirycznych jest minimalne. Przedstawiono wyniki empirycznego porównania pewnych teoretycznych modeli rozptyłu ruchu (rozptyw równomierny, grawitacyjny, wykładniczy oraz tzw. warunkowo asymetryczny) w odniesieniu do polskiej sieci teleksowej. Wskazano iteracyjną metodę wyznaczania wag węzłów opartą na znajomości ruchu wychodzącego i/lub przychodzącego do tych węzłów — dla poszczególnych modeli.

K. PALMOWSKA

ETUDE DES MODÈLES DE RÉPARTITION DU TRAFIC DANS LE RÉSEAU TÉLEX POLONAIS

Résumé

On a présenté la méthode de l'étude comparative des modèles théoriques de répartition du trafic télex en se basant sur les mesures du trafic routier. Dans cette méthode on calcule les valeurs optimales des paramètres d'un modèle, en particulier les poids optimaux des noeuds (paramètres se rapportant aux noeuds), pour lesquels l'écart entre les données théoriques et empiriques est minime.

On a présenté les résultats de comparaison empirique de certains modèles théoriques de la répartition du trafic (répartition uniforme, de gravitation, exponentielle et conditionnellement asymétrique) par rapport à données de trafic télex polonais. On a indiqué la méthode itérative de détermination des poids des noeuds à partir des valeurs du trafic de départ et d'arrivée.

K. PALMOWSKA

UNTERSUCHUNG DER VERKEHRSVERTEILUNGSMODELLE
HINSICHTLICH DES POLNISCHEN TELEXNETZES

Zusammenfassung

Die Methode der Vergleichsuntersuchung von theoretischen Verkehrsverteilungsmodellen in Anlehnung an Verkehrsmessungen in Leitungsbündeln wird dargestellt. Bei dieser Methode werden die optimalen Parameterwerte des Modells bestimmt, insbesondere die optimalen Knotengewichte (knotenbetreffende Parameter), bei denen die Abweichung der theoretischen von empirischen Daten minimal ist.

Es werden Ergebnisse eines empirischen Vergleichs mancher theoretischen Verkehrsverteilungsmodelle (Uniformverteilung, Gravitationsmodell, Exponentialmodell und ein sog. bedingt asymmetrisches Modell) in Bezug auf polnische Telexverkehrsdaten gebracht. Es wurde auf die iterative Methode der Bestimmung von Knotengewichten aufgrund abgehender oder ankommender Verkehrswerte hingewiesen.

К. ПАЛЬМОВСКА

ИССЛЕДОВАНИЕ МОДЕЛЕЙ РАСПРЕДЕЛЕНИЯ ПОТОКОВ В ПОЛЬСКОЙ СЕТИ АБОНЕНТСКОЙ ТЕЛЕГРАФНОЙ СВЯЗИ

Резюме

Представлен метод сравнительного анализа теоретических моделей распределения потоков сообщений на основании данных измерений нагрузки соединительных пучков каналов. В соответствии с этим методом вычисляются оптимальные величины параметров математических моделей в частности оптимальные веса узлов (параметров касающихся узлов) которые сводят к минимуму отклонение теоретических и эмпирических данных. Представляются результаты эмпирического сравнения некоторых теоретических моделей распределения потоков (равномерное распределение, гравитационная модель, экспоненциальная модель и так называемая условно асимметрическая модель) относительно к польской сети абонентского телеграфа). Показан итеративный метод определения весов узлов исходя из известных значений исходящего и входящего сообщения для этих узлов.

XX LAT POLSKIEGO TOWARZYSTWA ELEKTOTECHNIKI TEORETYCZNEJ I STOSOWANEJ

W dniu 26 stycznia 1981 roku minęło 20 lat od utworzenia PTETS. Rocznicą tą zbiegła się z doniosłymi wydarzeniami w naszym kraju. Odnowa na wszystkich odcinkach życia społecznego i gospodarczego kraju zapoczątkowana przez Robotników Wybrzeża i Śląska stawia przed Oddziałami i Członkami Towarzystwa nowe zadanie — jak najlepszego wykorzystania otwierających się możliwości dla dobra nauki i techniki polskiej oraz dla odrodzenia najlepszych tradycji szkolnictwa wyższego. PTETS przyłącza się w pełni do analogicznych opinii PAN (z 23.10.1980 r.) i Rady Towarzystw Naukowych przy Prezydium PAN (z 24.11.80 r.). Po obszernej dyskusji tych problemów na zebraniu plenarnym Zarządu Głównego PTETS w dniu 6.12.1980 r. podjęto jednomyślną uchwałę, iż „Plenum Zarządu Głównego PTETS w pełni popiera pozytywne przemiany społeczne zachodzące obecnie w Polsce i deklaruje włączenie się do działalności na rzecz tych przemian w zakresie społecznego oddziaływania na naukę”. Zebranie to przyjęło także wstępne sprawozdania Oddziałów i Zarządu Głównego za rok 1980 i zatwierdziło projekt planu na rok 1981.

Rocznicą dwudziestolecia zastaje nasze Towarzystwo w pełnym i ustabilizowanym rozwoju. Od roku 1961, kiedy to staraniem 41 Członków Założycieli zostało utworzone nasze Towarzystwo [1], liczba jego Oddziałów w końcu 1980 r. wzrosła do 10, a liczba jego członków do 572 osób, w tym w poszczególnych Oddziałach: Gdańsk — 58 osób, Gliwice — 49 osób, Kielce — 14 osób (Oddział powołany 8.01.80 r.), Kraków — 54, Łódź — 57, Poznań — 59, Szczecin — 24, Warszawa — 145, Wrocław — 98, Zielona Góra — 15 (Oddział powołany dn. 6.12.1980 r.).

W latach 1961—1980 w Oddziałach ogłoszono łącznie 1.278 referatów naukowych, zorganizowano 75 konkursów naukowych i 112 imprez zbiorowych w postaci konferencji, sympozjów i seminariów. Szczegółowa historia Towarzystwa ujęta jest w kolejnych komunikatach [1] ..., [5], poczynając od jubileuszu XV-lecia [1].

W dniach 25—26 kwietnia 1980 r. odbył się w Krakowie kolejny XVI Zjazd Delegatów PTETS połączony z Sesją Naukową.

Na wstępie zebrani uczcili chwilą ciszy pamięć zmarłych od ostatniego Zjazdu Członków Towarzystwa: Członka Honorowego PTETS prof. Romana Kurdziela (Wrocław), Członka Honorowego PTETS prof. dra Stefana Manczarskiego (Warszawa), prof. dra Felsa Andrzejewskiego (Wrocław), dra Krzysztofa Gajdy (Warszawa) i dra Henryka Karpowskiego (Łódź).

XVI Zjazd Delegatów przyjął sprawozdanie ustępującego Zarządu Głównego za lata 1978 i 1979 oraz dokonał wyboru przewodniczącego, członków Zarządu Głównego i członków Głównej Komisji Rewizyjnej na lata 1980—1982.

Na wniosek ustępującego Zarządu Głównego Zjazd nadał przez aklamację godność Członka Honorowego PTETS prof. dr Tadeuszowi Zagajewskiemu (Gliwice), prof. Bronisławowi Sochorowi (Łódź) i prof. dr Józefowi Żydanowiczowi (Warszawa).

W części naukowej Zjazdu referaty wygłosili: prof. dr M. Suski (Wrocław), doc. dr hab. S. Przeździecki (Warszawa), prof. dr hab. L. Szklarski (Kraków), prof. dr hab. K. Bisztyga (Kraków).

Następny Zjazd Delegatów odbędzie się na początku 1982 r.

Władze PTETS wybrane na okres do roku 1982

Zarząd Główny:

- | | |
|--|---|
| 1. Prof. dr hab. Janusz Turowski | — Przewodniczący (Łódź) |
| 2. Prof. dr Bolesław Dubicki | — Wiceprzewodniczący (Warszawa) |
| 3. Doc. dr hab. Stanisław Przeździecki | — Wiceprzewodniczący (Warszawa) |
| 4. Prof. dr hab. Eugeniusz Koziej | — Sekretarz Generalny (W-wa) |
| 5. Prof. dr Józef Żydanowicz | — Skarbnik (Warszawa) |
| 6. Dr inż. Andrzej Marusak | — Z-ca Sekretarza Generalnego i Przewodniczący O/Warszawa |
| 7. Doc. dr Jacek Przygodzki | — Z-ca Skarbnika i Przewodniczący O/Kielce |
| 8. Prof. dr hab. Kazimierz Bisztyga | — Członek Zarządu Głównego (Kraków) |
| 9. Doc. dr hab. Lech Fuliński | — Członek Zarządu Głównego (Wrocław) |
| 10. Doc. dr hab. Zbigniew Szczerba | — Członek Zarządu Głównego (Gdańsk) |

oraz Przewodniczący Oddziałów:

- | | |
|------------------------------------|------------------|
| 11. Prof. dr hab. Jerzy Sawicki | — O/Gdańsk |
| 12. Doc. dr hab. Stanisław Szpilka | — O/Gliwice |
| 13. Doc. dr hab. Józef Machowski | — O/Kraków |
| 14. Prof. dr hab. Zdzisław Korzec | — O/Łódź |
| 15. Doc. dr hab. Lech Różański | — O/Poznań |
| 16. Prof. dr Ryszard Sikora | — O/Szczecin |
| 17. Prof. dr hab. Marian Piekarski | — O/Wrocław |
| 18. Doc. dr Zdzisław Drozdowski | — O/Zielona Góra |

Zastępcy Członków Zarządu Głównego:

- | | |
|----------------------------------|------------|
| 19. Doc. dr hab. Romuald Nowicki | — Wrocław |
| 20. Doc. dr hab. Jerzy Mędrzycki | — Warszawa |

Główna Komisja Rewizyjna:

- | | |
|-------------------------------------|-----------------------------|
| 1. Prof. dr hab. Daniel Bem | — Przewodniczący (Wrocław) |
| 2. Prof. dr Jan Kożuchowski | — Warszawa |
| 3. Doc. dr hab. Kazimierz Mikołajuk | — Warszawa |
| 4. Prof. dr Marian Suski | — Wrocław |
| 5. Doc. dr Henryk Wierzb | — Z-ca Członka GKR (Gdańsk) |
| 6. Doc. dr Stefan Wojciechowski | — Z-ca Członka GKR (Łódź) |

Sprawozdanie za lata 1979 i 1980

Tablica 1

Liczbowe zestawienie działalności naukowej

	Rok	Zebrania organizacyjne	Zebrania naukowe				Sympozja i seminaria		
			zebrań	referatów		uczestników	liczba	referatów	uczestników
				z Polski	z zagr.				
Oddziały	1979	46	64	52	15	1.309	6	186	408
łącznie	1980	59	57	59	—	1.193	7	171	328
Zarząd	1979	7	—	—	—	—	—	—	—
Główny	1980	5	1	—	—	36	1	4	36

Liczba członków PTETS wyniosła w r. 1979—567 i w r. 1980—572.

Liczba nadanych tytułów Członka Honorowego PTETS wzrosła do 19.

W latach 1979—80 zorganizowano następujące sympozja i seminaria:

1. „Problemy wyładowań niezupełnych w układach elektroizolacyjnych”, III Sympozjum zorganizowane przez O/Kraków w Zakopanem 13—15.XII.79; 25 referatów, 70 uczestników.
2. „Zjawiska elektromagnetyczne w obwodach nieliniowych — Obwody z zaworami półprzewodnikowymi”, VI Sympozjum zorganizowane przez O/Poznań w Kołobrzegu 25—27.X.1979; 22 referaty, 5 komunikatów, 42 uczestników.
3. „Optymalizacja w zagadnieniach elektrotechniki”, II Sympozjum zorganizowane przez O/Warszawa w Zakopanem 5—8.VI.1979; 20 referatów, 40 uczestników.
4. „Metody matematyczne w elektrotechnice”, VIII Sympozjum zorganizowane przez O/Warszawa i WSI Opole w Pokrzywniej 21—26.V.1979; 81 referatów, 160 uczestników.
5. „Fizyka i elektronika w medycynie”, VII Seminarium zorganizowane przez O/Gdańsk 14—15.XII.1979; 34 referaty, 96 uczestników.
6. III Międzynarodowe Seminarium Podstaw Elektrotechniki i Teorii Obwodów zorganizowane przez O/Gliwice w Ustroniu 25—27.IV.1979; 4 referaty.
7. „Metody matematyczne w elektrotechnice”, MME-9 zorganizowane przez O/Warszawa wspólnie z WSI Opole w Pokrzywniej 26—31.V.1980; 97 referatów, 170 uczestników.
8. „Symulacja procesów dynamicznych”, SPD-1 Sympozjum zorganizowane przez O/Warszawa 9—14.06.1980 w Zakopanem; 34 referaty, 58 uczestników.
9. „Maszyny liniowe indukcyjne i ich zastosowanie”, Seminarium otwarte zorganizowane przez O/Wrocław 25 listopada, 1980 r.; 8 referatów.
10. „Systemy telekomunikacyjne i teoria obwodów”, Seminarium zorganizowane przez O/Poznań oraz Instytutu Elektroniki i Elektrotechniki Przemysłowej Politechniki Poznańskiej 15—16.05.1980; 7 referatów, 20 uczestników.
11. „Zastosowanie maszyn matematycznych w elektrotechnice”, VIII Sypozjum zorganizowane przez O/Łódź wspólnie z Ośrodkiem Elektrotechnicznej Techniki Obliczeniowej Politechniki Łódzkiej, luty 1980.

12. „IV Międzynarodowe Seminarium Podstaw Elektrotechniki i Teorii Obwodów” zorganizowane przez Instytut Podstawowych Problemów Elektrotechniki i Energoelektroniki wraz z O/Gliwice 16—18.04.1980 w Ustroniu; 3 referaty wygłoszone przez członków O/Gliwice.
13. „Fizyka i elektronika w medycynie”, VIII Seminarium zorganizowane przez O/Gdańsk 12—13.12.1980; 22 referaty, 80 uczestników.
14. Sesja Naukowa z okazji XVI Zjazdu Delegatów PTETS, Kraków 26.04.80; 4 referaty.

W konkursach naukowych w r. 1979 nagrody otrzymali:

1. W konkursie Zarządu Głównego pod hasłem „Zastosowanie modelowania i symulacji w urządzeniach elektrycznych”: I nagroda — dr inż. K. Stańkowski (Warszawa), II nagroda — dr inż. K. Januszkiewicz (Łódź), III nagroda — dr inż. W. Brociek (W-wa).
2. W konkursie Oddziału Łódzkiego pod hasłem „Mikroprocesory i ich zastosowanie”: I nagroda — mgr B. Ostojski, II nagroda — mgr inż. W. Szaniawski i dr inż. Z. Leszczyński, III nagroda — mgr inż. W. Górnisiewicz, mgr inż. W. Zatorski, mgr inż. T. Górnicki i mgr inż. M. Wyrzykowski.

W konkursach naukowych w r. 1980 nagrody otrzymali:

1. W konkursie naukowym Zarządu Głównego pod hasłem „Numeryczna analiza pól fizycznych w maszynach i urządzeniach elektrycznych i elektronicznych” nagrody otrzymali: II nagroda — dr K. Sachse (Wrocław), III nagroda — dr A. Karwowski (Wrocław), III nagroda dr Lucjan Węgrowicz i dr J. Radzki (Warszawa).
2. W konkursie Oddziału Gliwickiego pod hasłem „Energoelektronika w układach napędowych”: I nagroda — dr inż. A. Dębowski (Gliwice).

W ramach akcji wydawniczej w r. 1979 wydano 6 pozycji materiałów konferencyjnych.

Plan pracy na rok 1981 przewiduje wzrost liczby członków Towarzystwa do 606, liczbę zebrań naukowych 64 oraz następujące sympozja i seminaria:

1. VIII Seminarium „Fizyka i elektronika w medycynie” organizowane przez Oddział w Gdańsku wspólnie z PTFM w Gdańsku.
2. Sympozjum „Podstawowe problemy energoelektroniki” organizowane przez Oddział w Gliwicach.
3. VII Sympozjum „Zjawiska elektromagnetyczne w obwodach nieliniowych — Obwody z ferromagnetykami” organizowane przez Oddział w Poznaniu.
4. Seminarium „Filtry cyfrowe” organizowane przez Oddział we Wrocławiu.
5. Sympozjum „Metody matematyczne w elektrotechnice” MME-10 organizowane przez Oddział w Warszawie, Karpacz, 25—30 maja 1981.

W r. 1981 zostanie ogłoszony konkurs na najlepszą pracę naukową z dziedziny „Prace eksperymentalne w elektrotechnice”.

BIBLIOGRAFIA

1. *XV lat Polskiego Towarzystwa Elektrotechniki Teoretycznej i Stosowanej*. Komunikat 1, Rozprawy El., 1976, 22, Nr 3, ss. 749—750.
2. *Działalność PTETS w roku 1976 i plan na rok 1977*, Komunikat 2, Rozprawy El., 1977, 23, Nr 2, ss. 479—481.
3. *XV Zjazd Delegatów PTETS*, Komunikat 3, Rozprawy El., 1978, 24, Nr 3, ss. 763—766.
4. *Działalność PTETS w roku 1978 i plan na rok 1979*, Komunikat 4, Rozprawy El., 1979, 25, Nr 4, ss. 1105—1107.
5. J. Turowski: *XX lat PTETiS*, Prace X Sympozjum MME, Karpacz, 25—30.05.1981, s. 9—16.

J.T.

TREŚĆ

K. Adamiak: Synteza pola magnetycznego w obszarach dwuwymiarowych za pomocą uzwojeń bezrdzeniowych	311
A. Szatkowski: Remarks on the Existence of Periodic Solutions of Nonlinear Autonomous Systems and Networks	327
K. Schoepp, K. Makowski: Silnik indukcyjny z ekranowanymi biegunami o zwiększonym momencie rozruchowym	339
M. Uruski: Synteza bezstratnych układów sterowanych generatorem o zespolonej impedancji wewnętrznej i obciążonych wieloma impedancjami zespolonymi	355
I. Tataara: Wykorzystanie czułości drugiego rzędu do analizy filtrów aktywnych RC o zerowej czułości pierwszego rzędu	367
C. Michalik: Filtr dolnoprzepustowy o maksymalnie płaskiej charakterystyce tłumieniowej i maksymalnie tłumionej energii w pasmie zaporowym	379
J. Sykulski: Three-Dimensional Analysis of Electromagnetic Field Distribution in Solid Metallic Structures in Vicinity of Current Carrying Finite-Length Bar Conductors	387
W. Smołuch: Symulacja cyfrowa układu przekształtnikowego współpracującego z indukcyjną maszyną prądu przemiennego	403
E. Z. Pawlak: Szumy typu 1/f w cienkich polikrystalicznych warstwach półprzewodnikowych	419
D. Tollik, B. Wilamowski, W. Kordalski: Nieulotne pamięci półprzewodnikowe	431
J. Warzywoda: Wybór matematycznego modelu rozkładu prawdopodobieństwa napięć przekroju	443
M. Domański, R. Stasiński: Porównanie możliwości obliczeniowych mikroprocesorów i procesorów minikomputerów	453
H. Szymański: Optymalizacja parametrów mikrowiązki elektronowej	467
Z. Perkowski: Optymalizacja konstrukcji przewodów współosiowych wielkiej częstotliwości	471
A. Dziech, I. M. Smolewski: Ocena korelacyjnych własności dyskretnych sygnałów losowych	531
A. Ruciński, A. Klimontowicz: Estimation of Service Effectiveness of Overflow Traffic Component Flows in Overflow Systems with Full-Availability Trunk Groups	539
A. Kusy, A. Szpytma: Współczesne metody i kierunki badania właściwości szumów typu 1/f	553
A. Karwowski, M. Moszczyński: Właściwości impedancyjne anten klatkowych	577
K. Palmowska: A Study of Traffic Distribution Models with Reference to the Polish Telex Network	593
Komunikat: XX lat Polskiego Towarzystwa Elektrotechniki Teoretycznej i Stosowanej	607

CONTENTS — TABLES DES MATIÈRES — INHALT

K. Adamiak: Magnetic Field Synthesis in Two-Dimensional Areas by Means of Coreless Windings	324
Synthèse du champ magnétique dans les espaces bidimensionnels à l'aide des enroulements sans noyau	324
Synthese eines Magnetfeldes in zweidimensionalen Gebieten mit Hilfe kernloser Wicklungen	325
A. Szatkowski: Uwagi o istnieniu rozwiązań okresowych dla nieliniowych systemów i układów elektronicznych	336
Remarques sur les solutions périodiques des systèmes électroniques non linéaires	336
Bemerkungen über die Existenz der periodischen Lösungen für nichtlineare Systeme	337

K. Schoepp, K. Makowski: Induction Motor with Screened Poles and Increased Starting Torque	352
Moteur à induction à poles blindés au couple de démarrage augmenté	352
Induktionsmotor mit abgeschirmten Polen und vergrößertem Ablaufmoment	353
M. Uruski: Synthesis of a Lossless Network Excited by a Frequency-Dependent Generator and Terminated in Arbitrary Lumped Passive Loads	365
Synthèse des systèmes sans dissipation commandés par un générateur à impédance interne complexe et chargés de multiples impédances complexes	365
Synthese verlustloser, mit einem Generator mit Komplex-Innenimpedanz gesteuerter und mit mehreren Kompleximpedanzen belasteter Netze	365
I. Tatar: Application of Second Order Sensitivity to Analysis of Parameter Changes in Active Filters with Zero Sensitivity of First Order	377
Application de la sensibilité de second ordre à l'analyse des changements des paramètres des filtres actifs à sensibilité zéro de premier ordre	377
Ausnutzung der Empfindlichkeit 2. Ordnung zur Analyse von Parameterveränderungen aktiver Filter mit Nullempfindlichkeit 1. Ordnung	377
C. Michalik: Maximally Flat Low-Pass Filters with Maximum Integrated Power-Loss Ratio in the Stopband	385
Filtre passe-bas à caractéristique maximale plane et à l'énergie atténuée au maximum dans la zone d'atténuation	385
Tiefpassfilter mit maximaler Ebenung des Betrages im Durchlassbereich und mit maximal gedämpfter Energie im Sperrbereich	385
J. Sykalski: Analiza trójwymiarowa pola elektromagnetycznego przewodów szynowych o skończonej długości w sąsiedztwie masywnych półprzestrzeni przewodzących	400
Analyse tridimensionnelle du champ électromagnétique dans le système des conducteurs parallèles finis en voisinage des écrans conductifs massifs semi-unfinis	401
Dreidimensionale Analyse eines elektromagnetischen Feldes endlichlanger Schienenleiter in der Nähe massiver, leitender Halbräume	401
W. Smoluch: Digital Simulation of Thyristor Converter Collaborating with AC Machine	417
Simulation digital d'un convertisseur statique qui collabore avec le moteur asynchrone	418
Digitalsimulation eines mit einer induktiven Wechselstrommaschine zusammenarbeitenden Stromrichtersystems	418
E. Z. Pawlak: $1/f$ Noise in Thin Polycrystalline Semiconductor Layers	429
Bruits du type $1/f$ dans des couches polycristallines minces des semi-conducteurs	429
$1/f$ -Rauschen bei dünnen polykristallinen Halbleiter-Schichten	429
D. Tollik, B. Wilamowski, W. Kordalski: Nonvolatile Semiconductor Memories	441
Mémoires semi-conducteurs non-volatiles	441
Nonvolatile Halbleiterspeicher	441
J. Warzywoda: Selection of Mathematical Model of Distribution of Flash-Over Voltage Probability	451
Choix du modèle mathématique de répartition de la probabilité des tensions de contournement	451
Wählmethode für ein mathematisches Modell der Wahrscheinlichkeitsverteilung von Überschlagespannungen	451
M. Domański, R. Stasiński: Comparison of Computational Abilities of Microprocessors and Minicomputer Processors	465
Comparaison d'efficacité de quelques microprocesseurs et processeurs minicalculateurs	465
Vergleich der Effektivität der Mikroprozessoren und der Minikomputerprozessoren	465
H. Szymański: Optimization of Electron Microbeam Parameters	470
Optimisation des paramètres déterminant le faisceau d'électrons	470
Optimierung der Elektronenstrahlparameter	470

Z. Perkowski: Optimization of Construction of Flexible High Frequency Coaxial Cables	528
Optimisation de la construction des cables coaxiaux flexibles pour hautes fréquences	528
Konstruktionsoptimierung der koaxialen Hochfrequenzleitungen	528
A. Dziech, I. M. Smolewski: Evaluation of Correlation Properties in Rectangular Pulse Streams	537
Estimation des propriétés de corrélation des signaux aléatoires discrets	537
Schätzungsweise korelativer Eigenschaften diskreter stochastischer Signale	537
A. Ruciński, A. Klimontowicz: Ocena jakości obsługi strumieni składowych ruchu przelewowego w układach przelewowych z wiązkami doskonałymi	550
Estimation de l'efficacité de service des courants composants du trafic avec le débordement dans les systèmes avec le débordement à faisceaux à accessibilité totale	550
Abschätzung der Bedienungsgüte von Komponentenflüssen des Überlaufverkehrs bei Überlaufsystemen mit vollendeten Bündeln	551
A. Kusy, A. Szpytma: Modern Methods and Directions of Investigations of $1/f$ Type Noise Properties	574
Méthodes et directions actuelles de l'examen des propriétés des bruits du type $1/f$	574
Moderne Methoden und Untersuchungsrichtungen von Strukturen des $1/f$ -Rauschens	575
A. Karwowski, M. Moszczyński: Impedance Properties of Cage Antennas	590
Propriétés d'impédance des antennes à cage	590
Impedanzeigenschaften der Käfigantennen	590
K. Palmowska: Badanie modeli rozplywu ruchu w odniesieniu do polskiej sieci teleksowej	604
Etude des modèles de répartition du trafic dans le réseau télex polonais	604
Untersuchung der Verkehrsverteilungsmodelle hinsichtlich des polnischen Telexnetzes	605

СОДЕРЖАНИЕ

К. Адамьяк: Синтез магнитного поля в двумерном пространстве при помощи обмоток без сердечника	325
А. Шатковски: Замечания о существовании периодических решений нелинейных динамических систем	337
К. Шэпп, К. Маковски: Асинхронный двигатель с экранированными полюсами и увеличенным пусковым моментом	353
М. Уруски: Синтез цепей без потерь, управляемых генератором с комплексным внутренним сопротивлением и нагруженных многими комплексными двухполюсниками	366
И. Татара: Использование чувствительности второго порядка для анализа изменений параметров активных фильтров имеющих нулевую чувствительность первого порядка	378
Ч. Михалик: Фильтр нижних частот с максимально плоской характеристикой затухания и с максимальным затуханием энергии в полосе задерживания	386
Я. Сыкульски: Трехмерный анализ электромагнитного поля системы параллельных рельсовых проводов конечной длины вблизи проводящих полупространств	401
В. Смолух: Цифровое моделирование преобразовательной системы работающей совместно с машиной переменного тока	418
Э. З. Павляк: Шумы типа $1/f$ в тонких поликристаллических полупроводниковых пленках	430
Д. Толлик, Б. Виямовски, В. Кордадьски: Некратковременные полупроводниковые памяти	442
Я. Вазивода: Избрание математической модели распределения вероятности пробоя изоляции	452
М. Доманьски, Р. Стасиньски: Сравнение вычислительных возможностей микропроцессоров и процессоров линии — ЭВМ	466
Х. Шиманьски: Оптимализация параметров электронного микропучка	470
З. Перковски: Оптимализация конструкции гибких коаксиальных кабелей	529

А. Дзех, И. М. Смолевски: Оценка корреляционных свойств дискретных случайных сигналов	437
А. Рущински, А. Климонтович: Оценка качества обслуживания составных потоков систем с избыточной нагрузкой с полnodоступными пучками	551
А. Кусы, А. Шпытма: Современные методы и направления исследований структуры шумов типа $1/f$	575
А. Карвовски, М. Мощиньски: Импедационные свойства клеточных антенн	591
К. Пальмовска: Исследование моделей распределения потоков в польской сети абонентской телеграфной связи	605

Wytyczne dla Autorów

Komitet Redakcyjny prosi o przestrzeganie następujących wytycznych przy opracowywaniu maszynopisów artykułów nadsyłanych do opublikowania.

1. *Tematyka charakter i oryginalność artykułów.* „Rozprawy Elektrotechniczne” są przeznaczone do ogłaszania prac naukowych wchodzących w zakres szeroko pojętej elektrotechniki, zwłaszcza jeżeli traktują one omawiane zagadnienia w sposób możliwie kompletny i zmierzają do osiągnięcia celu informacyjno-dydaktycznego. W związku z tym publikowane prace mogą być w mniejszym lub większym stopniu pracami nieoryginalnymi, jeśli chodzi o treść w nich zawartą, co nie zwalnia Autorów od obowiązku ujęcia tematu w sposób oryginalny. Twórczy wkład Autora może polegać wówczas np. na: porównaniu opisywanych teorii, metod lub systemów, krytycznej ich ocenie, wyciągnięciu własnych wniosków co do celowości takiego lub innego działania, zastosowaniu własnej klasyfikacji zagadnień, prognoście itp. W szczególności „Rozprawy” ogłaszają takie prace naukowe, które zawierają: analizę porównawczą teorii, metod lub systemów, syntetyczne ujęcia określonego zagadnienia naukowego, omówienie obecnego stanu wybranego działu elektrotechniki, omówienie jego postępów lub tendencji rozwojowych itp.

2. *Wymagania ogólne.* Artykuły należy nadsyłać w maszynopisie, w dwóch egzemplarzach. Maszynopis powinien być napisany jednostronnie przez czarną taśmę, na maszynie do pisania z niezabrudzonymi i nieuszkodzonymi znakami. Dopuszcza się odręczne czytelne uzupełnienie lub pisanie atramentem albo długopisem kolorem czarnym lub ciemnoniebieskim znaków specjalnych oraz znaków w językach, których alfabetów nie ma na maszynach do pisania, np. znaków matematycznych, chemicznych, liter greckich. Maszynopis powinien być napisany na papierze do maszyn do pisania koloru białego, formatu A-4; numeracja ciągła na wszystkich kartkach. W artykułach należy stosować Międzynarodowy Układ Jednostek SI.

3. *Sposób pisania tekstu.* Tekst w maszynopisie powinien być pisany bez używania wyróżnień, a w szczególności nie dopuszcza się spacjiowania (rozstrzelania), podkreślania i pisania tekstów dużymi literami. Proponowane wyróżnienia Autor może wprowadzić do maszynopisu (zwykłym ołówkiem) za pomocą przyjętych znaków adiustacyjnych (podkreślenie linią przerywaną — spacjiowanie, podkreślenie linią ciągłą — pogrubienie, podkreślenie wężykiem — kursywa itp.). Na jednej kartce maszynopisu powinno być 30 wierszy po około 60 znaków łącznie z odstępami. Marginesy każdej kartki powinny mieć następujące wymiary: górny — ok. 25 mm, lewy — ok. 35 mm. Tekst maszynopisu powinien być napisany z podwójnym odstępem między wierszami; tytuły i podtytuły małymi literami. Akapity należy rozpoczynać z wcięciem równym trzem uderzeniom maszyny do pisania.

4. *Sposób pisania tablic.* Tablice powinny być napisane na oddzielnych arkuszach, w układzie zbliżonym do układu zecerskiego. Tytuły rubryk pionowych i poziomych powinny być napisane małymi literami z podwójnym odstępem między wierszami. Przypisy (notki) dotyczące tablic należy pisać bezpośrednio pod tablicą. Tablice należy numerować kolejno liczbami arabskimi; u góry każdej tablicy podać tytuł.

5. *Sposób pisania wzorów matematycznych.* Rozmieszczenie znaków, cyfr, liter i odstępów powinno być zbliżone do rozmieszczenia elementów druku. Wskaźniki i wykładniki powinny być napisane wyraźnie i być prawidłowo obniżone lub podwyższone w stosunku do linii wiersza podstawowego. Znaki nad literami i cyframi: strzałki, linie, kropki, daszki itp. powinny być napisane dokładnie nad tymi elementami, do których się odnoszą. Numery wzorów należy umieszczać z prawej strony.

6. *Przygotowanie materiału ilustracyjnego.* Rysunki, wykresy i fotografie należy wykonywać zgodnie z obowiązującymi normami, na oddzielnych arkuszach, z podaniem kolejnych numerów rysunków. W maszynopisie artykułu na marginesie, obok właściwego tekstu, należy podać jedynie odnośny numer rysunku, a na oddzielnym arkuszu wykaz podpisów pod rysunki. Wszystkie rysunki, wykresy i fotografie należy nazywać w tekście rysunkami (skrót: rys.). U samego dołu rysunku (a przy fotografiach na odwrocie) należy

wpisać czytelnie numer rysunku, tytuł pracy i nazwisko autora. Ostateczne wykonanie rysunków obowiązuje Redakcję.

7. *Streszczenia.* Do każdej nadsyłanej pracy należy dołączyć krótkie streszczenie (analizę) w języku polskim (w 5 egz.) oraz streszczenia w językach obcych: angielskim, francuskim, niemieckim i rosyjskim. W razie niemożności przygotowania streszczeń w językach obcych Autor powinien przynajmniej podać niezbędną do wykonania tłumaczenia obcojęzyczną terminologię.

8. *Literatura.* Po zakończeniu artykułu należy podać wykaz literatury wymieniając w następującej kolejności: pierwsze litery imion, nazwisko autora, po czym po przecinku pełny tytuł dzieła lub artykułu; dalej, w przypadku książki — wydawcę, miejsce wydania i rok, a w przypadku artykułu — tytuł czasopisma, numer zeszytu, rok wydania i ewent. numer strony. Pozycje wykazu powinny być ponumerowane.

Uwagi

1. Niezastosowanie się Autora do podanych wyżej wytycznych pociągnie za sobą konieczność potrącenia z honorarium autorskiego kosztów związanych z doprowadzeniem dostarczonych materiałów do postaci wymaganej przez Redakcję.

2. Autorowi przysługuje bezpłatnie 25 egz. odbitek pracy. Dodatkowe egzemplarze Autor może zamówić w Redakcji na własny koszt.

3. Autora obowiązuje korekta autorska, którą należy zwracać w ciągu 3 dni pod adresem Redakcji.

Rozprawy Elektrotechniczne

Kwartalnik

CZASOPISMO
ROZPRAWY ELEKTROTECHNICZNE

Warunki prenumeraty
Cena prenumeraty krajowej
półrocznie zł 80,—
rocznie zł 160,—

Prenumeratę na kraj przyjmują Oddziały RSW „Prasa—Książka—Ruch” oraz urzędy pocztowe i doręczyciele w terminach:

— do 25 listopada na I półrocze roku bieżącego i na cały rok następny,
— do 10 czerwca na II półrocze roku bieżącego.

Jednostki gospodarki uspołecznionej, instytucje, organizacje i wszelkiego rodzaju zakłady pracy zamawiają prenumeratę w miejscowych Oddziałach RSW „Prasa—Książka—Ruch”, w miejscowościach zaś, w których nie ma Oddziałów RSW — w urzędach pocztowych.

Czytelnicy indywidualni opłacają prenumeratę wyłącznie w urzędach pocztowych i u doręczycieli.

Prenumeratę ze zleceniem wysyłki za granicę przyjmują RSW „Prasa—Książka—Ruch”, Centrala Kolportażu Prasy i Wydawnictw, ul. Towarowa 28, 00-958 Warszawa, konto NBP XV Oddział w Warszawie Nr 1153-201045-139-11 — w terminach podanych dla prenumeraty krajowej.

Prenumerata ze zleceniem wysyłki za granicę jest droższa od prenumeraty krajowej o 50% dla zleceniodawców indywidualnych i o 100% dla zleceniodawców instytucji i zakładów pracy.

Bieżące i archiwalne numery można nabyć lub zamówić we Wzorcowni Wydawnictw Naukowych PAN-Ossolineum-PWN, Pałac Kultury i Nauki (wysoki parter) 00-901 Warszawa oraz w księgarniach naukowych „Domu Książki”.

A subscription order stating the period of time, along with the subscriber's name and address can be sent to your subscription agent or directly to Foreign Trade Enterprise Ars Polona — Ruch, 00-068 Warszawa, 7 Krakowskie Przedmieście, P.O. Box 1001, Poland. Please send payments to the account of Ars Polona — Ruch in Bank Handlowy S.A., 7 Traugutt Street, 00-067 Warszawa, Poland.

Rozpr. Elektrot. T. 27, z. 2, 309—616, Warszawa 1981

Index 37483

POLSKA AKADEMIA NAUK
KOMITET ELEKTROTECHNIKI

PL ISSN 0035-9386

ROZPRAWY ELEKTROTECHNICZNE

TOM XXVIII—ZESZYT 3

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1981

POLSKA AKADEMIA NAUK
KOMITET ELEKTROTECHNIKI

ROZPRAWY ELEKTROTECHNICZNE

TOM XXVII — ZESZYT 3

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1981

RADA REDAKCYJNA

Przewodniczący

prof. dr hab. JANUSZ GROSZKOWSKI — członek rzeczywisty PAN

Członkowie

prof. dr hab. LUDWIK BADIAN — członek korespondent PAN, prof. TADEUSZ CHOLEWICKI — członek rzeczywisty PAN, doc. dr ZDZISŁAW KACHLICKI, prof. dr hab. TADEUSZ KACZOREK, prof. dr LEONARD KNOCH, doc. dr hab. ANDRZEJ KUSY, prof. JAN MANITIUS, prof. dr WITOLD NOWICKI, prof. dr WITOLD ROSIŃSKI — członek korespondent PAN, prof. dr JERZY SEIDLER — członek korespondent PAN, prof. dr RYSZARD SIKORA, prof. dr hab. JANUSZ TUROWSKI, prof. dr STEFAN WĘGRZYN — członek rzeczywisty PAN

REDAKCJA

Redaktor naczelny

prof. dr WITOLD NOWICKI

Zastępca redaktora naczelnego

prof. dr hab. TADEUSZ KACZOREK

Sekretarz odpowiedzialny

red. JANINA MATHEISEL

ADRES REDAKCJI

00-661 Warszawa, Plac Jedności Robotniczej 1
Politechnika, Instytut Telekomunikacji, Gmach Elektroniki, pok. 483

Dyżury Redakcji: środy i piątki, godz. 11-13
tel. 28 89 81 lub 210 07-737

Telefony domowe: redaktor naczelny — 17 37 03
zastępca redaktora naczelnego — 44 88 34
sekretarz odpowiedzialny — 33 84 97

PAŃSTWOWE WYDAWNICTWO NAUKOWE

Warszawa, Miodowa 10

Nakład 650	Oddano do składania 12.XI.1981 r.
Ark. druk. 20,0	Podpisano do druku
Ark. wyd. 24,25	w listopadzie 1982 r.
Papier druk. b. sat. imp. kl III, 70 g 70×100	Druk ukończono w styczniu 1983 r.
Zamówienie nr 1026/12/81	Cena zł 40,— Z-82

Drukarnia im. Rewolucji Październikowej, Warszawa

Synteza macierzy rozproszenia układów bezstratnych metodami przestrzeni stanów

ANDRZEJ L. DOBRUCKI (OŁAWA), MARIAN S. PIEKARSKI (WROCŁAW)

Instytut Telekomunikacji i Akustyki Politechniki Wrocławskiej

Otrzymano 7.5.1979

W pracy przedstawiono nowy, bezpośredni, algebraiczny algorytm syntezy bezstratnych układów n -wrotowych opisanych macierzą rozproszenia. Wykorzystano ideę minimalnej realizacji macierzy metodami przestrzeni stanów. W przeciwieństwie do wcześniej znanych metod, podany algorytm jest prostszy obliczeniowo i polega na bezpośrednim obliczeniu jednej z możliwych macierzy opisujących bezstratną, bezinercyjną część realizowanego układu. Zaproponowano również nowy sposób realizacji bezstratnych układów bezinercyjnych opisanych macierzą rozproszenia.

1. WSTĘP

Problemem syntezy macierzy rozproszenia układów bezstratnych metodami przestrzeni stanów zajmowało się w ostatnich latach wielu autorów (patrz: bibliografia w pracy [1]). W większości tych prac nie uwzględniono jednak podstawowych ograniczeń, wynikających ze specyfiki technologii mikroelektronicznych, odnośnie struktury uzyskiwanego układu [8]. Ponadto metody syntezy proponowane w tych pracach cechuje duża ilość skomplikowanych obliczeń. Dla układów bezstratnych opisanych macierzą admitancyjną znana jest prosta obliczeniowo metoda syntezy uwzględniająca ograniczenia technologiczne [7], [8]. Kierując się podobną ideą, w poniższej pracy przedstawiono algorytm syntezy skupionych, liniowych, skończonych, stacjonarnych i bezstratnych (SLSSB) układów n -wrotowych opisanych macierzą rozproszenia $S(p)$ uwzględniający podstawowe ograniczenia technologii mikroelektronicznych. Algorytm ten był w skrócie po raz pierwszy prezentowany w pracy [3].

2. MACIERZ ROZPROSZENIA BEZSTRATNYCH n -WROTNIKÓW

Niech skupiony, liniowy, skończony, stacjonarny i bezstratny (SLSSB) n -wrotnik N będzie opisany znormalizowaną macierzą rozproszenia $S(p)$ o wymiarach $n \times n$ [6].

Twierdzenie 1. [6]. Macierz $S(p)$ jest realizowalna w układzie SLSSB wtedy i tylko wtedy, gdy:

1. $S(p)$ jest macierzą rzeczywistą, wymierną,
2. $S(p)$ jest macierzą Hurwitza, tzn. nie ma biegunów dla $\operatorname{Re} p \geq 0$,
3. $S(p)$ jest macierzą paraunitarną, tj.

$$S(p)S^*(-p) = \mathbf{1}_n \quad (1)$$

Dla układów bezinercyjnych macierz $S(p)$ jest macierzą stałą i warunkiem bezstratności jest ortogonalność macierzy S , tj.

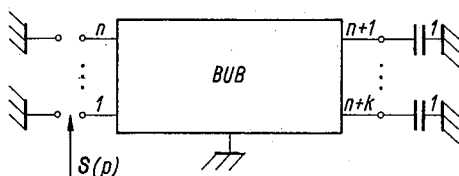
$$SS^t = \mathbf{1}_n. \quad (2)$$

Model syntezy macierzy $S(p)$ uwzględniający ograniczenia technologiczne [7], [8] przedstawiono na rys. 1.

Bezstratny układ bezinercyjny (BUB) opisany jest stałą ortogonalną macierzą rozproszenia S_a

$$S_a = \begin{bmatrix} S_{11} & S_{12} \\ S_{21} & S_{22} \end{bmatrix} \begin{matrix} n \\ k \end{matrix}. \quad (4)$$

Minimalna liczba kondensatorów niezbędnych do syntezy jest równa k , gdzie k jest stopniem kompleksowości macierzy $S(p)$ [1], [6]. Dla uproszczenia przyjęto jednakowe pojemności kondensatorów równe jedności.



Rys. 1. Model syntezy bezstratnych n -wrotników

Macierz rozproszenia $S(p)$ wyznaczona względem pierwszych n wrót BUB dana jest zależnością:

$$S(p) = S_{11} + S_{12} \left(\frac{1+p}{1-p} \mathbf{1}_k - S_{22} \right)^{-1} S_{21}. \quad (4)$$

Dokonując w zależności (4) podstawienia

$$s = \frac{1+p}{1-p}$$

otrzymuje się

$$S(p) = S_{11} + S_{12}(s\mathbf{1}_k - S_{22})^{-1}S_{21} = S\left(\frac{s-1}{s+1}\right). \quad (5)$$

Macierz $S(p)$ można wyrazić jako funkcję nowej zmiennej s

$$W(s) = S_{11} + S_{12}(s\mathbf{1}_k - S_{22})^{-1}S_{21}. \quad (6)$$

Z zależności (6) wynika, że czwórka macierzy S_{11} , S_{12} , S_{21} , S_{22} jest minimalną realizacją macierzy $W(s)$. Zatem macierz S_a można znaleźć przez obliczenie minimalnej realizacji macierzy $W(s)$. Jednakże znane dotychczas metody są skomplikowane obliczeniowo i nie zapewniają bezstratności macierzy S_a .

Istnieje inna droga prowadząca do znalezienia macierzy S_a . Wiadomo [1], że macierz $S(p)$ ma minimalną realizację o postaci

$$\begin{aligned} S(p) &= D + C(pI_k - A)^{-1}B, \\ k &= \delta[S(p)]. \end{aligned} \quad (7)$$

W celu znalezienia związków między minimalną realizacją (A, B, C, D) i macierzą S_a w zależności (7) podstawiono

$$p = \frac{s-1}{s+1}.$$

Po dokonaniu przekształceń otrzymuje się [4]:

$$S_{11} = D + C(I_k - A)^{-1}B, \quad (8a)$$

$$S_{12} = \sqrt{2} C(I_k - A)^{-1}, \quad (8b)$$

$$S_{21} = \sqrt{2} (I_k - A)^{-1}B, \quad (8c)$$

$$S_{22} = (I_k + A)(I_k - A)^{-1}. \quad (8d)$$

Korzystając ze związków pomiędzy macierzami S_{11} , S_{12} , S_{21} , S_{22} oraz A , B , C , D , z warunków bezstratności macierzy S_a można znaleźć odpowiednie warunki dla minimalnej realizacji (A, B, C, D) .

Twierdzenie 2. [1], [4]. Minimalna realizacja (A, B, C, D) bezstratnej macierzy rozproszenia $S(p)$ jest realizacją bezstratną wtedy i tylko wtedy, gdy:

$$A + A^t = -C^t C, \quad (9a)$$

$$B = -C^t D, \quad (9b)$$

$$DD^t = I_n. \quad (9c)$$

Zatem problem syntezy macierzy $S(p)$ można sprowadzić do obliczenia minimalnej realizacji (A, B, C, D) spełniającej warunki bezstratności.

Jednakże znane dotychczas metody nie pozwalają obliczyć w sposób bezpośredni bezstratnej realizacji (A, B, C, D) . Dla macierzy admitancyjnej znana jest metoda pozwalająca w sposób bezpośredni znaleźć macierz Y_a opisującą BUB [7], [8]. Kierując się podobną ideą opracowano algorytm pozwalający w sposób bezpośredni znaleźć minimalną realizację (A, B, C, D) , spełniającą warunki bezstratności (9), a co za tym idzie, dającą w efekcie ortogonalną macierz S_a .

3. ALGORYTM SYNTEZY

Macierz rozproszenia $S(p)$ można zapisać w postaci

$$S(p) = \frac{\sum_{i=0}^{i=r} B_i p^{r-i}}{g(p)}, \quad (10a)$$

gdzie B_i są rzeczywistymi macierzami o wymiarach $n \times n$, natomiast

$$g(p) = \sum_{i=0}^{i=r} b_i p^{r-i} \quad (10b)$$

jest najmniejszym wspólnym monicznym (tj. $b_0 = 1$) mianownikiem elementów macierzy $S(p)$. Na podstawie twierdzenia 1 $g(p)$ jest wielomianem Hurwitza [10] (tzn. $g(p) \neq 0$ dla $\text{Re } p \geq 0$).

Macierz $S(p)$ można również rozwinąć w szereg Laurenta w otoczeniu punktu $p = \infty$

$$S(p) = S_{-1} + \sum_{i=0}^{\infty} \frac{S_i}{p^{i+1}}, \quad (11)$$

gdzie macierze Markowa S_i są rzeczywistymi macierzami o wymiarach $n \times n$.

Porównując prawe strony zależności (10a) i (11) można napisać

$$\sum_{i=0}^{i=r} B_i p^{r-i} = g(p) \left(S_{-1} + \sum_{i=0}^{\infty} \frac{S_i}{p^{i+1}} \right). \quad (12)$$

Porównując następnie współczynniki przy jednakowych potęgach p otrzymuje się

$$B_j = \sum_{k=0}^{k=j} b_k S_{j-k-1}, \quad j = 0, \dots, r, \quad (13a)$$

oraz

$$0_n = \sum_{k=0}^{k=r} b_k S_{j-k-1}, \quad j \geq r+1. \quad (13b)$$

Zależności (13) pozwalają w sposób rekurencyjny wyznaczyć macierze S_i .

Podstawiając zależność (11) do (1) i porównując współczynniki przy jednakowych potęgach p otrzymuje się:

$$S_{-1} S_{-1}^t = I_n, \quad (14a)$$

$$\sum_{k=1}^{l+1} (-1)^{k-1} S_{l-k} S_{k-2}^t = 0_n, \quad l = 1, 2, \dots \quad (14b)$$

Zatem rzeczywista wymierna macierz $S(p)$ jest macierzą rozproszenia układu SLSSB wtedy i tylko wtedy, gdy są spełnione zależności (14) oraz $g(p)$ jest wielomianem Hurwitza.

Rozwijając w szereg Laurenta wokół punktu $p = \infty$ prawą stronę zależności (7)

$$S(p) = D + \sum_{i=0}^{\infty} \frac{CA^i B}{p^{i+1}} \quad (15)$$

i porównując z zależnością (11) otrzymuje się:

$$D = S_{-1}, \quad (16a)$$

$$CA^i B = S_i. \quad (16b)$$

Zależność (16a) jednoznacznie określa macierz D spełniającą warunek bezstratności (9c). Natomiast nieskończony ciąg równań (16b) nie daje jednoznacznego rozwiązania dla macierzy A , B i C spełniających warunki bezstratności (9a) i (9b).

W celu znalezienia macierzy A , B i C spełniających wspomniane warunki zdefiniowano macierz klatkową

$$T_l = [T_{ij}] = H_l X_l, \quad (17)$$

gdzie

$$H_l = \begin{bmatrix} S_0 & S_1 & \dots & S_l \\ S_1 & S_2 & \dots & S_{l+1} \\ \vdots & \vdots & \ddots & \vdots \\ S_l & S_{l+1} & \dots & S_{2l} \end{bmatrix} \quad (18)$$

jest uogólnioną macierzą henkelowską oraz

$$X_l = \begin{bmatrix} -S_{-1}^t & -S_0^t & \dots & -S_{l-1}^t \\ 0_n & S_{-1}^t & \dots & S_{l-2}^t \\ \vdots & \vdots & \ddots & \vdots \\ 0_n & 0_n & \dots & (-1)^{l+1} S_{-1}^t \end{bmatrix} \quad (19)$$

jest macierzą klatkową trójkątną.

Macierze T_{ij} można zapisać następująco

$$T_{ij} = \sum_{k=1}^{k=j} (-1)^{k+j-1} S_{i+j-k-1} S_{k-2}^t. \quad (20)$$

Korzystając z zależności (16) oraz warunków bezstratności (9) zależność (20) można przekształcić do następującej postaci

$$T_{ij} = CA^{i-1} (A^t)^{j-1} C^t. \quad (21)$$

Macierz T_l przyjmuje zatem postać

$$T_l = \begin{bmatrix} CC^t & CA^t C^t & \dots & C(A^t)^l C^t \\ CAC^t & CAA^t C^t & \dots & CA(A^t)^l C^t \\ \vdots & \vdots & \ddots & \vdots \\ CA^l C^t & CA^l A^t C^t & \dots & CA^l (A^t)^l C^t \end{bmatrix} = P_l P_l^t, \quad (22)$$

gdzie

$$P_l = \begin{bmatrix} C \\ CA \\ \vdots \\ CA^l \end{bmatrix}. \quad (23)$$

Zależność (22) sugeruje, że macierz C można wyznaczyć drogą faktoryzacji macierzy T_l . Jednakże z góry nie wiadomo, czy macierz T_l dana zależnością (17) może być faktoryzowana. Należy zatem zbadać własności macierzy T_l .

Macierz X_l jest nieosobliwa (bo, z uwagi na zal. (14a), nieosobliwa jest macierz S_{-1}), zatem rząd macierzy T_l jest równy rządowi macierzy H_l . Wiadomo [1], [10], że dla $l \geq r-1$

$$\text{rz } H_l = \text{rz } H_{l-1} = k = \delta[S(p)]. \quad (24)$$

Zatem

$$\text{rz } T_l = \text{rz } T_{l-1} = k. \quad (25)$$

W dalszych rozważaniach wystarczy zatem ograniczyć się jedynie do macierzy T_{r-1} o wymiarach $rn \times rn$.

Korzystając z zależności (14b) i (20) można pokazać, że $\mathbf{T}_{ij} = \mathbf{T}_{ji}^t$, a zatem macierz $\mathbf{T}_{r-1}$ jest symetryczna

$$\mathbf{T}_{r-1} = \mathbf{T}_{r-1}^t. \quad (26)$$

Uogólniona macierz stowarzyszona $\mathbf{\Omega}$ ma postać

$$\mathbf{\Omega} = \begin{bmatrix} \mathbf{0}_n & \mathbf{1}_n & \dots & \mathbf{0}_n \\ \vdots & \vdots & & \vdots \\ \mathbf{0}_n & \mathbf{0}_n & \dots & \mathbf{1}_n \\ -b_r \mathbf{1}_n & -b_{r-1} \mathbf{1}_n & \dots & -b_1 \mathbf{1}_n \end{bmatrix}. \quad (27)$$

Jej wielomian charakterystyczny jest równy

$$\det[s\mathbf{1}_{nr} - \mathbf{\Omega}] = [g(s)]^n. \quad (28)$$

Wszystkie wartości własne macierzy $\mathbf{\Omega}$ leżą zatem w lewej półpłaszczyźnie, tj.

$$\operatorname{Re} s_l < 0. \quad (29)$$

Z zależności (28) oraz twierdzenia Cayleya-Hamiltona wynika, że

$$g(\mathbf{\Omega}) = \mathbf{0}_{nr}. \quad (30)$$

Wprowadza się również macierz

$$\mathbf{T}_d = \mathbf{\Omega T}_{r-1} = \mathbf{\Omega H}_{r-1} \mathbf{X}_{r-1}. \quad (31)$$

W celu udowodnienia, że macierz $\mathbf{T}_{r-1}$ jest dodatnio półokreślona rozpatrzono wyrażenie

$$\mathbf{\Omega T}_{r-1} + \mathbf{T}_{r-1} \mathbf{\Omega}^t = \mathbf{\Omega T}_{r-1} + \mathbf{T}_{r-1}^t \mathbf{\Omega}^t = \mathbf{T}_d + \mathbf{T}_d^t.$$

Korzystając z zależności (31), (20), (27) i (14) otrzymuje się

$$\mathbf{\Omega T}_{r-1} + \mathbf{T}_{r-1} \mathbf{\Omega}^t = \begin{bmatrix} -S_0 S_0^t & -S_0 S_1^t & \dots & -S_0 S_{r-1}^t \\ -S_1 S_0^t & -S_1 S_1^t & \dots & -S_1 S_{r-1}^t \\ \vdots & \vdots & & \vdots \\ -S_{r-1} S_0^t & -S_{r-1} S_1^t & \dots & -S_{r-1} S_{r-1}^t \end{bmatrix} = - \begin{bmatrix} S_0 \\ S_1 \\ \vdots \\ S_{r-1} \end{bmatrix} [S_0^t \ S_1^t \ \dots \ S_{r-1}^t].$$

Oznaczając

$$\mathbf{L} = \begin{bmatrix} S_0 \\ S_1 \\ \vdots \\ S_{r-1} \end{bmatrix},$$

można otrzymać równanie, którego rozwiązaniem jest macierz $\mathbf{T}_{r-1}$

$$\mathbf{\Omega T}_{r-1} + \mathbf{T}_{r-1} \mathbf{\Omega}^t = -\mathbf{L L}^t. \quad (32)$$

Ponieważ wartości własne macierzy $\mathbf{\Omega}$ leżą w lewej półpłaszczyźnie (zal. (29)), równanie macierzowe (32) ma jedyne rozwiązanie, które jest macierzą symetryczną i dodatnio półokreśloną [2].

Pokazano zatem, że macierz $\mathbf{T}_{r-1}$ jest symetryczną, dodatnio półokreśloną macierzą rzędu k . Macierz taką można faktoryzować w postaci [5]

$$\mathbf{T}_{r-1} = \mathbf{M M}^t, \quad (33)$$

gdzie

$$\mathbf{M} = \begin{bmatrix} \mathbf{M}_0 \\ \mathbf{M}_1 \\ \vdots \\ \mathbf{M}_{r-1} \end{bmatrix} \quad (34)$$

jest rzeczywistą macierzą o wymiarach $nr \times k$.

Porównując zależności (33) i (34) z zależnościami (22) i (23) przyjmuje się:

$$\mathbf{P}_{r-1} = \mathbf{M}, \quad (35)$$

$$\mathbf{C} = \mathbf{M}_0, \quad (36)$$

oraz z zależnościami (9b)

$$\mathbf{B} = -\mathbf{C}^t \mathbf{D} = -\mathbf{M}_0^t \mathbf{S}_{-1}. \quad (37)$$

Macierze $\mathbf{B}$ i $\mathbf{C}$ spełniają zatem warunek bezstratności.

Do obliczenia macierzy $\mathbf{A}$ należy skorzystać z wcześniej wprowadzonych macierzy $\mathbf{T}_d$ i $\mathbf{\Omega}$.

Porównując zależności (17), (22) i (31) można napisać następującą równość [4]

$$\mathbf{T}_d = \mathbf{P}_{r-1} \mathbf{A} \mathbf{P}_{r-1}^t. \quad (38)$$

Z zależności (35) i (38) wynika równanie, którego rozwiązaniem jest macierz $\mathbf{A}$

$$\mathbf{M} \mathbf{M}^t = \mathbf{T}_d. \quad (39)$$

Równanie (39) ma rozwiązanie wtedy i tylko wtedy, gdy [9]

$$\mathbf{M} \mathbf{M}^+ \mathbf{T}_d (\mathbf{M} \mathbf{M}^+)^t = \mathbf{T}_d, \quad (40)$$

gdzie $\mathbf{M}^+$ jest macierzą pseudoodwrotną do $\mathbf{M}$ o następujących własnościach [5], [9]:

$$\mathbf{M} \mathbf{M}^+ \mathbf{M} = \mathbf{M}, \quad (41a)$$

$$\mathbf{M}^+ \mathbf{M} \mathbf{M}^+ = \mathbf{M}^+, \quad (41b)$$

$$(\mathbf{M} \mathbf{M}^+)^t = \mathbf{M} \mathbf{M}^+, \quad (41c)$$

$$(\mathbf{M}^+ \mathbf{M})^t = \mathbf{M}^+ \mathbf{M}. \quad (41d)$$

W rozważanym przypadku macierz $\mathbf{M}$ ma wymiar $nr \times k$ i rząd k (tzn. jest macierzą maksymalnego rzędu); zatem $\mathbf{M}^+$ dana jest zależnością [5]

$$\mathbf{M}^+ = (\mathbf{M}^t \mathbf{M})^{-1} \mathbf{M}^t. \quad (42)$$

Z zależności (42) wynika, że

$$\mathbf{M}^+ \mathbf{M} = \mathbf{I}_k. \quad (43)$$

Korzystając z zależności (17), (31) i (41), faktu, że $\mathbf{\Omega} \mathbf{H}_{r-1} = \mathbf{H}_{r-1} \mathbf{\Omega}^t$ [10] oraz nieosobliwości macierzy $\mathbf{X}_{r-1}$ można wykazać, że zależność (40) jest spełniona, a ponadto

$$\mathbf{M} \mathbf{M}^+ \mathbf{T}_d = \mathbf{T}_d \mathbf{M} \mathbf{M}^+ = \mathbf{T}_d. \quad (44)$$

Ponieważ macierz $\mathbf{M}$ jest macierzą maksymalnego rzędu, równanie (39) ma jedyne rozwiązanie, dane zależnością [4], [9]

$$\mathbf{A} = \mathbf{M}^+ \mathbf{T}_d (\mathbf{M}^+)^t. \quad (45)$$

Korzystając z zależności (31) oraz (43) można uzyskać inną postać rozwiązania

$$\mathbf{A} = \mathbf{M}^+ \mathbf{\Omega} \mathbf{M}. \quad (46)$$

Należy sprawdzić, czy tak obliczona macierz $\mathbf{A}$ spełnia warunek bezstratności (9a). Obliczając

$$\mathbf{A} + \mathbf{A}^t = \mathbf{M}^+ (\mathbf{T}_a + \mathbf{T}_d') (\mathbf{M}^+)^t,$$

na podstawie zależności (32), (33) i (17) oraz korzystając z własności macierzy pseudo-odwrotnej, otrzymuje się

$$\mathbf{A} + \mathbf{A}^t = -\mathbf{M}_0' \mathbf{M}_0.$$

Ponieważ $\mathbf{M}_0 = \mathbf{C}$, stąd

$$\mathbf{A} + \mathbf{A}^t = -\mathbf{C}^t \mathbf{C},$$

a zatem warunek bezstratności jest spełniony.

Należy ponadto wykazać, że tak wybrane macierze $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$ i $\mathbf{D}$ są rzeczywiście minimalną realizacją macierzy $\mathbf{S}(p)$, tzn. że spełniona jest zależność (7) lub też równoważny jej ciąg zależności (16b). Z zależności (34), (36) i (37) wynika, że równanie (16b) jest spełnione dla $i = 0, 1, \dots, r-1$. Należy wykazać słuszność tej zależności dla $i \geq r$. Korzystając z zależności (45), (46), (44), (43), (31) i (41) otrzymuje się

$$\mathbf{A}^2 = \mathbf{M}^+ \mathbf{\Omega}^2 \mathbf{M},$$

oraz ogólnie

$$\mathbf{A}^i = \mathbf{M}^+ \mathbf{\Omega}^i \mathbf{M}, \quad i = 0, 1, \dots \quad (47)$$

Ponieważ na podstawie zależności (30) i (46)

$$g(\mathbf{A}) = \sum_{i=0}^{i=r} b_i \mathbf{A}^{r-1} = \mathbf{M}^+ \left(\sum_{i=0}^{i=r} b_i \mathbf{\Omega}^{r-1} \right) \mathbf{M} = \mathbf{M}^+ g(\mathbf{\Omega}) \mathbf{M} = \mathbf{0}_k, \quad (48)$$

wynika stąd, że

$$\mathbf{A}^r = - \sum_{i=1}^{i=r} b_i \mathbf{A}^{r-i}. \quad (49)$$

Zatem, na podstawie zależności (13b)

$$\mathbf{S}_r = - \sum_{i=1}^{i=r} b_i \mathbf{S}_{r-1} = - \sum_{i=1}^{i=r} b_i \mathbf{C} \mathbf{A}^{r-1} \mathbf{B} = \mathbf{C} \left(- \sum_{i=1}^{i=r} b_i \mathbf{A}^{r-1} \right) \mathbf{B} = \mathbf{C} \mathbf{A}^t \mathbf{B}. \quad (50)$$

Można wykazać przez indukcję, że

$$\mathbf{S}_i = \mathbf{C} \mathbf{A}^i \mathbf{B}, \quad i \geq r.$$

Oznacza to, że zależność (16b) jest spełniona dla wszystkich $i \geq 0$.

Z przeprowadzonych rozważań wynika, że macierze $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$ i $\mathbf{D}$ dane zależnościami (16a), (36), (37) i (46) zostały wybrane prawidłowo. Zatem macierz $\mathbf{S}_a$ dana zależnością

$$\mathbf{S}_a = \begin{bmatrix} \mathbf{S}(1) & \sqrt{2} \mathbf{M}_0 (\mathbf{1}_k - \mathbf{M}^+ \mathbf{\Omega} \mathbf{M})^{-1} \\ -\sqrt{2} (\mathbf{1}_k - \mathbf{M}^+ \mathbf{\Omega} \mathbf{M})^{-1} \mathbf{M}_0' \mathbf{S}_{-1} & (\mathbf{1}_k + \mathbf{M}^+ \mathbf{\Omega} \mathbf{M}) (\mathbf{1}_k - \mathbf{M}^+ \mathbf{\Omega} \mathbf{M})^{-1} \end{bmatrix} \quad (51)$$

jest poszukiwaną macierzą rozproszenia bezstratnego układu bezinercyjnego.

Przykład

Dana jest macierz rozproszenia bezstratnego czwórnika

$$S(p) = \frac{1}{p^2 + p + 1} \begin{bmatrix} -p & p^2 + 1 \\ p^2 + 1 & -p \end{bmatrix}.$$

Korzystając z zależności (13) wyznacza się

$$S_{-1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix},$$

$$S_0 = \begin{bmatrix} -1 & -1 \\ -1 & -1 \end{bmatrix},$$

$$S_1 = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}$$

$$S_2 = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}.$$

Następnie tworzy się macierz T_1

$$T_1 = \begin{bmatrix} 1 & 1 & -1 & -1 \\ 1 & 1 & -1 & -1 \\ -1 & -1 & 2 & 2 \\ -1 & -1 & 2 & 2 \end{bmatrix}.$$

Rząd macierzy T_1 równa się 2. Po faktoryzacji otrzymuje się

$$M = \begin{bmatrix} -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ \sqrt{2} & 0 \\ \sqrt{2} & 0 \end{bmatrix},$$

oraz

$$M^+ = \begin{bmatrix} 0 & 0 & \frac{\sqrt{2}}{4} & \frac{\sqrt{2}}{4} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} & \frac{\sqrt{2}}{4} & \frac{\sqrt{2}}{4} \end{bmatrix}$$

Stąd

$$C = \frac{1}{\sqrt{2}} \begin{bmatrix} -1 & 1 \\ -1 & 1 \end{bmatrix},$$

$$B = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix}.$$

Macierz stowarzyszona Ω jest równa

$$\Omega = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ -1 & 0 & -1 & 0 \\ 0 & -1 & 0 & -1 \end{bmatrix},$$

zatem

$$A = \begin{bmatrix} -\frac{1}{2} & -\frac{1}{2} \\ \frac{3}{2} & -\frac{1}{2} \end{bmatrix}.$$

Macierz S_a jest więc równa

$$S_a = \begin{bmatrix} -\frac{1}{3} & \frac{2}{3} & 0 & \frac{2}{3} \\ \frac{2}{3} & -\frac{1}{3} & 0 & \frac{2}{3} \\ \frac{2}{3} & \frac{2}{3} & 0 & -\frac{1}{3} \\ 0 & 0 & 1 & 0 \end{bmatrix}.$$

Łatwo sprawdzić, że jest ona ortogonalna, a zatem jest ona macierzą rozproszenia bezstratnego układu bezinercyjnego.

4. REALIZACJA BEZSTRATNEGO UKŁADU BEZINERCYJNEGO

Zasadniczym problemem przy syntezie układów metodami przestrzeni stanów jest realizacja bezstratnego układu bezinercyjnego. Większość prac, dotyczących syntezy układów bezinercyjnych, dotyczy wielowrotników opisanych macierzą admitancyjną, dlatego też celowe wydaje się sprowadzenie problemu realizacji BUB opisanego macierzą rozproszenia, do realizacji BUB opisanego macierzą admitancyjną.

Wiadomo [5], że przez użycie transformacji ortogonalnej, można każdą macierz ortogonalną S_a o wymiarach $m \times m$ sprowadzić do postaci kanonicznej S_0

$$S_a = TS_0T^t, \quad (52)$$

gdzie:

$$S_0 = \left[-\mathbf{1}_p + \begin{pmatrix} \cos x_1 & -\sin x_1 \\ \sin x_1 & \cos x_1 \end{pmatrix} + \dots + \begin{pmatrix} \cos x_g & -\sin x_g \\ \sin x_g & \cos x_g \end{pmatrix} + \mathbf{1}_r \right], \quad (53)$$

$$p = m - rz(\mathbf{1}_m + S_a), \quad (54)$$

$$r = m - rz(\mathbf{1}_m - S_a), \quad (54b)$$

$$2g = rz(S_a - S_a^t) = m - p - r, \quad (54c)$$

oraz

$$\lambda_i = \cos x_i \pm j \sin x_i, \quad \sin x_i > 0, \quad i = 1, 2, \dots, g,$$

są zespolonymi wartościami własnymi macierzy S_a , natomiast kolumny macierzy T są znormalizowanymi wektorami własnymi macierzy S_a .

Zależność (52) opisuje kaskadowe połączenie multitransformatora ortogonalnego i układu opisanego macierzą S_0 składającego się z p wrót zwartych, g żyatorów i r wrót rozwartych. Jest to jednak rozwiązanie niepraktyczne, gdyż prowadzi do realizacji nie spełniającej ograniczeń technologicznych. Zależność (52) można zapisać w postaci

$$S_a = -T(-S_0)T^t = -TS_iT^t. \quad (55)$$

Odpowiada to kaskadowemu połączeniu multiżyrotora ortogonalnego o macierzy

$$S_b = \begin{bmatrix} 0_m & \mp T \\ \pm T^t & 0_m \end{bmatrix} \quad (56)$$

i układu opisanego macierzą S_i składającego się z:

1. p wrót zwartych,
2. g żyrotorów, opisanych macierzami rozproszenia S_i o wymiarach 2×2

$$S_i = \begin{bmatrix} -\cos x_i & \sin x_i \\ -\sin x_i & -\cos x_i \end{bmatrix}, \quad (57)$$

a stąd

$$Y_i = (1_2 - S_i)(1_2 + S_i)^{-1} = \frac{1}{1 - \cos x_i} \begin{bmatrix} 0 & -\sin x_i \\ \sin x_i & 0 \end{bmatrix} = \begin{bmatrix} 0 & y_i \\ -y_i & 0 \end{bmatrix}, \quad (58)$$

3. r wrót zwartych.

Dla multiżyrotora ortogonalnego istnieje zawsze zarówno macierz impedancyjna jak i admitancyjna [4]

$$Y_b = -S_b = -Z_b = \begin{bmatrix} 0_m & \pm T \\ \mp T^t & 0_m \end{bmatrix}. \quad (59)$$

Ponieważ admitancja wejściowa żyrotora, którego wyjście zostało zwarte, jest równa zeru, w układzie opisanym zależnością (59) można wyeliminować r ostatnich wrót multiżyrotora. Macierz admitancyjna przyjmuje wówczas postać

$$Y_c = \begin{bmatrix} 0_m & Y_1 \\ -Y_1^t & 0_{p+2g} \end{bmatrix}, \quad (60)$$

gdzie

$$Y_1 = T \cdot 1_{m \times (p+2g)}, \quad (61)$$

czyli Y_1 jest złożona z $p+2g$ pierwszych kolumn macierzy T .

Ostatnie $2g$ wrót układu opisanego macierzą Y_c można również wyeliminować przez włączenie g pojedynczych żyrotorów do wnętrza multiżyrotora. Nowy $(m+p)$ -wrotowy BUB jest opisany macierzą admitancyjną

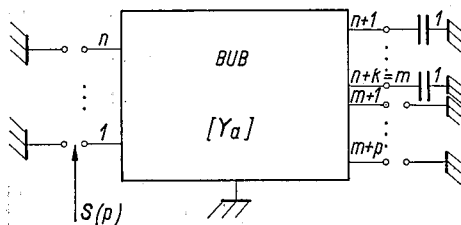
$$Y_a = \begin{bmatrix} -Y_3 Y_1^{-1} Y_3^t & Y_2 \\ -Y_2^t & 0_p \end{bmatrix} \begin{matrix} m \\ p \end{matrix} \quad (62)$$

gdzie: Y_2 składa się z p pierwszych kolumn macierzy T , a Y_3 z kolejnych $2g$ kolumn oraz

$$Y_1 = y_1 \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix} + \dots + y_g \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}. \quad (63)$$

Zatem realizację m -wrotowego BUB, opisanego macierzą rozproszenia S_a , można sprowadzić do realizacji $(m+p)$ -wrotowego BUB opisanego macierzą admitancyjną Y_a , którego ostatnie p wrót jest nieobciążonych.

Realizację BUB opisanego macierzą admitancyjną można przeprowadzić różnymi metodami, jednakże szczególnie korzystna wydaje się być realizacja przedstawiona w pracy [8]. Zastosowano tam pary źródeł sterowanych typu: źródło napięciowe sterowane napięciowo — różnicowe źródło prądowe sterowane prądowo (ZNSN/RZPSP). Podstawową zaletą takiej realizacji BUB jest mała wrażliwość parametrów układu na zmiany wartości elementów układu.



Rys. 2. Model syntezy bezstratnych n -wrotników opisanych macierzą rozproszenia $S(p)$

Przykład

Dana jest macierz rozproszenia

$$S(p) = \frac{1}{p^2 + p + 1} \begin{bmatrix} 0 & p^2 - p + 1 \\ p^2 + p + 1 & 0 \end{bmatrix}.$$

Macierz S_a obliczona wcześniej przedstawioną metodą jest równa

$$S_a = \begin{bmatrix} 0 & \frac{1}{3} & 0 & -\frac{2\sqrt{2}}{3} \\ 1 & 0 & 0 & 0 \\ 0 & \frac{2\sqrt{2}}{3} & 0 & \frac{1}{3} \\ 0 & 0 & -1 & 0 \end{bmatrix}.$$

Wartości własne S_a są równe

$$\lambda_1 = -1, \quad \lambda_2 = j, \quad \lambda_3 = -j, \quad \lambda_4 = 1;$$

stąd $p = 1, g = 1, r = 1$.

Stosując omówioną metodę obliczania macierzy Y_a , tzn. obliczając kolejno macierze S_0 i T (52), następnie Y_b (59), Y_c (60) otrzymuje się ostatecznie Y_a (62)

$$Y_a = \begin{bmatrix} 0 & \frac{1}{3} & 0 & \frac{\sqrt{2}}{3} & -\frac{1}{\sqrt{3}} \\ -\frac{1}{3} & 0 & \frac{\sqrt{2}}{3} & 0 & \frac{1}{\sqrt{3}} \\ 0 & -\frac{\sqrt{2}}{3} & 0 & -\frac{2}{3} & -\frac{1}{\sqrt{6}} \\ -\frac{\sqrt{2}}{3} & 0 & \frac{2}{3} & 0 & -\frac{1}{\sqrt{6}} \\ \frac{1}{\sqrt{3}} & -\frac{1}{\sqrt{3}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{6}} & 0 \end{bmatrix}.$$

5. PODSUMOWANIE

Algorytm syntezy przedstawiony w niniejszej pracy polega na bezpośrednim obliczeniu minimalnej, bezstratnej realizacji (A, B, C, D) macierzy rozproszenia $S(p)$. Po znalezieniu macierzy A, B, C i D obliczenie jednej z możliwych macierzy S_a opisującej bezstratny układ bezinercyjny jest bardzo proste. Stosując znane przekształcenia ekwiwalentne [1] można znaleźć wszystkie możliwe macierze S_a , co może być przydatne przy optymalizacji układu. W przeciwieństwie do metod podanych przez Youlę i Tissiego [10] oraz Andersona i Vongpanitlerda [1] nie jest konieczne rozwiązywanie równania macierzowego. Konieczna jest jedynie faktoryzacja rzeczywistej, symetrycznej, dodatnio półokreślonej macierzy T_{r-1} w postaci MM' .

Zaproponowano również nowy sposób realizacji m -wrotowego BUB opisanego macierzą rozproszenia, poprzez realizację $(m+p)$ -wrotowego BUB opisanego macierzą admitancyjną. Proponuje się realizację $(m+p)$ -wrotowego BUB przy użyciu źródeł sterowanych typu ZNSN/RZPSP. Cechuje ją mniejsza wrażliwość parametrów układu na zmiany wartości elementów w porównaniu z innymi znanymi metodami.

LITERATURA

1. B. D. O. Anderson, S. Vongpanitlerd, *Network analysis and synthesis. A modern systems theory approach*, Prentice-Hall, New York 1973.
2. R. W. Brockett, *Finite dimensional linear systems*, John Wiley, New York 1970.
3. A. L. Dobrucki, M. S. Piekarski, *Synteza bezstratnych układów mikroelektronicznych opisanych macierzą rozproszenia*, Materiały I Krajowej Konferencji „Teoria Obwodów i Układy Elektroniczne”, Podlesice 1977.
4. A. L. Dobrucki, *Synteza macierzy rozproszenia bezstratnych wielowrotników z uwzględnieniem wymagań technologii mikroelektronicznych*, Praca doktorska. Komunikat ITA, I-28/K-046/77, Wrocław 1977.
5. F. R. Gantmacher, *Teorija matric*, Izdatelstvo Nauka, Moskva 1966.
6. R. W. Newcomb, *Linear multiport synthesis*, McGraw-Hill, New York 1966.
7. M. S. Piekarski, *Synteza bezstratnych układów scalonych metodą przestrzeni stanów*, Archiwum Elektrotechniki, t. XXV, z. 2, 1976.
8. M. S. Piekarski, *Wybrane zagadnienia syntezy liniowych układów mikroelektronicznych*, Prace naukowe ITA, Seria Monografie, nr 9, Wrocław 1976.
9. M. Warmus, *Uogólnione odwrotności macierzy*, PWN, Warszawa 1972.
10. D. C. Youla, P. Tissi, *n-Port synthesis via reactance extraction — part I*, IEEE Intern. Conv. Rec., 1966.

A. L. DOBRUCKI, M. S. PIEKARSKI

A SCATTERING MATRIX SYNTHESIS OF LOSSLESS NETWORKS
VIA STATE-SPACE TECHNIQUES

Summary

In the paper a new direct algebraic algorithm of synthesis of lossless n -ports described by a scattering matrix is presented. An idea of minimal realization of a matrix via state-space technique is utilised. The algorithm of synthesis is much simpler than those known up to now. A new method of realization of lossless nondynamic network described by the scattering matrix is also presented.

A. L. DOBRUCKI, M. S. PIEKARSKI

SYNTHÈSE DE SCATTERING MATRICE DES CIRCUITS
SANS DISSIPATION DANS L'ESPACE D'ÉTATS

Résumé

Dans le travail on a présenté un nouveau, direct et algébrique algorithme de la synthèse des n -pôles circuits sans dissipation, décrits à l'aide de scattering matrice. On a utilisé l'idée de la réalisation minimale de la matrice à l'aide des états variables. Du point de vue de calcul, l'algorithme proposé est plus simple que les méthodes connus jusqu'à présent. On a donné aussi une nouvelle méthode de la réalisation des circuits non dynamiques et sans dissipation décrits à l'aide de scattering matrice.

A. L. DOBRUCKI, M. S. PIEKARSKI

SYNTHESE DER STREUMATRIX VERLUSTFREIER SCHALTUNGEN
MIT HILFE VON ZUSTANDSRAUMMETHODEN

Zusammenfassung

Es wird eine neue, direkte, algebraische Synthesenprozedur für verlustfreie, durch eine Streumatrix beschriebene N -Tor-Systeme gebracht. Der angegebene Synthesenalgorithmus läßt sich — im Gegensatz zu den früher bekannten Algorithmen — einfacher berechnen. Er beruht auf direkter Berechnung einer der zulässigen Matrizen, die den verlustfreien inertionslosen Teil der realisierten Schaltung beschreiben. Man hat auch eine neue Realisierungsmethode für eine verlustfreie, inerztionslose, durch die Streumatrix beschriebene Schaltung vorgeschlagen.

A. Л. ДОБРУЦКИ, М. С. ПЕКАРСКИ

СИНТЕЗ МАТРИЦЫ РАССЕЯНИЯ ЦЕПЕЙ БЕЗ ПОТЕРЬ
МЕТОДАМИ ПРОСТРАНСТВА СОСТОЯНИЙ

Резюме

Приведен новый, непосредственный, алгебраический алгоритм синтеза n -парополосников без потерь, описанных матрицей рассеяния. Использована идея минимальной реализации матрицы методами пространства состояний. Приведенный алгоритм синтеза на много проще чем известные до сих пор методы. Приведен также новый метод реализации нединамических цепей без потерь описанных матрицей рассеяния.

Potential Distribution Along Underground Conductors Connected to Station Earth Electrode

MACIEJ KRAKOWSKI (ŁÓDŹ)

Instytut Informatyki Politechniki Łódzkiej

GRZEGORZ SZYMAŃSKI (POZNAŃ)

Instytut Elektrotechniki Przemysłowej Politechniki Poznańskiej

ALFRED ZIĘTKOWIAK (POZNAŃ)

Biuro Studiów i Projektów Energetycznych „Energoprojekt” w Poznaniu

The paper received 9.7.1979

A method of computing of the currents, potentials and voltages between sheath and cable conductor along a bundle of underground conductors connected to a station earth electrode is presented. General matrix equation for the currents and potentials along the parallel underground conductors is given. The equivalent circuit and differential equations for the conductors are derived. The results of calculations are compared with the experimental data.

1. INTRODUCTION

The underground conductors such as cables, pipes etc. are often connected to station earth electrodes. Thus the potential due to the fault to ground is imposed on the underground conductors from the station to remote points. The problems concerned with such effects are discussed in a number of publications e.g. [1, 2, 7].

In the paper, a bundle of underground conductors connected to a station earth electrode is considered. If the conductors are placed in a trench or interconnected for instance by cable boxes, it can be assumed that the potentials of all the conductors are identical. Such an assumption is accepted in the following work.

List of principal symbols

a_{mn} — horizontal distance between parallel underground conductors m and n ,

d_m — depth at which underground conductor m is buried,

G_{im} — unit-length leakage conductance of insulation or coating of underground conductor m ,

$I_m(x)$ — current which flows through underground conductor m ,

$$k^2 = j\omega\mu_0\gamma,$$

r_m — external radius of underground conductor m ,

$U(x)$ — voltage between sheath and cable conductor,

$V_m(x)$ — potential along underground conductor m ,

Z_{im} — unit-length external-surface impedance of underground conductor m ,

Z_o — characteristic impedance of a single underground conductor,

Γ — propagation coefficient of underground conductor,

γ — earth conductivity,

$\mu_0 = 4\pi \cdot 10^{-7}$ H/m — permeability of vacuum,

ρ — earth resistivity,

ω — angular frequency.

2 TWO PARALLEL UNDERGROUND CONDUCTORS

A pair of parallel underground conductors is discussed at length in [4]. Now the fundamental concepts concerning such conductors will be presented, which permits to obtain an approximate solution of the problem, in the shape similar as for the transmission line.

The propagation coefficient Γ_{om} of the conductor denoted by m ($m = 1, 2, \dots$) is the solution of the transcendental equation

$$\Gamma_{om}^2 \left[G_{im}^{-1} + \frac{1}{\pi\gamma} \ln \frac{1,12}{\sqrt{2r_m d_m (\Gamma_{om}^2 + k^2)}} \right] = Z_{im} + \frac{j\omega\mu_0}{\pi} \ln \frac{1,12}{\sqrt{2r_m d_m (\Gamma_{om}^2 + k^2)}}, \quad (1)$$

with the requirement $\text{Re}(\Gamma_m) > 0$, which can be solved by method of successive approximations as in [3].

Denoting

$$Y_{mm}^{-1} = G_{im}^{-1} + \frac{1}{\pi\gamma} \ln \frac{1,12}{\sqrt{2r_m d_m (\Gamma_{om}^2 + k^2)}}, \quad (2)$$

$$Z_{mm} = Z_{im} + \frac{j\omega\mu_0}{\pi} \ln \frac{1,12}{\sqrt{2r_m d_m (\Gamma_{om}^2 + k^2)}}, \quad (3)$$

the propagation coefficient becomes

$$\Gamma_{om} = \sqrt{Z_{mm} Y_{mm}}, \quad \text{Re}(\Gamma_{om}) > 0, \quad (4)$$

whereas the characteristic impedance of the conductor m is defined by

$$Z_{om} = \sqrt{\frac{Z_{mm}}{Y_{mm}}}, \quad \text{Re}(Z_{om}) > 0. \quad (5)$$

The propagation coefficient Γ_{om} concerns a single conductor m without interaction of the other conductors i.e. if the conductor m is remote from the other ones.

Furthermore

$$Y_{mn}^{-1} = \frac{1}{\pi\gamma} \ln \frac{1,12}{\sqrt{s_{mn} s'_{mn} (\Gamma_{mn}^2 + k^2)}}, \quad (6)$$

$$Z_{mn} = \frac{j\omega\mu_0}{\pi} \ln \frac{1,12}{\sqrt{s_{mn} s'_{mn} (\Gamma_{mn}^2 + k^2)}},$$

where

$$s_{mn} = \sqrt{a_{mn}^2 + (d_m - d_n)^2}, \quad s'_{mn} = \sqrt{a_{mn}^2 + (d_m + d_n)^2}, \quad \Gamma_{mn} = \sqrt{\Gamma_{om}\Gamma_{on}}$$

and $m, n = 1, 2, \dots$. The quantities Y_{mn}^{-1} and Z_{mn} express the interaction of the conductors m and n .

The equations due to the pair of underground conductors 1 and 2 which are parallel to the x -axis of a rectangular coordinate system can be accepted in the form

$$\begin{aligned} \frac{dV'_1}{dx} + Z_{11}I_1(x) + Z_{12}I_2(x) &= E_1^0(x), \\ \frac{dV'_2}{dx} + Z_{21}I_1(x) + Z_{22}I_2(x) &= E_2^0(x), \\ V'_1(x) &= -Y_{11}^{-1} \frac{dI_1}{dx} - Y_{12}^{-1} \frac{dI_2}{dx}, \\ V'_2(x) &= -Y_{21}^{-1} \frac{dI_1}{dx} - Y_{22}^{-1} \frac{dI_2}{dx}, \end{aligned} \quad (7)$$

where $Z_{mn} = Z_{nm}$ and $Y_{mn}^{-1} = Y_{nm}^{-1}$ are taken from eqns. (6). The $E_m^0(x)$ and $V_m'(x)$ appearing in the last equations are the electric intensity of the primary and secondary components of the electric field along conductor m , respectively [3, 4]. The primary component of the electric field is impressed by the external factors which generate the field in general (6) the assumption that no currents flow in the conductors. On the other hand, the secondary component is excited by the currents flowing through the conductors.

The elimination of the potentials $V'_1(x)$, $V'_2(x)$ from eqns. (7) gives

$$\begin{aligned} Y_{11}^{-1} \frac{d^2 I_1}{dx^2} - Z_{11}I_1(x) + Y_{12}^{-1} \frac{d^2 I_2}{dx^2} - Z_{12}I_2(x) &= -E_1^0(x), \\ Y_{21}^{-1} \frac{d^2 I_1}{dx^2} - Z_{21}I_1(x) + Y_{22}^{-1} \frac{d^2 I_2}{dx^2} - Z_{22}I_2(x) &= -E_2^0(x). \end{aligned} \quad (8)$$

It is easy to solve the last equations using the Fourier integral if the conductors are assumed to be infinitely long. The currents and the electric intensities appearing in eqns. (8) are expressed by the Fourier integrals

$$I_1(x) = \int_{-\infty}^{+\infty} i_1(u) e^{jxu} du, \quad I_2(x) = \int_{-\infty}^{+\infty} i_2(u) e^{jxu} du, \quad (9a)$$

$$E_1^0(x) = \int_{-\infty}^{+\infty} e_1^0(u) e^{jxu} du, \quad E_2^0(x) = \int_{-\infty}^{+\infty} e_2^0(u) e^{jxu} du, \quad (9b)$$

where

$$e_1^0(u) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} E_1^0(x) e^{-jux} dx, \quad e_2^0(u) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} E_2^0(x) e^{-jux} dx. \quad (10)$$

Inserting eqns. (9) and (10) into eqns. (8) we obtain

$$\begin{aligned} \Delta_{11}i_1(u) + \Delta_{12}i_2(u) &= e_1^0(u), \\ \Delta_{12}i_1(u) + \Delta_{22}i_2(u) &= e_2^0(u), \end{aligned} \quad (11)$$

where

$$\Delta_{11} = Y_{11}^{-1}(u^2 + \Gamma_{01}^2), \quad \Delta_{22} = Y_{22}^{-1}(u^2 + \Gamma_{02}^2), \quad \Delta_{12} = Y_{12}^{-1}(u^2 + Z_{12}Y_{12}). \quad (12)$$

Substitution of $i_1(u)$ and $i_2(u)$ which are the solutions of eqn. (11) into eqn. (9a) yields

$$I_1(x) = \int_{-\infty}^{+\infty} \frac{\Delta_{22} e_1^0 - \Delta_{12} e_2^0}{\Delta_{11} \Delta_{22} - \Delta_{12}^2} e^{jxu} du, \quad (13)$$

$$I_2(x) = \int_{-\infty}^{+\infty} \frac{\Delta_{11} e_2^0 - \Delta_{12} e_1^0}{\Delta_{11} \Delta_{22} - \Delta_{12}^2} e^{jxu} du.$$

The last expressions are identical with the results derived in another way in [4].

3. BUNDLE OF n PARALLEL UNDERGROUND CONDUCTORS

The approach presented in the preceding section can be extended to a bundle of n parallel conductors. Thus the equations can be represented in the matrix form

$$\frac{d}{dx} [V'_m] = -[Z_{mn}][I_n] + [E_m^0], \quad (14a)$$

$$[V'_m] = -[Y_{mn}^{-1}] \frac{d}{dx} [I_n], \quad (14b)$$

where the square matrices

$$[Z_{mn}] = \begin{bmatrix} Z_{11} & Z_{12} & \dots & Z_{1n} \\ Z_{21} & Z_{22} & \dots & Z_{2n} \\ \dots & \dots & \dots & \dots \\ Z_{n1} & Z_{n2} & \dots & Z_{nn} \end{bmatrix}, \quad [Y_{mn}^{-1}] = \begin{bmatrix} Y_{11}^{-1} & Y_{12}^{-1} & \dots & Y_{1n}^{-1} \\ Y_{21}^{-1} & Y_{22}^{-1} & \dots & Y_{2n}^{-1} \\ \dots & \dots & \dots & \dots \\ Y_{n1}^{-1} & Y_{n2}^{-1} & \dots & Y_{nn}^{-1} \end{bmatrix}$$

are symmetrical; the elements of the last matrices are given by eqns. (6). Furthermore

$$[V'_m] = \begin{bmatrix} V'_1(x) \\ V'_2(x) \\ \vdots \\ V'_n(x) \end{bmatrix}, \quad [I_m] = \begin{bmatrix} I_1(x) \\ I_2(x) \\ \vdots \\ I_n(x) \end{bmatrix}, \quad [E_m^0] = \begin{bmatrix} E_1^0(x) \\ E_2^0(x) \\ \vdots \\ E_n^0(x) \end{bmatrix}$$

where $V'_m(x)$, $I(x)$ and $E_m^0(x)$ are the potential of the secondary electric field, the current and the electric intensity of the impressed electric field, respectively, which are considered along the conductor m .

It is assumed that the impressed electric field is equal to zero along each conductor in the bundle, which frequently occurs in technical applications (such a situation occurs, for instance, if the conductors are energized by the current). Thus $[E_m^0] = 0$ and $[V'_m] = [V_m]$, where $V_m(x)$ is the potential of the conductor m , should be inserted into eqns. (14). Moreover it is assumed that the potential of all the conductors in the bundle is the same and is denoted by $V(x)$.

Hence,

$$[V'_m] = [1] V(x) \quad (15)$$

where $[1]$ is the column matrix, whose all elements are equal to unity.

Thus eqns. (14) take the form

$$[1] \frac{dV}{dx} = -[Z_{mn}][I_n], \quad (16a)$$

$$[Y_{mn}^{-1}] \frac{d}{dx} [I_n] = -[1]V(x). \quad (16b)$$

From eqn. (16a) we obtain

$$[I_n] = -[y_{mn}][1] \frac{dV}{dx} \quad (17)$$

where

$$[y_{mn}] = [Z_{mn}]^{-1}. \quad (18)$$

The total current flowing in the bundle is given by

$$I(x) = [1]^T [I_n] \quad (19)$$

where $[1]^T$ as the transpose of $[1]$ is the row matrix, whose all elements are equal to unity.

Inserting eqn. (17) into eqn. (19) we get

$$\frac{dV}{dx} = -ZI(x) \quad (20)$$

where

$$\frac{1}{Z} = [1]^T [y_{mn}][1] = [1]^T [Z_{mn}]^{-1} [1] \quad (21)$$

giving

$$Z = \left(\sum_{k,l=1}^n y_{kl} \right)^{-1} \quad (22)$$

as the result.

On the ground of eqn. (14b) the equation

$$\frac{dI}{dx} = -YV(x) \quad (23)$$

can be found, where

$$Y = [1]^T [Y_{mn}^{-1}]^{-1} [1] \quad (24)$$

whence

$$Y = \sum_{k,l=1}^n \bar{y}_{kl} \quad (25)$$

and $\bar{y}_{kl}$ is an element of the matrix $[Y_{kl}^{-1}]^{-1}$.

It is easy to show that the differential equations (20) and (23) follow from Fig. 1. Thus the circuit shown in Fig. 1 is the equivalent circuit for a bundle of underground conductors in absence of the impressed external field on the assumption that the potential of all the conductors is the same.

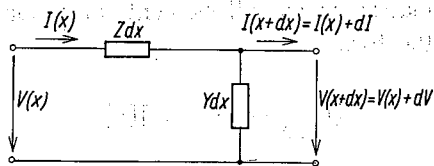


Fig. 1. Equivalent circuit for a bundle of underground conductors

Consider a bundle of underground conductors of the length l energized by the voltage U_0 and loaded by the impedance Z_k . Obviously, U_0 is the potential of the station earth electrode and is impressed by the short-circuit current. The current and potential along the bundle can be derived as the solutions of eqns. (20) and (23) with boundary conditions

$$V(0) = U_0, \quad V(l) = Z_k I(l).$$

As a result, we obtain

$$V(x) = U_0 \frac{\text{sh} \Gamma_0(l-x) + \frac{Z_k}{Z_0} \text{ch} \Gamma_0(l-x)}{\text{sh} \Gamma_0 l + \frac{Z_k}{Z_0} \text{ch} \Gamma_0 l}, \quad (26a)$$

$$I(x) = \frac{U_0}{Z_0} \frac{\text{ch} \Gamma_0(l-x) + \frac{Z_k}{Z_0} \text{sh} \Gamma_0(l-x)}{\text{sh} \Gamma_0 l + \frac{Z_k}{Z_0} \text{ch} \Gamma_0 l}, \quad (26b)$$

where

$$Z_0 = \sqrt{\frac{Z}{Y}}, \quad \text{Re}(Z_0) > 0,$$

$$\Gamma_0 = \sqrt{ZY}, \quad \text{Re}(\Gamma_0) > 0$$

are the characteristic impedance and propagation coefficient of the bundle.

The input impedance of the bundle is given by

$$Z_{in} = \frac{V(0)}{I(0)} = Z_0 \frac{Z_k + Z_0 \text{th} \Gamma_0 l}{Z_0 + Z_k \text{th} \Gamma_0 l}. \quad (27)$$

Eqns. (20) and (23) can be also applied to systems including various segments of different parameters.

4. VOLTAGE BETWEEN SHEATH AND CABLE CONDUCTOR

The current $I_s(x)$ flowing through a cable sheath produces the voltage between sheath and each cable conductor, which is given by [5]

$$U(x) = R_s \int_x^{x_0} I_s(\tau) d\tau \quad (28)$$

where R_s is the unit-length d.c. resistance of the cable sheath. It is assumed that $U(x) = 0$ at $x = x_0$.

Consider a bundle of cables 1, 2, ..., n . The current flowing in the sheath of the cable k can be derived from eqn. (17), giving,

$$I_k(x) = - \sum_{l=1}^n y_{kl} \frac{dV}{dx}$$

and substitution of eqn. (20) yields

$$I_k(x) = Z \sum_{l=1}^n y_{kl} I(x) \quad (29)$$

where $I(x)$ is the total current flowing in the bundle and Z is taken from eqn. (22). The relationships presented concern a bundle which is assumed to be homogeneous along its length. The voltage between sheath and cable conductor k can be computed by eqn. (28) with I_s replaced by I_k from eqn. (29).

5. EXAMPLE

In a system including a pair of cables, which is shown in Fig. 2, a short-circuit occurs on the high voltage side of the station. The potentials of the cable line have been measured in excavations shown in Fig. 2.

A scheme of the system from Fig. 2 is presented in Fig. 3. The unit-length d.c. resistance and leakage conductance of insulation of the cable sheath are $R_s = 0.737 \cdot 10^{-3} \Omega/\text{m}$ and $G_i^{-1} = 0$, respectively. The radius of the cable and depth at which that is buried are $r = 0.057 \text{ m}$ and $d = 0.8 \text{ m}$, respectively. The horizontal distance between both cables $a = 0.35 \text{ m}$, whereas the frequency is of 50 Hz. The remaining data are presented in Fig. 3.

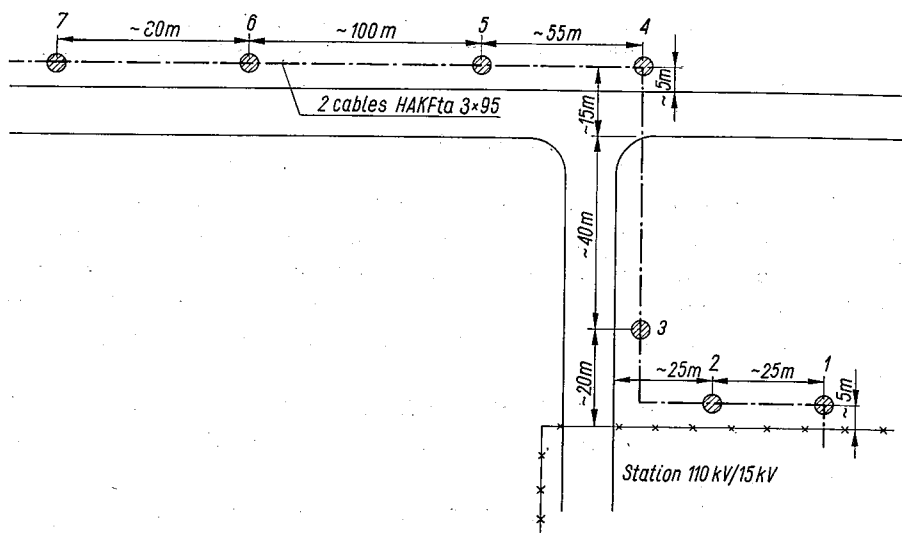


Fig. 2. Cable line 15 kV led out from a station 110 kV/15 kV
The places of excavations for measuring purposes are denoted by 1, 2, ..., 7

The measurements and calculations by computer ODRA-1305 were carried out in Biuro Studiów i Projektów Energetycznych „Energoprojekt” in Poznań. The results of calculations and measurements presented in Fig. 4 are in a rather good agreement.

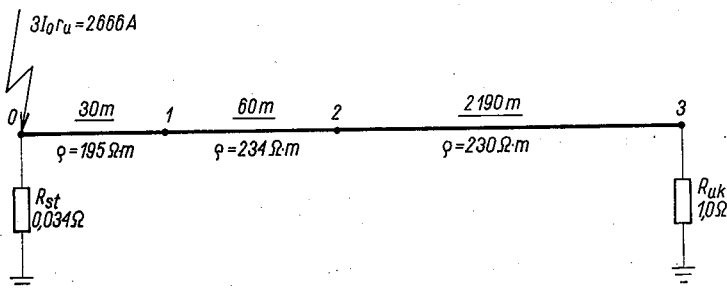


Fig. 3. Scheme of the system shown in Fig. 2

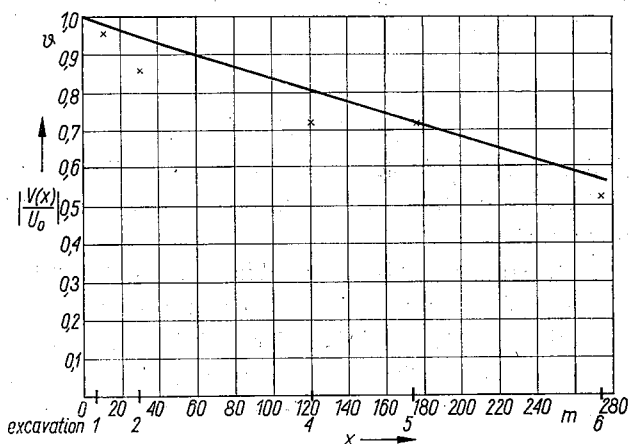


Fig. 4. Potential distribution along the cable line from Fig. 2
The values obtained from experimental data are denoted by crosses

6. CONCLUSIONS

A method of computing of the currents and potentials as well as the voltages between sheath and cable conductor along a bundle of the underground conductors connected to a station earth electrode is presented. The results derived permit also to compute the input impedance of the bundle. The method considered can be also applied to the system including various segments of different parameters. The calculations can be carried out by digital computers.

The considerations are based on idealized assumption that the ground is uniformly conducting. The result of calculation can be improved if the earth conductivity is taken from experimental data.

It is shown that the result of calculations and measurements are in a good agreement. Thus the method presented can be used in many technical problems dealing with the currents and potentials in underground conductors.

REFERENCES

1. K. H. Feist, *Über den Reduktionsfaktor von Starkstromkabeln*, Siemens Zeitschrift, 1965, No. 1, pp. 61—67.
2. K. H. Feist, *Erdung in Industrie- und Kabelnetzen. Trennung von Erdpotentialen*, ETZ-b, 1977, No. 22, pp. 711—713.
3. M. Krakowski, *Currents and potentials along extensive underground conductors*, Proceedings IEE, vol. 115, 1968, No. 9, pp. 1299—1304.
4. M. Krakowski, *Currents and potentials along parallel underground conductors*, Archiwum Elektrotechniki, vol. XX, 1971, No. 4, pp. 937—957.
5. M. Krakowski, A. Jałocha, *Napięcie na izolacji kabla znajdującego się w pobliżu uziomu elektroenergetycznego*, Przegląd Telekomunikacyjny, 1967, No. 4, pp. 119—122.
6. M. Krakowski, *Mathematical study of voltages between sheath and cable conductor due to currents from nearby earth electrodes*, Proceedings IEE, vol. 116, 1969, No. 10, pp. 1701—1705.
7. S. M. Rožavskii, N. A. Korz, *Rosprostranenie potencjalov po jestestvennym zazemlitelam*, Energetika, 1971, No. 4, pp. 3—7.
8. E. D. Sunde, *Earth conduction effects in transmission systems*, D. Van Nostrand, New York 1949.

M. KRAKOWSKI, G. SZYMAŃSKI, A. ZIĘTKOWIAK

ROZKŁAD POTENCJAŁU WZDŁUŻ PRZEWODÓW PODZIEMNYCH
DOŁĄCZONYCH DO UZIOMU STACYJNEGO

Streszczenie

Przedstawiono metodę obliczania prądów, potencjałów i napięć między powłoką i żyłą kabla wzdłuż wiązki przewodów podziemnych dołączonych do uziomu stacyjnego. Podano ogólne równanie macierzo-we dla prądów i potencjałów wzdłuż równoległych przewodów podziemnych oraz schemat zastępczy i równanie różniczkowe. Wyniki obliczeń porównano z danymi eksperymentalnymi.

M. KRAKOWSKI, G. SZYMAŃSKI, A. ZIĘTKOWIAK

RÉPARTITION DU POTENTIEL LE LONG DES CONDUCTEURS
ENTERRÉS RACCORDÉS À LA PRISE DE TERRE DE LA STATION

Résumé

On a présenté la méthode de calcul des courants, des potentiels et des tensions, se produisant entre la gaine et le fil du câble, le long du faisceau des conducteurs enterrés, raccordés à la prise de terre de la station. On a donné l'équation matricielle générale pour les courants et les potentiels le long des conducteurs parallèles enterrés et le schéma de remplacement ainsi que l'équation différentielle. On a comparé les résultats de calcul avec les données expérimentales.

M. KRAKOWSKI, G. SZYMAŃSKI, A. ZIĘTKOWIAK

POTENTIALVERTEILUNG IN ERDVERLEGTEM, AN STATIONSERDER
ANGESCHALTETEN LEITUNGEN

Zusammenfassung

In der Arbeit wurde eine Methode für die Berechnung von Strömen, Potentialen und Spannungen zwischen Kabeladern in einem erdverlegten, an Stationserder angeschalteten Leitungsbündel dargestellt. Es wurde die allgemeine Matrixgleichung für Ströme und Potentiale in erdverlegten Leitungen angegeben, auch ein Ersatzbild und Differentialgleichungen wurden gebracht. Die Berechnungsergebnisse wurden mit experimentellen Daten verglichen.

М. КРАКОВСКИ, Г. ШИМАНЬСКИ, А. ЗЕНТКОВЯК

РАСПРЕДЕЛЕНИЕ ПОТЕНЦИАЛА В ПОДЗЕМНЫХ ПРОВОДАХ
ПРИСОЕДИНЕННЫХ К ЗАЗЕМЛЕНИЮ СТАНЦИИ

Р е з ю м е

Рассмотрен метод вычисления токов, потенциалов и напряжений между жилами и оболочками кабелей в подземном пучке проводов присоединенном к заземлению станции. Приведено общее матричное уравнение для токов и потенциалов в параллельных подземных проводах, а также эквивалентная схема и дифференциальные уравнения. Результаты вычислений сравнены с экспериментальными данными.

Analiza i realizacja aktywnych obwodów z okresowo zmiennymi parametrami

MACIEJ SIWCZYŃSKI (OPOLE), LESŁAW TOPÓR-KAMIŃSKI (GLIWICE)

Instytut Podstawowych Problemów Elektrotechniki Politechniki Śląskiej

Otrzymano 17.7.1979

W pracy przedstawiono sposób interpolacji operatora całkowitego obwodu z okresowo zmiennymi parametrami, ciągiem operatorów dyskretnych. Opisano metodę numeryczną określania jądra operatora dyskretnego za pomocą zamkniętego rezystancyjnego obwodu drabinkowego. Prócz tego rozpatrzono zagadnienie realizacji dwójników czasowo zmiennych w klasie aktywnych sterowanych elementów.

1. WSTĘP

Artykuł ten dotyczy analizy liniowych obwodów elektrycznych o okresowo zmiennych parametrach. Ograniczono się do dwójników, ale podaną metodę można łatwo uogólnić na obwody wielowejściowe.

Jak wiadomo, odpowiedź napięciową u dwójnika na pobudzenie prądowe i określa operator $\hat{z}$:

$$u(t) = \hat{z}i(t) = \int_{-\infty}^{\infty} z(t, r)i(r)dr, \quad (1)$$

gdzie $z(t, r) = 0$ dla $t < r$ jest odpowiedzią napięciową dwójnika na pobudzenie prądowe $\delta(t-r) \in \mathcal{D}'_+$ [16]. Ilustruje to rysunek 1a.

Weźmy pod uwagę dwójnik o okresowo zmiennych parametrach, pobudzany prądem o wspólnym okresie T i założmy, że istnieje ustalona odpowiedź napięciowa o okresie T . Wówczas

$$u(t) = \hat{Z}'i(t) = \int_0^T Z'(t, r)i(r)dr, \quad (2)$$



Rys. 1. Schematy do określenia jąder $z(t, r)$ i $Z'(t, r)$ operatorów impedancji

gdzie $Z'(t, r)$ jest odpowiedzią napięciową dwójnika na pobudzenia prądowe $III(t-r) \in \mathcal{P}'_T$ [16], przy czym

$$Z'(t+T, r+T) = Z'(t, r).$$

Ilustruje to rys. 1b. Nietrudno pokazać, że jądro $Z'(t, r)$ operatora $\hat{Z}'$ jest związane z jądrem $z(t, r)$ operatora $\hat{z}$ wyrażeniem

$$Z'(t, r) = \begin{cases} \sum_{n=1}^{\infty} z(t, r-nT), & \text{gdy } -T < t-r < 0, \\ \sum_{n=0}^{\infty} z(t, r-nT), & \text{gdy } 0 \leq t-r < T. \end{cases} \quad (3)$$

Analiza stanu ustalonego dwójnika pobudzonego okresowo polega na określeniu jądra $Z'(t, r)$ na podstawie równania

$$\hat{i}u = \hat{u}i, \quad (4)$$

gdzie $\hat{i}$, $\hat{u}$ są operatorami różniczkowymi o postaci

$$\begin{aligned} \hat{i}x &= \sum_{n=0}^{N_i} f_n x^{(n)}, \\ \hat{u}x &= \sum_{n=0}^{N_u} g_n x^{(n)}, \end{aligned} \quad (5)$$

a f_n i g_n są funkcjami czasu o okresie T .

Na ogół jądro $Z'(t, r)$ nie da się przedstawić za pomocą formuł elementarnych, dlatego metoda przedstawiona w tym artykule polega na poszukiwaniu wartości funkcji $Z'(t, r)$ na pewnej przeliczalnej siatce płaszczyzny (t, r) .

2. INTERPOLACJA OPERATORA CAŁKOWEGO

Przez $\Pi_T^{(N)}$ oznaczmy przestrzeń funkcji rzeczywistych o okresie T , N -krotnie różniczkowalnych o skończonej energii za okres. Podzielmy okres T na M równych odcinków o długości h (rys. 2).

Przez P_M oznaczmy przestrzeń wszystkich funkcji rzeczywistych, ograniczonych, określonych na zbiorze liczb całkowitych $\mathcal{N}$, z okresem M . W przestrzeniach $\Pi_T^{(N)}$ i P_M okreśmy iloczyny skalarne; dla $x, y \in \Pi_T^{(N)}$ oraz $X, Y \in P_M$, mamy

$$(x, y)_{\Pi_T^{(N)}} = \int_0^T x(t)y(t)dt, \quad (6)$$

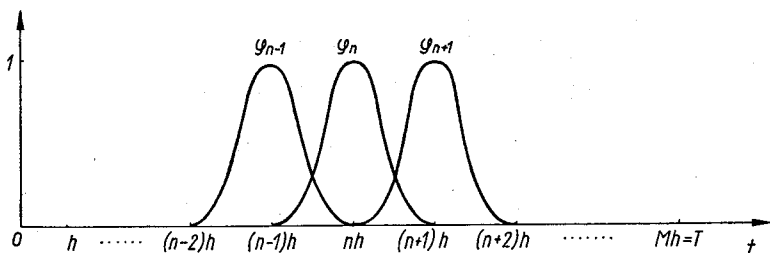
$$(X, Y)_{P_M} = \sum_{n=0}^{M-1} X(n)Y(n). \quad (7)$$

Określmy teraz ciąg $\{\varphi_n\}$ funkcji N -krotnie różniczkowalnych. Kilka elementów tego ciągu pokazano na rys. 2. Poszczególne funkcje φ_n utworzone są z wielomianów argumentu t w następujący sposób

$$\varphi_n(t) = \varphi\left(\frac{t-nh}{h}\right), \quad (8)$$

gdzie

$$\varphi(\vartheta) = \begin{cases} 1 + c_N|\vartheta|^N + c_{N+1}|\vartheta|^{N+1} + \dots + c_{2N-1}|\vartheta|^{2N-1}, & \vartheta \in [-1, 1], \\ 0, & \vartheta \notin [-1, 1]. \end{cases} \quad (9)$$



Rys. 2. Kilka elementów ciągu $\{\varphi_n\}$

Z warunku N -krotnej różniczkowalności funkcji φ wynika układ równań liniowych [12]:

$$\begin{bmatrix} 1 & 1 & \dots & 1 \\ N & N+1 & \dots & 2N-1 \\ N(N-1) & (N+1)N & \dots & (2N-1)(2N-2) \\ N(N-1)(N-2) & (N+1)N(N-1) & \dots & (2N-1)(2N-2)(2N-3) \\ \dots & \dots & \dots & \dots \end{bmatrix} \begin{bmatrix} C_N \\ C_{N+1} \\ C_{N+2} \\ \vdots \\ C_{2N-1} \end{bmatrix} = \begin{bmatrix} -1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad (10)$$

z którego można określić poszczególne współczynniki c występujące w równości (9). Na przykład dla N równego kolejno 1, 2, 3 otrzymuje się

$$\begin{aligned} \varphi(\vartheta) &= \begin{cases} 1 - |\vartheta|, & \text{gdy } \vartheta \in [-1, 1], \\ 0, & \text{gdy } \vartheta \notin [-1, 1], \end{cases} \\ \varphi(\vartheta) &= \begin{cases} 1 - 3\vartheta^2 + 2|\vartheta^3|, & \text{gdy } \vartheta \in [-1, 1], \\ 0, & \text{gdy } \vartheta \notin [-1, 1], \end{cases} \\ \varphi(\vartheta) &= \begin{cases} 1 - 10|\vartheta^3| + 15\vartheta^4 - 6|\vartheta^5|, & \text{gdy } \vartheta \in [-1, 1], \\ 0, & \text{gdy } \vartheta \notin [-1, 1]. \end{cases} \end{aligned}$$

Nietrudno zauważyć, że funkcje φ_n spełniają warunek

$$\bigwedge_{N_1, N_2 \leq N} \int_{-\infty}^{\infty} \varphi_{n_1}^{(N_1)}(t) \varphi_{n_2}^{(N_2)}(t) dt = 0, \quad \text{gdy } |n - m| > 1. \quad (11)$$

Niech $P_M^{(N)}$ będzie podprzestrzenią przestrzeni $\Pi_T^{(N)}$, utworzoną ze wszystkich elementów postaci

$$x = \sum_{n \in \mathcal{N}} X(n) \varphi_n \quad (12)$$

gdzie $X \in P_M$. Ze względu na własności funkcji φ_n , liczby $X(n)$ są wartościami funkcji x w węzłach aproksymacji, to znaczy

$$X(n) = x(nh). \quad (13)$$

Widać, że dla ustalonego N pomiędzy elementami przestrzeni P_M i $P_M^{(N)}$ istnieje odpowiedniość wzajemnie jednoznaczna. Tym samym zostaje określone wzajemne jednoznaczne odwzorowanie

$$\chi: P_M \rightarrow P_M^{(N)}.$$

W rozdziale 1 wspomniano, że określenie operatora $\hat{Z}$ wymaga odwrócenie operatora $\hat{i}$. Dokonamy tego z zastosowaniem metody Ritza.

Utwórzmy funkcjonał F nieujemnie określony na $\Pi_T^{(N)}$ taki, że

$$F(x) = (\hat{i}x - \hat{u}i, \hat{i}x - \hat{u}i)_{\Pi_T^{(N)}} = (\hat{i}^\dagger \hat{i}x, x)_{\Pi_T^{(N)}} - 2(\hat{i}^\dagger \hat{u}i, x)_{\Pi_T^{(N)}} + (\hat{u}i, \hat{u}i)_{\Pi_T^{(N)}}, \quad (14)$$

gdzie i jest ustalonym elementem przestrzeni $\Pi_T^{(N)}$, $N = N_T^{-1}$, a $\hat{i}^\dagger$ oznacza operator sprzężony względem $\hat{i}$.

Funkcjonał ten ma minimum równe zeru dla elementu $u \in \Pi_T^{(N)}$ spełniającego równanie

$$\hat{i}^\dagger \hat{u}i - \hat{i}^\dagger \hat{u}i = 0,$$

które jest równoważne równaniu (4), gdyż operacja $\hat{i}^\dagger$ jest odwracalna.

Utwórzmy dalej ciąg $\{P_M^{(N)}\}$ (N ustalone) podprzestrzeni przestrzeni $\Pi_T^{(N)}$ izomorficzny z ciągiem $\{P_M\}$ (izomorfizm χ). W każdej z przestrzeni P_M znajdujemy punkt U_M , w którym funkcjonał $F \circ \chi: P_M \rightarrow \mathcal{R}_+$ przyjmuje wartość minimalną, przy czym

$$F \circ \chi(U_M) > 0.$$

Jeżeli ciąg $\{F \circ \chi(U_M)\}$ wartości funkcjonału $F \circ \chi$ w punktach U_M dąży do zera, to ciąg $\{\chi^{-1}(U_M)\}$ jest zbieżny do rozwiązania $u \in \Pi_T^{(N)}$ równania (4). Funkcjonał $F \circ \chi: P_M \rightarrow \mathcal{R}_+$ jest dany wyrażeniem

$$F \circ \chi(X) = \sum_{n,m=0}^M a(n,m)X(n)X(m) - 2 \sum_{n,m=0}^M b(n,m)X(n)I(m) + F_0, \quad (15)$$

gdzie:

$$a(n,m) = (\hat{i}^\dagger \hat{i} \varphi_n, \varphi_m)_{P_M^{(N)}}, \quad (16)$$

$$b(n,m) = (\varphi_n, \hat{i}^\dagger \hat{u} \varphi_m)_{P_M^{(N)}}, \quad (17)$$

a F_0 jest stałą liczbą.

Z wyrażeń (11), (16) i (17) wynika, że macierz $a(n,m)$ jest symetryczna i trójdzielna ($a(n,m) = 0$ dla $|n-m| > 1$), natomiast macierz $b(n,m)$ jest trójdzielna. Funkcjonał $F \circ \chi$ jest dodatnio określony na P_M , a jego minimum przypada w punkcie $U \in P_M$ spełniającym równanie

$$\sum_{m=0}^M a(n,m)U(m) = \sum_{m=0}^M b(n,m)I(m) = J(n). \quad (18)$$

¹⁾ Należy założyć, że $N_T > N_a$, w przeciwnym razie $z(t,r)$ nie jest funkcją, lecz dystrybucją nieregularną.

Funkcje U , I są okresowe, zachodzą więc równości $U(0) = U(M)$, $I(0) = I(M)$, dzięki którym układ równań (18) można „zwinąć”, czyli dodać ostatnie równanie do pierwszego. Po zwinieniu układu równań (18) otrzymamy równanie operatorowe w P_M

$$\hat{I}U = \hat{U}I = J'; \quad (19)$$

operatory $\hat{I}, \hat{U}: P_M \rightarrow P_M$ mają postać

$$\hat{I}X(n) = \sum_{m=0}^{M-1} a'(n, m)X(m), \quad (20)$$

$$\hat{U}X(n) = \sum_{m=0}^{M-1} b'(n, m)X(m), \quad (21)$$

a ich jądra są następującymi macierzami

$$a'(n, m) = \begin{bmatrix} a(0, 0) + a(M, M), & a(0, 1), & 0, & \dots & 0, & a(M, M-1) \\ a(n, m) & \text{dla } 2 \leq n \leq M-2, & 0 \leq m \leq M-1 \\ a(M-1, M), & 0, & \dots, & a(M-1, M-2), & a(M-1, M-1) \end{bmatrix}, \quad (22)$$

$$b'(n, m) = \begin{bmatrix} b(0, 0) + b(M, M), & b(0, 1), & 0, & \dots, & 0, & b(M, M-1) \\ b(n, m) & \text{dla } 2 \leq n \leq M-2, & 0 \leq m \leq M-1 \\ b(M-1, M), & 0, & \dots, & b(M-1, M-2), & b(M-1, M-1) \end{bmatrix}. \quad (23)$$

Równania (18) i (19) mają prostą interpretację elektryczną. Układ (18) jest analogiczny do równań drabinki rezystancyjnej pokazanej na rys. 3. Potencjały węzłów tej drabinki równe są odpowiednio $U(0), U(1), \dots, U(M)$, a źródła prądowe $J(0), J(1), \dots, J(M)$.

Jeżeli „skleimy” początek z końcem drabinki widocznej na rys. 3, to otrzymamy drabinkę zwiniętą opisaną równaniem (19) (rys. 4 i 5).

Źródła prądowe drabinki zwiniętej określa element $J' \in P_M$:

$$J'(n) = \begin{cases} J(0) + J(M), & \text{gdy } n = 0, \\ J(n), & \text{gdy } 1 \leq n \leq M-1. \end{cases} \quad (24)$$

Oznaczając przez $\hat{R}$ operator rozwiązujący drabinkę zwiniętą, zapiszemy rozwiązanie równania (19) w postaci

$$U(n) = \hat{I}^{-1}J'(n) = \hat{R}J'(n) = \sum_{m=0}^{M-1} r(n, m)J'(m). \quad (25)$$

Macierz $r(n, m)$ będącą jądrem operatora $\hat{R}$ można łatwo wyznaczyć poprzez przyłączenie do drabinki zwiniętej jednostkowych źródeł próbnych. Poszczególne potencjały drabinki będą wówczas równe elementom macierzy $r(n, m)$ (rys. 6).

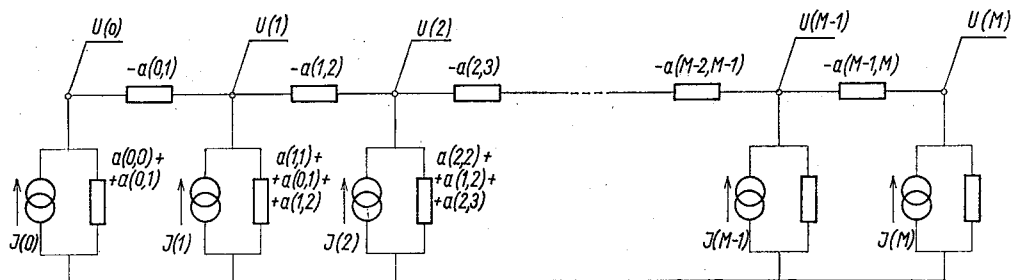
Z równości (19) i (25) określa się operator impedancji dwójnika

$$U(n) = \hat{I}^{-1}\hat{U}I(n) = \hat{Z}I(n) = \sum_{m=0}^{M-1} Z(n, m)I(m), \quad (26)$$

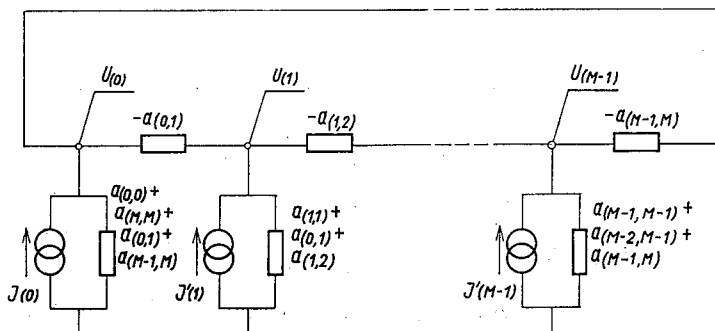
którego jądro jest określone wyrażeniem

$$Z(n, m) = \sum_{k=0}^{M-1} r(n, k)b'(k, m). \quad (27)$$

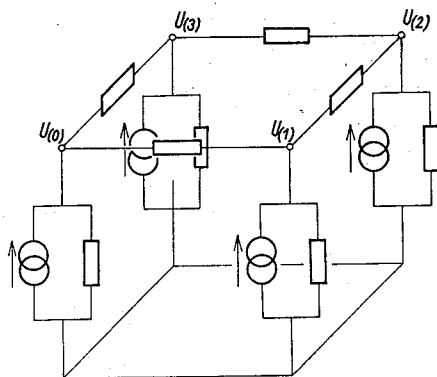
Funkcja Z określona wzorem (27) jest okresowa względem n i m . Pozwala ona obliczyć przybliżone wartości jądra $Z'(t, r)$ na przeliczalnej siatce płaszczyzny (t, r) .



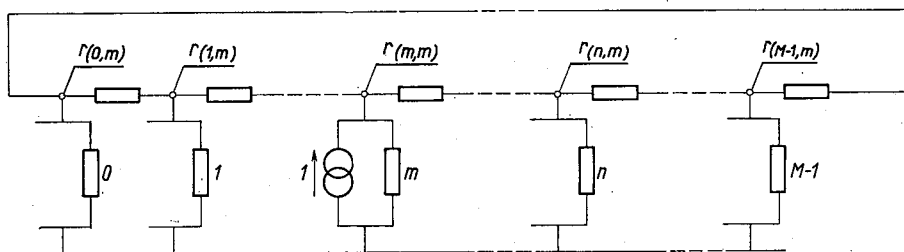
Rys. 3. Drabinka rezystancyjna odpowiadająca układowi równań (18)



Rys. 4. Drabinka zwinięta odpowiadająca równaniu (19)



Rys. 5. Przykład drabinki zwiniętej odpowiadającej równaniu (19)



Rys. 6. Schemat drabinki do wyznaczania macierzy $r(n, m)$

3. ZBIEŻNOŚĆ I DOKŁADNOŚĆ METODY

Do celów badania zbieżności procedury, bez utraty ogólności można przyjąć, że operator $\hat{u}$ w równaniu (4) jest tożsamościowy. Wówczas równanie (4) przyjmie postać

$$\hat{ix} = i. \quad (28)$$

W tym przypadku funkcjonal (14) będzie miał postać

$$F(x) = (\hat{ix} - i, \hat{ix} - i)_{\Pi_T^{(N)}} = (\hat{i}^T \hat{ix}, x)_{\Pi_T^{(N)}} - 2(\hat{i}^T i, x)_{\Pi_T^{(N)}} + (i, i)_{\Pi_T^{(N)}}. \quad (29)$$

Funkcjonał $F \circ \chi$ określony na P_M przyjmuje postać

$$F \circ \chi(X) = \sum_{n,m=0}^M a(n, m) X(n) X(m) - 2 \sum_{n,m=0}^M b(n, m) X(n) I(m) + \sum_{n,m=0}^M c(n, m) I(n) I(m), \quad (30)$$

gdzie

$$c(n, m) = (\varphi_n, \varphi_m)_{F_M^{(N)}}. \quad (31)$$

Funkcjonał (30) można zapisać w postaci

$$F \circ \chi(X) = (\hat{IX}, X)_{P_M} - 2(X, J')_{P_M} + (\hat{CI}, I)_{P_M}, \quad (32)$$

gdzie J' jest elementem przestrzeni P_M zdefiniowanym równością (19) (zob. też (24)), $\hat{C}$ jest operatorem o jądrze określonym macierzą $c'(n, m)$

$$c'(n, m) = \begin{bmatrix} c(0, 0) + c(M, M), & c(0, 1), & 0, & \dots, & 0, & c(M, M-1) \\ c(n, m) & \text{dla} & 2 \leq n \leq M-2, & 0 \leq m \leq M-1 \\ c(M-1, M), & 0, & \dots, & c(M-1, M-2), & c(M-1, M-1) \end{bmatrix},$$

a I jest ustalonym elementem przestrzeni P_M .

Zwiększając wymiar M tworzymy ciąg przestrzeni P_M . W każdej przestrzeni P_M tego ciągu istnieje element U_M minimalizujący funkcjonal (32). Zbieżność procedury wymaga, aby ciąg liczbowy $\{F \circ \chi(U_M)\}$ dążył do zera.

Jak wiadomo, element $U_M \in P_M$ minimalizujący funkcjonal (32) jest rozwiązaniem drabinki rezystancyjnej ze źródłami $J' \in P_M$, pokazanej na rys. 4. Bilans mocy tej drabinki wymaga, aby

$$S_M = (U_M, J')_{P_M} = (\hat{IU}_M, U_M)_{P_M}. \quad (33)$$

Liczba S_M jest mocą dyssypacji drabinki pokazanej na rys. 4. Z wyrażeń (32) i (33) wynika, że

$$F \circ \chi(U_M) = (\hat{CI}, I)_{P_M} - S_M = S_M^c - S_M. \quad (34)$$

Liczba S_M^c równa jest mocy dyssypacji drabinki zwiniętej o macierzy $c'(n, m)$ i potencjałach $I \in P_M$.

Z wyrażenia (34) wynika więc, że zbieżność procedury wymaga, aby

$$S_M^c - S_M \rightarrow 0. \quad (35)$$

Warunek (35) ma następującą interpretację. Zwiększając M tworzymy dwa ciągi drabinek o operatorach odpowiednio $\hat{I}$, $\hat{C}$ oraz potencjałach U_M , I . Każdej drabince wchodzącej w skład tych ciągów odpowiada moc dyssypacji S_M oraz S_M^c . Procedura interpolacji ope-

ratora impedancji dwójnika jest zbieżna, gdy ciąg $\{S_M^c - S_M\}$ różnic mocy dyssypacji jest zbieżny do zera.

Wartość funkcjonału $F \circ \chi$ w punkcie $U_M \in P_M$ określona wzorem (34) może być wskaźnikiem dokładności interpolacji i służyć do oceny błędu.

4. ALGORYTM ANALIZY DRABINKI ZWINIĘTEJ

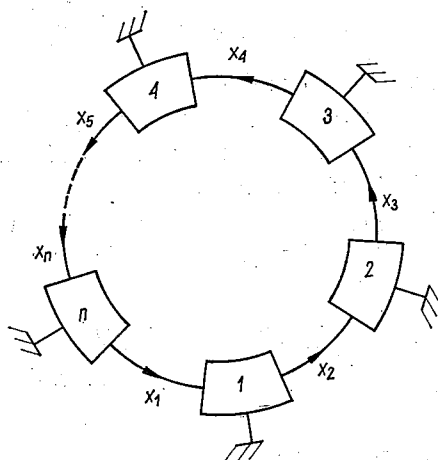
W rozdziale 2 pokazano, że analiza stanu ustalonego dwójnika z okresowo zmiennymi parametrami sprowadza się do analizy zwiniętej drabinki rezystancyjnej. Drabinkę taką można traktować jak zamknięty łańcuch czwórników aktywnych złożonych z rezystorów i źródeł niesterowanych. Łańcuch taki pokazano na rys. 7. Mimo, że znanych jest wiele algorytmów analizy aktywnych obwodów drabinkowych, podamy tutaj jeszcze jeden, szczególnie przydatny do analizy drabinek zwiniętych.

Każdy czwórnik widoczny na rys. 8, zawierający źródła niesterowane, można opisać równaniem

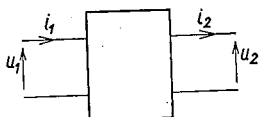
$$\begin{bmatrix} u_1 \\ i_1 \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \begin{bmatrix} u_2 \\ i_2 \end{bmatrix} + \begin{bmatrix} u_{10} \\ i_{10} \end{bmatrix}. \quad (36)$$

Macierz $[u_{10}, i_{10}]^T$ obserwujemy na wejściu czwórnika wówczas, gdy do wyjścia zostanie dołączony nulator (rys. 9). Oznaczając macierze kolumnowe w równości (36) odpowiednio przez x_1, x, s , a macierz kwadratową przez A , można napisać

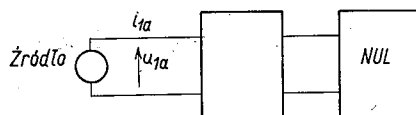
$$x_1 = Ax + s. \quad (37)$$



Rys. 7. Zamknięty łańcuch czwórników służący do analizy drabinki zwiniętej



Rys. 8. Schemat czwórnika aktywnego opisanego równaniem (36)



Rys. 9. Schemat do wyznaczenia wektora (u_{10}, i_{10}) czwórnika aktywnego

Równość (37) można traktować jako odwzorowanie przekształcające dwuwymiarowy wektor kolumnowy x w inny wektor kolumnowy x_1 . Odwzorowanie takie jest scharakteryzowane parą (A, s) . Fakt ten zanotujemy następująco

$$(A, s)x = Ax + s. \quad (38)$$

Element A pary nazwiemy macierzą, a element s — źródłem.

Zbiór takich par staje się przestrzenią liniową, jeżeli zdefiniować dodawanie par i mnożenie ich przez liczbę α :

$$\begin{aligned} (A_1, s_1) + (A_2, s_2) &= (A_1 + A_2, s_1 + s_2), \\ \alpha(A, s) &= (\alpha A, \alpha s). \end{aligned} \quad (39)$$

Złożenie odwzorowań (A_1, s_1) i (A_2, s_2) daje

$$(A_2, s_2) [(A_1, s_1)x] = (A_2, s_2) [A_1 x + s_1] = A_2 A_1 x + A_2 s_1 + s_2. \quad (40)$$

Z wyrażenia (40) wynika sposób określenia mnożenia w zbiorze par

$$(A_2, s_2)(A_1, s_1) = (A_2 A_1, A_2 s_1 + s_2). \quad (41)$$

Nietrudno sprawdzić, że mnożenie par jest łączne i że istnieje element neutralny $(E, 0)$, gdzie E jest macierzą jednostkową, a 0 jest zerowym dwuwymiarowym wektorem kolumnowym. Zatem zbiór par tworzy półgrupę. Jeżeli ponadto macierz każdej z par zbioru jest nieosobliwa, to zbiór par tworzy grupę z elementem odwrotnym

$$(A, s)^{-1} = (A^{-1}, A^{-1}s).$$

Można też przekonać się, że mnożenie par nie jest rozdzielne z dodawaniem, co oznacza, że zbiór par nie jest pierścieniem.

Mnożenie par możemy wykorzystać do określenia napięcia i prądu w dowolnym miejscu zamkniętego łańcucha czwórników widocznego na rys. 7. Przypuśćmy, że chcemy obliczyć wektor $x_1 = [u_1, i_1]^T$ na wejściu czwórnika 1. Dokonując rozcięcia łańcucha zamkniętego między czwórnikami n i 1, obliczamy parę łańcucha otwartego ze wzoru

$$(A, s)^1 \equiv (A^1, s^1) = (A, s)_1 \dots (A, s)_n \equiv \prod_{k=1}^n (A, s)_k, \quad (42)$$

gdzie $(A, s)_k$ oznacza parę k -tego czwórnika.

Dla łańcucha zamkniętego wektor x_1 jest punktem stałym odwzorowania (A^1, s^1) (rys. 10), czyli

$$(A^1, s^1)x_1 = x_1,$$

skąd

$$x_1 = (E - A^1)^{-1} s^1. \quad (43)$$

Postępując analogicznie można określić wektor x_m w dowolnym miejscu łańcucha. Wektor ten jest punktem stałym odwzorowania

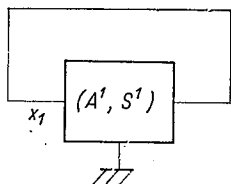
$$(A, s)^m = (A, s)_m (A, s)_{m+1} \dots (A, s)_{m-1} = \prod_{k=m}^{m-1} (A, s)_k. \quad (44)$$

Przyjmując $n = M-1$ można za pomocą powyższego algorytmu obliczyć napięcie $V(m)$ oraz prąd i_m w dowolnym miejscu drabinki zwiniętej i tym samym znaleźć rozkład napięcia i prądu w drabince.

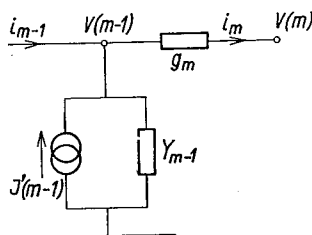
Algorytm taki wymaga wielokrotnego użycia wzoru (44) dla kolejnych wartości m , co jest bardzo pracochłonne. W celu uproszczenia procedury proponuje się, aby wzór (44) zastosować jednokrotnie do określenia dwóch wybranych wartości $V(m)$ i i_m , a pozostałe napięcia i prądy drabinki wyznaczyć ze wzorów rekurencyjnych

$$\begin{aligned} V(m-1) &= V(m) + \frac{1}{g_m} i_m, \\ i_{m-1} &= i_m + y_{m-1} V(m-1) - J'(m-1). \end{aligned} \quad (45)$$

Wzory te wynikają bezpośrednio z rys. 11.



Rys. 10. Schemat do wyznaczenia punktu stałego odwzorowania (A^1, S^1)



Rys. 11. Ogniwo drabinki zwiniętej

5. REALIZACJA ELEKTRONICZNA OBWODÓW O ZMIENNYCH PARAMETRACH

Z równania (4) wynika, że każdy dwójnik jest scharakteryzowany parą operatorów $\{\hat{u}, \hat{i}\}$. W zbiorze takich operatorów można w zwykły sposób określić działania dodawania i mnożenia

$$\begin{aligned} \hat{u}_1 + \hat{u}_2: \bigwedge_i (\hat{u}_1 + \hat{u}_2)i &= \hat{u}_1 i + \hat{u}_2 i, \\ \hat{u}_1 \hat{u}_2: \bigwedge_i \hat{u}_1 \hat{u}_2 i &= \hat{u}_1 (\hat{u}_2 i). \end{aligned}$$

Z powyższych określeń wynikają następujące własności zbioru operatorów.

Własność 1. Zbiór operatorów tworzy grupę abelową ze względu na działanie dodawania oraz co najmniej półgrupę ze względu na działanie mnożenia, zatem zbiór operatorów tworzy pierścień.

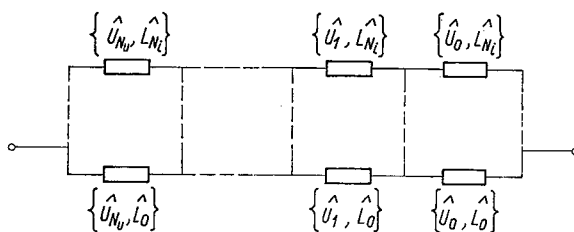
Własność 2. Dodawanie operatorów wiąże się z łączeniem dwójników. Dwójnik $\{\hat{u}_1 + \hat{u}_2, \hat{i}\}$ jest równoważny szeregowemu połączeniu dwójników $\{\hat{u}_1, \hat{i}\}$, $\{\hat{u}_2, \hat{i}\}$, a dwójnik $\{\hat{u}, \hat{i}_1 + \hat{i}_2\}$ jest równoważny równoległemu połączeniu dwójników $\{\hat{u}, \hat{i}_1\}$, $\{\hat{u}, \hat{i}_2\}$.

Z własności 2 wynika metoda modelowania dwójnika poprzez rozkład na szeregowo-równoległe połączenie dwójników prostszych. Rozkład taki przeprowadza się w celu otrzymania sieci (zwanej dalej podstawową) w kształcie kratownicy (rys. 12) zawierającej $(N_{\hat{u}} + 1) \times (N_{\hat{i}} + 1)$ dwójników elementarnych. Elementarnym nazywamy dwójnik opisany operatorami jednomianowymi, to znaczy

$$f_n u^{(n)} = g_m i^{(m)}.$$

Każde przekształcenie sieci podstawowej nie zmieniające pary $\{\hat{u}, \hat{i}\}$ operatorów dwójnika nazywamy dopuszczalnym. Z własności 1 i 2 wynikają następujące przekształcenia dopuszczalne:

- zamiana miejscami dwóch dowolnych wierszy sieci podstawowej,
- zamiana miejscami dwóch dowolnych kolumn sieci podstawowej,
- usunięcie dowolnego wewnętrznego połączenia między dwoma wierszami sieci podstawowej.

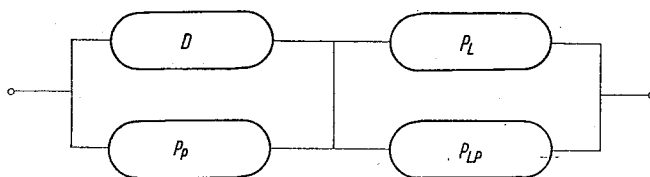


Rys. 12. Schemat sieci podstawowej

Określony w ten sposób zbiór przekształceń sieci podstawowej pozwala tworzyć sieci pochodne, wygodniejsze do praktycznej realizacji. Często w praktyce tylko niektóre współczynniki operatorów (5) zależą od czasu. Umożliwia to rozdzielenie dwójnika na części czasowo niezmienną i czasowo zmienną.

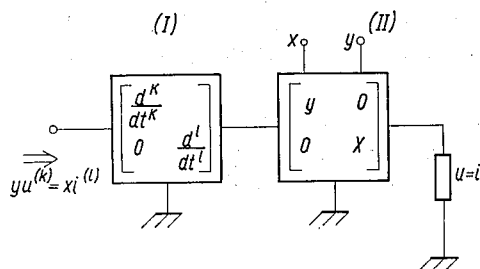
Rozpatrzmy dwójnik opisany równaniem typu (4), w którym tylko dwa współczynniki $f_k(t)$ i $g_l(t)$ są czasowo zmienne. Zatem w sieci podstawowej współczynniki zależne od czasu zawierać będą dwójniki w k -tej kolumnie i l -tym wierszu.

Stosując wyżej podane przekształcenia dopuszczalne można dokonać podziału sieci podstawowej na cztery dwójniki (rys. 13). Jeden z nich jest czasowo niezmienny (D), dwa jednostronnie czasowo zmienne (P_L i P_P), a jeden dwustronnie czasowo zmienny. W dwójniku jednostronnie czasowo zmiennym współczynniki tylko jednego z operatorów $\hat{u}$ lub $\hat{i}$ zależą od czasu.

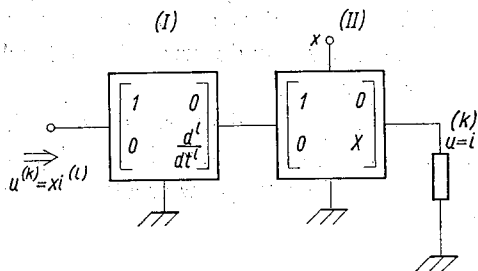


Rys. 13. Sposób podziału sieci podstawowej na części czasowo niezmienną i czasowo zmienną

Ostateczne realizacje elektroniczne dwójników jednostronnie i dwustronnie czasowo zmiennych pokazano na rysunkach 14 i 15. Blok (I) może być mutatorem [4] lub uogólnionym konwertorem impedancyjnym [2], [3], natomiast blok (II) można zrealizować w postaci konwertora impedancyjnego sterowanego [14] lub inwertora impedancyjnego sterowanego [15].



Rys. 14. Realizacja elektroniczna dwójnika jednostronnie czasowo zmiennego



Rys. 15. Realizacja elektroniczna dwójnika dwustronnie czasowo zmiennego

LITERATURA

1. M. Bogucki, *Analiza i synteza pewnej klasy liniowych obwodów zmiennych okresowo*, Instytut Podstaw Elektrotechniki i Elektrotechnologii, Politechnika Wrocławska, Komunikat K-141, 1978.
2. L. T. Bruton, *Non-ideal performance of a class of two-amplifier positive impedance converters*, IEEE Trans. Circuit Theory, Nov. 1969.
3. L. T. Bruton, *Biquadratic Sections using Generalized Impedance Converters*, The Radio and Electronic Engineer, Nov. 1971.
4. L. Q. Chua, *Synthesis of New Nonlinear Network Elements*, Proc. IEEE, August 1968.
5. R. G. Cooke, *Infinite Matrices and Sequence Spaces*, Macmillan and Company Ltd., London 1950.
6. J. B. Cruz, *Sensitivity considerations for time-varying sampled-data feedback systems*, IRE Trans. on Automatic Control, vol. 6, pp. 228—236, May 1961.
7. B. Friedland, *A technique for the analysis of time-varying sampled-data systems*, Trans AIEE, vol. 75, Applications and Industry, part II, pp. 407—413, January 1957.
8. S. K. Godunow, W. S. Riabenkij, *Raznostnyje schemy*, Nauka, Moskwa 1977.
9. A. W. Naylor, *Generalized Frequency Response Concepts for Time-Varying, Discrete-Time Linear Systems*, IEEE Trans. on Circuit Theory, pp. 428—440, September 1963.
10. A. A. Samarskij, E. S. Nikołajew, *Metody reszenija sietocznych urawnienij*, Nauka, Moskwa 1978.
11. I. W. Sandberg, *Realization of Class of Periodically Variable Systems*, IRE Trans. on Circuit Theory, v. CT-9, Dec. 1962, pp. 416—447.
12. M. Siwczyński, L. Topór-Kamiński, *Analiza obwodów o okresowo zmiennych parametrach metodą wariacyjną*, Prace Naukowe Instytutu Telekomunikacji i Akustyki Politechniki Wrocławskiej, No 40, Konferencje No 10, Wrocław 1978.
13. G. Strang, G. J. Fix, *An Analysis of the Finite Element Method*, Prentice-Hall, Inc., 1973.
14. L. Topór-Kamiński, *Konwertor impedancyjny sterowany*, Zeszyty Naukowe Politechniki Śląskiej Elektryka, z. 60, 1978.
15. L. Topór-Kamiński, *Inwertor impedancyjny sterowany*, Zeszyty Naukowe Politechniki Śląskiej Automatyka, w druku.
16. A. H. Zemanian, *Teoria dystrybucji i analiza transformat*, PWN, Warszawa 1969.

M. SIWCZYŃSKI, L. TOPÓR-KAMIŃSKI

ANALYSIS AND REALIZATION OF LINEAR PERIODICALLY TIME-VARYING NETWORKS

Summary

The paper contains an analysis of linear periodically time-varying networks by digital interpolation of integral network's operator, using a ring chain resistive model. Moreover, the problem of time-varying networks realization in class of active electronics elements is described.

M. SIWCZYŃSKI, L. TOPÓR-KAMIŃSKI

ANALYSE ET SIMULATION DES CIRCUITS ACTIFS
À PARAMÈTRES PÉRIODIQUEMENT VARIABLES

Résumé

Dans l'article est présentée une méthode d'analyse des circuits actifs périodiquement variables, par l'interpolation de l'opérateur de circuit avec l'usage de circuit résistif en échelle. En outre est décrit le problème de simulation des circuits à l'aide des éléments électroniques actifs.

M. SIWCZYŃSKI, L. TOPÓR-KAMIŃSKI

ANALYSE UND REALISATION LINEARER KREISE
MIT ZAITVERÄNDERLICHEN PARAMETERN

Zusammenfassung

In der Arbeit wird die Methode der Interpolation des integralen Operators des Kreises mit periodisch veränderlichen Parametern dargelegt. Diese Interpolation wird mittels entsprechender Folge der digitalen Operatoren umschrieben. Eine numerische Methode für die Kernbestimmung des digitalen Operators wird angegeben. Außerdem betrachtet man das Problem eines parametrischen Zweipols auf Basis aktiver gesteuerter Elemente.

М. СИВЧЫŃСКИ, Л. ТОПУР-КАМИŃСКИ

АНАЛИЗ И РЕАЛИЗАЦИЯ ЛИНЕЙНЫХ ЦЕПЕЙ
С ПЕРИОДИЧЕСКИ ИЗМЕНЯЮЩИМИСЯ ПАРАМЕТРАМИ

Резюме

Интегральный оператор линейного периодически — параметрического двухполосника интерполирован последовательностью импульсных операторов. Представлен численный метод определения ядра импульсного оператора с помощью замкнутой резистивной лестничной цепи. Кроме того рассматривается задача реализации параметрических двухполосников в классе активных управляемых элементов.

Odwrotne zagadnienia brzegowe typu Dirichleta dla równania Laplace'a

KAZIMIERZ ADAMIAK (KIELCE)

Instytut Elektrotechniki Politechniki Świętokrzyskiej

Otrzymano 10.1.1980

W pracy podano metodę syntezy pól potencjalnych. Zadany rozkład pola w pewnym obszarze uzyskuje się dobierając w odpowiedni sposób wartości potencjału pola na założonej powierzchni. Zagadnienie sprowadza się do rozwiązania całkowitego równania Fredholma I rodzaju lub układu równań całkowitych. Zaprezentowano wyniki obliczeń numerycznych kilku przykładów w układzie płaskorównoległym. Obliczenia wykonano metodą regularyzacji oraz metodami iteracyjnymi.

1. WSTĘP

Najczęściej rozpatrywanymi w teorii pola wewnętrznymi zagadnieniami brzegowymi dla równania Laplace'a są zagadnienia proste, nazywane również zagadnieniami analizy. Polegają one na poszukiwaniu w obszarze D rozwiązania równania Laplace'a

$$\nabla^2 V = 0 \quad (1)$$

spełniającego na brzegu Γ obszaru D określone warunki. Są to warunki brzegowe:

— I rodzaju (Dirichleta)

$$V|_{\Gamma} = f(x), \quad x \in \Gamma, \quad (2a)$$

— II rodzaju (Neumanna)

$$\left. \frac{\partial V}{\partial n} \right|_{\Gamma} = g(x), \quad x \in \Gamma, \quad (2b)$$

— III rodzaju

$$\frac{\partial V}{\partial n} + p(x) V|_{\Gamma} = h(x), \quad x \in \Gamma. \quad (2c)$$

Niekiedy rozpatruje się też mieszane warunki brzegowe

$$\begin{aligned} V|_{\Gamma_1} &= f(x), \quad x \in \Gamma_1, \\ \left. \frac{\partial V}{\partial n} \right|_{\Gamma_2} &= g(x), \quad x \in \Gamma_2, \\ \Gamma &= \Gamma_1 \cup \Gamma_2. \end{aligned} \quad (2d)$$

Istnieje szereg metod pozwalających rozwiązywać zagadnienie proste [1].

Można również rozpatrywać zagadnienie odwrotne dla równania Laplace'a. Zadać mianowicie w obszarze D_0 rozwiązanie, a poszukiwać warunków brzegowych (I, II lub III rodzaju) na brzegu obszaru D , $D_0 \subset D$. Zagadnienie takie jest też nazywane syntezą. Odmianą zagadnienia odwrotnego jest poszukiwanie brzegu obszaru D , na którym zadane zostały warunki brzegowe, tak aby rozwiązanie było równe założonemu. Zagadnienia odwrotne były tematem nielicznych prac [2]. Częściej rozwiązywano problemy pokrewne, a mianowicie zadania syntezy pól magnetycznych [3, 4]. Polegają one na wyznaczeniu, dla zadanego rozwiązania, prawej strony równania Poissona.

Praktyczne znaczenie zagadnień odwrotnych jest bardzo duże. Konstrukcji wytwarzających pola o zadanym rozkładzie poszukuje się w szeregu dziedzin techniki [5, 6] i fizyki stosowanej [2].

Van der Merve [2] zastosował do rozwiązywania zagadnienia odwrotnego dla równania Laplace'a w układzie cylindrycznym szeregi Fouriera-Bessela. Metoda taka pozwoliła dla zadanego rozwiązania na osi znaleźć rozwiązanie w jej otoczeniu. Uzyskane wyniki mają duże znaczenie w optyce elektronowej. Metoda Fouriera-Bessela posiada istotną wadę. Mianowicie pozwala przeprowadzać syntezę za pomocą warunków brzegowych w obszarach nieograniczonych. W praktyce częściej przeprowadza się syntezę pola w obszarach ograniczonych. Metoda opisana w pracy pozwoli rozwiązać zagadnienie odwrotne przy założeniu, że brzeg Γ obszaru D , na którym zadaje się warunki brzegowe, jest ograniczony.

2. ODWROTNE ZAGADNIENIE DIRICHLETA DLA RÓWNIANIA LAPLACE'A

Niech dany będzie obszar $D_0 \subset R^n$, oraz obszar $D \subset R^n$, $D_0 \subset D$, ograniczony brzegiem Σ . Rozwiązaniem równania (1) w D_0 jest potencjał warstwy pojedynczej [7]

$$V(P) = \int_{\Sigma} \mu(Q) E(P, Q) dS, \quad P \in D_0, \quad Q \in \Sigma, \quad (3)$$

gdzie [8]:

$$E(P, Q) = \frac{1}{r_{PQ}^{n-2}} \quad \text{dla} \quad n > 2,$$

$$E(P, Q) = \ln \frac{1}{r_{PQ}} \quad \text{dla} \quad n = 2.$$

Jeśli znana jest gęstość powierzchniowa $\mu(Q)$, równanie (3) pozwala rozwiązać zagadnienie proste dla równania Laplace'a w obszarze D_0 . Traktując natomiast $V(P)$ jako daną, a poszukując $\mu(Q)$ rozwiązuje się zagadnienie odwrotne. Po znalezieniu $\mu(Q)$ znaleźć można rozkład potencjału na powierzchni Σ . W tym celu należy powtórnie wykorzystać wzór (3), lecz biorąc $P \in \Sigma$. Otrzymuje się więc, dla zadanego rozwiązania równania Laplace'a, warunek brzegowy na powierzchni Σ . A więc w ten sposób rozwiązane zostało odwrotne zagadnienie Dirichleta.

W szczególności dla układu płaskorównoległego, gdy obszar D_0 leży na płaszczyźnie $y = 0$

$$D_0 = \{(x, y); \quad y = 0 \wedge x \in (-a, a)\},$$

a brzeg Σ obszaru D jest dany zależnością

$$\Sigma = \{y; y = f(x), x \in (-a, a)\},$$

wzór (3) można sprowadzić do

$$V(x) = \int_{-a}^a \mu(s) \ln \frac{1}{\sqrt{(x-s)^2 + f^2(s)}} ds. \quad (4)$$

Potencjał na krzywej $y = f(x)$, po wyznaczeniu $\mu(s)$, określać będzie wzór

$$V(x) = \int_{-a}^a \mu(s) \ln \frac{1}{\sqrt{(x-s)^2 + [f(x)-f(s)]^2}} ds. \quad (5)$$

W sposób analogiczny jak syntezę rozwiązania przeprowadza się syntezę jego pochodnych. Sytuację ułatwia fakt, iż pochodną potencjału (3) można obliczyć przez różniczkowanie pod znakiem całki [8].

W praktyce zagadnienia odwrotne rozwiązuje się bardzo często w układach symetrycznych. Np. dla przypadku płaskorównoległego

$$\begin{aligned} \Sigma &= \Sigma_1 \cup \Sigma_2, \\ \Sigma_1 &= \{y; y = f(x), \quad x \in (-a, a)\}, \\ \Sigma_2 &= \{y; y = -f(x), \quad x \in (-a, a)\}. \end{aligned} \quad (6)$$

W takim przypadku dopuścić można symetrię lub antysymetrię rozwiązania. Jeśli rozwiązanie jest symetryczne, wówczas

$$\mu_1(s) = \mu_2(s) = \mu(s),$$

gdzie $\mu_1(s)$ — gęstość na powierzchni Σ_1 , a $\mu_2(s)$ — na powierzchni Σ_2 .

Na powierzchni $y = 0$, przy takim założeniu, pochodna normalna rozwiązania jest równa zero. Można więc przeprowadzać syntezę potencjału lub jego pochodnej względem osi x

$$V(x) = \int_{-a}^a \mu(s) \ln \frac{1}{\sqrt{(x-s)^2 + f^2(s)}} ds, \quad (7a)$$

$$\frac{\partial V(x)}{\partial x} = \int_{-a}^a \mu(s) \frac{x-s}{f^2(s) + (x-s)^2} ds. \quad (7b)$$

Podobnie zakładając rozwiązanie antysymetryczne

$$\mu_1(s) = -\mu_2(s) = \mu(s),$$

otrzymuje się równanie dla pochodnej normalnej

$$e(x) = -\frac{\partial V(x)}{\partial y} = \int_{-a}^a \frac{\mu(s) ds}{f^2(s) + (x-s)^2}. \quad (8)$$

Potencjał oraz jego pochodna styczna są w układzie antysymetrycznym równe zero dla $y = 0$.

Równania (3), (4), (5), (7) i (8) są całkowitymi równaniami Fredholma I rodzaju. Równania takie są zagadnieniami źle uwarunkowanymi w sensie Hadamarda, dlatego do ich rozwiązywania należy zastosować metodę regularyzacji [9, 10, 3, 5].

3. POSZUKIWANIE POWIERZCHNI PRZY ZADANYCH WARUNKACH BRZEGOWYCH

Przy syntezie pól elektromagnetycznych metoda podana w p. 2 jest niewygodna, ponieważ powierzchnie, na których zadaje się warunki brzegowe, są najczęściej powierzchniami ekwipotencjalnymi. W polach elektrostatycznych i przepływowych prądów stałych są to powłoki przewodzące, a w magnetostatyce — ferromagnetyczne. Realizowanie zmiennej wartości potencjału na takiej powierzchni powoduje konieczność jej dzielenia i oddzielnego zasilania. Dlatego duże znaczenie może mieć rozwiązanie następującego problemu odwrotnego.

Znaleźć brzeg Σ obszaru D , na którym zadany został warunek brzegowy

$$V(P) = V_0 = \text{const}, \quad P \in \Sigma, \quad (9)$$

aby w obszarze $D_0 \subset D$ rozwiązanie równania Laplace'a było równe zadanemu.

Problem ten można rozwiązać wykorzystując rozwiązanie zadania poprzedniego. Znając rozkład gęstości $\mu(Q)$ na brzegu Σ obszaru D określić można w otoczeniu obszaru D_0 rozkład linii ekwipotencjalnych. Dowolna z nich spełnia postawione wymagania. Szybciej jednak prowadzi do celu metoda bezpośrednia. Polega ona na wyznaczeniu Σ z układu równań

$$\begin{aligned} V(P) &= \int_{\Sigma} \mu(Q) E(P, Q) dS, \quad P \in D_0, \quad Q \in \Sigma, \\ V_0 &= \int_{\Sigma} \mu(Q) E(R, Q) dS, \quad R \in \Sigma. \end{aligned} \quad (10)$$

Przypadki szczególne w układzie płaskorównoległym wyprowadza się podobnie jak w p. 2. Dla przykładu, przy syntezie pochodnej normalnej rozwiązania na odcinku $(-a, a)$ płaszczyzny $y = 0$ w układzie antysymetrycznym układ równań (10) sprowadza się do:

$$e(x) = \frac{\partial V(x)}{\partial y} = \int_{-a}^a \frac{\mu(s)f(s)}{f^2(s) + (x-s)^2} ds, \quad (11a)$$

$$V_0 = \int_{-a}^a \mu(s) \ln \sqrt{\frac{(x-s)^2 + [f(x)+f(s)]^2}{(x-s)^2 + [f(x)-f(s)]^2}} ds. \quad (11b)$$

Rozpatrywany problem odwrotny sprowadza się więc do rozwiązania układu równań całkowitych I rodzaju z niewiadomymi $f(s)$ i $\mu(s)$; $e(x)$ oraz V_0 są wielkościami znanymi. Otrzymane równania są równaniami liniowymi ze względu na $\mu(s)$ oraz nieliniowymi ze względu na $f(s)$. Ich rozwiązanie jest możliwe bądź to metodą iteracyjną opartą na linearyzacji i regularyzacji [11], bądź modyfikacją klasycznej metody iteracyjnej [12].

4. PRZYKŁAD ROZWIĄZANIA ODWROTNEGO ZAGADNIENIA DIRICHLETA

Równanie (8) rozwiązane zostało, metodą różnic skończonych, dla następujących danych: $f(s) = 1,0$, $e(x) = 1,0$, $a = 1,0$. Ze względu na parzystość $e(x)$, $\mu(s)$ będzie również funkcją parzystą, wobec czego równanie (8) można zapisać w postaci

$$e(x) = \int_0^a \left[\frac{1}{1+(x-s)^2} + \frac{1}{1+(x+s)^2} \right] \mu(s) ds. \quad (8a)$$

Z powodu złego uwarunkowania konieczna jest regularyzacja równania (8a). Stosując metodę podaną przez Ławrentiewa [9] oraz dokonując całkowania przybliżonego według wzoru Simpsona, otrzymuje się układ równań liniowych

$$\sum_{i=1}^N \left[\frac{1}{1+(x_j-s_i)^2} + \frac{1}{1+(x_j+s_i)^2} \right] \mu(s_i)h + \alpha \mu(s_j) = 1,0, \quad (12)$$

$$j = 1, \dots, N,$$

gdzie:

$$\begin{aligned} s_i &= (i-0,5)h, \\ x_j &= (j-0,5)h, \\ h &= a/N. \end{aligned}$$

Jeśli natomiast wykorzystać metodę Tichonowa [10], otrzymuje się

$$-\alpha \left[\frac{\mu(s_{j+1}) - 2\mu(s_j) + \mu(s_{j-1}))}{h^2} - \mu(s_j) \right] + \sum_{i=1}^N L_{ij} \mu(s_i)h = b_j, \quad (13)$$

$$L_{ij} = \sum_{l=1}^N K_{il} K_{lj} h, \quad b_j = \sum_{i=1}^N K_{ij} h,$$

$$K_{ij} = \frac{1}{1+(x_j-s_i)^2} + \frac{1}{1+(x_j+s_i)^2}.$$

Parametr α występujący w obu równaniach jest tzw. parametrem regularyzacji. Jego wartość wyznacza się metodą iteracyjną [10, 4].

Rozwiązania równań (12) i (13) dla $N = 50$ zestawiono w tabl. 1. W wielu punktach wartości $\mu(s_i)$ są równe zero. Dla metody Tichonowa różne od zera są jedynie $\mu(s_1)$ oraz $\mu(s_{50})$. Rozwiązania zerowe w tabl. 1 pominięto.

Wartość parametru regularyzacji przy zastosowaniu metody Tichonowa może być równa zero. Dla metody Ławrentiewa największą dokładność uzyskuje się dla $\alpha = 10^{-5}$. Zmniejszanie α poniżej tej wartości prowadzi do niestabilności obliczeń. Zwiększanie — powoduje pogorszenie dokładności wyników (rys. 1).

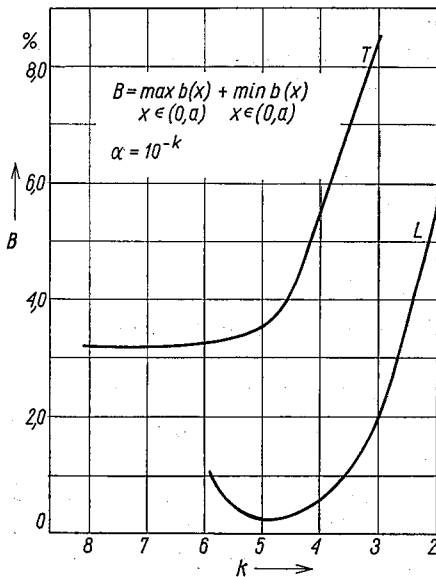
Otrzymane wyniki, w celu ich sprawdzenia, podstawiono do równania (8a). Procentową odchyłkę uzyskanych w ten sposób funkcji $e(x)$ od zadanej wartości $e_o(x) = 1,0$

$$b(x) = \frac{e(x) - e_o(x)}{e_o(x)} 100\% \quad (14)$$

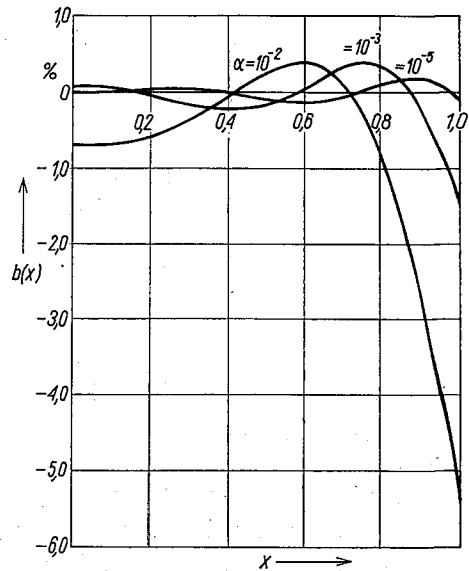
Tablica 1

Wyniki rozwiązania równania (8a)

	Metoda Ławrentiewa	Metoda Tichonowa
i	$\mu(s_i)$	$\mu(s_i)$
1	8,400	7,735
14	-6,460	—
15	-5,497	—
16	-2,227	—
17	-3,024	—
18	-4,892	—
19	-3,100	—
20	7,593	—
21	6,132	—
22	7,208	—
29	39,162	—
35	17,574	—
36	8,859	—
37	11,624	—
40	-91,562	—
41	-29,136	—
43	-6,380	—
44	-18,646	—
50	108,827	33,939



Rys. 1. Maksymalna procentowa odchyłka rozwiązań uzyskanych metodami Ławrentiewa (L) i Tichonowa (T) od rozwiązania założonego

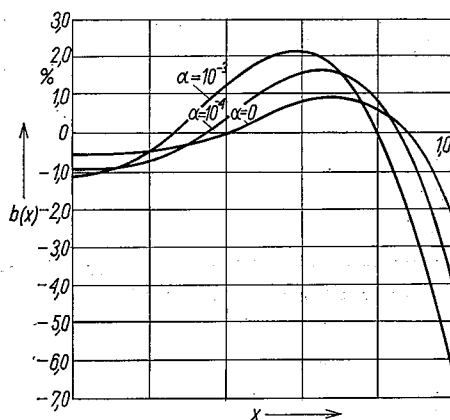


Rys. 2. Odchyłka rozwiązań zagadnienia odwrotnego uzyskanych metodą Ławrentiewa od rozwiązania założonego

dla wyników przedstawionych w tabl. 1 i innych (dla różnych wartości α) pokazano na rys. 2 i 3. Biorąc pod uwagę dokładność uzyskanych wyników, zdecydowaną przewagę ma metoda Ławrentiewa. Daje ona dokładność wyników rzędu 0,2%, podczas gdy metoda Tichonowa ok. 3,5%.

Znajomość gęstości $\mu(x)$ pozwala na wyznaczenie potencjału $V(x)$ na rozpatrywanej powierzchni. Należy się w tym celu posłużyć wzorem

$$V(x) = \int_0^a \mu(s) \ln \sqrt{\frac{[(x-s)^2 + 4][(x+s)^2 + 4]}{(x-s)^2(x+s)^2}} ds. \quad (15)$$



Rys. 3. Odchyłka rozwiązań zagadnienia odwrotnego od rozwiązania założonego dla metody Tichonowa

Funkcja podcałkowa we wzorze (15) zawiera osobliwość dla $x = s$. Dlatego całkowanie we wzorze (15) według metody Simpsona odbywa się zgodnie z zależnością

$$V(x) = \sum_{i=1}^N \mu(s_i) \left\{ \ln \sqrt{\frac{[(x-s_i)^2 + 4][(x+s_i)^2 + 4]}{(x+s_i)^2}} + A \right\} h, \quad (16)$$

gdzie:

$$h = a/N,$$

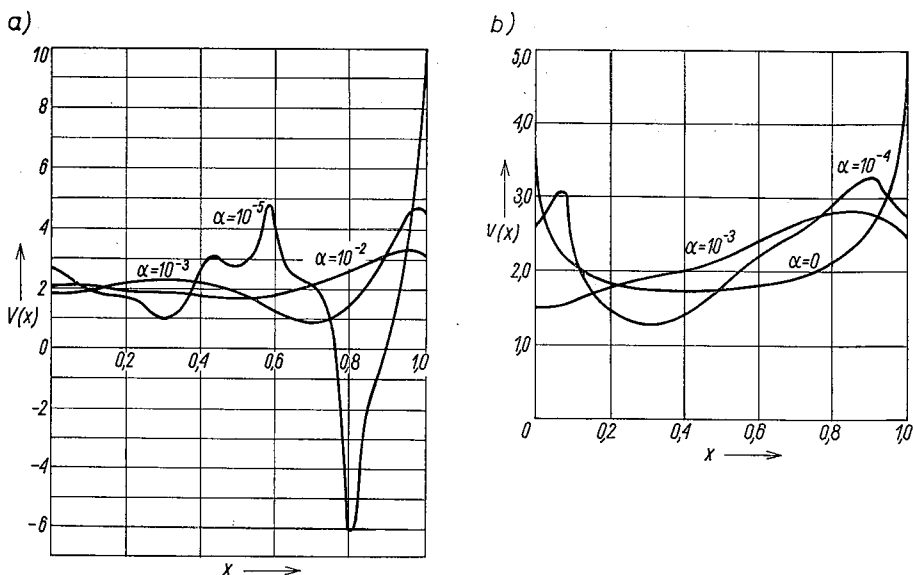
$$A = \begin{cases} \ln \frac{1}{x-s_i} & \text{gdy } x \neq s_i, \\ 1 - \ln 0,5h & \text{gdy } x = s_i. \end{cases}$$

Dla rozwiązań podanych w tabl. 1 rozkłady potencjału na płaszczyźnie $y = 1,0$ pokazane zostały na rys. 4.

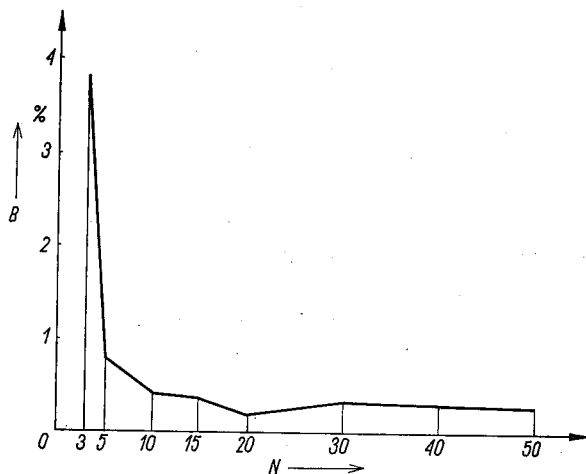
Równanie (12) rozwiązane zostało również dla innych wartości N . Maksymalną procentową odchyłkę

$$B = \max_{x \in (0, b)} b(x) - \min_{x \in (0, b)} b(x) \quad (17)$$

dla kilku wartości N zaprezentowano na rys. 5. Zwiększanie N powyżej 10 jest niecelowe, ponieważ nie przynosi to poprawy wyników.



Rys. 4. Wykresy potencjałów na powierzchni $y = 1,0$
 a — metoda Ławrentiewa, b — metoda Tichonowa



Rys. 5. Wartości maksymalnej procentowej odchyłki dla różnych N

5. PRZYKŁAD ZAGADNIENIA ODWROTNEGO Z NIEWIADOMYM RÓWNANIEM POWIERZCHNI, NA KTÓREJ ZADANY ZOSTAŁ ROZKŁAD POTENCJAŁU

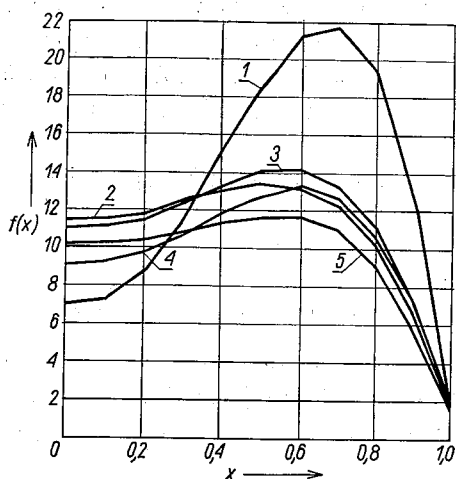
Układ równań (11) również rozwiązany zostanie metodą różnic skończonych. Ze względu na $\mu(s)$ oba równania są liniowe; wobec tego zastosować można metodę podaną w p. 4. Ze względu na $f(s)$ równania (11) są równaniami nieliniowymi. Do rozwiązania nieliniowego równania I rodzaju

$$g(x) = 0 \quad (18)$$

zastosować można odmianę klasycznej metody iteracyjnej [12]

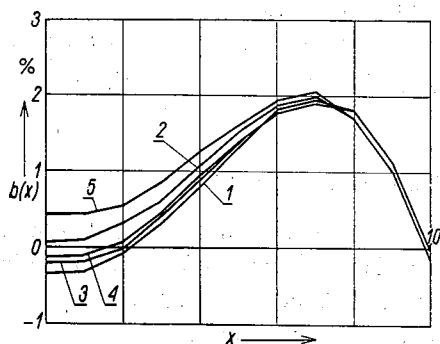
$$x_{n+1} = x_n + \beta g(x_n). \quad (19)$$

Metoda ta różni się od tradycyjnie stosowanej wartością parametru β . Jego wartość nie może być zbyt duża, ze względu na możliwość niestabilności obliczeń. Dla całkowego równania Fredholma I rodzaju maksymalną wartość β oszacowano w zależności od wartości własnych jądra całkowego [12]. Dla równań nieliniowych znalezienie takiego osza-

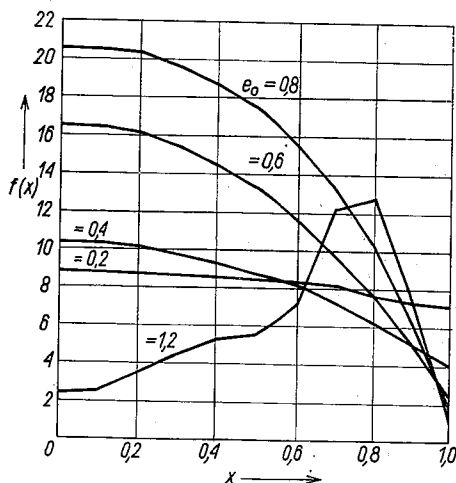


Rys. 6. Równania powierzchni będące rozwiązaniem odwrotnego zagadnienia Laplace'a przy założeniu $V(x) = V_0 = \text{const}$

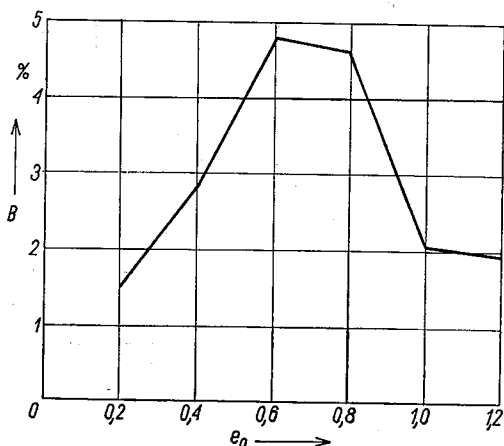
krzywa 1: $\beta = 1,0$, $L = 20$, $K = 50$, $F = 2,5$; krzywa 2: $\beta = 0,5$, $L = 30$, $K = 30$, $F = 5,0$; krzywa 3: $\beta = 0,5$, $L = 20$, $K = 50$, $F = 2,5$; krzywa 4: $\beta = 0,5$, $L = 20$, $K = 50$, $F = 5,0$; krzywa 5: $\beta = 0,5$, $L = 20$, $K = 40$, $F = 5,0$



Rys. 7. Procentowa odchyłka rozwiązań z rys. 6



Rys. 8. Równania powierzchni dla różnych e_0



Rys. 9. Wartości maksymalnej procentowej odchyłki dla rozwiązań z rys. 7

cowania jest problemem bardzo trudnym, dlatego wartość β znajdowano metodą prób.

Układ równań (11) rozwiązano metodą iteracji. Zakładając znajomość n -tego przybliżenia funkcji $f_n(s)$, $n+1$ przybliżenie $\mu_{n+1}(s)$ wyznacza się z (11b), a następnie $f_{n+1}(s)$ z (11a). Jako zerowe przybliżenie $f_0(s)$ przyjmowano funkcję stałą

$$f_0(s) = F = \text{const.}$$

Wyniki obliczeń dla $e_0(x) = 1,0 = \text{const}$, $V_0 = 10$, $N = 10$ i $a = 1,0$ oraz dla różnych K (ilości iteracji przy rozwiązywaniu równania (19)), L (ilości iteracji przy rozwiązywaniu układu równań (11)), F i β zestawiono na rys 6. Rozwiązania dla różnych L , K i F różnią się znacznie między sobą. Wszystkie spełniają jednak układ równań (11) z podobną dokładnością (rys. 7).

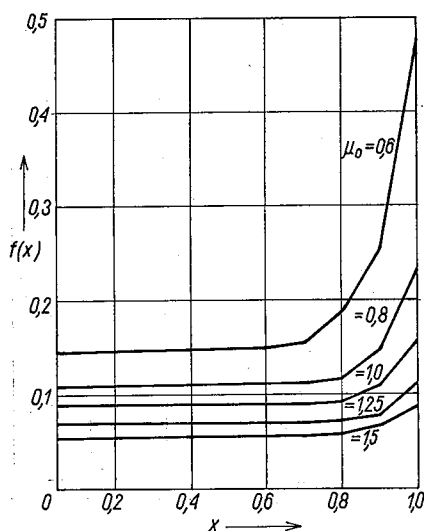
Wyniki dla innych wartości $e_0(x)$ oraz maksymalną procentową odchyłkę uzyskanych wyników od rozwiązania założonego pokazano na rys. 8 i 9.

6. O PEWNYM PROBLEMIE ODWROTNYM Z NIEWIADOMYM RÓWNANIEM POWIERZCHNI

W wielu dziedzinach techniki duże znaczenie ma rozwiązanie problemu następującego: znaleźć równanie powierzchni, na której potencjał $V(x)$ i gęstość $\mu(x)$ są funkcjami zadanymi. Zadanie to w antysymetrycznym układzie płaskorównoległym sprowadza się do rozwiązania równania

$$V(x) = \int_{-a}^a \ln \sqrt{\frac{(x-s)^2 + [f(x)+f(s)]^2}{(x-s)^2 + [f(x)-f(s)]^2}} \mu(s) ds \quad (20)$$

z niewiadomą $f(x)$.



Rys. 10. Równania powierzchni, na których dla $V(x) = V_0 = 10,0$ $\mu(x) = \mu_0$

Problem taki występuje między innymi w technice wysokich napięć, przy konstrukcji elektrod do badania dielektryków [13]. Aby natężenie pola elektrycznego na powierzchni elektrod było możliwie jednakowe, stosuje się elektrody o odpowiednim kształcie tzw. profil Rogowskiego lub elektrody Bruce'a [14]. Jednorodność natężenia pola można zapewnić przez jednorodność gęstości powierzchniowej ładunku, a więc rozwiązując równanie (20) profil Rogowskiego wyznaczyć można w sposób ścisły. Wyniki rozwiązania równania (20) dla $a = 1,0$, $N = 20$, $V_0 = 10,0 = \text{const}$ oraz różnych $\mu(s) = \mu_0 = \text{const}$, uzyskane metodą iteracyjną (19) pokazano na rys. 10. Rozwiązania te spełniają równanie (20) z bardzo dużą dokładnością (wyższą od $10^{-3} \%$).

7. ZAKOŃCZENIE

Opisana metoda rozwiązywania zagadnień odwrotnych może znaleźć zastosowanie wszędzie tam, gdzie zjawiska są opisywane równaniami Laplace'a. W teorii pola elektromagnetycznego można ją zastosować do syntezy pól elektrostatycznych, przepływowych i magnetostatycznych. W analogiczny sposób rozwiązuje się problemy bardziej ogólne niż przykłady podane w pracy, a więc przy dowolnych powierzchniach Σ i w obszarach trójwymiarowych.

W pracy nie poruszano zagadnienia istnienia rozwiązania zagadnienia odwrotnego dla równania Laplace'a. Z technicznego punktu widzenia nie ma to istotnego znaczenia. Jeśli bowiem nawet rozwiązanie w sensie ścisłym nie istnieje, poszukuje się quasi-rozwiązania, tzn. elementu spełniającego równanie z możliwie największą dokładnością. Poszukiwanie quasi-rozwiązania może być problemem wieloznacznym. Wyniki obliczeń uzyskane różnymi metodami, lub nawet tą samą metodą, ale dla odmiennych założeń np. co do dokładności obliczeń, mogą się znacznie różnić między sobą. W zastosowaniach wystarcza znalezienie jednego, dowolnego, z wielu istniejących quasi-rozwiązań.

LITERATURA

1. P. Moon, D. E. Spencer, *Field Theory for Engineers*, New York 1961 (tłum. polskie, Warszawa 1966).
2. J. P. van der Merwe, *The Inverse Interior Dirichlet Problem*, J. Appl. Phys., 50, 5120—5126 (1979).
3. K. Adamiak, *Tichonov's Regularization Method Applied to the Magnetic Field on a Solenoid Axis Synthesis*, Arch. f. Elektr., 59, 19—22 (1977).
4. K. Adamiak, *Synthesis of Homogeneous Magnetic Field in Internal Region of Cylindrical Solenoid*, Arch. f. Elektr., 62, 75—79 (1980).
5. K. Adamiak, *Synteza pola magnetycznego w obszarach dwuwymiarowych za pomocą uzwojeń bezrdzeniowych*, Rozpr. Elektr., z. 2, 1981, 311—325.
6. O. Coufal, *Optimum Cross-Section of the Superconducting Coils for MHD Generators*, Acta Technika ČSAV, 16, 541—552 (1971).
7. H. Marcinkowska, *Wstęp do teorii równań różniczkowych cząstkowych*, Warszawa 1972.
8. W. M. Babicz i inni, *Linijnyje urawnienia matematycznej fiziki*, Moskwa 1964 (tłum. polskie, Warszawa 1970).
9. M. M. Ławrentiew, *Ob niekotorych niekorektnych zadaczach matematycznej fiziki*, Nowosibirsk 1962.

10. A. N. Tichonov, W. J. Arsenin, *Metody reszenia niekorektnych zadacz*, Moskwa 1974.
11. A. N. Tichonow, *O reszenii nelinejnych integralnych urawnienij pierwszego rodzaju*, Dokl. AN SSSR, 156, 1236—1239 (1964).
12. W. M. Fridman, *Metod posledowatielnych priblizenij dla integralnogo urawnienia Fredgolma I roda*, Uspiechi mat. nauk, 11, 233—234 (1967).
13. Z. Siciński, *Badanie materialów elektroizolacyjnych*, Warszawa 1968.
14. F. M. Bruce, *Calibration of Uniform-Field Spark-Gaps for High-Voltage Measurement at Power Frequencies*, Journal of IEE, 94, 138—154 (1947).

K. ADAMIAK

REVERSE EDGE PROBLEMS OF DIRICHLET TYPE FOR LAPLACE EQUATION

Summary

The paper presents a method of potential field synthesis. The necessary field distribution is obtained by appropriately selecting the values of field in the area assumed. The problem can be reduced to solving a Fredholm integral equation of the first type or a set of integral equations. The results of numerical calculations of several examples in a plano-parallel arrangement are presented. The calculations were performed by means of the regularization and iterative methods.

K. ADAMIAK

PROBLÈMES DES LIMITES INVERSES DU TYPE DE DIRICHLET POUR L'ÉQUATION DE LAPLACE

Résumé

Dans le travail on a donné la méthode de synthèse des champs potentiels. On obtient la répartition du champ demandée en choisissant de la manière adéquate les valeurs du potentiel du champ sur la surface admise. Le problème se réduit à la résolution de l'équation intégrale de Fredholm du premier genre ou à celle du système des équations intégrales. On a présenté les résultats des calculs numériques de quelques exemples du système plat-parallèle. Les calculs ont été faits par méthode de régularisation et méthodes itératives.

K. ADAMIAK

REZIPROKES RANDPROBLEM DIRICHLET-TYPS FÜR LAPLACE-GLEICHUNGEN

Zusammenfassung

Es wurde die Synthesenmethode für Potentialfelder dargestellt. Die gewünschte Feldverteilung auf einer bestimmten Fläche wird dadurch errichtet, daß in einer entsprechenden Weise die Potentialwerte des Feldes auf der gegebenen Fläche zusammengestellt werden. Das Problem reduziert sich zur Lösung der Fredholm I-Differentialrechnung der Art bzw. des Differentialrechnungssystems. Die Ergebnisse numerischer Berechnungen einiger Beispiele in flach-paralleler Anordnung wurden gebracht. Die Berechnungen wurden mittels der Regularisationsmethode und den Iterationsmethoden durchgeführt.

К. АДАМЯК

ОБРАТНАЯ КРАЕВАЯ ЗАДАЧА ТИПА ДИРИХЛЕ ДЛЯ УРАВНЕНИЯ ЛАПЛАССА

Р е з ю м е

Приведен метод синтеза потенциальных полей. Заданное распределение поля в некоторой области получается путем соответствующего подбора значений потенциалов поля на принятой поверхности. Вопрос сводится к решению интегрального уравнения Фредгольма первого порядка или системы интегральных уравнений. Представлены результаты числовых расчетов нескольких примеров в плоскопараллельной системе. Расчеты проведены регуляризационными и итерационными методами.

Algorytm sterowania dyskretnego układem Leonarda

KRZYSZTOF MIĄCZYŃSKI (KRAKÓW)

Akademia Górniczo-Hutnicza

Otrzymano 19.7.1979

W artykule przedstawiono koncepcję bezpośredniego sterowania cyfrowego rozruchem i nawrotem układu Leonarda ze wzbudzeniem przekształtnikowym. Stwierdzono, że dla napędu tego typu może okazać się wystarczający linowy algorytm działania regulatora prędkości. Zaproponowany algorytm należy do klasy dających przebiegi przejściowe o skończonym czasie trwania i zerowym błędzie ustalonym zadanej prędkości. Pokazano, że zarówno maksymalna pochodna jak i pulsacja prądu rozruchowego mieszczą się w dopuszczalnych granicach. Przedyskutowano ogólnie stosowane metody realizacji dyskretniej transmittancji regulatora.

1. WSTĘP

Zagadnienie sterowania dyskretnego silnikiem prądu stałego nie jest zagadnieniem nowym, lecz w większości przypadków rozważano serwomechanizmy lub napędy bardzo małej mocy, co w istotny sposób rzutowało na wybór algorytmu regulatora. W rozważanym przypadku przyjęto napęd o cechach charakterystycznych układu Leonarda średniej mocy.

Rozważania przeprowadzono przy następujących założeniach:

- 1) przyjęto liniowy model układu napędowego o rzeczywistych wartościach własnych,
- 2) wzbudzenie generatora zasila nawrotny układ przekształtnikowy, sterowany poprzez człon liniaryzujący charakterystykę przekształtnika w przedziale stosowanych napięć wzbudzenia,
- 3) pominięto opóźnienie przekształtnika, traktując go jako człon proporcjonalny,
- 4) na wyjściu regulatora cyfrowego znajduje się ekstrapolator rzędu zerowego.

Pozostałe założenia, celem skupienia uwagi, zamieszczono w dalszym tekście. Zadaniem jest wyznaczenie regulatora cyfrowego zapewniającego możliwie szybkie osiągnięcie żądanej prędkości od stanu zerowego, bądź nawrót od prędkości ustalonej. Uwzględnione przy tym winny być ograniczenia związane z prędkością narastania oraz maksymalną wartością prądu twornika. Ponadto przyjęto, że rozruch (lub nawrót) odbywa się bez obciążenia. Obliczenia przeprowadzono dla modelu w układzie jednostek względnych.

2. UKŁAD JEDNOSTEK I WIELKOŚCI WZGLĘDNYCH

Przyjęty układ jednostek:

jednostka napięcia U_N — napięcie znamionowe silnika,

jednostka prądu I_N — prąd znamionowy silnika,

jednostka prędkości ω_{ON} — prędkość biegu luzem silnika przy napięciu znamionowym,

jednostka czasu B — elektromechaniczna stała czasu układu Leonarda, określona

$$\text{jako } B = J \frac{R}{c_s^2}.$$

Wielkości względne:

$$\nu = \frac{\omega}{\omega_{ON}} — \text{prędkość silnika,}$$

$$i = \frac{I}{I_N} — \text{prąd silnika,}$$

$$e_g = \frac{E_g}{U_N} — \text{siła elektromotoryczna generatora,}$$

$$x_1 = \frac{HI}{HI_N} = i — \text{sygnał przetwornika prądu (} H \text{ jest współczynnikiem wzmocnienia przetwornika),}$$

$$\mu = \frac{\mathfrak{M}}{c_s I_N} — \text{moment mechaniczny (} c_s \text{ jest stałą silnika),}$$

$$\varrho = \frac{RI_N}{U_N} — \text{uogólniona rezystancja obwodu twornika (} R \text{ jest rzeczywistą rezystancją uogólnioną),}$$

$$\tau = \frac{t}{B} — \text{czas względny,}$$

$$x_2 = \frac{\omega k_{Tg}}{\omega_{ON} k_{Tg}} = \nu — \text{sygnał przetwornika prędkości (} k_{Tg} \text{ jest współczynnikiem wzmocnienia przetwornika prędkości).}$$

$$u_w = \frac{U_w}{U_N} — \text{napięcie wzbudzenia generatora,}$$

$$i_w = \frac{I_w r_w}{U_N} — \text{prąd wzbudzenia generatora,}$$

$$u_s = \frac{U_s}{U_N} — \text{napięcie sterujące układu zasilającego wzbudzenia generatora,}$$

$$k_u = \frac{E_{gN}}{r_w I_{wN}} — \text{współczynnik wzmocnienia napięciowego generatora,}$$

$$k_p — \text{współczynnik wzmocnienia układu zasilającego wzbudzenie generatora,}$$

$$p = \frac{B}{T_g} — \text{odwrotność względnej stałej czasowej generatora,}$$

$q = \frac{B}{T_e}$ — odwrotność względnej elektromagnetycznej stałej czasowej obwodu głównego,
 r_w — rezystancja rzeczywista obwodu wzbudzenia generatora.

3. DYSKRETNY MODEL UKŁADU

Przy tak przyjętym układzie wielkości względnych otrzymamy następujący zespół równań stanu części ciągłej rozważanego systemu dynamicznego:

$$\begin{bmatrix} \dot{i}_w \\ \dot{x}_1 \\ \dot{y} \end{bmatrix} = \begin{bmatrix} -p & 0 & 0 \\ \frac{q}{\varrho} k_u & -q & -\frac{q}{\varrho} \\ 0 & \varrho & 0 \end{bmatrix} \begin{bmatrix} i_w \\ x_1 \\ y \end{bmatrix} + \begin{bmatrix} p & 0 \\ 0 & 0 \\ 0 & -\varrho \end{bmatrix} \begin{bmatrix} u_w \\ \mu \end{bmatrix}, \quad (1)$$

$$y = \nu,$$

lub w postaci zwartej:

$$\begin{aligned} \dot{\mathbf{x}} &= \mathbf{F} \cdot \mathbf{x} + \mathbf{G} \cdot \mathbf{u}, \\ \mathbf{y} &= \mathbf{C} \cdot \mathbf{x}. \end{aligned} \quad (2)$$

W równaniach tych nie uwzględniono wzmocnienia przekształtnika, co wobec założenia 3 jest bez znaczenia, bowiem uzyskane napięcie wzbudzenia należy tylko podzielić przez k_p dla uzyskania napięć sterujących.

W dalszym ciągu zrezygnowano z prowadzenia obliczeń na liczbach ogólnych, wprowadzając dane liczbowe rzeczywistego układu napędowego. Na podstawie danych znamionowych napędu, których tutaj nie zamieszczono, mamy:

$$p = 0,09; \quad q = 4,5; \quad \varrho = 0,048; \quad k_u = 6,1385.$$

Wartości własne wynoszą:

$$s_1 = -\lambda_1 = -1,5; \quad s_2 = -\lambda_2 = -3,0; \quad s_3 = -p = -0,09.$$

Układ równań (1) poddano dyskretyzacji wyznaczając macierz fundamentalną rozwiązań $e^{\mathbf{F}\tau}$, skąd

$$\begin{aligned} \mathbf{x}(k+1) &= e^{\mathbf{F}T} \mathbf{x}(k) + \int_0^T e^{\mathbf{F}(T-\xi)} \mathbf{G} d\xi \mathbf{u}(k), \\ \mathbf{y}(k) &= \mathbf{C} \cdot \mathbf{x}(k). \end{aligned} \quad (3)$$

Należy jednak zwrócić uwagę na pewien znamieny fakt. O ile oczywiste jest zastosowanie ekstrapolacji rzędu zerowego odnośnie sygnału sterującego u_w , o tyle wyznaczenie modelu dyskretnego przy zerowej ekstrapolacji momentu obciążającego ma sens tylko dla momentów stałych lub o bardzo małej pochodnej. Przy zmiennym momencie popełniony błąd może być zbyt duży, a wyniki obliczeń mogą znacznie różnić się od przebiegów rzeczywistych.

Należałoby wyznaczyć model stosując dla momentów interpolację (jeżeli przebieg momentu jest znany a priori) lub ekstrapolację rzędu pierwszego. W rozważanym przypadku

rozpatruje się jednak układ nieobciążony (lub z bardzo niewielkim momentem stałym, praktycznie równym zeru). Okres próbkowania ustalono drogą prób, przyjmując $T = 1$ (czyli w czasie rzeczywistym $= B$).

Uzasadnienie tego wyboru zostanie podane dalej. Nie wyklucza się przy tym krótszego okresu próbkowania, np. $T = 0,5$, jednak dla tego rodzaju układu wydaje się to niecelowe, prowadzi bowiem do znacznych napięć forsujących wzbudzenie, zaś zysk czasowy na całości przebiegu jest niewielki. Przy powyższych założeniach uzyskano:

$$\begin{bmatrix} i_w(k+1) \\ x_1(k+1) \\ v(k+1) \end{bmatrix} = \begin{bmatrix} 0,91393, & 0,0, & 0,0 \\ 59,84078, & -0,12356, & -10,83394 \\ 3,55367, & 0,00555, & 0,39647 \end{bmatrix} \begin{bmatrix} i_w(k) \\ x_1(k) \\ v(k) \end{bmatrix} + \begin{bmatrix} 0,08607, & 0,0 \\ 6,66289, & 0,60353 \\ 0,15106, & -0,03452 \end{bmatrix} \begin{bmatrix} u_w(k) \\ \mu(k) \end{bmatrix}, \quad (4)$$

co w postaci ogólnej zapiszemy

$$\mathbf{x}(k+1) = \mathbf{A}\mathbf{x}(k) + [\mathbf{b}_1, \mathbf{b}_2]\mathbf{u}(k); \quad (5)$$

pomijając kolumnę $\mathbf{b}_2$ zgodnie z założeniem o $\mu = 0$ i transformując równanie (5) otrzymamy równanie

$$\mathbf{x}(z) = (z\mathbf{I} - \mathbf{A})^{-1}\mathbf{b}_1 u_w(z), \quad (6)$$

z którego uzyskamy transmitancje współrzędnych stanu, ze względu na sterowanie $u_w(z)$. W badanym przypadku otrzymamy:

$$G_{vu}(z) = 0,15106 \frac{z^{-1}(1+1,40290z^{-1})(1+0,07316z^{-1})}{(1-0,91393z^{-1})(1-0,22294z^{-1})(1-0,04997z^{-1})}, \quad (7)$$

$$G_{x_1u}(z) = 6,66289 \frac{z^{-1}(1-z^{-1})(1+0,21698z^{-1})}{(1-0,91393z^{-1})(1-0,22294z^{-1})(1-0,04997z^{-1})}, \quad (8)$$

$$G_{i_wu}(z) = 0,08607 \frac{z^{-1}}{1-0,91393z^{-1}}. \quad (9)$$

4. OGÓLNE ZASADY SYNTEZY REGULATORA

Transmitancje (7), (8) i (9) stanowią punkt wyjściowy dla syntezy dyskretnego regulatora prędkości $D(z)$. Jako rozwiązanie zadania proponuje się nieco zmodyfikowaną wersję algorytmu, z klasy dających przebiegi przejściowe o skończonym czasie trwania [1, 2, 3]. Transmitancja układu zamkniętego jest wówczas wielomianem [1, 2, 3]

$$K(z) = \sum_{k=0}^{N-1} a_k z^{-k}, \quad (10)$$

zaś

$$y[nT] = \sum_{k=0}^{N-1} a_k v[(n-k)T], \quad (11)$$

gdzie:

- $y[nT]$ — wielkość wyjściowa,
 $v[(n-k)T]$ — wielkość sterująca,
 $N-1$ — liczba „przeszłych” wartości sygnału sterującego $v[kT]$, od których, jak i od aktualnych sygnałów $v[nT]$, zależy wielkość wyjściowa.

Jeżeli przy tym

$$\sum_{k=0}^{N-1} a_k = 1, \quad (12)$$

to otrzymamy zerowy uchyb ustalony odpowiedzi skokowej. Innymi słowy, układy o takich transmitancjach posiadają tzw. „skończoną pamięć” [3]. Są one stabilne ze względu na brak biegunów, lub raczej ze względu na wszystkie bieguny równe zero, jeżeli przyjmiemy

$$K(z) = \frac{K(z)z^{N-1}}{z^{N-1}}.$$

Algorytmy tego typu są niewygodne w realizacji przy dużej liczbie okresów próbkowania w czasie przebiegu przejściowego, tj. dla układów wysokiego rzędu. Natomiast w przypadku układów tak niskiego rzędu, jak omawiany, algorytm nie obciąża nadmiernie cyfrowego urządzenia liczącego.

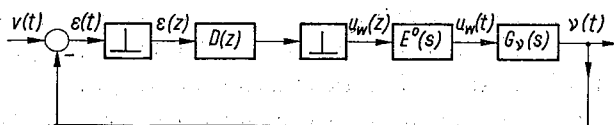
Zaletą takiej koncepcji sterowania jest możliwość uzyskania pożądanych stanów układu po m okresach próbkowania, gdzie m jest rzędem układu. Wymagane jest przy tym dysponowanie dokładnym, dobrze zidentyfikowanym modelem sterowanego układu dynamicznego. Wydaje się jednak, że w przypadku układu Leonarda, postulat ten może być spełniony z dostateczną dokładnością.

Przyjęto, ze względu na duży stosunek $\frac{B}{T_e}$, że sygnał zadający generowany przez EMC winien być dyskretnym odpowiednikiem kombinacji funkcji liniowych i skokowych, przy czym te ostatnie stanowią tylko niezbędne uzupełnienia o niewielkiej amplitudzie. W odróżnieniu od rozpowszechnionej dla układów ciągłych metody sterowania za pomocą regulatorów prądu i prędkości, zakłada się jedynie sprzężenie prędkościowe. Sprzężeniu prądowemu można by przeznaczyć rolę zabezpieczenia lub korekcji, co jednak nie jest przedmiotem rozważań w niniejszym artykule.

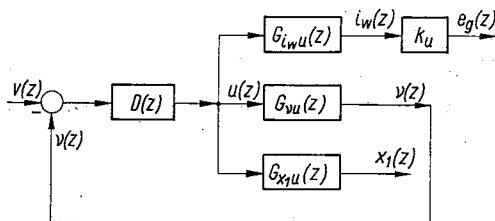
W wyniku takich założeń otrzymamy strukturę układu, jak na rys. 1 lub rys. 2, gdzie pokazano układ po dyskretyzacji. Na rysunkach tych oznaczono:

$E^0(s)$ — sktrapolator rzędu zerowego,

$G_v(s)$ — transmitancja prędkość/napięcie wzbudzenia generatora układu ciągłego, pozostałe oznaczenia podano uprzednio.



Rys. 1. Schemat blokowy sterowanego dyskretnie układu Leonarda
 $D(z)$ — dyskretny regulator prędkości, $E^0(s)$ — ekstrapolator rzędu zerowego



Rys. 2. Schemat blokowy sterowanego dyskretnie układu Leonarda w postaci zdyskretyzowanej

Na podstawie schematu z rys. 2 napiszemy

$$K(z) = \frac{D(z)G_{vu}(z)}{1 + D(z)G_{vu}(z)}, \quad (13)$$

gdzie transmitancja $G_{vu}(z)$ wyznaczona jest z uwzględnieniem ekstrapolatora rzędu zerowego. Stąd

$$D(z) = \frac{1}{G_{vu}(z)} \frac{K(z)}{1 - K(z)}. \quad (14)$$

Pozwala to wyznaczyć $D(z)$, jeżeli założy się $K(z)$. Wybór $K(z)$ nie może być zupełnie dowolny [1, 3], lecz musi spełniać pewne warunki, pominięcie których mogłoby doprowadzić do fizycznie nierealizowalnej transmitancji $D(z)$, tj. takiej, w której przy dodatnich potęgach z stopień licznika przewyższa stopień mianownika. Warunki te według [1, 3] są następujące:

- 1) jeżeli obiekt zawiera opóźnienie z^{-p} , to musi ono wystąpić w transmitancji $K(z)$ układu zamkniętego (nie może być skompensowane przez fizycznie realizowalną transmitancję $D(z)$),
- 2) jeżeli $G_{vu}(z)$ zawiera zera leżące na zewnątrz koła jednostkowego, to muszą one wystąpić jako zera w $K(z)$,
- 3) bieguny $G_{vu}(z)$ leżące na zewnątrz koła jednostkowego muszą wystąpić jako zera w transmitancji $1 - K(z)$,
- 4) jeżeli układ zamknięty ma mieć zerowy uchyb ustalony na wymuszenie $v(t) = t^q$ (gdzie q — liczba całkowita), to transmitancja $1 - K(z)$ musi zawierać czynnik $(1 - z^{-1})^{q+1}$.

Niespełnienie punktów 2 i 3 może prowadzić do niestabilności układu. Projektowanie układu ze względu na przebieg przejściowy o skończonym czasie trwania wymaga założenia transmitancji zgodnie z zależnością (10).

Na to, aby efekt skończonego czasu przebiegu przejściowego uzyskać nie tylko dla dyskretnych chwil czasu, lecz dla całego przebiegu [1, 3] trzeba, aby:

- 5) wielomian $K(z)$ zawierał wszystkie zera transmitancji $G_{vu}(z)$.

Jest to tzw. silna wersja syntezy na skończony czas przebiegu przejściowego. Warunek ten jest jednak niewystarczający. Aby można go było spełnić w naszym przypadku dla sygnału $v(t) = at$, musiałby w transmitancji układu ciągłego $G_r(s)$ istnieć przynajmniej

jeden integrator. Sygnał $u_w(t)$ jest bowiem sygnałem schodkowym, zatem przy braku integratora muszą między chwilami próbkowania wystąpić różnice pomiędzy $v(t)$ a $v(t)$.

Z napędowego punktu widzenia pulsacje prądu i prędkości pomiędzy chwilami próbkowania są bardzo niepożądane, stąd wydać się może wątpliwe zastosowanie zaproponowanego algorytmu. Należy jednak zwrócić uwagę na dużą stałą czasu generatora (o rząd większą od okresu próbkowania), dzięki której w poszczególnych okresach próbkowania uzyskamy efekt zbliżony do takiego, jaki dałby integrator. Zatem należy oczekiwać, że pulsacje będą niewielkie.

Postępując według wskazanych wyżej zasad, przy równoczesnym założeniu, że prędkość powinna narastać liniowo, przyjmiemy

$$K(z) = z^{-1}(1 + k'_2 z^{-1} + k'_3 z^{-2})(a_0 + a_1 z^{-1}), \quad (15)$$

gdzie

$$k'_2 = \frac{k_2}{k_1}; \quad k'_3 = \frac{k_3}{k_1},$$

zaś

$$L_{Gv}(z) = z^{-1}(k_1 + k_2 z^{-1} + k_3 z^{-2}) \quad (16)$$

jest licznikiem transmitancji $G_{vu}(z)$.

Równocześnie założymy

$$1 - K(z) = (1 - z^{-1})^2(1 + b_1 z^{-1} + b_2 z^{-2}). \quad (17)$$

Współczynniki a_0 , a_1 , b_1 , b_2 , nie znane z góry, należy wyznaczyć porównując współczynniki przy jednakowych potęgach z transmitancji $K(z)$, wyznaczonej w zależności (15) i (17), a następnie z zależności (14) wyznaczyć transmitancję $D(z)$.

W rozważanym przypadku uzyskamy transmitancję typu

$$D(z) = \frac{d_0 + d_1 z^{-1} + d_2 z^{-2} + d_3 z^{-3} + d_4 z^{-4}}{1 + c_1 z^{-1} + c_2 z^{-2} + c_3 z^{-3} + c_4 z^{-4}}. \quad (18)$$

Transmitancja ta jest fizycznie realizowalna i niezbyt złożona. Nie wprowadza ona zbyt wielkiego obciążenia cyfrowego urządzenia liczącego, a większość operacji zostaje dokonana na wartościach „przeszłych” w początku okresu próbkowania.

Jedyną operacją, która musi być wykonana w przedziale czasu $(kT - \delta, kT)$ pomijalnym w stosunku do okresu próbkowania, jest operacja dodawania aktualnej wartości błędu $\varepsilon_v(k)$ do zawczasu wyznaczonej wartości $[u_w(k) - \varepsilon_v(k)]$. Jest to zresztą niedogodność wszystkich algorytmów, w których wykorzystuje się bieżącą wartość błędu. Transmitancja (18) zawiera czynnik: $\frac{1 - 0,91393 z^{-1}}{(1 - z^{-1})^2}$, w którym $z = 0,91393$ jest dyskretnym

biegunem generatora. Przyjęcie mniejszego okresu próbkowania, tj. zbliżenie bieguna do jedynki, mogłoby pozwolić na pewne przybliżenie upraszczające transmitancję. Jednakże transmitancja (18) jest nie do przyjęcia z napędowego punktu widzenia. Daje ona mianowicie silne przeregulowanie wartości prądu silnika w pierwszej fazie przebiegu przejściowego. Jest to zrozumiałe, gdyż pochodna prędkości w tej fazie musi być wyższa niż pochodna sygnału zadającego, za którym silnik winien nadążać bez błędu statycznego.

5. SYNTEZA ALGORYTMU STEROWANIA

Zdecydowano się zatem na uproszczenie algorytmu sterowania, zakładając jedynie brak błędu ustalonego po osiągnięciu zadanej prędkości. Natomiast nie zrobiono żadnych założeń odnośnie błędu statycznego w trakcie nadążania za liniowo narastającym przebiegiem sygnału zadającego. Modyfikując zatem poprzednią koncepcję sterowania założymy

$$K(z) = a_0(z^{-1} + k'_2 z^{-2} + k'_3 z^{-3}), \quad (19)$$

$$1 - K(z) = (1 - z^{-1})(1 + b_1 z^{-1} + b_2 z^{-2}). \quad (20)$$

Wobec rzeczywistych pierwiastków licznika transmitancji $G_{vu}(z)$ możemy położyć $z_1 + z_2 = k'_2$, $z_1 z_2 = k'_3$, skąd otrzymamy

$$\underbrace{\begin{bmatrix} 1, & 1, & 0 \\ z_1 + z_2, & -1, & 1 \\ z_1 z_2, & 0, & -1 \end{bmatrix}}_{= F} \begin{bmatrix} a_0 \\ b_1 \\ b_2 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}. \quad (21)$$

Zatem

$$\text{DET } F = 1 + z_1 + z_2 + z_1 z_2 \neq 0 \quad (22)$$

jest warunkiem istnienia rozwiązania układu (21).

W badanym przypadku

$$K(z) = a_0(z^{-1} + 1,47921z^{-2} + 0,10287z^{-3}), \quad (23)$$

$$\text{DET } F = 2,58208,$$

w wyniku czego otrzymano

$$a_0 = 0,38728,$$

$$b_1 = 0,57288,$$

$$b_2 = 0,03983.$$

Transmitancja układu zamkniętego będzie zatem

$$\begin{aligned} K(z) &= 0,38728z^{-1}(1 + 1,47921z^{-1} + 0,10287z^{-2}) = \\ &= 0,38728z^{-1}(1 + 1,40290z^{-1})(1 + 0,07316z^{-1}). \end{aligned} \quad (24)$$

Na podstawie schematu (rys. 2) wyznaczamy

$$u_w(z) = \frac{K(z)}{G_{vu}(z)} v(z), \quad (25)$$

$$x_1(z) = K(z) \frac{G_{xu}(z)}{G_{vu}(z)} v(z). \quad (26)$$

Biorąc pod uwagę $G_{xu}(z) = \frac{L_{xu}(z)}{M(z)}$; $G_{vu}(z) = \frac{L_{vu}(z)}{M(z)}$ otrzymamy transmitancję układu zamkniętego, ze względu na prąd

$$G_{xv}(z) = K(z) \frac{L_{xu}(z)}{L_{vu}(z)}. \quad (27)$$

Wobec tego, że $K(z)$ zawiera wszystkie zera transmitancji $G_{vu}(z)$, transmitancja G_{xv} będzie również transmitancją o „skończonej pamięci”. Zatem

$$G_{xv} = 17,08218z^{-1}(1-z^{-1})(1+0,21698z^{-1}). \quad (28)$$

Można teraz ustalić współczynnik nachylenia α sygnału zadającego, generowanego przez cyfrowe urządzenia liczące

$$\mathcal{L}[\alpha t] = \frac{z^{-1}}{(1-z^{-1})^2}; \quad (29)$$

stąd

$$x_1(z) = \alpha \cdot 17,08218z^{-2} \frac{1+0,21698z^{-1}}{1-z^{-1}}, \quad (30)$$

co ostatecznie daje

$$x_1(z) = \alpha \cdot 17,08218(z^{-2} + 1,21698z^{-3} + 1,21698z^{-4} + \dots). \quad (30a)$$

Zakładając, że „dyskretny pułap” prądu winien wynosić $2I_N$, otrzymamy $\alpha = 0,09621$ skąd

$$x_1(z) = 1,64341z^{-2} + 1,99999z^{-3} + 1,99999z^{-4} + \dots; \quad (30b)$$

wówczas

$$v(z) = (0,03726z^{-1} + 0,05512z^{-2} + 0,00383z^{-3}) \frac{z^{-1}}{(1-z^{-1})^2}. \quad (31)$$

Napięcie wzbudzenia generatora wyznaczmy z zależności

$$u_w(z) = \frac{K(z)}{G_{vu}(z)} v(z), \quad (32)$$

skąd

$$u_w(z) = 2,56375(1 - 1,18684z^{-1} + 0,26056z^{-2} - 0,01018z^{-3}) 0,09621 \frac{z^{-1}}{(1-z^{-1})^2}. \quad (33)$$

Biorąc pod uwagę, że w jednostkach względnych $u_{wN} = 0,16291$ (gdzie u_{wN} napięcie dające U_N silnika), otrzymamy

$$u_w(z) = u_{wN} \cdot 1,51408(1 - 1,18684z^{-1} + 0,26056z^{-2} - 0,01018z^{-3}) \frac{z^{-1}}{(1-z^{-1})^2}. \quad (33a)$$

Podobnie jak w przypadku prądu silnika wyznacza się

$$i_w(z) = G_{i_w u}(z) \frac{K(z)}{G_{vu}(z)} v(z). \quad (34)$$

Ostatecznie

$$i_w(z) = 0,22066(z^{-1} - 0,27291z^{-2} + 0,01114z^{-3}) 0,09621 \frac{z^{-1}}{(1-z^{-1})^2}. \quad (35)$$

Sygnał sterujący $v(z) = 0,09621 \frac{z^{-1}}{(1-z^{-1})^2}$ obowiązuje w przedziale k , (1, 10).

W dziesiątym okresie próbkowania $v = 0,96210$, zaś w jedenastym $v = 1,05831$. Aby

uniknąć niepożądanych przeregulowań, cyfrowe urządzenie liczące winno wygenerować dla przedziału k , $(10, \infty)$ funkcję zadającą o postaci

$$v(z) = 0,09621 \frac{z^{-1}}{(1-z^{-1})^2} - 0,09621 \frac{z^{-1}z^{-10}}{(1-z^{-1})^2} + 0,0379 \frac{z^{-11}}{1-z^{-1}}. \quad (36)$$

Ostatni składnik stanowi uzupełnienie $v(z)$ do wartości zadanej.

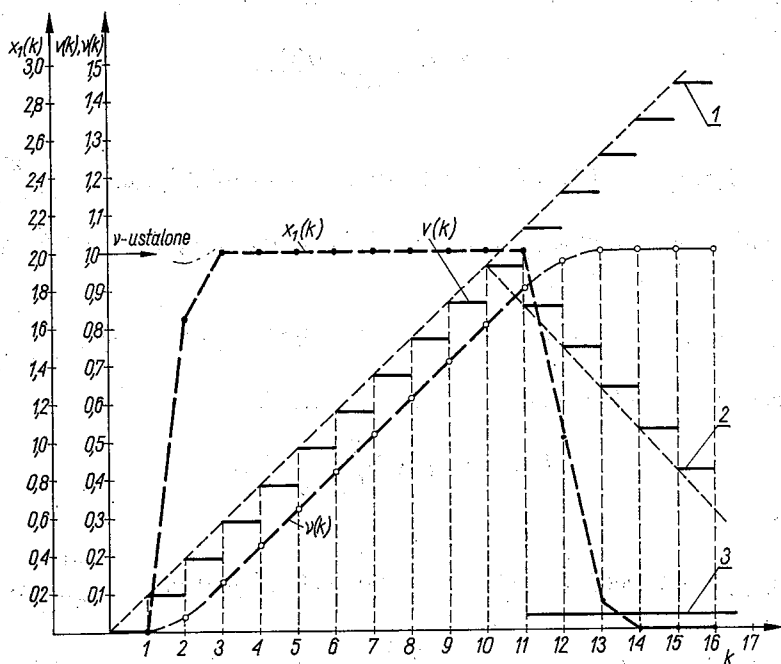
W ogólności przyjęto, że gdy $v(k+1) - v(k) > \delta(k)$, gdzie $\delta(k) = v_z - v(k)$, powinno nastąpić sterowanie według zależności (36) (v_z — wartość zadana prędkości). Założenie to pozwala uniknąć znacznych dyskretnych przeregulowań prędkości, zaś w omawianym przypadku daje dyskretny przebieg aperiodyczny. Uwaga ta dotyczy wyłącznie transmiancji przyjętego w pracy typu.

Dyskretny przebieg $v(k)$, $v(k)$, $x_1(k)$ przedstawia rysunek 3, zaś $\frac{u_w}{u_{wN}}$ oraz $\frac{i_w}{i_{wN}}$ rysunek 4. Linie przerywane łączące dyskretnie punkty przebiegu nie odtwarzają tego co się dzieje w układzie ciągłym między chwilami próbkowania; służą one jedynie do nadania rysunkom przejrzystości.

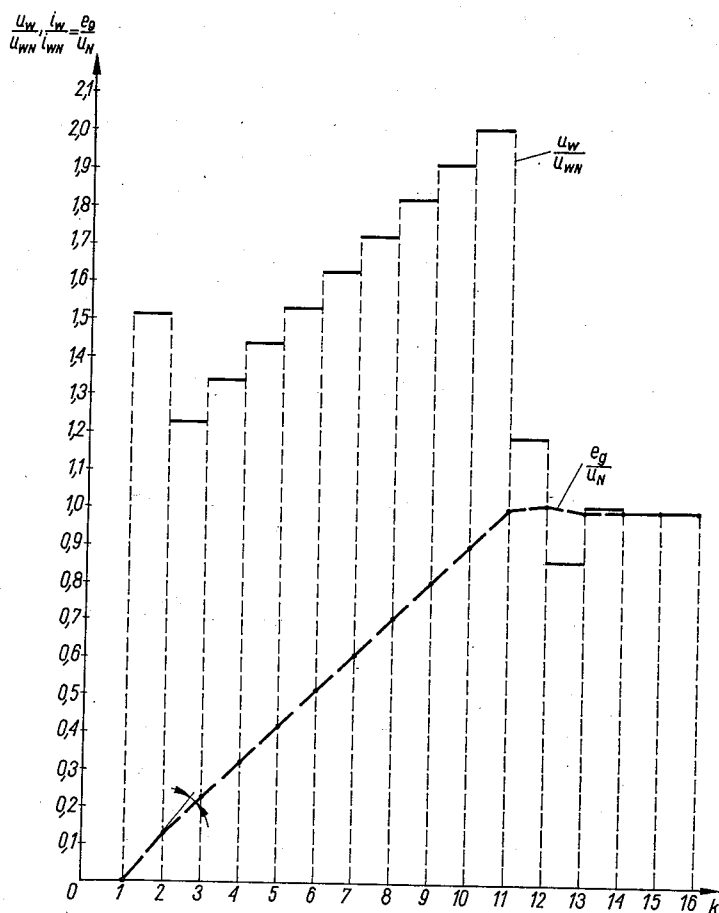
Prąd wzbudzenia generatora wygodnie wyznacza się na podstawie ciągu sterowań $u_w(k)$, wprost z równania różnicowego

$$i_w(k+1) = 0,91393 i_w(k) + 0,08607 u_w(k); \quad i_w(0) = 0. \quad (37)$$

Siła elektromotoryczna generatora przekracza wartość U_N silnika o 1,19%, co stanowi sprawdzian, że przebiegi mieszczą się w założonym przedziale liniowości układu.



Rys. 3. Dyskretny przebieg sygnału zadającego v , prądu x i prędkości v



Rys. 4. Dyskretnie przebiegi względnego napięcia wzbudzenia $\frac{u_w}{u_{wN}}$ oraz względnej siły elektromotorycznej generatora $\frac{e_g}{u_N}$

Transmitancję regulatora wyznaczmy z zależności (14)

$$D(z) = 2,56378 \frac{(1 - 0,91393z^{-1})(1 - 0,22294z^{-1})(1 - 0,04997z^{-1})}{(1 - z^{-1})(1 + 0,49189z^{-1})(1 + 0,08099z^{-1})}. \quad (38)$$

Po wymnożeniu wyrażeń w nawiasach otrzymamy postać dogodną do realizacji metodą bezpośrednią lub kanoniczną

$$D(z) = 2,56378 \frac{1 - 1,18684z^{-1} + 0,26082z^{-2} - 0,01018z^{-3}}{1 - 0,42712z^{-1} - 0,53304z^{-2} - 0,03984z^{-3}}. \quad (38a)$$

Ogólna postać otrzymanej transmitancji jest znacznie prostsza, niż opisana zależnością (18)

$$D(z) = \frac{u_w(z)}{\varepsilon(z)} = d_0 \frac{1 + d'_1 z^{-1} + d'_2 z^{-2} + d'_3 z^{-3}}{1 + c_1 z^{-1} + c_2 z^{-2} + c_3 z^{-3}}, \quad (39)$$

gdzie: $\varepsilon(z) = v(z) - v(z)$, $d'_1 = \frac{d_1}{d_0}$.

6. PROGRAMOWANIE TRANSMITANCJI REGULATORA

Programowanie transmitancji (39) najdogodniej jest przeprowadzić metodą bezpośrednią lub kanoniczną [3]. Pierwsza z nich wymaga więcej komórek pamięci oraz więcej operacji przesunięć, lecz mniej operacji po wprowadzeniu aktualnej wartości błędu $\varepsilon(k)$, zaś druga odwrotnie, mniej obciąża pamięć, lecz wymaga więcej operacji po wprowadzeniu $\varepsilon(k)$.

Metoda bezpośrednia

Przyjmując $u'_w(z) = \frac{u_w(z)}{d_0}$ otrzymamy

$$u'_w(z) (1 + c_1 z^{-1} + c_2 z^{-2} + c_3 z^{-3}) = (1 + d'_1 z^{-1} + d'_2 z^{-2} + d'_3 z^{-3}) \varepsilon(z); \quad (40)$$

stąd

$$u'_w(z) = \varepsilon(z) + \sum_{i=1}^3 d'_i z^{-i} \varepsilon(z) - \sum_{j=1}^3 c_j z^{-j} u'_w(z), \quad (40a)$$

zatem

$$\mathcal{Z}^{-1} [u'_w(z)] = u'_w(k) = \varepsilon(k) + \sum_{i=1}^3 d'_i \varepsilon(k-i) - \sum_{j=1}^3 c_j u'_w(k-j), \quad (41)$$

przy czym

$$\varepsilon(-1) = \varepsilon(-2) = \varepsilon(-3) = u'_w(-1) = u'_w(-2) = u'_w(-3) = 0.$$

Operację mnożenia $u_w(k) = d_0 u'_w(k)$ można przeprowadzić w programie, lecz także celem oszczędności czasu można ją wykonać hardware'owo przy konwersji cyfra-analog. Różnicę drugiego i trzeciego składnika można wyliczyć podczas przerwy między chwilami próbkowania.

Ostatnią operacją jest dodanie $\varepsilon(k)$, które musi być wykonane w bardzo krótkim przedziale czasu $[kT - \delta, kT]$ (przy czym $\delta \ll T$). Schemat blokowy realizacji regulatora metodą bezpośrednią, pozwalający napisać program maszynowy przedstawia rys. 5.

Metoda kanoniczna

Podobnie jak poprzednio, przyjmiemy $u'_w(z) = \frac{u_w(z)}{d_0}$

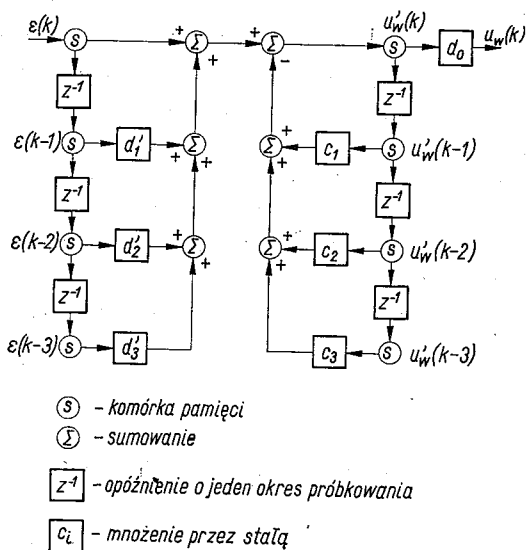
$$\frac{u'_w(z)}{1 + \sum_{i=1}^3 d'_i z^{-i}} = \frac{\varepsilon(z)}{1 + \sum_{j=1}^3 c_j z^{-j}} = F(z), \quad (42)$$

gdzie $F(z)$ jest wielkością pomocniczą.

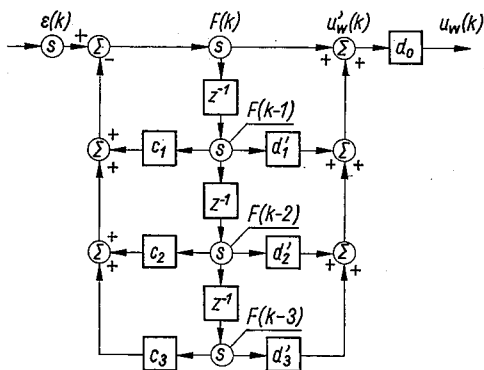
Z zależności (42) otrzymamy

$$F(z) = \varepsilon(z) - \sum_{j=1}^3 c_j z^{-j} F(z), \quad (43)$$

$$u'_w(z) = F(z) + \sum_{i=1}^3 d'_i z^{-i} F(z). \quad (44)$$



Rys. 5. Schemat blokowy regulatora zaprogramowanego metodą bezpośrednią



Rys. 6. Schemat blokowy regulatora zaprogramowanego metodą kanoniczną

Wykonując transformację $\mathcal{Z}^{-1}$ równań (43), (44) otrzymamy

$$F(k) = \varepsilon(k) - \sum_{j=1}^3 c_j F(k-j), \quad (45)$$

$$u'_w(k) = F(k) + \sum_{i=1}^3 d'_i F(k-i) \quad (46)$$

oraz

$$u_w(k) = d_0 u'_w(k),$$

gdzie $F(-1) = F(-2) = F(-3) = 0$.

W tym przypadku, po pobraniu $\varepsilon(k)$, wykonać należy dwie operacje sumowania. Zależności (45) i (46) nie można połączyć w jedną, ponieważ należy zapamiętać oddzielnie wartość $F(k)$, a po wygenerowaniu sterowania przesunąć ją do komórki pamięci odpowiadającej indeksowi $(k-1)$. Uwaga odnośnie mnożenia przez d_0 z opisu poprzedniej metody zachowuje swoją moc.

Schemat blokowy realizacji regulatora, pozwalający napisać program maszynowy przedstawia rysunek 6. Wydaje się, że metoda kanoniczna jest w omawianym przypadku mniej korzystna niż metoda bezpośrednia. Ścisłą analizę porównawczą metod należałoby przeprowadzić w oparciu o dane konkretnego cyfrowego urządzenia sterującego, przy czym uwzględnić należy metody: kaskadową, równoległą oraz rozważyć rozwiązania kombinowane.

7. POCHODNE I PULSACJE PRZEBIEGU PRĄDU UKŁADU

W wyniku generacji przez cyfrowe urządzenie liczące funkcji $v(z)$ określonej przez zależność (36) powstanie „martwy” okres próbkowania na początku przebiegu sterowania $u_w(k)$. W związku z tym należy generować funkcję zadającą $v_1(z) = zv(z)$. Układ napędowy z silnikiem prądu stałego winien spełniać postulat odnośnie maksymalnej pochodnej prądu mniejszej od dopuszczalnej oraz pulsacji prądu ustalonego, mniejszych od dopuszczalnych. Maksymalna pochodna prądu pojawi się w pierwszym okresie próbkowania. Otrzymamy zatem

$$\frac{dx_1}{d\tau} = 12,75540[-0,02193e^{-0,09\tau} + 0,70922e^{-1,5\tau} - 0,68730e^{-3,0\tau}]. \quad (47)$$

Stąd przechodząc na jednostki bezwzględne wyznaczymy

$$\left. \frac{dI}{dt} \right|_{\max} \approx 11,5 \frac{I_N}{[s]} < 100 \frac{I_N}{[s]}.$$

Pulsacje prądu (amplitudę) ustalonego można oszacować na podstawie $e(k)$, $v(k)$, $x_1(k)$, q , dla $k > 3$. Amplituda ich jest mniejsza niż 1,62% I_N czyli 0,81% I_r , gdzie I_r jest prądem rozruchowym. Można zatem przyjąć, że przebiegi prądu spełniają warunki stawiane napędom prądu stałego.

Rzeczywisty okres próbkowania wynosi 0,18 [s], wydaje się więc, że osiągnięcie zadanego prądu rozruchu w czasie 0,36 [s] (2 okresy próbkowania) nie jest zbyt dobrym wynikiem. Należy jednak pamiętać, że czas ten wynosi tylko około 14% całego czasu rozruchu oraz że po pierwszym okresie próbkowania prąd osiąga 82% wartości prądu rozruchowego.

Istnieje możliwość skrócenia czasu narastania prądu poprzez zmniejszenie okresu próbkowania, jednakże będzie to wymagało znacznie wyższych napięć forsujących. Autor sądzi, że postępowanie takie jest niecelowe, a ponadto utrudnia dostosowanie proponowanych algorytmów do wolniejszych cyfrowych urządzeń sterujących.

8. PODSUMOWANIE I WNIOSKI

Przedstawiony algorytm sterowania należy do klasy algorytmów suboptymalnych, o minimalnej liczbie okresów próbkowania, w których osiąga się pożądany stan układu. W omawianym przypadku liczba tych okresów wynosi 3 (to jest tyle, ile rząd układu). Zwraca uwagę fakt, że prąd ustalony zostaje osiągnięty po dwóch okresach próbkowania. Uzyskano to poprzez brak założeń odnośnie błędu statycznego układu w chwili przejścia do ustalonej fazy rozruchu, uznając, że błąd ten nie będzie miał istotnego wpływu na czas osiągania prędkości ustalonej.

Okazuje się, że dla omawianej klasy układów i zadań stawianych przed nimi, tj. dla układów napędowych o dominującej stałej czasowej, można znaleźć rozwiązanie z proponowanej klasy algorytmów w obszarze liniowej teorii układów dyskretnych, bez uciekania się do metod uwzględniających nasycenie sygnału sterującego [2].

Wybór okresu próbkowania pozostaje nadal sprawą otwartą, ale można wnioskować, że jeżeli nie pojawią się ograniczenia związane ze zbyt małą szybkością cyfrowego urządzenia liczącego, to z napędowego punktu widzenia można dobierać okres próbkowania rzędu elektromechanicznej stałej czasowej. Wydłużenie go może spowodować nadmierne pulsacje prądu w czasie ustalonej fazy rozruchu, zaś skrócenie może doprowadzić do nadmiernej pochodnej prądu.

Nie rozważano problemu stabilizacji prędkości przy obciążeniu momentem. Badaniu takiemu należy algorytm poddać, wszelako negatywny wynik nie dewaluuje przedstawionej koncepcji sterowania. Rozwiązanie bowiem pomyślane było dla układów, w których rozruch odbywa się na biegu luzem, zaś obciążenie pojawia się po osiągnięciu prędkości ustalonej lub praktycznie nie występuje, jak w wielu przypadkach napędów hutniczych. W przypadku negatywnym, można przejść na uprzednio zaprogramowany inny algorytm sterowania, lepiej realizujący stabilizację układu ze względu na moment obciążający. W technice analogowej mogło to powodować istotne trudności, lecz w technice cyfrowej adaptację przeprowadza się software'owo, co zupełnie zmienia postać rzeczy.

Kontynuując tę problematykę należałoby rozważyć możliwość adaptacji przedstawionej koncepcji na układy przekształtnikowe. Zdaniem autora możliwość taka istnieje, jednakże zadanie jest znacznie trudniejsze. Sugeruje się, że należy jako łącznika pomiędzy częścią cyfrową i analogową użyć ekstrapolatora rzędu pierwszego, mimo niekorzystnego przesunięcia fazowego jakie taki ekstrapolator wprowadza do układu, powodując zmniejszenie zapasu stabilności.

LITERATURA

1. J. T. Tou, *Digital and Sampled-data Control Systems*, Mc Graw Hill Book Company Inc., New York, Toronto, London 1959.
2. J. T. Tou, *Nowoczesna teoria sterowania*, WNT, Warszawa 1967.
3. A. Manitius, J. Kostro, *Bezpośrednie sterowanie cyfrowe*, NOT, Warszawa 1971.
4. T. Kaczorek, *Teoria układów regulacji automatycznej*, WNT, Warszawa, 1974.
5. J. Ackermann, *Regulacja impulsowa*, WNT, Warszawa, 1976.

K. MIĄCZYŃSKI

SAMPLE DATA ALGORITHM FOR LEONHARD DRIVE SYSTEM

Summary

This paper presents the conception of the direct digital control for the starting and the reverse of Leonhard drive system with SCR excitation circuit. It was found that for this type of drive system the linear algorithm of controller action may prove sufficient. The proposed algorithm belongs to the class giving a finite time of transient process and null fixed error of reference speed. It is shown that both the maximum derivative and the pulsation of the starting current are within allowable limits. Convenient methods of realisation of the controller transmittance were discussed generally.

K. MIĄCZYŃSKI

ALGORITHME DE LA COMMANDE NUMÉRIQUE DIRECTE
DU SYSTÈME DE LEONARD

Résumé

Dans l'article on a présenté la conception de la commande numérique directe du démarrage et du détour du système de Leonard avec excitation à l'aide d'un convertisseur statique. On a constaté que, pour la transmission de ce type, l'algorithme linéaire d'action du régulateur de vitesse peut se montrer suffisant. L'algorithme proposé appartient à la catégorie qui donne les régimes transitoires de la durée finie et l'erreur zéro fixée pour la vitesse donnée. On a montré que la dérivée maximale ainsi que les pulsations du courant de démarrage se renferment dans les limites admissibles. On a discuté les méthodes, généralement appliquées, de la réalisation discrète de la transmittance du régulateur.

K. MIĄCZYŃSKI

ALGORITHMUS DER ABTASTSTEUERUNG EINES LEONARD-ANTRIEBES

Zusammenfassung

Ein Leonard-Antrieb mittlerer Leistung mit Gleichrichter-erregung des Generators und einem digitalen Drehzahlregler wurde untersucht. Unter Voraussetzung, daß die Geschwindigkeit des Stromanstiegs und der Stromamplitude beim Anlauf und Reversieren eines leerlaufenden Motors begrenzt werden, und der Annahme, daß ein Halteglied 0. Ordnung ist, wurde der Steueralgorithmus gewonnen. Es wurde ebenfalls angenommen, daß die Abtastzeit der mechanischen Zeitkonstante B gleicht. Der gewonnene und rechnerisch nachgeprüfte Algorithmus gehört der suboptimalen Klasse mit einer minimalen Zahl der Abtastperioden an. Die interessanten Schlußfolgerungen können beim Entwerfen digitaler Regler für Antriebsregelkreise behilflich sein.

К. МІЇЧЫНСЬКІ

АЛГОРИТМ ДИСКРЕТНОГО УПРАВЛЕННЯ СИСТЕМОЮ ГЕНЕРАТОР-ДВИГАТЕЛЬ

Резюме

Представлена концепция непосредственного цифрового управления пуском в ход и реверсом системы генератор-двигатель с преобразовательным возбуждением. Констатировано, что для привода этого типа может оказаться достаточным линейный алгоритм действия регулятора скорости. Представленный алгоритм принадлежит к классу дающих конечное время переходных процессов и нулевую ошибку заданной скорости. Показано, что максимальная скорость роста пускового тока и его пульсации находятся в допустимых границах. Продискутировано применимые методы реализации дискретной передаточной функции регулятора.

Analysis of Transformer Tank Shielding Effectiveness

JAN SYKULSKI (ŁÓDŹ)

*Instytut Transformatorów, Maszyn i Aparatów Elektrycznych
Politechniki Łódzkiej*

The paper received 3.8.1979

The electromagnetic and magnetic transformer tank shielding effectiveness is investigated on the basis of two- and three-dimensional leakage field analyses. The shielding effect coefficients are defined for the leakage flux and eddy-current losses. The influence on these coefficients of geometric and matter parameters is also discussed.

1. INTRODUCTION

The preceding papers [2], [3] and [4] describe the vector potential method of two- and three-dimensional analyses of electromagnetic field due to currents in parallel conductors with two-sided solid metallic structures. Although the configuration of the problem is simple, actually the simplest possible, its analysis is not. However, numerical methods were found helpful to reach the systematic and flexible solutions which can be adapted easily to practical problems, examples of which were also given. Particularly, the system mentioned above may be considered as a basic element of transformer winding.

In this paper the shielding effectiveness of electromagnetic and magnetic shields of transformer tank is investigated theoretically on the basis of electromagnetic field analyses given in [3] and [4]. The results obtained from two- and three-dimensional approach were checked against experiments conducted on the physical models and both the methods were found helpful in the leakage field analysis in transformer. The goal of this paper is the extension of calculations to such practical arrangements for which experimental results are not available or difficult to obtain. Once verified the methods are hoped to be efficient in any case. Results presented in this paper may be treated as an auxiliary tool in the design of large transformers.

2. SHIELDING EFFECT COEFFICIENTS

Much useful information on leakage field distribution and eddy-current losses may be presented in a very convenient form by using the coefficients expressing the shielding effect for magnetic flux and power losses in transformer tank ([6], [7], [8], [9], [10]).

Consider the system presented in Fig. 1 as a basic model of transformer winding. It can be easily seen that the electromagnetic field is symmetrical to the $y = 0$ plane, if the co-ordinate system is employed as in the Fig. 1. The solution for the field in all regions examined, due to currents in bar conductors, has been given in [3] and [4] (two- and three-dimensional solutions respectively). It is convenient here to enter the synthetic treatment of shielding effectiveness not only by making a comparison of leakage field distributions in various arrangements but using the *shielding effect coefficients* based upon the foregoing field analyses.

The following coefficients were found helpful in the shielding problem investigation.

1. *Magnetic flux shielding effect coefficient*

$$p_{\Phi} = \frac{\Phi_k(y=0)_{\text{with shield}}}{\Phi_k(y=0)_{\text{without shield}}}, \quad (1)$$

which is seen to describe the decreased value of magnetic flux in tank due to the application of the shield. In a 3-dimensional field p_{Φ} is a function of z .

2. *Magnetic flux local shielding effect coefficient*

$$p_{\Phi \text{ local}} = \frac{\{\Phi_k(y, z)_{\max}\}_{\text{with shield}}}{\{\Phi_k(y, z)_{\max}\}_{\text{without shield}}}, \quad (2)$$

which expresses the ratio of the maximum value of unitary magnetic flux in a tank with shield to the one without shield. In a general case the points of maximum values of unitary flux may be different for both configurations.

3. *Power loss shielding effect coefficient*

$$p_m = \frac{P_k(z)_{\text{with shield}}}{P_k(z)_{\text{without shield}}}, \quad (3)$$

which shows the variation of unitary power loss in tank after any shield has been used. In a 3-dimensional field p_m is a function of z .

4. *Power loss local shielding effect coefficient*

$$p_{m \text{ local}} = \frac{\{\Delta P_k(y, z)_{\max}\}_{\text{with shield}}}{\{\Delta P_k(y, z)_{\max}\}_{\text{without shield}}}, \quad (4)$$

which shows that the maximum value of power density in tank changes after one has used a shield. The maximum of power density need not occur in the same place in various configurations.

5. *Total power loss change coefficient*

$$p_{m \text{ total}} = \frac{\{P_{\text{total}}\}_{\text{with shield}}}{\{P_{\text{total}}\}_{\text{without shield}}}, \quad (5)$$

which expresses the change of power losses in the whole system (tank and shield) after employment of a shield.

The set of coefficients defined by eqns (1) ÷ (5) enables to compare the effectiveness of various shields in very synthetic and convenient form. The application of this treatment to practical arrangements is given in the next chapter.

3. NUMERICAL RESULTS

In the following discussion the influence of matter and geometric parameters on the shielding effect coefficients is presented. Consider the basic configuration and background data of the system as shown in Fig. 1.

The calculations were made by using the field solutions and computer programs described in [3] and [4]. The effect of a number of parameters was examined separately. The numerical results are presented in Fig. 2, 3, 4, 5, 6 and 7 (for which we assume: $a - \gamma_{\text{emgnt}} = 2,2 \cdot 10^6 \text{ S/m}$, $b - \gamma_{\text{emgnt}} = 0$).

First it is to be said, that both two-dimensional ([3]) and three-dimensional ([4]) analyses give approximate results. The comparison between two- and three-dimensional approaches given in [4] led to the conclusion, that in the analysis of leakage field in the cylindrical part of transformer tank (but not in the flat part), appropriate results may be obtained from any of these methods. The maximum values of magnetic field quantities occur in this region of tank, thus the values of shielding effect coefficients calculated from both methods should be and are equal.

We note now the following properties of results. It follows from [3] that if the laminated magnetic shield of infinite extent in y and z directions is assumed in calculations, the normal component of magnetic field strength H on the shield surface and magnetic flux in the shield are in good agreement with measurements, while the tangential component of H and power losses in the shield are not correct. The results obtained in this paper confirm the latter conclusion. Fig. 3 shows the strong effect of the assumed conductivity of the shunt on the total power loss change coefficient (eqn. 5) and neither $\gamma_e = 2,2 \cdot 10^6 \text{ S/m}$ (since $p_{m\text{total}} > 1$) nor $\gamma_e = 0$ (since $p_{m\text{total}} \approx 0$) gives correct power loss in the magnetic shield. Now, it can only be said, that the effective conductivity of magnetic shield should be selected for calculations between these two limit values. As far as we do not know

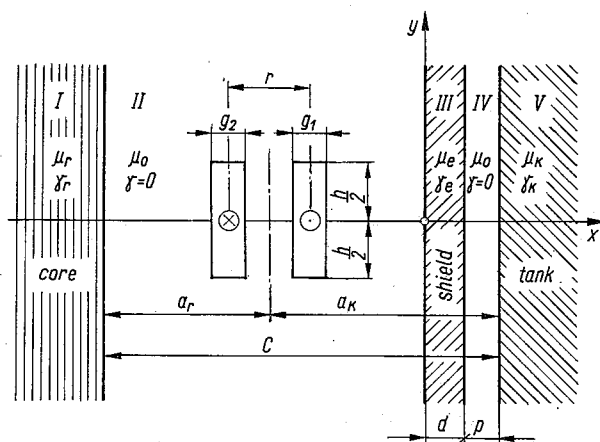


Fig. 1. The basic model of transformer windings for leakage field evaluation [3] (or the three-dimensional plane model — see [2] and [4])

Background data for calculations: $h = 2 \text{ m}$, $g_1 = 0,15 \text{ m}$, $g_2 = 0,2 \text{ m}$, $r = 0,3 \text{ m}$, $C = 0,8 \text{ m}$, $a_k = 0,45 \text{ m}$, $a_r = 0,35 \text{ m}$, $\gamma_{Fe} = 7 \cdot 10^6 \text{ S/m}$, $\gamma_r = \gamma_{\text{emgnt}} = 2,2 \cdot 10^6 \text{ S/m}$ or $\gamma_r = \gamma_{\text{emgnt}} = 0$, $\gamma_{Cu} = 56 \cdot 10^6 \text{ S/m}$, $\gamma_{Al} = 35 \cdot 10^6 \text{ S/m}$, $\mu_r = 3000 \cdot 4\pi \cdot 10^{-7} \text{ H/m}$, $\mu_{\text{emgnt}} = 4000 \cdot 4\pi \cdot 10^{-7} \text{ H/m}$, $\mu_k = 200 \cdot 4\pi \cdot 10^{-7} \text{ H/m}$, $I_{m1} = I_{m2}$; for electromagnetic shield $d = 3 \cdot 10^{-3} \text{ m}$, $p = 0$, for magnetic shield $d = 2,1 \cdot 10^{-3} \text{ m}$ (6 plates), $p = 0$

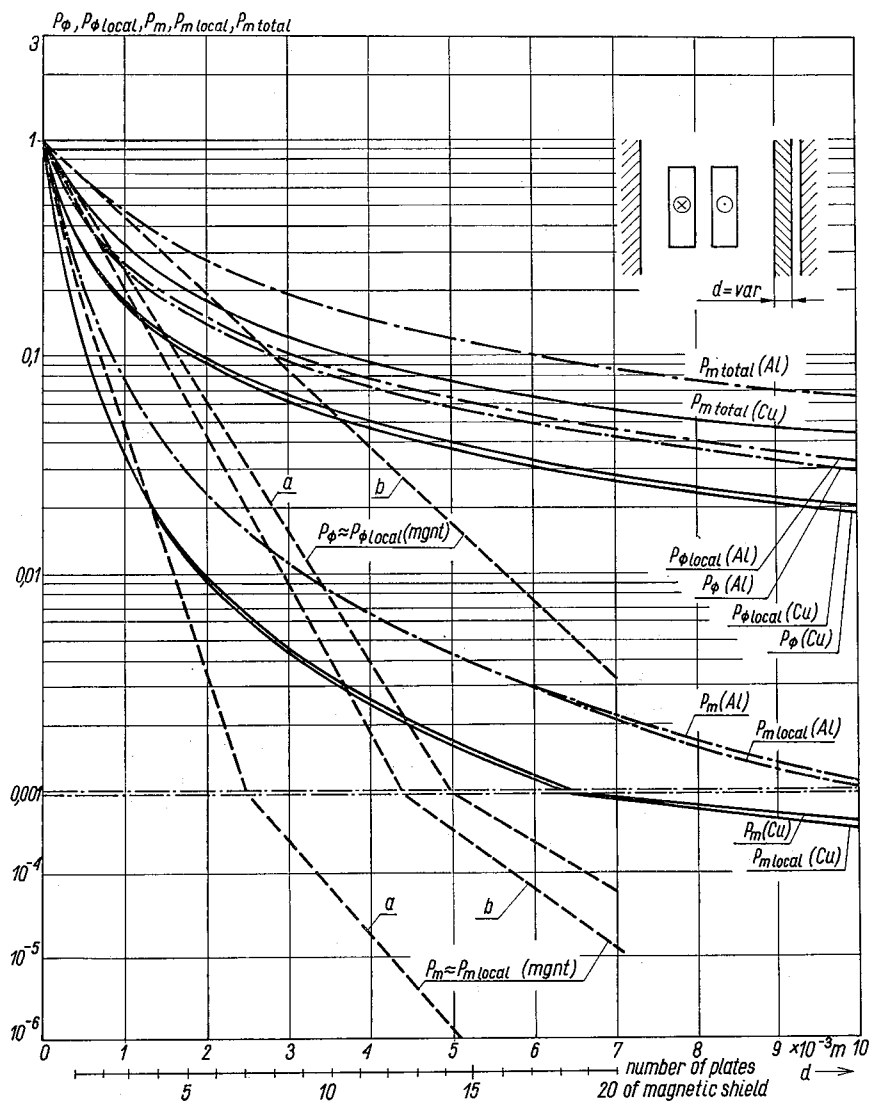


Fig. 2. The variation of shielding effect coefficients with the thickness of shield

how to estimate this effective conductivity, the coefficient $p_{m total}$ for magnetic shield must not be used and is not presented in other figures.

Several additional comments and detailed suggestions follow from these results, which may be summarized in a following conclusions.

1. The shielding effectiveness increases with the growing value of thickness of the shield which mainly affects the power loss and magnetic flux in transformer tank.

2. The local shielding effect coefficients $p_{\phi local}$ and $p_{m local}$ assume the values not much lower than coefficients p_ϕ and p_m , respectively (divergences less than 5 percent). Therefore no relatively increasing power density or unitary magnetic flux occur in the tank after employment of a full shield (of infinite extent in y and z directions).

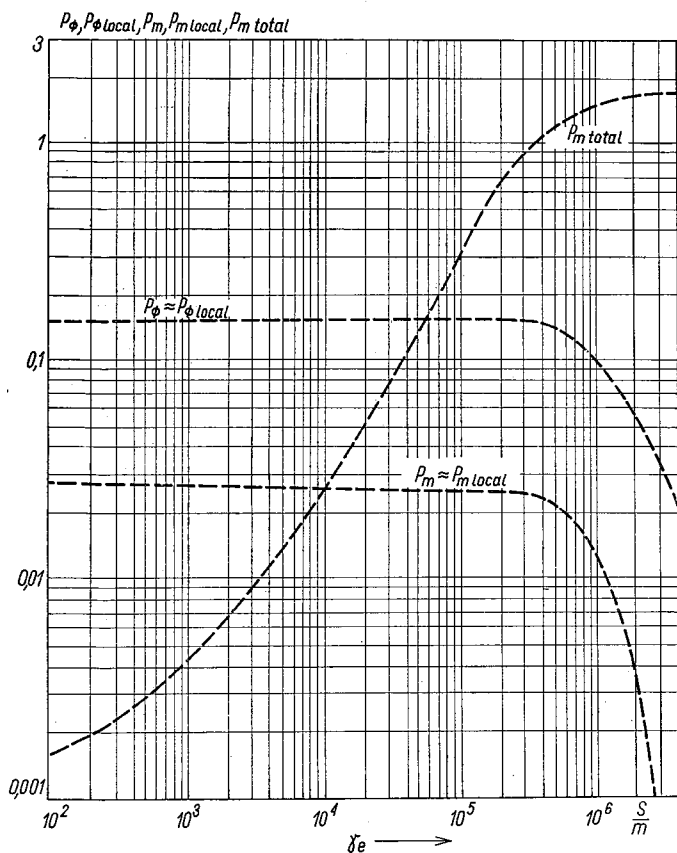


Fig. 3. The variation of shielding effect coefficients with the conductivity of magnetic shield

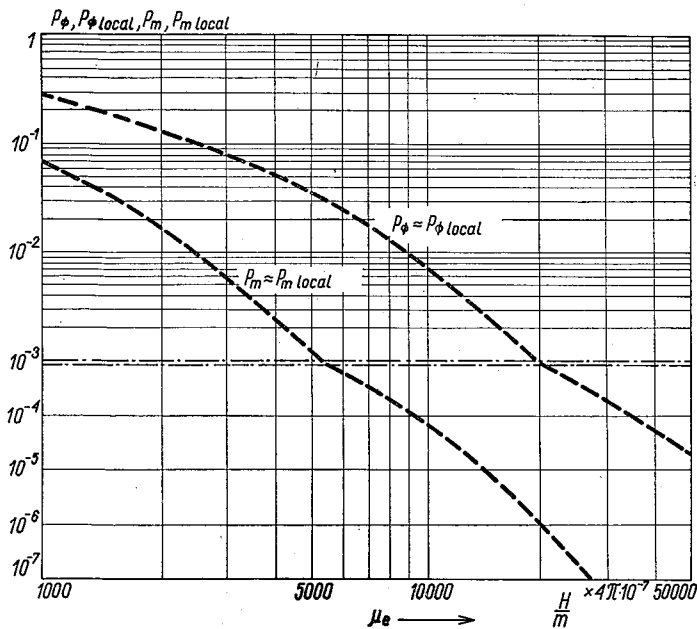


Fig. 4. The variation of shielding effect coefficients with the permeability of magnetic shield
 $\gamma_{\text{emgnt}} = 2.2 \cdot 10^6 \text{ S/m}$

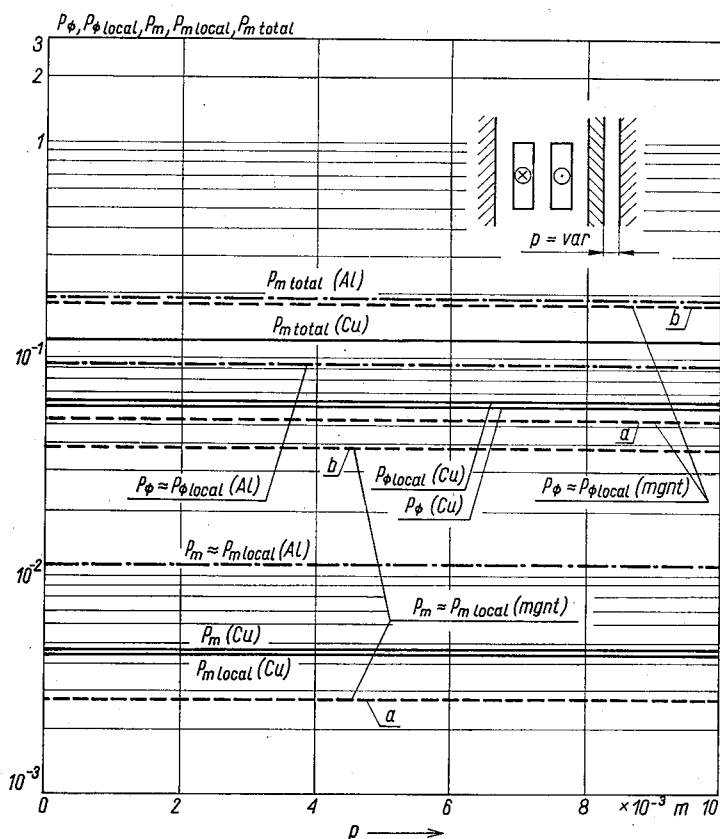


Fig. 5. The variation of shielding effect coefficients with the thickness of insulation gap between shield and tank

3. The magnetic shield functions more effectively than electromagnetic ones (Cu or Al), if the same shield thickness is assumed, except very small values of d . The larger shield thickness, the better magnetic shield as against Cu or Al ones.

4. To compare Cu shield to Al one, assume the same value of d . The Cu shield ensures about 1,6-fold reduction of power loss and 2,5-fold reduction of power density in tank. Then, the variation of shield thickness influences the efficiency of both shields in the same extent.

5. The shielding effect coefficients decrease with increasing permeability of magnetic shield. Therefore the shielding effectiveness varies over a very wide range due to the saturation of magnetic shield.

6. The effect of the insulation gap between shield and tank on the power losses and magnetic flux in tank is very slight and may be neglected.

7. The shielding effectiveness varies both with the distance between tank and core as well as with the height of exciting windings. However this decreasing (for magnetic shield) or increasing (for Cu shield) tendency can be observed only if the tank is relatively

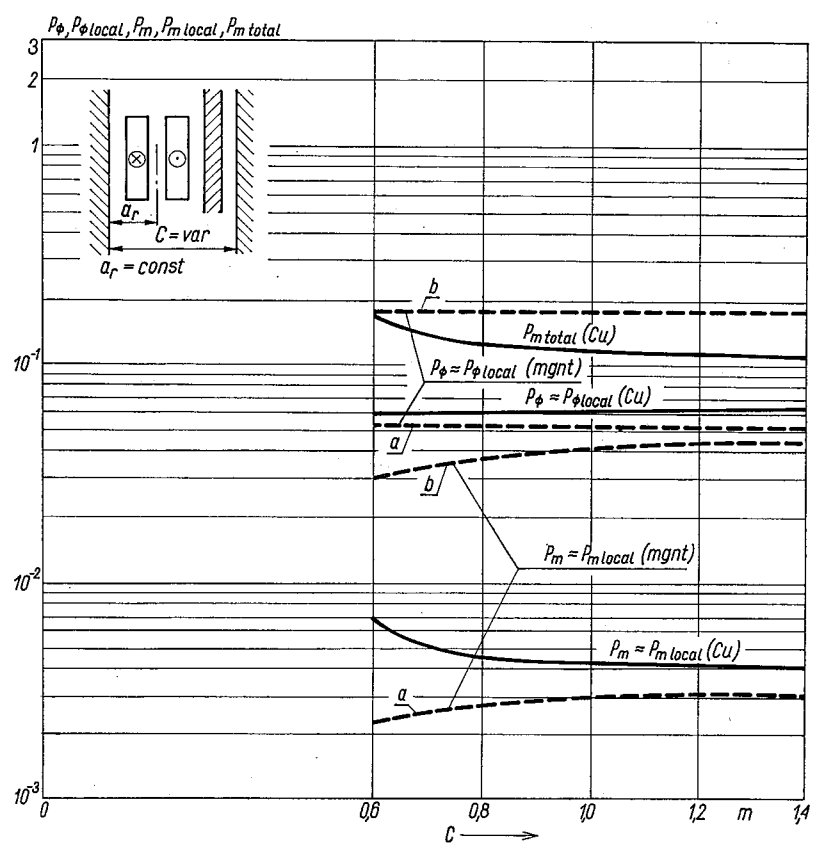


Fig. 6. The variation of shielding effect coefficients with the distance between tank and core

close to windings i.e. the C or h dimensions are small enough, which is beyond the limits of practical interest.

8. The effects of many additional parameters were also investigated. It was found however, that they were negligible. Nevertheless in any practical application the effects of all system variables may be examined by using this numerical methods presented in this and preceding papers. These methods were used here to obtain the large number of data for the comparisons. The relationships shown in this paper are to hint to the designer, which parameters mainly affect the power loss in transformer tank and how to estimate the effectiveness of magnetic and electromagnetic shielding.

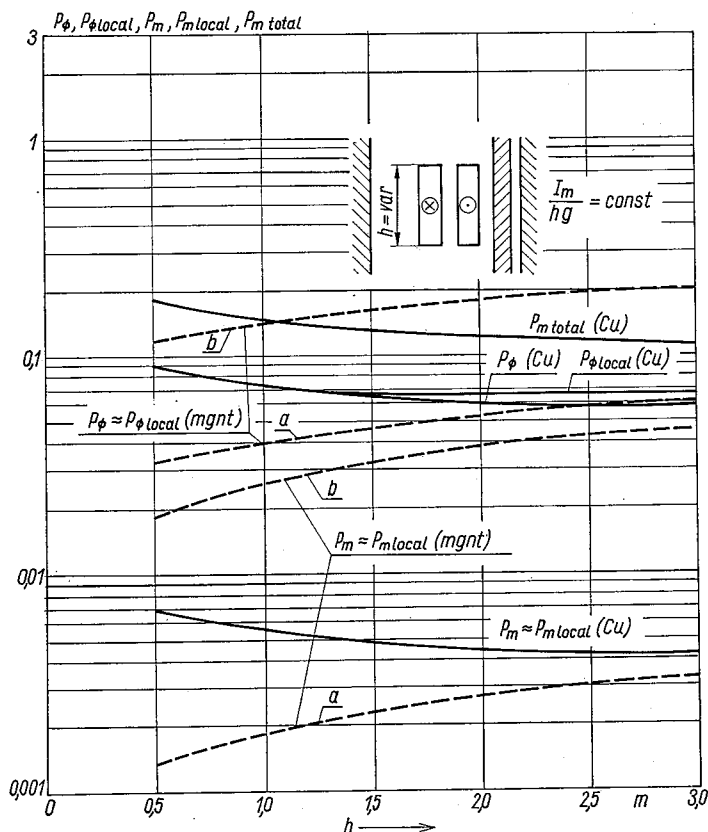


Fig. 7. The variation of shielding effect coefficients with the height of exciting windings

REFERENCES

1. J. Sykulski, *Obliczanie pola elektromagnetycznego i strat w układach transformatorowych ekranowanych*, praca doktorska, ITMA PL, Łódź, 1978.
2. J. Sykulski, *Electromagnetic Field Due to A.C. Flowing Through Finite-Length Conductors in the Presence of Semi-Infinite Solid Shield*, *Rozprawy Elektrotechniczne*, t. XXV, z. 1, 1979.
3. J. Sykulski, *Two-Dimensional Analysis of Electromagnetic Field Distribution in Solid Metallic Structures in Vicinity of Current Carrying Bar Conductors*, *Rozprawy Elektrotechniczne*, 27, z. 1, 1981.
4. J. Sykulski, *Three-Dimensional Analysis of Electromagnetic Field Distribution in Solid Metallic Structures in Vicinity of Current Carrying Finite-Length Bar Conductors*, *Rozprawy Elektrotechniczne*, 27, z. 2, 1981.
5. J. Turowski, *Elektrodynamika techniczna*, WNT, Warszawa 1968.
6. K. Zakrzewski, J. Turowski, *Metody analizy ekranowania magnetycznego kadzi transformatorów*, *Rozprawy Elektrotechniczne*, t. XVIII, z. 3, 1972.
7. P. Jezierski, *Analiza skuteczności ekranowania elektromagnetycznego silnych prądumiennych pól magnetycznych*, *Rozprawy Elektrotechniczne*, t. XVI, z. 1, 1970.
8. M. Kaźmierski, *Przenikalność zastępcza masywnego żelaza przy częstotliwości technicznej*, *Rozprawy Elektrotechniczne*, t. XXI, z. 2, 1975.
9. S. Apanasewicz, M. Kaźmierski, *Rozkład pola elektromagnetycznego i strat na powierzchni częściowo ekranowanej kadzi transformatora*, *Archiwum Elektrotechniki*, t. XXVI, z. 2, 1977.
10. S. Apanasewicz, *Metody analizy elektromagnetycznego pola rozproszenia w transformatorach i strat mocy w różnych przypadkach częściowego ekranowania ich kadzi*, *Prace IEL*, z. 104.

J. SYKULSKI

ANALIZA SKUTECZNOŚCI EKRANOWANIA KADZI TRANSFORMATORÓW

Streszczenie

W pracy zbadano skuteczność ekranowania elektromagnetycznego i magnetycznego kadzi transformatorów wykorzystując analizę dwu- i trójwymiarową pola rozproszenia. Zdefiniowano współczynniki zaekranowania strumienia magnetycznego i strat mocy oraz przedstawiono ich zależność od parametrów fizycznych i geometrycznych układu.

J. SYKULSKI

ANALYSE DE L'EFFICACITÉ DU BLINDAGE DE LA CUVE
DES TRANSFORMATEURS

Résumé

On a présenté l'efficacité du blindage électromagnétique de la cuve des transformateurs, en se servant de l'analyse bi- et tridimensionnelle du champ de fuite. On a défini les coefficients du blindage du flux magnétique et des pertes, en démontrant leur dépendance des paramètres physiques et géométriques du système.

J. SYKULSKI

ZUR ANALYSE DER EFFEKTIVITÄT DER ABSCHIRMUNG
EINES TRANSFORMATOR-KESSELS

Zusammenfassung

Es wurde die Effektivität einer elektromagnetischen und magnetischen Abschirmung geprüft unter Anwendung einer zwei- bzw. dreidimensionalen Analyse des Streufeldes. Die Abschirmungskoeffizienten des magnetischen Flusses und der Leistungsverluste in der Funktion der physikalischen und geometrischen System-Parameter werden beschrieben.

Я. СЫКУЛЬСКИ

АНАЛИЗ ЭФФЕКТИВНОСТИ ЭКРАНИРОВАНИЯ БАКОВ ТРАНСФОРМАТОРОВ

Резюме

Обсуждается эффективность магнитного и электромагнитного экранирования баков трансформаторов при использовании двух- или трехмерного анализа поля рассеяния. Определены коэффициенты экранирования для магнитного потока и потерь мощности. Исследована зависимость этих коэффициентов от физических параметров и геометрической конфигурации системы.

Składowe aktywne i zerowe prądów w obwodach elektromagnetycznych

KRZYSZTOF KLUSZCZYŃSKI (GLIWICE)

Instytut Maszyn i Urządzeń Elektrycznych Politechniki Śląskiej

Otrzymano 8.11.1979

Na podstawie teorii przekształceń liniowych w przestrzeniach wektorowych wykazano, że w liniowych obwodach elektromagnetycznych można wyodrębnić składowe aktywne i zerowe prądów o istotnym znaczeniu z punktu widzenia przemian energetycznych (przemiany energii elektrycznej, magnetycznej i mechanicznej). Przedstawiono sposób wyznaczania składowych zerowych w obwodach elektromagnetycznych o dowolnie złożonej strukturze topologicznej, a w szczególności w maszynach asynchronicznych. Sformułowano warunki eliminacji, bądź minimalizacji składowej zerowej prądów. Wprowadzono nowy, k -osiowy układ współrzędnych, w którym równania różniczkowe stanu nieustalonego obwodów elektromagnetycznych znacznie się upraszczają.

1. WSTĘP

Rdzeń liniowego obwodu elektromagnetycznego, którego stan magnetyczny związany ze strumieniem głównym analizować można w oparciu o I i II prawo Kirchhoffa, nazywać będziemy siecią magnetyczną. Uzwojenia elektryczne, znajdujące się na rdzeniu obwodu magnetycznego, mogą być galwanicznie niezależne lub też połączone, tworząc sieć elektryczną. W ogólnym przypadku obwód elektromagnetyczny stanowi więc układ dwóch wzajemnie powiązanych sieci: elektrycznej, odpowiedzialnej za wywoływanie strumienia magnetycznego oraz magnetycznej, przewodzącej strumień magnetyczny (rys. 1). Stan nieustalony obwodu elektromagnetycznego, którego rdzeń stanowi spójną sieć magnetyczną, a n uzwojeń rozmieszczonych jest pojedynczo na dowolnych gałęziach, opisuje w naturalnym układzie współrzędnych układ równań różniczkowych:

$$\begin{bmatrix} u_1 \\ u_2 \\ \dots \\ u_n \end{bmatrix} = \begin{bmatrix} R_1 & & 0 \\ & R_2 & \\ & & \ddots \\ 0 & & R_n \end{bmatrix} \begin{bmatrix} i_1 \\ i_2 \\ \vdots \\ i_n \end{bmatrix} + \frac{d}{dt} \begin{bmatrix} L_{s1} & & 0 \\ & L_{s2} & \\ & & \ddots \\ 0 & & L_{sn} \end{bmatrix} \begin{bmatrix} i_1 \\ i_2 \\ \vdots \\ i_n \end{bmatrix} + \frac{d}{dt} \begin{bmatrix} \psi_1 \\ \psi_2 \\ \vdots \\ \psi_n \end{bmatrix}, \quad (1)$$

$$\begin{bmatrix} \psi_1 \\ \psi_2 \\ \vdots \\ \psi_n \end{bmatrix} = \begin{bmatrix} M_{11} & \dots & M_{1n} \\ M_{21} & \dots & M_{2n} \\ \dots & \dots & \dots \\ M_{n1} & \dots & M_{nn} \end{bmatrix} \begin{bmatrix} i_1 \\ i_2 \\ \vdots \\ i_n \end{bmatrix}, \quad (2)$$

gdzie:

u_i, i_i, ψ_i — napięcie, prąd i strumień główny skojarzony z i -tym uzwojeniem,

R_i, L_{si} — rezystancja i indukcyjność rozproszenia i -tego uzwojenia,

M_{ij} — indukcyjność wzajemna i -tego i j -tego uzwojenia.

Obliczając rząd macierzy indukcyjności głównych $[M]$ w równaniu (2) dla różnych obwodów elektromagnetycznych można stwierdzić, że $\text{rz}[M] \leq n$. Nierówność ta ma fundamentalne znaczenie dla zapewnienia optymalnych warunków przetwarzania energii w obwodach elektromagnetycznych oraz w przetwornikach elektromechanicznych i określa warunki wzajemnego dopasowania sieci elektrycznej i magnetycznej, a w szczególności — doboru struktury topologicznej obu sieci, rozkładu uzwojeń na gałęziach sieci magnetycznej oraz liczby zwojów poszczególnych uzwojeń. Odpowiedź na to, od czego zależy rząd macierzy indukcyjności głównych obwodu elektromagnetycznego oraz jaką rolę odgrywa w przemianach energetycznych, można sformułować na gruncie teorii przekształceń liniowych w przestrzeniach wektorowych.

2. WYMIAR PRZESTRZENI STRUMIENI SKOJARZONYCH Ψ SIECI MAGNETYCZNEJ

n -wierszowe macierze jednokolumnowe wartości chwilowych przebiegów elektrycznych lub magnetycznych n -uzwojeniowego obwodu elektromagnetycznego o postaci $[v] = [v_1(t) \dots v_n(t)]^T$ ($v_i(t)$ — wartość chwilowa przebiegu związana z i -tym uzwojeniem) traktować możemy jako wektory v rzeczywistej przestrzeni liniowej V . Fizykalnie, wektory przestrzeni V reprezentować mogą napięcia na zaciskach uzwojeń, prądy w uzwojeniach lub też strumienie skojarzone z poszczególnymi uzwojeniami, wyznaczając odpowiednio przestrzenie wektorowe: napięć U , prądów I oraz strumieni skojarzonych Ψ . W przestrzeni wektorowej V wprowadźmy iloczyn skalarny wektorów $v^{(1)} \cdot v^{(2)} \in V$

$$v^{(1)} \cdot v^{(2)} = \sum_{i=1}^n v_i^{(1)} v_i^{(2)} = [v^{(1)}]^T [v^{(2)}] \quad (3)$$

oraz normę $\|v\|$ wektora $v \in V$

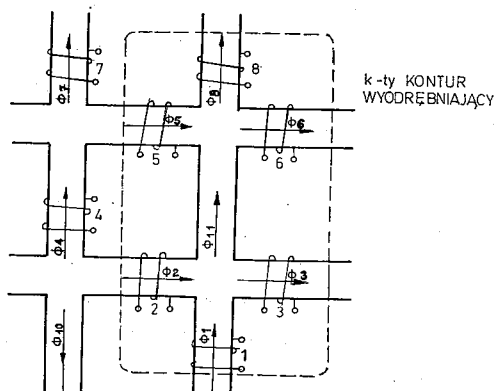
$$\|v\| = \sum_{i=1}^n v_i v_i = [v]^T [v]. \quad (4)$$

Sens fizykalny iloczynu skalarnego jest oczywisty, gdy jednym z elementów jest wektor napięcia $[u]$ a drugim — wektor prądu $[i]$. Wartość iloczynu skalarnego $p(t) = \sum_{i=1}^n u_i(t) i_i(t) = [u]^T [i]$ jest wówczas równa mocy chwilowej, pobieranej przez n uzwojeń obwodu elektromagnetycznego.

Wymiary przestrzeni wektorowych: napięć U i prądów I są określone przez graf sieci elektrycznej. Wykażemy, że podobnie wymiar przestrzeni strumieni skojarzonych Ψ wyznaczony jest przez graf sieci magnetycznej i dodatkowo — przez rozłożenie uzwojeń. Rozważmy płaską i spójną sieć magnetyczną o g_m gałęziach, w_m węzłach i n uzwojeniach, rozmieszczonych pojedynczo na dowolnych gałęziach sieci. Strumienie magnetyczne w gałęziach rdzenia z nawiniętymi uzwojeniami nie mogą być dowolne: przyjmować muszą

wartości zgodne z równaniami ułożonymi na podstawie I prawa Kirchhoffa dla wszystkich dających się wyodrębnić części sieci magnetycznej, łączących się z resztą obwodu magnetycznego wyłącznie poprzez gałęzie z nawiniętymi uzwojeniami. Oznaczając przez z_i liczbę zwojów i -tego uzwojenia otrzymujemy dla wyodrębnionej części sieci magnetycznej z rys. 1 równanie

$$\left[\frac{1}{z_1} \frac{1}{z_2} 0 - \frac{1}{z_4} 0 - \frac{1}{z_6} \frac{1}{z_7} 0 \frac{1}{z_9} 0 \dots 0 \right] [\psi_1 \psi_2 \dots \psi_n]^T = 0. \quad (5)$$



Rys. 1. Fragment obwodu elektromagnetycznego z k -tym konturem wyodrębniającym

Równanie (5) wyznacza hiperpłaszczyznę w przestrzeni V^n , która jest zbiorem wszystkich wektorów strumieni skojarzonych ortogonalnych do wektora $\left[\frac{1}{z_1} \frac{1}{z_2} 0 - \frac{1}{z_4} 0 - \frac{1}{z_6} \frac{1}{z_7} 0 \frac{1}{z_9} 0 \dots 0 \right]^T$. Jeśli liczba niezależnych liniowo równań, ułożonych dla kolejnych, dających się wyodrębnić części sieci magnetycznej wynosi $n-k$, to wymiar przestrzeni Ψ , będącej przecięciem wszystkich $n-k$ hiperpłaszczyzn, odpowiadających $n-k$ równaniom wynosi k

$$\dim \Psi = n - (n - k) = k. \quad (6)$$

Zasadniczą rolę odgrywa liczba uzwojeń w stosunku do liczby strun sieci magnetycznej s_m . Można wykazać, że jeśli:

$$\begin{aligned} n < s_m & \quad \dim \Psi \leq n, \\ n = s_m & \quad \dim \Psi \leq n, \\ s_m > n > s_m & \quad \dim \Psi \leq s_m. \end{aligned} \quad (7)$$

Wymiar przestrzeni Ψ nigdy nie przewyższa liczby strun s_m . Nierówności (7) przechodzą w równości tylko wówczas, gdy istnieje drzewo sieci magnetycznej, na którego strunach eżą wszystkie uzwojenia sieci elektrycznej.

3. RZĄD MACIERZY INDUKCYJNOŚCI GŁÓWNYCH

Macierz indukcyjności głównych $[M]$ w układzie równań algebraicznych (2) traktować można jako macierz pewnego przekształcenia liniowego M , przeprowadzającego wektory przestrzeni prądów I w wektory przestrzeni Ψ

$$M(i) = \psi, \quad i \in I, \quad \psi \in \Psi. \quad (8)$$

Wymiar przestrzeni Ψ , będącej obrazem przekształcenia M , nazywamy rzędem przekształcenia M ($\text{rz } M$), a więc $\dim \Psi = \text{rz } M$. Można wykazać, że rząd przekształcenia liniowego M równy jest rzędowi macierzy przekształcenia $[M]$, czyli

$$\dim \Psi = \text{rz } [M]. \quad (9)$$

Równość (9) ma ważne znaczenie w teorii liniowych obwodów elektromagnetycznych. Wynika z niej, że rząd macierzy indukcyjności głównych $[M]$ jest cechą topologiczną sieci magnetycznej, a więc cechą zależną od grafu sieci magnetycznej i rozłożenia uzwojeń, a niezależną od reluktancji magnetycznych gałęzi sieci magnetycznej i liczby zwojów poszczególnych uzwojeń. Rząd macierzy $[M]$ jest określony nierównościami (7). Zbiór wektorów, których obrazem w przekształceniu M jest wektor zerowy, a więc takich wektorów i_0 , dla których $M(i_0) = 0$, nosi nazwę przestrzeni zerowej I_0 przekształcenia M . Elementy przestrzeni zerowej I_0 nazywamy wektorami zerowymi. Wymiar przestrzeni zerowej I_0 , zwany zerowością, oznaczamy przez $\text{def } M$. Z twierdzenia, obowiązującego dla przestrzeni liniowej o skończonym wymiarze n : $\text{def } M + \text{rz } M = n$ wynika, że: $\dim I_0 = n - \text{rz } [M] = n - k$. Stąd zaś już bezpośrednio uzyskujemy warunek istnienia przestrzeni zerowej I_0^{n-k} : $\text{rz } [M] < n$.

4. ROZKŁAD ORTOGONALNY WEKTORA PRĄDU NA SKŁADOWE: AKTYWNA I ZEROWA

Nierówności (7) są podstawą do podziału liniowych obwodów elektromagnetycznych na dwie rozłączne klasy o odmiennych własnościach energetycznych. Do pierwszej klasy zaliczamy obwody, w których $\text{rz } [M] = n$, do drugiej zaś te, w których $\text{rz } [M] < n$. W obwodach klasy drugiej przestrzeń prądów I^n można przedstawić w postaci sumy prostej dwóch ortogonalnych podprzestrzeni: k -wymiarowej przestrzeni aktywnej $I_\perp^k$ oraz $(n-k)$ -wymiarowej przestrzeni zerowej I_0^{n-k}

$$I^n = I_\perp^k \oplus I_0^{n-k}, \quad I_\perp^k \perp I_0^{n-k}. \quad (10)$$

Przestrzeń aktywna $I_\perp^k$ jest ortogonalnym dopełnieniem przestrzeni zerowej I_0^{n-k} do I^n . Co więcej, z własności symetrii macierzy indukcyjności głównych wynika, że przestrzeń aktywna $I_\perp^k$ jest tożsama z przestrzenią strumieni skojarzonych Ψ

$$I_\perp^k \equiv \Psi^k. \quad (11)$$

Relacja (10) oznacza, że każdy wektor $i \in I^n$ można jednoznacznie przedstawić w postaci sumy dwóch wzajemnie ortogonalnych składowych — aktywnej $i_\perp$ i zerowej i_0 :

$$i = i_\perp + i_0, \quad i_\perp \in I_\perp^k, \quad i_0 \in I_0^{n-k}. \quad (12)$$

Zachodzi: $\|i\|^2 = \|i_\perp\|^2 + \|i_0\|^2$. Stąd: $\|i_\perp\|^2 \leq \|i\|^2$. Ponadto z własności przestrzeni zerowych wynika, że: $\psi = M(i) = M(i_\perp + i_0) = M(i_\perp) + M(i_0) = M(i_\perp)$. W świetle tych relacji oczywisty staje się sens fizyczny składowej aktywnej $i_\perp$.

Składowa aktywna $i_{\perp}$ wektora i jest to wektor prądu o najmniejszej normie, wywołujący w obwodzie elektromagnetycznym identyczny rozkład strumieni skojarzonych głównych jak rzeczywisty wektor prądu i . Składowa zerowa i_0 prądu i jest to wektor prądu, którego przepływ w n uzwojeniach obwodu elektromagnetycznego nie wytwarza wektora strumienia skojarzonego.

Tak więc dany rozkład strumieni skojarzonych wywołuje w obwodach elektromagnetycznych pierwszej klasy jeden i tylko jeden wektor prądu, zaś w obwodach drugiej klasy — istnieje ich nieskończenie wiele. Wytwarzanie jednak wektora ψ nie za pomocą wektora prądu o najmniejszej normie $i_{\perp}$, lecz wektorem prądu $i = i_{\perp} + \lambda i_0$ ($\lambda \in R$, $\lambda \neq 0$), wywołuje dodatkowo straty mocy czynnej na rezystancjach uzwojeń i dodatkowy strumień rozproszenia na indukcyjnościach rozproszonych. Z energetycznego punktu widzenia celowe jest więc wykorzystywanie do przetwarzania energii elektrycznej w magnetyczną i na odwrót obwodów elektromagnetycznych klasy pierwszej. W obwodach elektromagnetycznych drugiej klasy składową zerową można eliminować, bądź minimalizować.

5. WYZNACZANIE SKŁADOWYCH AKTYWNYCH I ZEROWYCH PRĄDU

Wyznaczenie składowych aktywnych i zerowych wektora prądu wymaga wprawdzie znalezienia baz przestrzeni: aktywnej oraz zerowej. Bazy przestrzeni $I_{\perp}^k$ i I_0^{n-k} można wyznaczyć dwoma metodami, bądź to na podstawie znajomości macierzy indukcyjności głównych $[M]$, bądź też — grafu sieci magnetycznej, rozkładu uzwojeń i liczby zwojów poszczególnych uzwojeń.

Jeżeli znamy macierz indukcyjności głównych $[M]$, bazę przestrzeni zerowej znajdujemy poprzez rozwiązanie jednorodnego układu równań

$$\begin{bmatrix} M_{11} & \dots & M_{1n} \\ M_{21} & \dots & M_{2n} \\ \dots & \dots & \dots \\ M_{n1} & \dots & M_{nn} \end{bmatrix} \begin{bmatrix} i_1 \\ i_2 \\ \vdots \\ i_n \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad \text{rz}[M] = k. \quad (13)$$

Można wykazać, że bazę przestrzeni aktywnej stanowi k niezależnych kolumn macierzy $[M]$.

Druga metoda polega na ułożeniu wszystkich $n-k$ niezależnych liniowo równań typu (5). Wówczas wektor $\left[\frac{1}{z_1} \frac{1}{z_2} 0 - \frac{1}{z_4} 0 - \frac{1}{z_6} \frac{1}{z_7} 0 - \frac{1}{z_9} 0 \dots \right]$ wraz z pozostałymi niezależnymi liniowo wektorami normalnymi do przestrzeni Ψ^k wyznaczają bazę przestrzeni zerowej I_0^{n-k} , co wynika z relacji (10) i (11).

Ortogonalną i unormowaną bazę przestrzeni aktywnej oznaczmy przez: $[K_1] = [k_{11} \ k_{12} \dots k_{1n}]^T \dots [K_k] = [k_{k1} \ k_{k2} \dots k_{kn}]^T$, a bazę przestrzeni zerowej przez: $[K_{k+1}] = [k_{k+11} \ k_{k+12} \dots k_{k+1n}]^T \dots [K_n] = [k_{n1} \ k_{n2} \dots k_{nn}]^T$. Połączone bazy przestrzeni $I_{\perp}^k$ oraz I_0^{n-k} stanowią ortonormalną bazę przestrzeni I^n . Układ współrzędnych związany z powyższą bazą będziemy nazywać k -osiowym, a współrzędne wektorów w tym układzie — wyróżniać górnym indeksem (k) . Macierz transformacji z naturalnego do k -osiowego układu współrzędnych ma postać

$$[K] = \begin{bmatrix} k_{11} & \dots & k_{1n} \\ k_{21} & \dots & k_{2n} \\ \dots & \dots & \dots \\ k_{n1} & \dots & k_{nn} \end{bmatrix}. \quad (14)$$

Ortogonalność osi k -osiowego układu współrzędnych zapewnia niezmienniczość mocy chwilowej oraz pozwala w prosty sposób obliczać macierz odwrotną transformacji, a mianowicie: $[K]^{-1} = [K]^T$. Rozkład ortogonalny (10) przyjmuje w k -osiowym układzie współrzędnych postać

$$[i^{(k)}] = \begin{bmatrix} i_1^{(k)} \\ \vdots \\ i_k^{(k)} \\ i_{k+1}^{(k)} \\ \vdots \\ i_n^{(k)} \end{bmatrix} = [i_1^{(k)}] + [i_0^{(k)}] = \begin{bmatrix} i_1^{(k)} \\ \vdots \\ i_k^{(k)} \\ 0 \\ \vdots \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ \vdots \\ 0 \\ i_{k+1}^{(k)} \\ \vdots \\ i_n^{(k)} \end{bmatrix}. \quad (15)$$

Widać, że pierwsze k współrzędnych wyznacza składową aktywną wektora prądu i_1 zaś pozostałe $n-k$ współrzędnych — składową zerową i_0 . W wyniku transformacji odwrotnej otrzymujemy składowe aktywne i zerowe w naturalnym układzie współrzędnych

$$\begin{aligned} [i_1] &= [K]^T [i^{(k)}], \\ [i_0] &= [K]^T [i_0^{(k)}]. \end{aligned} \quad (16)$$

6. ZASTOSOWANIE TRANSFORMACJI k -OSIOWEJ DO ANALIZY OBWODÓW ELEKTROMAGNETYCZNYCH

Analizę obwodów elektromagnetycznych, spełniających nierówność $\text{rz}[M] < n$ wygodnie jest prowadzić w k -osiowym układzie współrzędnych, w którym formalny matematyczny zapis układu równań różniczkowych stanu nieustalonego znacznie upraszcza się. Można wykazać, że na skutek niezmienniczości przestrzeni I_1^k oraz I_0^{n-k} równania macierzowe (2) przyjmuje w k -osiowym układzie współrzędnych postać

$$[\psi^{(k)}] = [M^{(k)}][i^{(k)}], \quad (17)$$

gdzie:

$$[M^{(k)}] = [K][M][K]^T = \left[\begin{array}{c|c} \begin{matrix} M_{11}^k & \dots & M_{1k}^k \\ \dots & \dots & \dots \\ M_{k1}^k & \dots & M_{kk}^k \end{matrix} & \begin{matrix} [0] \\ k(n-k) \end{matrix} \\ \hline \begin{matrix} [0] \\ (n-k)k \end{matrix} & \begin{matrix} [0] \\ (n-k)(n-k) \end{matrix} \end{array} \right].$$

Jeśli k -osiowy układ współrzędnych jest zgodny z układem współrzędnych o bazie złożonej z wektorów własnych, wówczas macierz $[M^{(k)}]$ ulega dalszemu uproszczeniu, a mianowicie jej pełna podmacierz w lewym górnym rogu przyjmuje postać diagonalną. W wielofazowych symetrycznych obwodach elektromagnetycznych, w których rezystancje i indukcyjności rozprożeń poszczególnych uzwojeń są równe, takie uproszczenie struktury macierzy indukcyjności głównych prowadzi do rozsprężenia układu równań różniczkowych stanu nieustalonego.

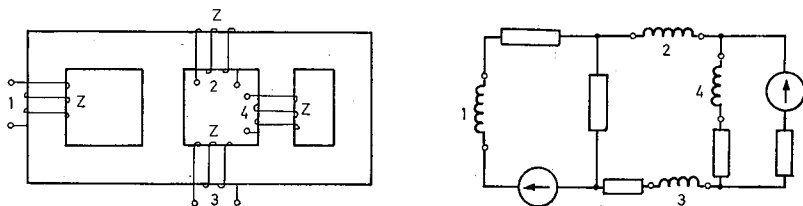
7. METODY ELIMINACJI SKŁADOWEJ ZEROWEJ WEKTORA PRĄDU

Układ równań różniczkowych (1) i (2) przyjmuje w k -osiowym układzie współrzędnych postać:

$$\begin{bmatrix} u_{\perp}^{(k)} \\ u_0^{(k)} \end{bmatrix} = \begin{bmatrix} R^{(k)} \end{bmatrix} \begin{bmatrix} i_{\perp}^{(k)} \\ i_0^{(k)} \end{bmatrix} + \frac{d}{dt} \begin{bmatrix} L_s^{(k)} \end{bmatrix} \begin{bmatrix} i_{\perp}^{(k)} \\ i_0^{(k)} \end{bmatrix} + \frac{d}{dt} \begin{bmatrix} \psi_{\perp}^{(k)} \\ 0 \end{bmatrix}, \quad (18)$$

$$\begin{bmatrix} \psi_{\perp}^{(k)} \\ 0 \end{bmatrix} = \begin{bmatrix} M_{\perp}^{(k)} & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} i_{\perp}^{(k)} \\ i_0^{(k)} \end{bmatrix}. \quad (19)$$

Warunki eliminacji składowej zerowej wektora prądu wynikają wprost z układu równań (18) i (19). W symetrycznych n -fazowych obwodach elektromagnetycznych, w których macierze $[R^{(k)}]$ i $[L_s^{(k)}]$ są diagonalne wystarcza, aby składowa u_0 była równa zero. W niesymetrycznych obwodach elektromagnetycznych wartość składowej zerowej określona jest zarówno przez składową zerową u_0 , jak i aktywną $u_{\perp}$, a więc może pojawić się i wówczas, gdy u_0 będzie równe zero. Jej wartość można wtenczas ograniczyć przez odpowiedni dobór parametrów obwodu.



Rys. 2. Obwód elektromagnetyczny i sposób połączenia jego uzwojeń, zapewniający eliminację prądów zerowych

W szczególnie prosty sposób można eliminować składową zerową wektora prądu z uzwojeń obwodów elektromagnetycznych o jednakowej liczbie zwojów, łącząc niezależne galwanicznie uzwojenia w sieć elektryczną o grafie, identycznym z grafem sieci magnetycznej. Przynależność wektorów prądów do przestrzeni aktywnej $I_{\perp}^k$ wymuszają wówczas węzły sieci elektrycznej. W obwodzie elektromagnetycznym przedstawionym na rys. 2 wymiar przestrzeni strumieni skojarzonych wynosi 3, co oznacza, że przestrzeń aktywna jest 3-wymiarowa, zaś przestrzeń zerowa 1-wymiarowa. Wektor bazowy przestrzeni zerowej ma postać: $\begin{bmatrix} 0 - \frac{1}{\sqrt{2}} \frac{1}{\sqrt{2}} 0 \end{bmatrix}^T$. Eliminację składowej zerowej prądu zapewnia przykładowo połączenie uzwojeń w sieć, przedstawioną na rys. 2, przy czym bez znaczenia jest to, jakie elementy wchodzi w skład poszczególnych gałęzi obwodu oraz jakimi źródłami napięcia obwód jest zasilany.

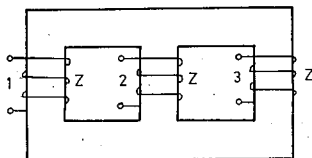
8. INTERPRETACJA GEOMETRYCZNA SKŁADOWYCH AKTYWNYCH I ZEROWYCH PRĄDU

Interpretację geometryczną wprowadzonych pojęć przedstawimy na przykładzie symetrycznego 3-fazowego obwodu elektromagnetycznego z rys. 3 w 3-wymiarowej przestrzeni afinicznej. Zakładamy jednakową przewodność magnetyczną Λ wszystkich kolumn

obwodu elektromagnetycznego oraz jednakową liczbę zwojów z poszczególnych uzwojeń. Macierz indukcyjności głównych jest symetryczna i cykliczna

$$[M] = z^2 A \begin{bmatrix} 1 & -\frac{1}{2} & -\frac{1}{2} \\ -\frac{1}{2} & 1 & -\frac{1}{2} \\ -\frac{1}{2} & -\frac{1}{2} & 1 \end{bmatrix}. \quad (20)$$

Rząd $[M]$ wynosi 2. Bazę przestrzeni aktywnej $I_{\perp}^2$ stanowią dwa dowolne wektory kolumny macierzy (20): $\left[1 - \frac{1}{2} - \frac{1}{2}\right]^T$, $\left[-\frac{1}{2} 1 - \frac{1}{2}\right]^T$. Bazę przestrzeni zerowej I_0^1 wyznaczamy, rozwiązując układ równań algebraicznych (13). W wyniku otrzymujemy:

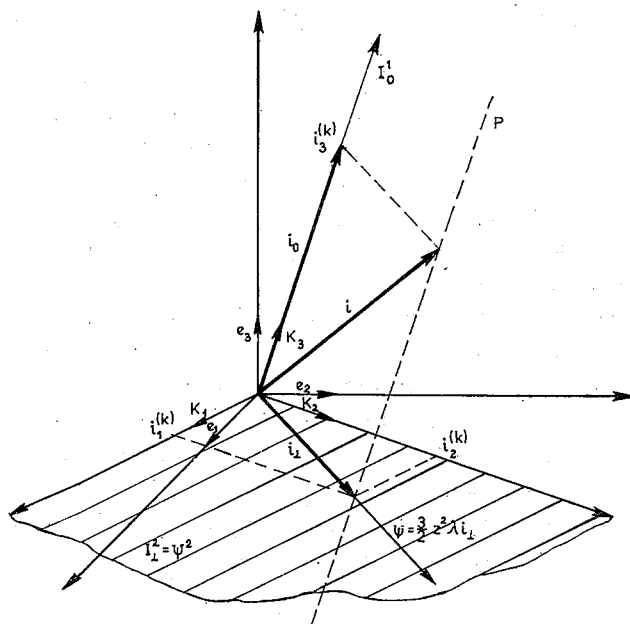


Rys. 3. 3-fazowy obwód elektromagnetyczny

$[i_0] = [1 \ 1 \ 1]^T$. Połączone, zortogonalizowane i unormowane bazy przestrzeni $I_{\perp}^2$ oraz I_0^1 wyznaczają bazę k -osiowego układu współrzędnych: $[K_1] = \sqrt{\frac{2}{3}} \left[1 - \frac{1}{2} - \frac{1}{2}\right]^T$, $[K_2] = \sqrt{\frac{2}{3}} \left[0 \frac{\sqrt{3}}{2} - \frac{\sqrt{3}}{2}\right]^T$, $[K_3] = \sqrt{\frac{2}{3}} \left[\frac{1}{\sqrt{2}} \frac{1}{\sqrt{2}} \frac{1}{\sqrt{2}}\right]^T$. Interpretację geometryczną przedstawia rys. 4. Wektory e_1, e_2, e_3 stanowią bazę naturalnego, a wektory K_1, K_2, K_3 — k -osiowego układu współrzędnych. Wektory K_1 i K_2 rozpinają przestrzeń (płaszczyznę) $I_{\perp}^2$, tożsamą z przestrzenią Ψ^2 , zaś wektor K_3 — przestrzeń (prostą) I_0^1 . Składowa aktywna $i_{\perp}$ jest rzutem wektora i na płaszczyznę $I_{\perp}^2$ i jest określona współrzędnymi $i_1^{(k)}, i_2^{(k)}$. Współrzędna $i_3^{(k)}$ wyznacza składową zerową. Taki sam strumień skojarzony ψ wytwarzają wszystkie wektory i , których hodografem jest prosta p o równaniu: $i = i_{\perp} + \lambda i_0$, $\lambda \in R$.

Macierz transformacji $[K]$ ma postać

$$[K] = \sqrt{\frac{2}{3}} \begin{bmatrix} 1 & -\frac{1}{2} & -\frac{1}{2} \\ 0 & \frac{\sqrt{3}}{2} & -\frac{\sqrt{3}}{2} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix}. \quad (21)$$



Rys. 4. Interpretacja geometryczna składowej aktywnej i zerowej

Macierz indukcyjności głównych jest w k -osiowym układzie współrzędnych diagonalna

$$[M^{(k)}] = \begin{bmatrix} \frac{3}{2} z^2 \Lambda & 0 & 0 \\ 0 & \frac{3}{2} z^2 \Lambda & 0 \\ 0 & 0 & 0 \end{bmatrix}. \quad (22)$$

Składową zerową prądu można wyeliminować poprzez połączenie uzwojeń w gwiazdę bez przewodu zerowego.

9. SKŁADOWE ZEROWE PRĄDU W NIESYMETRYCZNYCH MASZYNACH ASYNCHRONICZNYCH

Dla niesymetrycznej wewnętrznie maszyny asynchronicznej o n -fazowym stojanie i m -fazowym wirniku przyjmijmy następujące założenia upraszczające: liniowy obwód magnetyczny, sinusoidalny układ prądowy uzwojeń oraz równomierną szczelinę powietrzną. Założenie o sinusoidalnym układzie prądowym, równoznaczne ograniczeniu analizy do I harmonicznej przestrzennej indukcji magnetycznej w szczelinie maszyny, uzasadnione jest w maszynach, w których niesymetria powtarza się regularnie pod każdym biegunem. W szczególności osie faz stojana mogą być poprzysuwane na obwodzie

maszyny o kąt różny od $\frac{2\pi}{n}$ i odpowiednio osie faz wirnika — o kąt różny od $\frac{2\pi}{m}$ (niesymetria kątowa maszyny). W równaniach stanu nieustalonego maszyny asynchronicznej występują 4 różne macierze indukcyjności głównych: stojana $[M_{ss}]$, wirnika $[M_{ww}]$, stojana względem wirnika $[M_{sw}]$ oraz wirnika względem stojana $[M_{ws}]$. Współczynniki indukcyjności głównych macierzy $[M_{sw}]$ i $[M_{ws}]$ zależą od kąta obrotu wirnika φ . Można wykazać, że: $\text{rz}[M_{ss}] = \text{rz}[M_{ww}] = \text{rz}[M_{sw}] = \text{rz}[M_{ws}] = \dim I_{\perp} = k = 2$ oraz że istnieje wspólna 2-wymiarowa przestrzeń aktywna i wspólna $(n-2)$ -wymiarowa przestrzeń zerowa dla macierzy $[M_{ss}]$ i $[M_{ws}]$ oraz wspólna 2-wymiarowa przestrzeń aktywna i wspólna $(m-2)$ -wymiarowa przestrzeń zerowa dla macierzy $[M_{ww}]$ i $[M_{sw}]$. Istnienie wspólnych przestrzeni $I_{\perp}$ i I_0 dla par macierzy indukcyjności głównych maszyny asynchronicznej ma prostą interpretację fizyczną: wektor zerowy prądu stojana, niewytwarzający pola magnetycznego w szczelinie maszyny, nie wytwarza również strumienia skojarzonego z uzwojeniami wirnika. I odwrotnie. Składowa zerowa prądu stojana istnieje więc niezależnie od składowej zerowej prądu wirnika.

Moment elektromagnetyczny maszyny asynchronicznej można wyrazić jako formę dwuliniową wektora prądu stojana $[i_s]$ i wektora prądu wirnika $[i_w]$

$$M_e = \frac{\partial}{\partial \varphi} [i_s]^T [M_{sw}] [i_w] =$$

$$= \frac{\partial}{\partial \varphi} \begin{bmatrix} i_{s1} \\ i_{s2} \\ \vdots \\ i_{sn} \end{bmatrix}^T \begin{bmatrix} M_{sw11} & \dots & M_{sw1m} \\ \dots & \dots & \dots \\ M_{swn1} & \dots & M_{swnm} \end{bmatrix} \begin{bmatrix} i_{w1} \\ i_{w2} \\ \vdots \\ i_{wm} \end{bmatrix} \quad (23)$$

Przestrzeń zerowa stojana jest lewą przestrzenią zerową, a przestrzeń zerowa wirnika — prawą przestrzenią zerową formy dwuliniowej (23). W przetwornikach elektromechanicznych ulega więc rozszerzeniu sens fizyczny składowych zerowych. Składowe zerowe wektorów prądów stojana i wirnika nie tylko nie wytwarzają pola magnetycznego w szczelinie maszyny, ale również nie partycypują w wytwarzaniu momentu elektromagnetycznego maszyny.

Stały, niezależny od kąta φ rząd wszystkich macierzy indukcyjności głównych równy 2 tłumaczy, dlaczego w teorii wirujących maszyn elektrycznych, zamiast o k -osiowym układzie współrzędnych i k -osiowej transformacji, mówimy o 2-osiowym układzie współrzędnych i 2-osiowej transformacji. Korzyści, wynikające z zastosowania transformacji 2-osiowej są już znaczne przy analizie 3-fazowych maszyn asynchronicznych, a stają się jeszcze wyraźniejsze w przypadku maszyn o wyższej liczbie faz. Ma to istotne znaczenie przy analizie maszyn o budowie niesymetrycznej, bądź też stanów awaryjnych maszyn symetrycznych, kiedy poszczególne części uzwojeń, a nawet poszczególne zezwoje lub pręty klatki wirnika traktuje się jako odrębne fazy.

Istnienie różnych podprzestrzeni aktywnych i zerowych dla par macierzy indukcyjności głównych wymaga wprowadzenia innego, 2-osiowego układu współrzędnych dla wielkości elektrycznych i magnetycznych stojana oraz innego — dla wirnika. Bazę 2-osiowego układu

współrzędnych dla stojana wyznacza macierz $[M_{ss}]$, a dla wirnika — macierz $[M_{ww}]$. W nowym układzie współrzędnych dla szerokiej klasy niesymetrycznych maszyn asynchronicznych, w których rezystancje i indukcyjności rozproszeń poszczególnych faz po sprowadzeniu na fazę odniesienia są równe, układ równań różniczkowych stanu nieustalonego ulega rozprężeniu na: układ 2 równań dla składowych aktywnych stojana i wirnika, układ $n-2$ równań dla składowej zerowej stojana oraz układ $m-2$ równań dla składowej zerowej wirnika. Układ równań różniczkowych dla składowych aktywnych ma taką samą budowę formalną, jak układ równań maszyny 2-fazowej, jednakże współczynniki indukcyjności wzajemnej w nowym układzie współrzędnych w ogólnym przypadku nie spełniają zasady wzajemności. Zależność współczynników indukcyjności głównych od kąta obrotu wirnika dla maszyny niesymetrycznej kątowo jest zaś tego rodzaju, że nie może być usunięta na drodze transformacji obrotu o kąt φ .

10. SKŁADOWE ZEROWE W SYMETRYCZNYCH MASZYNACH ASYNCHRONICZNYCH

Pojęcie składowej zerowej wywodzi się od R. H. Parka, który na drodze rozważań intuicyjno-fizycznych wprowadził transformację 2-osiową do analizy stanów nieustalonych symetrycznych maszyn synchronicznych. Składową zerową w maszynach 3-fazowych przyjęto ogólnie wyróżniać dolnym indeksem 0. W nawiązaniu do tradycyjnego sposobu notacji będziemy kolejne współrzędne składowej zerowej stojana oznaczać indeksami: 01, 02, ..., 0n-2, zaś wirnika: 01, 02, ..., 0m-2. Zachowamy również tradycyjny sposób wyróżniania współrzędnych aktywnych przez α, β na tzw. płaszczyźnie stojana i przez d, q na tzw. płaszczyźnie wirnika. Transformacje, prowadzące do takich układów współrzędnych z osiami wirującymi w przestrzeni aktywnej, są złożeniem transformacji 2-osiowej z odpowiednimi transformacjami obrotu. Podsumowując, wektor prądu, zamiast przez $[i_1^{(k)} i_2^{(k)} i_3^{(k)} i_4^{(k)} \dots]^T$ będziemy oznaczać przez $[i_\alpha i_\beta i_{01} i_{02} \dots]^T$, bądź $[i_d i_q i_{01} i_{02} \dots]^T$.

Dla symetrycznego n -fazowego stojana, w którym kąt pomiędzy osiami kolejnych faz wynosi $\alpha_s = \frac{2\pi}{n}$, składowa zerowa ma w naturalnym układzie współrzędnych, zgodnie z (16), postać

$$[i_0] = \sqrt{\frac{2}{n}} \begin{bmatrix} \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \vdots \\ \frac{1}{\sqrt{2}} \end{bmatrix} i_{01} + i_{02} \begin{bmatrix} \cos 0 \\ \cos 2\alpha_s \\ \cos 4\alpha_s \\ \cos 6\alpha_s \\ \vdots \\ \cos(n-1)2\alpha_s \end{bmatrix} + i_{03} \begin{bmatrix} \sin 0 \\ \sin 2\alpha_s \\ \sin 4\alpha_s \\ \sin 6\alpha_s \\ \vdots \\ \sin(n-1)2\alpha_s \end{bmatrix} +$$

$$+i_{04} \begin{bmatrix} \cos 0 \\ \cos 3\alpha_s \\ \cos 6\alpha_s \\ \cos 9\alpha_s \\ \vdots \\ \cos(n-1)3\alpha_s \end{bmatrix} + \dots + i_{0n-3} \begin{bmatrix} \cos 0 \\ \cos \frac{n-1}{2}\alpha_s \\ \cos 2\frac{n-1}{2}\alpha_s \\ \cos 3\frac{n-1}{2}\alpha_s \\ \vdots \\ \cos(n-1)\frac{n-1}{2}\alpha_s \end{bmatrix} + i_{0n-2} \begin{bmatrix} \sin 0 \\ \sin \frac{n-1}{2}\alpha_s \\ \sin 2\frac{n-1}{2}\alpha_s \\ \sin 3\frac{n-1}{2}\alpha_s \\ \vdots \\ \sin(n-1)\frac{n-1}{2}\alpha_s \end{bmatrix} \quad (24)$$

gdy liczba faz n jest nieparzysta, lub

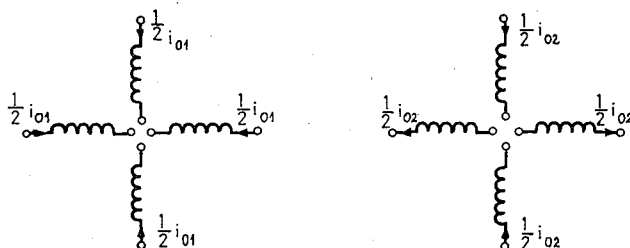
$$[i_0] = \sqrt{\frac{2}{n}} (i_{01} \begin{bmatrix} \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \vdots \\ \frac{1}{\sqrt{2}} \end{bmatrix} + i_{02} \begin{bmatrix} \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ \vdots \\ -\frac{1}{\sqrt{2}} \end{bmatrix} + i_{03} \begin{bmatrix} \cos 0 \\ \cos 2\alpha_s \\ \cos 4\alpha_s \\ \cos 6\alpha_s \\ \vdots \\ \cos(n-1)2\alpha_s \end{bmatrix} + i_{04} \begin{bmatrix} \sin 0 \\ \sin 2\alpha_s \\ \sin 4\alpha_s \\ \sin 6\alpha_s \\ \vdots \\ \sin(n-1)2\alpha_s \end{bmatrix} + \dots) \quad (25)$$

$$+i_{05} \begin{bmatrix} \cos 0 \\ \cos 3\alpha_s \\ \cos 6\alpha_s \\ \cos 9\alpha_s \\ \vdots \\ \cos(n-1)3\alpha_s \end{bmatrix} + \dots + i_{0n-3} \begin{bmatrix} \cos 0 \\ \cos \left(\frac{n}{2}-1\right)\alpha_s \\ \cos 2\left(\frac{n}{2}-1\right)\alpha_s \\ \cos 3\left(\frac{n}{2}-1\right)\alpha_s \\ \vdots \\ \cos(n-1)\left(\frac{n}{2}-1\right)\alpha_s \end{bmatrix} + i_{0n-2} \begin{bmatrix} \sin 0 \\ \sin \left(\frac{n}{2}-1\right)\alpha_s \\ \sin 2\left(\frac{n}{2}-1\right)\alpha_s \\ \sin 3\left(\frac{n}{2}-1\right)\alpha_s \\ \vdots \\ \sin(n-1)\left(\frac{n}{2}-1\right)\alpha_s \end{bmatrix} \quad (26)$$

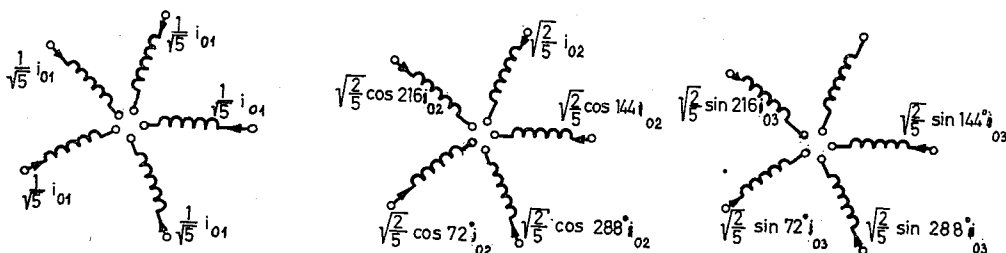
gdy liczba faz n jest parzysta.

Współczynnik $\sqrt{\frac{2}{n}}$ wynika z unormowania wektorów bazy. Podobnie można wyrazić składową zerową prądu symetrycznego m -fazowego wirnika o kącie pomiędzy osiami faz $\alpha_w = \frac{2\pi}{n}$, podstawiając we wzorach (24) i (25): $n = m$ i $\alpha_s = \alpha_w$.

Przykładowo, składowe zerowe 5-fazowego i 6-fazowego obwodu elektromagnetycznego maszyny przedstawiono na rys. 5 i 6. Prądy składowej zerowej stojana można ograniczać przez pomniejszanie składowej zerowej w napięciach fazowych maszyny. Jeśli składowa zerowa napięcia fazowego jest równa zero, wówczas prądy składowej zerowej nie płyną. W zwartym wirniku prądy zerowe mogą powstać jako wynik włączenia niejednakowej impedancji dodatkowej w poszczególne fazy wirnika. W wirniku zwartym bezpośrednio, bądź przez impedancję dodatkową włączoną symetrycznie, składowa zerowa prądu nie wystąpi.



Rys. 5. Składowe zerowe prądu 4-fazowego obwodu elektromagnetycznego maszyny

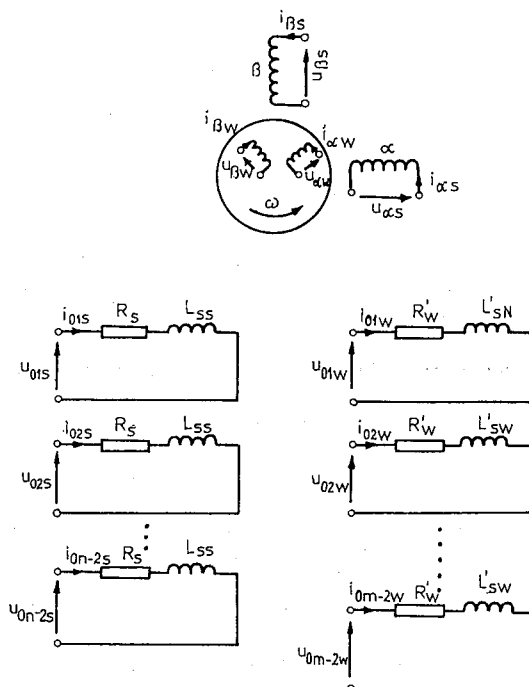


Rys. 6. Składowe zerowe prądu 5-fazowego obwodu elektromagnetycznego maszyny

Inna metoda ograniczenia składowej zerowej — to połączenie faz maszyny w gwiazdę bez przewodu zerowego, co uniemożliwia przepływ prądu i_{01} . Przewód zerowy nie ma natomiast wpływu na pozostałe prądy zerowe: i_{02} , i_{03} , ponieważ ich suma algebraiczna jest zawsze równa zero.

W stanach ustalonych można dokonać rozkładu wektorów prądów stojana i wirnika na składowe: współbieżną (wytworzącą moment napędzający), przeciwbieżną (wytworzącą moment hamujący) i zerową (nie wytworzącą momentu elektromagnetycznego). O ile w maszynach symetrycznych składowa zerowa jest niezależna od składowych: współbieżnej i przeciwbieżnej, to w maszynach niesymetrycznych parametry optymalnego obwodu elektromagnetycznego są wynikiem kompromisu pomiędzy wartościami wszystkich trzech składowych. Schematy zastępcze symetrycznej maszyny wielofazowej w 2-osiowym układzie współrzędnych są szczególnie proste. Równaniom składowej aktywnej stojana i wirnika odpowiada 2-fazowa maszyna o osiach α i β wzajemnie ortogonalnych, równoważna energetycznie maszynie rzeczywistej, zaś równaniom składowej zerowej stojana i wirnika — $n-2$ prostych obwodów elektrycznych, złożonych z rezystancji i indukcyj-

ności rozproszeń stojana i $m-2$ obwodów elektrycznych, złożonych z rezystancji i indukcyjności rozproszeń wirnika (rys. 7). W symetrycznych maszynach asynchronicznych celowe jest stosowanie transformacji k -osiowej i wtenczas, gdy w analizie uwzględniamy wyższe harmoniczne przestrzenne pola magnetycznego w szczelinie maszyny. W równaniach różniczkowych stanu nieustalonego maszyny w miejsce macierzy indukcyjności głównych występują wówczas szeregi macierzy indukcyjności głównych, odpowiadające



Rys. 7. Schematy zastępcze dla składowych aktywnych i zerowych symetrycznej maszyny wielofazowej w 2-osiowym układzie współrzędnych

kolejnym harmonicznym przestrzennym pola magnetycznego. Można wykazać, że w symetrycznych maszynach 3-fazowych istnieje wspólna 2-wymiarowa przestrzeń aktywna dla macierzy indukcyjności głównych, odpowiadających 1, 5, 7 ... harmonicznej przestrzennej i że jest ona identyczna z 2-wymiarową przestrzenią zerową dla macierzy indukcyjności głównych, odpowiadających 3, 6, 9 ... harmonicznej przestrzennej pola magnetycznego. Odpowiednio też 1-wymiarowa przestrzeń aktywna macierzy indukcyjności głównych, związanych z 3, 6, 9 ... harmoniczną przestrzenną jest taka sama, jak 1-wymiarowa przestrzeń zerowa macierzy indukcyjności głównych, odpowiadających 1, 5, 7 ... harmonicznej przestrzennej pola magnetycznego. Fizycznie oznacza to, że każdy wektor prądu, który wytwarza 1 — wytwarza także 5, 7, 11 ... harmoniczną przestrzenną pola magnetycznego w szczeli nie, a nie powoduje powstania pól magnetycznych 3, 6, 9 ... harmonicznej przestrzennej.

W wyniku transformacji k -osiowej układ równań różniczkowych stanu nieustalonego ulega formalnemu rozsprzężeniu na: układ równań różniczkowych dla 1, 5, 7, 11 ... harmonicznej przestrzennej oraz układ równań dla 3, 6, 9 ... harmonicznej przestrzennej

pola magnetycznego w maszynie. Uzyskanie autonomicznych równań różniczkowych dla kolejnych harmonicznych przestrzennych pola magnetycznego jest niemożliwe. Znaczne uproszczenie układu równań można jednak uzyskać przez przyjęcie założenia, że każda harmoniczna pola magnetycznego w szczelinie wywołuje wyłącznie sinusoidalną reakcję magnetyczną tego samego rzędu w wirniku.

LITERATURA

1. N. W. Jefimow, E. R. Rozendorn, *Algebra liniowa wraz z geometrią wielowymiarową*, PWN, Warszawa 1974.
2. K. Kluszczyński, *Podstawy teoretyczne transformacji k -osiowej i jej zastosowanie do analizy rozgałęzionych obwodów elektromagnetycznych*. Zeszyty Naukowe Pol. Śl. „Elektryka”, z. 61, Gliwice 1978.
3. W. Paszek, *Wpływ stałych rozłożonych w klatkach wirników głębokożłobkowych na nieustalone przebiegi elektrodynamiczne silników indukcyjnych*. Zeszyty Naukowe Pol. Śl. „Elektryka”, z. 61, Gliwice 1978.
4. A. Puchała, *Dynamika maszyn i układów elektromechanicznych*, PWN, Warszawa 1977.

K. KLUSZCZYŃSKI

ACTIVE AND ZERO COMPONENTS IN ELECTRO-MAGNETIC CIRCUITS

Summary

The orthogonal decomposition of the winding currents in the electro-magnetic circuits into active and zero components is presented. The method of calculating zero components in the electro-magnetic circuits, particularly in asynchronous machines is described. The conditions for eliminating or minimizing the zero components are discussed. The new k -axis coordinate system obtained simplifies the differential equations of the electro-magnetic circuits.

K. KLUSZCZYŃSKI

COMPOSANTES ACTIVES ET NEUTRES DES COURANTS DANS LES CIRCUITS ÉLECTROMAGNÉTIQUES

Résumé

En se basant sur la théorie des transformations linéaires dans les espaces vectoriels, on a montré que, dans les circuits linéaires électromagnétiques, on peut isoler des composantes actives et neutres des courants, ce qui a une importance effective du point de vue des conversions énergétiques (conversions d'énergie électrique, magnétique et mécanique). On a donné le mode de désignation des composantes neutres dans les circuits électromagnétiques à structure topologique composée aléatoirement et avant tout dans les machines asynchrones. On a formulé les conditions d'élimination ou de réduction au minimum de la composante neutre des courants. On a introduit un nouveau k -axial système de coordonnées, dans lequel les équations différentielles de l'état instable des circuits électromagnétiques se simplifient considérablement.

K. KLUSZCZYŃSKI

AKTIV- UND NULLKOMPONENTEN DER ELEKTROMAGNETISCHEN
STROMKREISE

Zusammenfassung

Es wird nachgewiesen, daß die Theorie linearer Transformationen im Vektorraum, angewandt auf lineare, elektromagnetisch gekoppelte Stromkreise, eine Aufteilung der Vektoren in Aktivkomponenten und Nullkomponenten ermöglicht. Diese Komponenten haben eine wesentliche physikalische Bedeutung in Bezug auf die Energieumwandlungsprozesse. Die Absonderung der Nullkomponenten bei elektromagnetisch gekoppelten Stromkreisen, insbesondere bei Asynchronmaschinen wird diskutiert. Bedingungen werden formuliert, die zur Ausscheidung bzw. zur Minimalisierung der Stromkomponenten führen.

К. КЛЮЩИНСКИ

АКТИВНЫЕ И НУЛЕВЫЕ СОСТАВЛЯЮЩИЕ ТОКОВ
В ЭЛЕКТРОМАГНИТНЫХ ЦЕПЯХ

Резюме

Доказано что для широкого класса разветвлённых электромагнитных цепей можно провести ортогональное разложение вектора тока в обмотках на активную и нулевую составляющие. Представлен метод определения нулевых составляющих в электромагнитных цепях с произвольной сложной топологической структурой, в особенности в электрических машинах. Сформулированы условия исключения или минимализации нулевой составляющей токов. Введена новая k -осевая координатная система, в которой дифференциальные уравнения неустойчившихся режимов значительно упрощаются.

Próba opisu charakterystyk napięciowo-prądowych wyładowania łukowego stabilizowanego strugą gazu przy wykorzystaniu równań magnetohydrodynamiki

CZESŁAW KRÓLIKOWSKI (POZNAŃ), ANIELA KAMIŃSKA-PRANKE (POZNAŃ),
RYSZARD NIEWIEDZIAŁ (POZNAŃ)

Instytut Elektroenergetyki Politechniki Poznańskiej

Otrzymano 10.12.1979

Praca zawiera analizę procesu nagrzewania gazu przez łuk elektryczny, przy uwzględnieniu równań magnetohydrodynamiki. W wyniku ustalono równanie podające związek między napięciem i prądem łuku, własnościami fizycznymi przepływającego gazu oraz geometrią kanału wyładowczego. Równania charakterystyk napięciowo-prądowych wyznaczono dla łuku palącego się w gazach jednoatomowych przyjmując, że ciepło właściwe tych gazów jest wartością stałą oraz dla gazów dwuatomowych uwzględniając zależność ciepła właściwego od temperatury. Obliczenia wykonano dla łuków palących się w argonie i azocie a wyniki porównano z wynikami otrzymanymi eksperymentalnie.

1. WPROWADZENIE

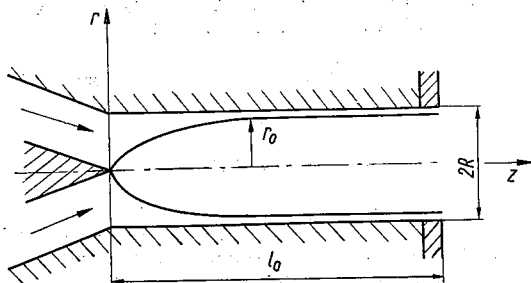
Proces nagrzewania przez łuk elektryczny strumienia przepływającego gazu był analizowany w wielu pracach. W pracy [1] była przeprowadzona analiza równania energii, w której głównym mankamentem było wprowadzenie do obliczeń rozkładu prędkości, wyliczonego na podstawie przypuszczalnego rozkładu temperatury w łuku elektrycznym. H. A. Stine i V. R. Watson w pracy [2] przeprowadzili analizę procesu nagrzewania gazu przy założeniu, że prędkość masy w polu przepływu jest stała i przy pominięciu równania pędu. Inne opublikowane rozwiązania równania energii dotyczyły zjawisk w obszarze przemian i w pełni rozwiniętych obszarach przepływu. Na uwagę zasługuje praca N. B. Murthy'a [3], który przeprowadził analizę zjawisk zachodzących w palniku plazmowym dla całego obszaru wyładowania, dzieląc go na obszar wlotu, obszar przejścia i całkowicie rozwinięty obszar przepływu.

W niniejszej pracy rozpatrzono proces nagrzewania gazu przez łuk elektryczny, uwzględniając równania magnetohydrodynamiki i wcześniejsze badania nad modelem wymiany ciepła między łukiem a gazem stabilizującym wyładowanie (gazem roboczym) [4]. W wyniku uzyskano równanie podające związek między napięciem łuku, prądem łuku, własnościami fizycznymi przepływającego gazu oraz geometrią kanału wyładowczego.

2. CHARAKTERYSTYKA NAPIĘCIOWO-PRĄDOWA WYŁADOWANIA ŁUKOWEGO STABILIZOWANEGO STRUGĄ GAZU

W przestrzeni cylindrycznego przepływu można wydzielić dwa obszary: obszar wewnętrzny, czyli plazmy łuku i obszar zewnętrzny — obszar gazu w otoczeniu. Oddzielone one są powierzchnią graniczną o promieniu $r_0 = r(z)$, co przedstawiono na rys. 1.

W obszarze wewnętrznym, czyli dla $r(z) < r_0$ konduktancja gazu $\sigma > 0$, natomiast w obszarze zewnętrznym dla $r(z) \geq r_0$, $\sigma = 0$. W pracy dokonano analizy obszaru wewnętrznego przepływu, rozważając wyładowanie łukowe poddane wolnemu przepływowi gazu, równoległe do osi łuku.



Rys. 1. Model pola przepływu

Proces wyładowania łukowego rozpatrzono przy następujących założeniach:

- przepływ gazu jest laminarny, a liczba Macha dużo mniejsza od jednośi,
- przepływ jest osiowo-symetryczny — nie ma składowej obwodowej prędkości, a więc nie ma także rotacji przepływu,
- składowa promieniowa prędkości jest dużo mniejsza od składowej osiowej, w wyniku czego ciśnienie jest jednakowe w płaszczyźnie prostopadłej do osi łuku,
- gaz znajduje się w równowadze termodynamicznej,
- zależność własności termodynamicznych (oprócz masy właściwej gazu) i własności transportu od ciśnienia jest nieznacząca,
- osiowe przenoszenie ciepła i energia promieniowania są do pominięcia.

Procesy zachodzące w wewnętrznym obszarze przepływu opisują równania magneto-hydrodynamiki, które z uwagi na symetrię osiową przepływu, rozpatrzono we współrzędnych cylindrycznych.

Równanie ciągłości przepływu przedstawia się następująco

$$\frac{\partial \varrho}{\partial t} + \frac{1}{r} \frac{\partial}{\partial r} (r \varrho \vartheta_r) + \frac{\partial}{\partial z} (\varrho \vartheta_z) = 0. \quad (1)$$

Ponieważ dla stanu ustalonego $\frac{\partial \varrho}{\partial t} = 0$, a na podstawie założenia przyjętego modelu wyładowania łukowego, że $\vartheta_r \ll \vartheta_z$ można pominąć drugi człon równania (1), zatem równanie ciągłości przepływu, w rozpatrywanym przypadku, przyjmie następującą postać

$$\frac{\partial}{\partial z} (\varrho \vartheta_z) = 0. \quad (2)$$

Równanie ruchu dla składowych wektorów będzie:

$$\varrho \left(\frac{\partial \vartheta_r}{\partial t} + \vartheta_r \frac{\partial \vartheta_r}{\partial r} + \frac{\vartheta_\varphi}{r} \frac{\partial \vartheta_r}{\partial \varphi} + \vartheta_z \frac{\partial \vartheta_r}{\partial z} - \frac{\vartheta_\varphi^2}{r} \right) + \frac{\partial p}{\partial r} = j_\varphi H_z - j_z H_\varphi, \quad (3)$$

$$\varrho \left(\frac{\partial \vartheta_\varphi}{\partial t} + \vartheta_r \frac{\partial \vartheta_\varphi}{\partial r} + \frac{\vartheta_\varphi}{r} \frac{\partial \vartheta_\varphi}{\partial \varphi} + \vartheta_r \frac{\partial \vartheta_\varphi}{\partial z} + \frac{\vartheta_r \vartheta_\varphi}{r} \right) + \frac{\partial p}{\partial \varphi} = j_z H_r - j_r H_z, \quad (4)$$

$$\varrho \left(\frac{\partial \vartheta_z}{\partial t} + \vartheta_r \frac{\partial \vartheta_z}{\partial r} + \frac{\vartheta_\varphi}{r} \frac{\partial \vartheta_z}{\partial \varphi} + \vartheta_z \frac{\partial \vartheta_z}{\partial z} \right) + \frac{\partial p}{\partial z} = j_r H_\varphi - j_\varphi H_r. \quad (5)$$

Na podstawie wcześniej podanych własności łuku elektrycznego stabilizowanego strumieniem gazu, równanie (5) przyjmie prostą postać

$$\varrho \vartheta_z \frac{\partial \vartheta_z}{\partial z} + \frac{\partial p}{\partial z} = 0. \quad (6)$$

Równanie zachowania energii ma postać

$$\begin{aligned} \frac{\partial}{\partial t} \left[\varrho \left(C_p T + \frac{\vartheta^2}{2} \right) \right] + \frac{1}{r} \frac{\partial}{\partial r} \left[r \varrho \vartheta_r \left(C_p T + \frac{\vartheta^2}{2} \right) \right] + \frac{\partial}{\partial z} \left[\varrho \vartheta_z \left(C_p T + \frac{\vartheta^2}{2} \right) \right] + \\ - \frac{1}{r} \frac{\partial}{\partial r} \left(r \lambda \frac{\partial T}{\partial r} \right) - \frac{\partial}{\partial z} \left(\lambda \frac{\partial T}{\partial z} \right) = j_r E_r + j_\varphi E_\varphi + j_z E_z, \end{aligned} \quad (7)$$

gdzie

$$\vartheta = \sqrt{\vartheta_r^2 + \vartheta_\varphi^2 + \vartheta_z^2}.$$

W równaniu (7) nie uwzględniono energii traconej z łuku drogą promieniowania, gdyż jak wykazały badania [5, 6] przy gęstościach mocy łuku rzędu tysięcy kW/cm² energia promieniowania nie przekracza 3%. Ponieważ w rozpatrywanych warunkach można uważać przepływ za laminarny, nie uwzględniono również ciepła odprowadzonego z łuku drogą konwekcji [3]. W równaniu (7) pominięto pierwsze dwa człony po lewej stronie, z uwagi na rozpatrywanie zjawisk w stanie ustalonym.

Przyjmując, że $\vartheta_r \ll \vartheta_z$ można także nie brać pod uwagę ciepła przenoszonego w kierunku osi z drogą przewodnictwa cieplnego gazu, jako że jest ono do pominięcia, w porównaniu z ciepłem przenoszonym w kierunku osi r .

Wykorzystując powyższe uproszczenia równanie (7) można zapisać w następującej postaci

$$\frac{\partial}{\partial z} \left[\varrho \vartheta_z \left(C_p T + \frac{\vartheta_z^2}{2} \right) \right] - \frac{1}{r} \frac{\partial}{\partial r} \left(r \lambda \frac{\partial T}{\partial r} \right) = j_z E_z, \quad (8)$$

równanie stanu gazu

$$p = \varrho R_g T. \quad (9)$$

Dla pełnego opisu zjawisk wykorzystano także równania elektromagnetyczne Maxwella oraz prawo Ohma dla pola przepływowego:

$$-\frac{\partial H_\varphi}{\partial z} = j_r, \quad (10)$$

$$\frac{\partial H_r}{\partial z} - \frac{\partial H_z}{\partial r} = j_\varphi, \quad (11)$$

$$\frac{1}{r} \frac{\partial}{\partial r} (r H_\varphi) = j_z, \quad (12)$$

$$\frac{\partial E_\varphi}{\partial z} = \mu \frac{\partial H_r}{\partial t}, \quad (13)$$

$$\frac{\partial E_r}{\partial z} - \frac{\partial E_z}{\partial r} = -\mu \frac{\partial H_\varphi}{\partial t}, \quad (14)$$

$$\frac{1}{r} \frac{\partial}{\partial r} (r E_\varphi) = -\mu \frac{\partial H_z}{\partial t}, \quad (15)$$

$$j_r = \sigma(E_r + \mu \partial_\varphi H_z - \mu \partial_z H_\varphi), \quad (16)$$

$$j_\varphi = \sigma(E_\varphi + \mu \partial_z H_r - \mu \partial_r H_z), \quad (17)$$

$$j_z = \sigma(E_z + \mu \partial_r H_\varphi - \mu \partial_\varphi H_r). \quad (18)$$

Z założeń b i c wynika, że równanie (4) przyjmie postać $j_z H_r - j_r H_z = 0$. Ponieważ w kierunku osi z występuje przepływ prądu, więc wektor gęstości prądu ma kierunek osi z czyli $j_z \neq 0$. Wynika stąd wniosek, że składowe wektora pola magnetycznego H_r i H_z muszą być równe zeru.

Rozpatrując równania Maxwella w postaci (10)–(15) stwierdzamy, że $j_\varphi = 0$ i dalej $E_r = E_\varphi = 0$. Prawo Ohma, dla przyjętych warunków wyładowania przyjmie postać

$$j_z = \sigma E_z. \quad (19)$$

Korzystając z równania ruchu (6), po prostych przekształceniach otrzymano następującą postać równania zachowania energii, dotyczące jednostki objętości

$$\frac{\partial}{\partial z} (\varrho \partial_z C_p T) - \partial_z \frac{\partial p}{\partial z} - \frac{1}{r} \frac{\partial}{\partial r} \left(r \lambda \frac{\partial T_r}{\partial r} \right) = j_z E_z. \quad (20)$$

W równaniu (20) występuje współczynnik wymiany ciepła λ , który dla wysokich temperatur łuku jest nieznany i zależy od temperatury i ciśnienia. Opierając się na przeprowadzonych badaniach w [9], przyjęto, że ilość ciepła odprowadzanego drogą przewodnictwa cieplnego z obszaru wewnętrznego do obszaru zewnętrznego przepływu jest proporcjonalna do wartości prądu łuku, a więc

$$-\int_0^{r_0} \int_0^{l_1} \frac{1}{r} \frac{\partial}{\partial r} \left(r \lambda \frac{\partial T}{\partial r} \right) 2\pi r dr dz = KI, \quad (21)$$

gdzie:

K — stała proporcjonalności,

l_1 — długość łuku.

Przepływ gazu przez kanał cylindryczny odbywa się pod wpływem różnicy ciśnienia na wlocie do kanału p_1 i ciśnienia na wylocie kanału p_0 . Założono, że rozkład ciśnienia wzdłuż kanału wyładowczego jest liniowy i wyrażony zależnością $p(z) = Bz + p_1$ gdzie

$$B = \frac{p_1 - p_0}{l}. \quad (22)$$

Przepływający gaz ogrzewa się w obszarze wyładowania łukowego i po przebyciu drogi w łuku osiąga temperaturę łuku T_l . Przyjmując, że wzrost temperatury jest liniowy, można zapisać

$$T_g = T_o + Az, \quad (23)$$

gdzie:

T_o — temperatura początkowa gazu,

$A = \frac{T_l - T_o}{l_l}$ — współczynnik wzrostu temperatury.

Ponieważ temperatura początkowa gazu $T_o \ll T_l$, można ostatecznie przyjąć, że

$$T_g = Az. \quad (24)$$

Wykorzystując równanie ciągłości przepływu w postaci (2) stwierdzamy, że masa gazu przepływająca w ciągu sekundy przez kolumnę łukową jest stała i wynosi

$$\pi r_o^2 \rho \vartheta_z = G = \text{const.}$$

Zależy ona od masy gazu stabilizującego wyładowanie łukowe G^* i może być opisana zależnością

$$G = G^* \frac{SR_g \rho_l T_l}{F}, \quad (25)$$

gdzie:

G^* — całkowity przepływ gazu w kanale wyładowczym,

R_g — stała gazowa J/g · deg,

$S = \pi r_o^2$ — przekrój poprzeczny kolumny łukowej,

$F = \pi R^2$ — przekrój poprzeczny kanału wyładowczego.

W oparciu o powyższe rozważania równanie gęstości mocy w łuku można zapisać w następującej postaci

$$\frac{G^* R_g \rho_l T_l}{F(Bz + p_1)} C_p A - \frac{G^* R_g T_l}{F(Bz + p_1)} B - \frac{1}{r} \frac{\partial}{\partial r} \left(r \lambda \frac{\partial T_l}{\partial r} \right) = j_z E_z. \quad (26)$$

Rozwiązanie równania (26) jest możliwe pod warunkiem, że ciepło właściwe C_p , gęstość ρ i konduktywność σ gazu są znanymi funkcjami temperatury i ciśnienia.

Korzystając z aproksymowanego w odpowiednio małych przedziałach przebiegu $\sigma(T)$, wyznaczono temperaturę łuku przy założeniu, że jest ona zależna od odległości rozpatrywanego przekroju od końca sworzniowej elektrody i nie zależy od promienia łuku. Stąd otrzymano

$$T_l = a + \frac{bz}{S} \frac{I}{U}, \quad (27)$$

gdzie a i b — współczynniki liczbowe uzyskane z aproksymacji liniowej przebiegu $\sigma(T)$.

Zmiany gęstości gazu można opisać równaniem

$$\rho_l = \alpha + \beta T_l, \quad (28)$$

gdzie α i β — współczynniki liczbowe uzyskane z aproksymacji liniowej przebiegu $\rho(T)$.

Równanie opisujące zależność napięcia na łuku w funkcji prądu łuku, uwzględniające własności fizyczne gazu oraz geometrię kanału wyładowczego otrzymano, całkując rów-

nanie gęstości mocy (26) po objętości łuku. Jako długość łuku l_1 przyjęto odległość równą połowie długości kanału wyładowczego. Określenie rzeczywistej długości łuku jest niemożliwe, ze względu na występujące w nim procesy gazo- i elektrodynamiczne. Natomiast graniczną wartość promienia łuku przyjmowano w zależności od prądu łuku i ciśnienia według [7]

$$r_o = \frac{kI^n}{p^m}. \quad (29)$$

Współczynnik proporcjonalności k określano dla danych warunków stabilizacji wyładowania.

Stosując gazy jednoatomowe, dla których ciepło właściwe jest wartością stałą, niezależną od temperatury, równanie całkowe opisujące charakterystykę $U-I$ ma postać

$$\frac{\pi k^2}{4} \frac{G^* R_g}{F} I^{2n} \left\{ (C_p A \alpha - \beta) \int_0^{l_1} \Phi_1 \Phi_2 dz + C_p A \beta \int_0^{l_1} \Phi_1 \Phi_2^2 dz \right\} + KI = UI, \quad (30)$$

gdzie

$$\Phi_1 = \frac{1}{(Bz + p_1)^{1+2m}}; \quad \Phi_2 = a + \frac{4bz}{\pi k^2} \frac{I^{1-2n}}{U} (Bz + p_1)^{2m}.$$

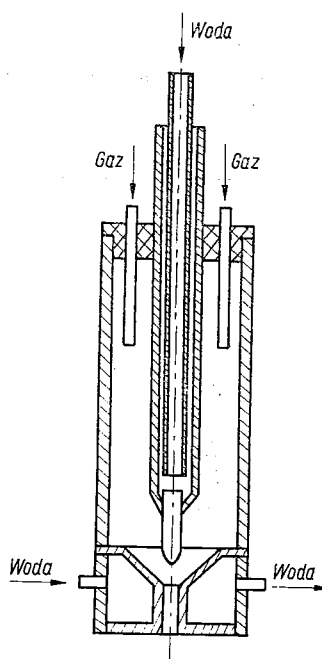
Całki wchodzące w skład równania (30) wyznaczono analitycznie i po prostych przekształceniach algebraicznych otrzymano równanie stopnia czwartego ze względu na U , które rozwiązano metodą Newtona.

3. OPIS CHARAKTERYSTYKI NAPIĘCIOWO-PRĄDOWEJ WYŁADOWANIA ŁUKOWEGO STABILIZOWANEGO ARGONEM

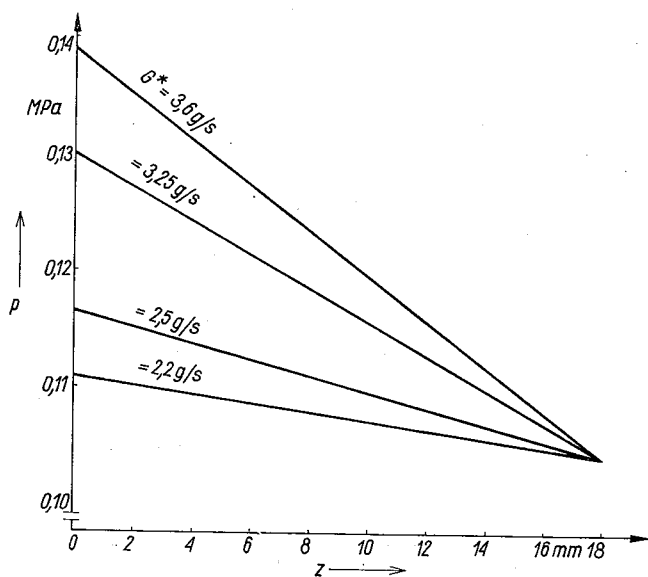
Przedstawiona w rozdziale 2 analiza może posłużyć do opisu charakterystyk napięciowo-prądowych łuku elektrycznego w plazmotronie, w którym całkowity przepływ można podzielić na dwa obszary scharakteryzowane uprzednio. Obliczenia charakterystyk $U-I$ wykonano wykorzystując równanie (30) dla łuku palącego się w plazmotronie, przedstawionym na rysunku 2. Łuk palił się w atmosferze argonu, tzn. gazu jednoatomowego. Obliczenia wykonano dla czterech wartości wydatku gazu i rozkładów ciśnień wzdłuż dyszy przedstawionych na rysunku 3. Do rozważań przyjęto następujące wartości współczynników stałych: $K = 10,8 \text{ V}$; $k = 4,14 \cdot 10^{-3} \text{ Nm/A}^n \text{m}^{2m-1}$; $n = 0,4$; $m = 0,25$; $R_g = 0,21 \text{ J/g} \cdot \text{deg}$; $C_p = 0,52 \text{ J/g} \cdot \text{deg}$. Otrzymane wyniki obliczeń porównano z wynikami pomiarów zamieszczonymi w [8]. Wartości pomierzonych i obliczonych napięć łuku zestawiono w tablicy 1, a na rysunku 4 przedstawiono obliczone charakterystyki $U-I$ w postaci graficznej.

Z tablicy 1 wynika, że różnice między wielkościami pomierzonymi i obliczonymi nie przekraczają $\pm 8\%$, co stanowi już błąd znaczący. Analizując równanie opisujące charakterystykę napięciowo-prądową łuku i założenia, które przyjęto dla jego otrzymania stwierdzono, że przyczyną znacznego błędu może być przyjęcie liniowych zmian ciśnienia w dyszy. Przyjęto więc dokładniejsze przybliżenie rozkładu ciśnień wzdłuż elektrody-dyszy potwierdzone również badaniami w [9], a mianowicie

$$p = Cz^2 + Dz + p_1, \quad (31)$$



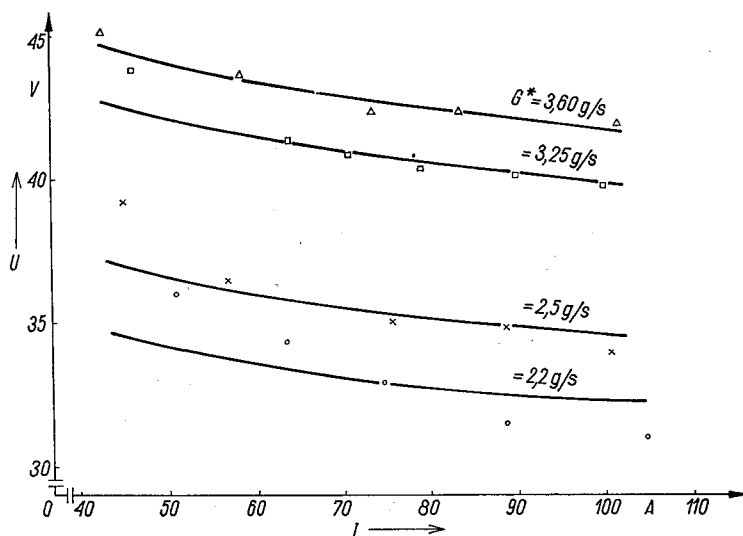
Rys. 2. Przekrój plazmotronu



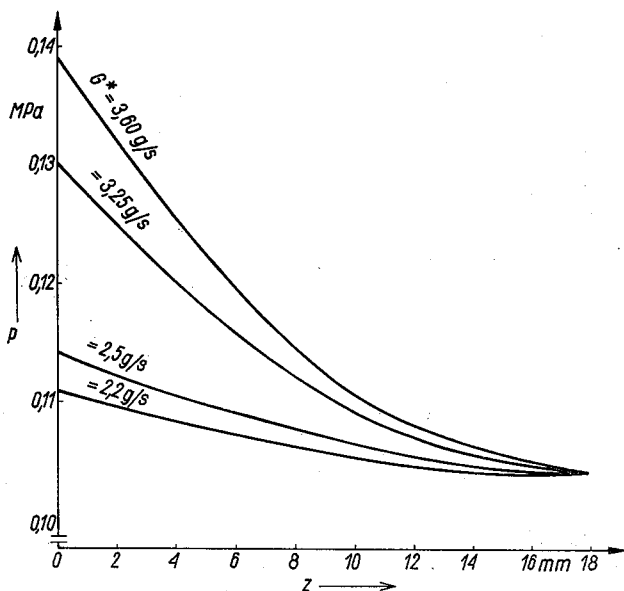
Rys. 3. Liniowy rozkład ciśnień wzdłuż elektrody-dyszy
Gaz roboczy argon

Pomierzone i obliczone zgodnie z równaniem (30) napięcia łuku
 $L = 18 \text{ mm}$, $R = 3,5 \text{ mm}$, gaz roboczy — argon

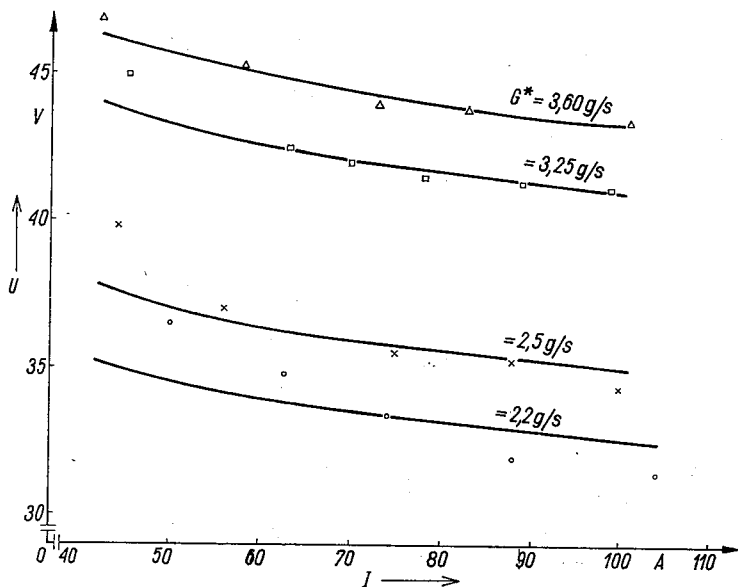
Lp.	U_{pom}	U_{obl}	I	G^*	p_1	$U\% = \frac{U_{obl} - U_{pom}}{U_{pom}} \cdot 100\%$
—	V	V	A	g/s	MPa	%
1.	36,5	34,1	50,0	2,2	0,111	-6,57
2.	34,8	33,5	62,5			-3,73
3.	33,0	33,0	74,0			0,00
4.	31,9	32,6	88,0			+2,19
5.	31,4	32,1	104,0	2,2	0,111	+2,23
1.	39,8	36,8	44,0	2,5	0,114	-7,54
2.	37,0	36,2	56,0			-2,16
3.	35,5	35,3	75,0			-0,56
4.	35,3	35,1	88,0			-0,56
5.	34,3	34,5	100,0	2,5	0,114	+0,58
1.	45,0	42,4	45,0	3,25	0,129	-5,78
2.	42,6	41,3	63,0			-3,05
3.	42,0	40,9	70,0			-2,62
4.	41,5	40,6	78,0			-2,17
5.	41,3	40,1	89,0			-2,90
6.	41,1	39,8	99,0	3,25	0,129	-3,16
1.	46,9	44,7	42,0	3,60	0,138	-4,69
2.	44,5	43,5	58,0			-2,25
3.	44,0	42,7	73,0			-2,95
4.	43,7	42,3	83,0			-3,20
5.	43,5	41,6	101,0	3,60	0,138	-4,37



Rys. 4. Obliczone zgodnie z równaniem (30) napięciowo-prądowe charakterystyki łuku elektrycznego
 Gaz roboczy argon. Punktami oznaczono rezultaty pomiarów



Rys. 5. Rozkład ciśnień wzdłuż elektrody-dyszy opisany równaniem (31)
Gaz roboczy argon



Rys. 6. Obliczone zgodnie z równaniem (32) napięciowo-prądowe charakterystyki łuku elektrycznego
Gaz roboczy argon. Punktami oznaczono rezultaty pomiarów

gdzie:

$$C = \frac{p_1 - p_0}{l^2}; \quad D = -2 \frac{p_1 - p_0}{l}.$$

W tym przypadku równanie opisujące charakterystykę napięciowo-prądową łuku elektrycznego przyjmie postać

$$\frac{\pi k^2}{4F} G^* R_g I^{2n} \left\{ \int_0^{l_1} [C_p A \alpha - (2Cz + D)] \Phi_3 \Phi_4 dz + C_p A \beta \int_0^{l_1} \Phi_3 \Phi_4^2 dz \right\} + KI = UI, \quad (32)$$

gdzie:

$$\Phi_3 = \frac{1}{(Cz^2 + Dz + p_1)^{1+2m}}, \quad (33)$$

$$\Phi_4 = a + \frac{4bz}{\pi k^2} \frac{I^{1-2n}}{U} (Cz^2 + Dz + p_1)^{2m}.$$

Obliczenia napięcia łuku według równania (32) wykonano dla tych samych warunków, dla których obliczono napięcia według równania (30), dla rozkładów ciśnień wzdłuż dyszy podanych na rysunku 5. Wyniki obliczeń i pomiarów zestawiono w tablicy 2, a na ry-

Tablica 2

Pomierzone i obliczone zgodnie z równaniem (32) napięcia łuku

$L = 18$ mm, $R = 3,5$ mm, gaz roboczy — argon

Lp.	U_{pom}	U_{obt}	I	G^*	p_1	$U_{\%} = \frac{U_{obl} - U_{pom}}{U_{pom}} \cdot 100\%$
—	V	V	A	g/s	MPa	%
1.	36,5	34,4	50,0	2,2	0,111	-5,75
2.	34,8	33,8	62,5			-2,87
3.	33,0	33,4	74,0			+1,21
4.	31,9	32,9	88,0			+3,13
5.	31,4	32,5	104,0	2,2	0,111	+3,50
1.	39,8	37,5	44,0	2,5	0,114	-5,78
2.	37,0	36,5	56,0			-1,35
3.	35,5	35,8	75,0			+0,84
4.	35,3	35,4	88,0			+0,28
5.	34,3	35,0	100,0	2,5	0,114	+2,04
1.	45,0	43,6	45,0	3,25	0,129	-3,11
2.	42,6	42,5	63,0			-0,23
3.	42,0	42,1	70,0			+0,24
4.	41,5	41,8	78,0			+0,72
5.	41,3	41,4	89,0			+0,24
6.	41,1	41,0	99,0	3,25	0,129	+0,24
1.	46,9	46,3	42,0	3,60	0,138	-1,28
2.	45,5	45,1	58,0			-0,88
3.	44,0	44,4	73,0			+0,91
4.	43,9	43,7	83,0			-0,45
5.	43,5	43,3	101,0	3,60	0,138	-0,46

sunku 6 przedstawiono charakterystyki napięciowo-prądowe w postaci graficznej. Błąd maksymalny względem wielkości pomierzonej nie przekracza $\pm 6\%$.

4. OPIS CHARAKTERYSTYKI NAPIĘCIOWO-PRĄDOWEJ WYŁADOWANIA ŁUKOWEGO STABILIZOWANEGO AZOTEM

W równaniu (32) ciepło właściwe gazu przyjmowano jako stałe niezależne od temperatury, stąd nie opisuje ono charakterystyk $U-I$ łuku w gazach dwuatomowych. Ciepło właściwe gazów dwuatomowych jest funkcją temperatury. Według Sirotinskiego i Burckharda zależność ta ma postać [10]

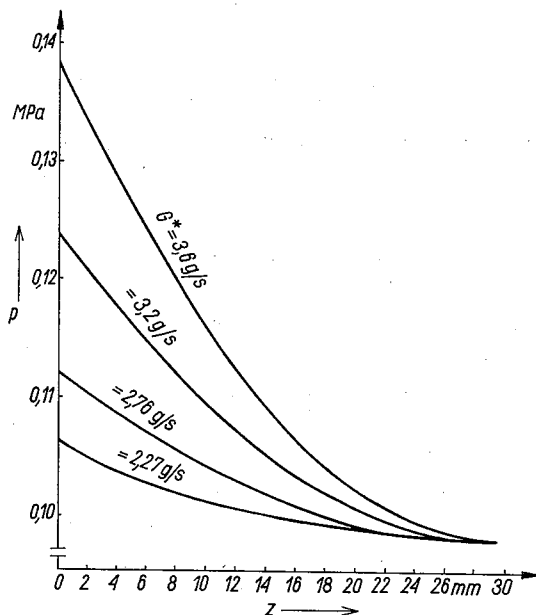
$$C_p = \frac{7}{2} R_g \left(\frac{T}{273} \right)^{1/2}. \quad (34)$$

Równanie opisujące charakterystykę napięciowo-prądową łuku palącego się w atmosferze gazów dwuatomowych można zapisać w postaci

$$\begin{aligned} \frac{7}{8} \frac{k^2}{F} G^* R_g^2 I^{2n} A 273^{-\frac{1}{2}} \left\{ \alpha \int_0^{l_1} \Phi_3 \Phi_4^{\frac{3}{2}} dz + \beta \int_0^{l_1} \Phi_3 \Phi_4^{\frac{5}{2}} dz \right\} + \\ - \frac{\pi k^2}{4F} G^* R_g I^{2n} \int_0^{l_1} (2Cz + D) \Phi_3 \Phi_4 dz + KI = UI, \quad (35) \end{aligned}$$

gdzie: Φ_3, Φ_4 określone są zależnościami (33).

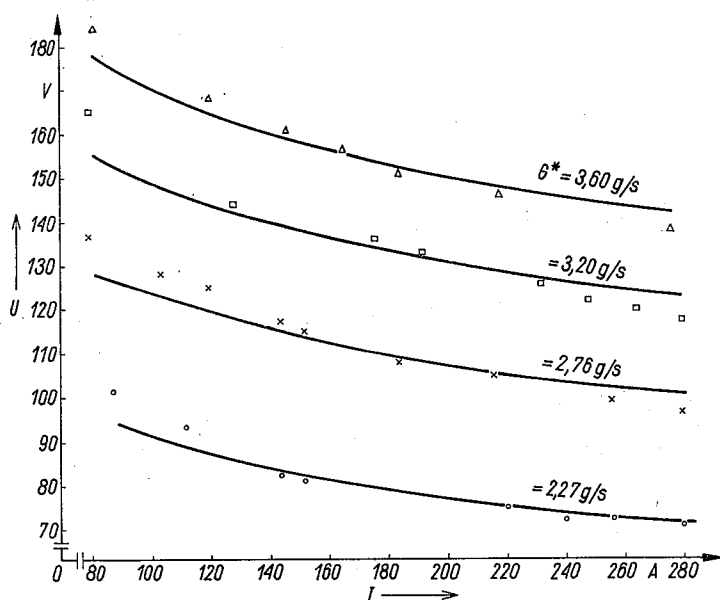
Całki występujące w równaniu (35) rozwiązano analitycznie i w wyniku prostych przekształceń algebraicznych otrzymano równanie piątego stopnia ze względu na U , które



Rys. 7. Rozkład ciśnień wzdłuż elektrody-dyszy opisany równaniem (31)
Gaz roboczy azot

rozwiązywano metodą przybliżeń Newtona. Obliczenia przeprowadzono dla łuku palącego się w azocie dla czterech wartości wydatku gazu i rozkładów ciśnień wzdłuż dyszy podanych na rysunku 7. Do obliczeń przyjęto następujące wartości współczynników stałych: $K = 10,8 \text{ V}$, $k = 4,14 \cdot 10^{-3} \text{ N}^m/\text{A}^n \text{m}^{2m-1}$, $n = 0,4$, $m = 0,25$, $R_g = 0,21 \text{ J/g} \cdot \text{deg}$.

Wyniki obliczeń podano w tablicy 3, a w postaci graficznej na rysunku 8. W celu porównania obliczonych charakterystyk $U-I$ z wynikami eksperymentalnymi wykonano konstrukcję plazmotronu ze wzdłużnym doprowadzeniem gazu i wykonano pomiary napięcia łuku w zależności od prądu łuku dla identycznych czterech wartości wydatku gazu. Wyniki pomiarów przedstawiono na rysunku 8 i zamieszczono w tablicy 3. Różnice między wartościami pomierzonymi i obliczonymi nie przekraczają $\pm 7\%$.



Rys. 8. Obliczone zgodnie z równaniem (35) napięciowo-prądowe charakterystyki łuku elektrycznego
Gaz roboczy azot. Punktami oznaczono rezultaty pomiarów

5. PODSUMOWANIE

W pracy dokonano próby wyznaczenia charakterystyk napięciowo-prądowych łuku palącego się w strumieniu przepływającego gazu w oparciu o równania magnetohydrodynamiki i przyjęty model wymiany ciepła między łukiem i przepływającym gazem. Równania charakterystyk napięciowo-prądowych wyznaczono dla dwóch różnych rozkładów ciśnień wzdłuż dyszy. Obliczenia wykazały, że przebieg charakterystyk jest zależny nie tylko od ciśnienia na wlocie i wylocie z dyszy, ale także od profilu jego zmian wzdłuż długości dyszy. Przyjęcie rozkładu ciśnienia wzdłuż dyszy opisanego równaniem (31) spowodowało obniżenie maksymalnego błędu obliczeń o około 2%, w stosunku do obliczeń przy założeniu liniowych zmian ciśnienia.

Główną przyczyną rozbieżności między wielkościami pomierzonymi i obliczonymi wydaje się być przyjęcie stałej długości łuku, niezależnej od wartości natężenia prądu

Tablica 3

Pomierzone i obliczone zgodnie z równaniem (35) napięcia łuku

 $L = 30$ mm, $R = 4$ mm, gaz roboczy — azot

Lp.	U_{pom}	U_{obl}	I	G^*	$U\% = \frac{U_{obl} - U_{pom}}{U_{pom}} \cdot 100\%$
—	V	V	A	g/s	%
1.	101,4	94,4	88	2,27	-6,90
2.	93,5	89,3	112		-4,49
3.	82,2	84,1	144		2,26
4.	81,0	83,0	152		2,47
5.	75,2	75,9	220		0,93
6.	72,1	74,0	240		2,63
7.	72,0	72,7	256		0,97
8.	70,4	71,0	280	2,27	0,85
1.	137,2	128,8	80	2,76	-6,12
2.	128,0	123,0	104		-3,91
3.	124,9	119,5	120		-8,04
4.	117,5	115,1	144		-1,31
5.	114,7	113,2	152		-1,30
6.	108,2	109,5	184		1,20
7.	104,9	106,0	216		1,05
8.	99,2	102,4	256		3,22
9.	96,5	100,5	280	2,76	4,14
1.	165,4	155,8	80	3,2	-5,80
2.	144,2	142,4	128		-1,25
3.	136,0	134,0	176		-1,47
4.	132,9	131,8	192		-0,83
5.	125,9	127,1	232		1,52
6.	122,0	125,5	248		2,87
7.	120,2	124,0	264		3,16
8.	117,5	122,6	280	3,2	4,34
1.	188,2	178,0	80	3,6	-5,42
2.	168,0	164,9	120		-1,84
3.	161,2	159,0	146		-1,36
4.	156,4	155,5	164		-0,57
5.	150,7	152,2	184		0,99
6.	145,9	147,4	218		1,03
7.	138,0	141,1	276	3,6	2,25

i wydatku gazu. Problem ten wymaga dalszych badań i analiz, które winny doprowadzić do ustalenia modelu wyładowania łukowego ze zmienną długością łuku przy zmiennych gazo- i elektrodynamicznych parametrach.

LITERATURA

1. K. S. W. Champion, *The Energy Balance Equation for the Positive Columns of High Pressure Arcs*, Proc. Phys. Soc. London, Vol. BLXVI, 1953.
2. H. A. Stine, V. R. Watson, *The Theoretical Enthalpy Distribution of Air in Steady Flow Along the Axis of a Direct Current Electric Arc*, NASATND, 1331, 1962.

3. S. N. B. Murthy, *Analysis arc heating phenomena in a tube*, IEEE Transactions on Nuclear Science, 1964.
4. R. Niewiedział, *Model elektrycznej dęgi stabilizowanej prądowym gazowym potokiem*, Mat. I Międzynarodowego Sympozjum pt. „Elektryczny łuk łączeniowy”, Łódź 1970.
5. C. Królikowski, *Próba analitycznego ujęcia podstawowych zjawisk występujących przy przetwarzaniu energii elektrycznej w wysokotemperaturową energię cieplną w palniku plazmowym prądu stałego*, Praca doktorska, AGH, Kraków 1964.
6. W. Finkelburg, H. Maecker, *Elektrische Bögen und thermisches Plasma*, Springer Verlag, Berlin 1956.
7. A. Au, Z. Ciok, *Aparaty elektryczne*, WPN, 1975.
8. C. Królikowski, *Analiza statecznej pracy palnika plazmowo-łukowego prądu stałego z łukiem pośrednim przy wykorzystaniu charakterystyk napięciowo-prądowych*, Rozprawy nr 22, Wyd. Uczelniane PP, Poznań 1967.
9. C. Królikowski, *Analityczno-eksperymentalna metoda wyznaczania podstawowych parametrów roboczych palnika plazmowego prądu stałego*, Zeszyty Naukowe Politechniki Poznańskiej, Elektryka, nr 6 (34), 1964.
10. G. Burckhard, *Zeitschrift Naturforschung*, 3a-603, 1948.

C. KRÓLIKOWSKI, A. KAMIŃSKA-PRANKE, R. NIEWIEDZIAŁ

AN ATTEMPT TO DESCRIBE VOLTAGE-CURRENT CHARACTERISTICS OF AN ARC WITH GAS JET STABILISATION WHEN UTILIZING EQUATIONS OF MAGNETOHYDRODYNAMICS

Summary

In this paper on the basis of magnetohydrodynamics equations and of the model of electrical arc with gas jet stabilisation the authors propose the equation determining voltage-current characteristics of the d.c. arc. The analytically calculated characteristics described by the proposed equation are presented and compared with the experimental ones.

C. KRÓLIKOWSKI, A. KAMIŃSKA-PRANKE, R. NIEWIEDZIAŁ

ESSAI POUR DÉTERMINER LES CARACTÉRISTIQUES TENSION-COURANT DE L'ARC ÉLECTRIQUE STABILISÉ PAR LE GAZ DE TRAVAIL EN UTILISANT LES ÉQUATIONS MAGNÉTOHYDRODYNAMIQUES

Résumé

En utilisant les équations magnétohydrodynamiques et le modèle de l'arc électrique à courant continu qui est stabilisé avec le gaz de travail, on a présenté l'équation pour déterminer les caractéristiques tension-courant.

Les caractéristiques théoriques calculées, à partir de l'équation proposée, ont été présentées et comparées avec les caractéristiques déterminées expérimentalement.

C. KRÓLIKOWSKI, A. KAMIŃSKA-PRANKE, R. NIEWIEDZIAŁ

ANALYTISCHER BESTIMMUNGSVERSUCH VON STROM-SPANNUNGS-KENNLINIEN DES MIT ARBEITSGAS STABILISIERTEN LICHTBOGENS UNTER AUSNUTZUNG MAGNETOHYDRODYNAMISCHER GLEICHUNGEN

Zusammenfassung

In der Arbeit wurden Strom-Spannungs-Kennlinien aufgrund magnetohydrodynamischer Gleichungen und eines Modells des elektrischen, mit Arbeitsgas stabilisierten Gleichstromlichtbogens dargestellt. Auf der vorgeschlagenen Gleichung basierend, wurden die theoretischen Kennlinien berechnet, und mit Versuchsergebnissen versehen.

Ч. КРУЛИКОВСКИ, А. КАМИНЬСКА-ПРАНКЕ, Р. НЕВВЕДЗЯЛ

ПОПЫТКА ОПИСАНИЯ ВОЛЬТАМПЕРНЫХ ХАРАКТЕРИСТИК ДУГИ
СТАБИЛИЗИРОВАННОЙ ПОТОКОМ ГАЗА
НА ОСНОВАНИИ УРАВНЕНИЙ МАГНИТОГИДРОДИНАМИКИ

Р е з ю м е

Проведен анализ нагрева газа электрической дугой на основании уравнений магнитогиродинамики. В результате получено уравнение описывающее зависимость между напряжением и током дуги, физическими свойствами газа, геометрией разрядного канала. Уравнения вольтамперных характеристик выведены для дуги, горящей в одноатомных газах, принимая теплоемкость этих газов постоянной и для двуатомных газов, принимая что их теплоемкость зависит от температуры. Расчеты проведены для дуг, горящих в аргоне и азоте. Полученные результаты сравнены с результатами эксперимента.

Influence des barrières sur l'amorçage des moyens intervalles d'air pointe-plan

AHMED BOUBAKEUR (ALGER)

Instytut Wysokich Napięć Politechniki Warszawskiej

Reçu le 17.12.1979

L'article présente les résultats des essais concernant l'influence des barrières propres et polluées sur la tension disruptive du système pointe-plan. Tensions appliquées: $+1,2/50 \mu\text{s}$, $+170/2500 \mu\text{s}$, 50 Hz. Distances d'électrodes $30 \div 300 \text{ cm}$. On a défini les conditions dans lesquelles le canal de la décharge contourne la barrière sans perforation et sans décharges glissantes.

1. INTRODUCTION

Le recours, pour des raisons économiques, à des tensions de plus en plus élevées pour le transport de l'énergie électrique a nécessité, et nécessitera encore, des distances d'isolement dans l'air de plus en plus grandes.

Dans des conditions atmosphériques données, la disruption de l'air dépend largement de la forme des électrodes et du genre de tension appliquée. Dans la pratique, la majorité des configurations rencontrées peut être caractérisée, à partir d'une certaine distance entre électrodes, par les géométries à champ non-uniforme *pointe-pointe* ou *pointe-plan* ou bien par celles qui possèdent la rigidité électrique intermédiaire. De toutes ces configurations seule la géométrie *pointe-plan* est la plus défavorable en ce qui concerne la rupture diélectrique de l'air qui se manifeste dans ce cas aux plus faibles tensions quand la pointe est positive. Cela traduit l'importance que revêt l'utilisation du système d'électrodes *pointe-plan* dans plusieurs recherches liées à la rigidité diélectrique de l'air. En effet, sur la base des résultats obtenus avec ce système d'électrodes, il devient facile d'estimer les distances minimales dans l'air à respecter pour une tension de tenue donnée [1]. Ceci devient de plus en plus nécessaire dans le domaine de très hautes tensions où les distances d'isolement ne restent plus à peu près proportionnelles à la tension de tenue comme dans le cas des moyennes tensions. Le choix des distances d'isolement est en général basé sur les résultats contenus avec les intervalles d'air *pointe-plan* soumis à la contrainte de tension la plus défavorable. Il fut vérifié [2, 3, 4, 5] que cette contrainte est représentée, pour les longs intervalles d'air, par les chocs positifs de manoeuvres à durée de front relativement longue.

Pour augmenter la tension disruptive des intervalles d'air, on a proposé l'emploi de barrières isolantes. Cet emploi doit tenir compte de la forme géométrique des électrodes et de la nature de la barrière, la pollution atmosphérique y incluse. L'avantage fourni

par les barrières est surtout l'augmentation appréciable de la tension disruptive des intervalles d'air *pointe positive-plan*. Cette augmentation fut remarquée par plusieurs chercheurs.

Pour les intervalles relativement petits (inférieurs à 30 cm), l'influence des barrières sur la rupture diélectrique de l'air fut déjà remarquée vers les années 1920 par C. P. Steinmetz [6]. Les principales recherches furent entreprises ensuite par E. Marx [7, 8] et H. Roser [9, 10]. Ces auteurs ont étudié l'influence de différents types de barrières, dans différents intervalles d'air à champ nonuniforme, aux tensions de chocs à front raide (type chocs de foudre), à tension continue et à fréquence industrielle. L'amélioration de la tension disruptive de l'air fut notable seulement pour le système d'électrodes *pointe positive-plan*. Roser [9] considère que cette amélioration serait due surtout à l'uniformisation du champ électrique, dans l'espace *barrière-plan*, provoquée par la charge positive déposée sur la surface de la barrière en face de la pointe. Cette constatation fut vérifiée ensuite par N. F. Wolochtkhenko [11] et par M. P. Verma [12]. Les études qui suivirent, après celles de Marx et Roser, ont été orientées surtout vers l'influence de la forme des électrodes, la distribution du champ électrique dans l'arrangement *pointe-barrière-plan*, l'existence de différentes phases de rupture (décharge *pointe-barrière* suivie par la décharge *barrière-plan* et le contournement de la barrière), la géométrie et la nature de la barrière, et l'influence des barrières polluées.

Pour les intervalles d'air longs de quelques mètres, les recherches sont plutôt rares. Les premières études de V. S. Komelkov & A. M. Lifchits [14] et de O. Salka & H. Norinder [15], vers les années 1950, concernaient surtout le mécanisme des décharges dans l'air. Des recherches plus détaillées et concernant la rupture diélectrique de l'air furent réalisées en 1973 par R. Finsterwalder [16]. Cet auteur étudia l'influence d'une barrière isolante sur la tension 100% de décharge disruptive aux chocs de manoeuvres des intervalles d'air *tige-sphère* longs jusqu'à 1 m. Il a défini les variations de cette tension disruptive en fonction de la position de la barrière dans l'intervalle et du temps jusqu'à la crête des chocs de manoeuvres. La tension disruptive augmente quand la barrière s'approche de la pointe, tout en restant minimale pour les chocs ayant une durée de front comprise entre 100 et 200 μ s. Pour les intervalles d'air *pointe-plan* longs de 3 m, M. Darveniza et B. Holcombe présentèrent au Congrès Mondial de l'Electrotechnique de 1977, leurs travaux concernant l'influence d'une barrière isolante sur la rupture diélectrique de l'air aux chocs de manoeuvres et de foudre [17, 18]. Ils constatèrent que, aussi pour les longs intervalles, l'augmentation de la tension disruptive a lieu quand la pointe est positive. Cette augmentation a été obtenue aussi à l'aide d'une barrière métallique à bords arrondis pour les intervalles *pointe-plan* de 80 cm par A. Sevigny et autres [19].

Contrairement à celui des petits intervalles d'air, le domaine de longs intervalles reste encore à explorer et plusieurs recherches sont à réaliser dont certaines d'entre elles sont présentées dans ce rapport.

2. PROGRAMME DE RECHERCHES

Notre étude touchera les trois genres de tension suivants: chocs de foudre $+1,2/50 \mu$ s, chocs de manoeuvres $+170/2500 \mu$ s et tension à fréquence industrielle 50 Hz, et sera

plus détaillée pour les deux derniers genres de tension car ils correspondent aux contraintes les plus sévères pour les intervalles d'air *pointe-plan* qui nous intéressent. Ces intervalles étaient compris entre 40 cm et 200 cm et nommés *moyens intervalles*. S'il s'agit de l'application de la tension de choc de manoeuvres à durée de front égale à 170 μ s, il fut tenu compte du fait que pour les intervalles d'air *pointe positive-plan* compris entre 100 cm et 300 cm, la tension disruptive reste très voisine d'une valeur minimale pour les chocs ayant une durée de front comprise entre 100 μ s et 250 μ s [4, 5]. D'autre part, d'après Finsterwalder [16], la tension disruptive des intervalles d'air avec barrière isolante devient minimale pour les chocs ayant une durée de front comprise entre 100 μ s et 200 μ s indépendamment de la position de la barrière dans l'intervalle.

Aux tensions de choc, l'étude de la polarité négative ne fut pas effectuée car il a été déjà démontré qu'avec la pointe négative la barrière n'entraîne aucune augmentation notable de la tension disruptive [18].

Comme il a été vérifié auparavant [10, 18], l'influence de la barrière isolante est surtout liée à l'accumulation des charges sur la surface de la barrière du côté de la pointe positive. Les principaux facteurs dont dépendrait l'influence de la barrière seraient alors, pour une forme de tension donnée et une même configuration d'électrodes, l'état de surface de la barrière, sa géométrie propre (largeur) et sa position dans l'intervalle entre électrodes. En conséquence, l'étude concernera les barrières à surfaces propres et les barrières recouvertes d'une couche semi-conductrice caractérisant ainsi la barrière polluée.

3. ESSAIS

Pour effectuer nos recherches, nous avons utilisé les principales sources de tension qui existent au Laboratoire H. T. de l'Ecole Polytechnique de Varsovie. La tension à fréquence industrielle était fournie par un transformateur d'essais 750 kV—750 kVA. La mesure de tension a été réalisée grâce à un diviseur de tension capacitif (15 pF). Les tensions de choc de foudre et de manoeuvres ont été fournies par un générateur de choc correspondant à un montage multiplicateur de Marx de 16 étages avec une capacité de choc par étage de 125 nF. L'énergie nominale de ce générateur est de 32 kJ et sa tension nominale de 2800 kV pour les chocs de foudre et environ 1400 kV pour les chocs de manoeuvres. Pour les chocs de foudre, la capacité de charge du générateur était de 500 pF, la résistance série de 800 Ω et la résistance parallèle de 6,4 k Ω . Cette dernière résistance a été utilisée comme branche H. T. d'un diviseur de tension résistif pour la mesure de la tension de choc. Pour les chocs de manoeuvres, la capacité de charge était de 800 pF, la résistance de front égale à 45 Ω et la résistance parallèle égale à 382,5 k Ω . La capacité de charge fut utilisée comme branche H. T. d'un diviseur capacitif de tension pour la mesure de la tension de choc de manoeuvres.

Les essais ont été réalisés conformément aux recommandations de la C.E.I. [20] en ce qui concerne la détermination de la tension moyenne de décharge disruptive et les corrections atmosphériques. La tension 50% de décharge disruptive de choc (foudre et manoeuvres) était déterminée par la méthode „montées et descentes” [21, 22]. Le nombre de chocs par série était de 50 et le temps entre deux chocs successifs était de l'ordre de

60 s. Ce temps était suffisant pour que la charge déposée sur la barrière n'aie pas d'influence sur la détermination statistique de la tension 50% de décharge disruptive.

L'arrangement *pointe-barrière-plan* utilisé est présenté sur la figure 1. L'électrode H.T. était constituée par un cylindre en fer de diamètre égal à 156 mm, terminé par une pointe conique 30° ayant un bout hémisphérique de 2 cm de diamètre. L'électrode plate, mise à la terre, était constituée par une plaque métallique circulaire en fer de 4 m de diamètre,

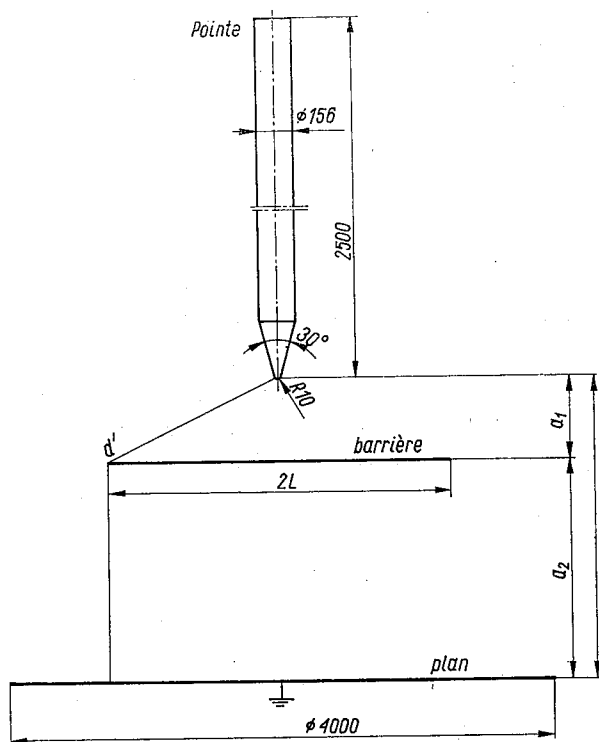


Fig. 1. Arrangement *pointe-barrière-plan*

ayant une épaisseur de 1,5 mm. Les barrières isolantes utilisées étaient en papier bakéliné ($\epsilon = 5,8$) de forme carrée (largeur maximale 130 cm) et avaient une épaisseur de 6 mm. Cette épaisseur empêcha la perforation des barrières par la décharge disruptive et assura une bonne rigidité mécanique de ces barrières (très faible flèche). Les barrières étaient suspendues par quatre fils en nylon (diamètre égal à 4 mm) à un support isolant.

4. RÉSULTATS

4.1. Barrière non-polluée

Il fut déjà observé [8, 9, 10, 11, 12, 13] que la tension disruptive des petits intervalles d'air *pointe-positive-plan* varie essentiellement en fonction de la position de la barrière par rapport aux électrodes. Cette tension passe par un maximum quand la barrière est près de la pointe. Cette constatation reste aussi valable pour les moyens intervalles d'air

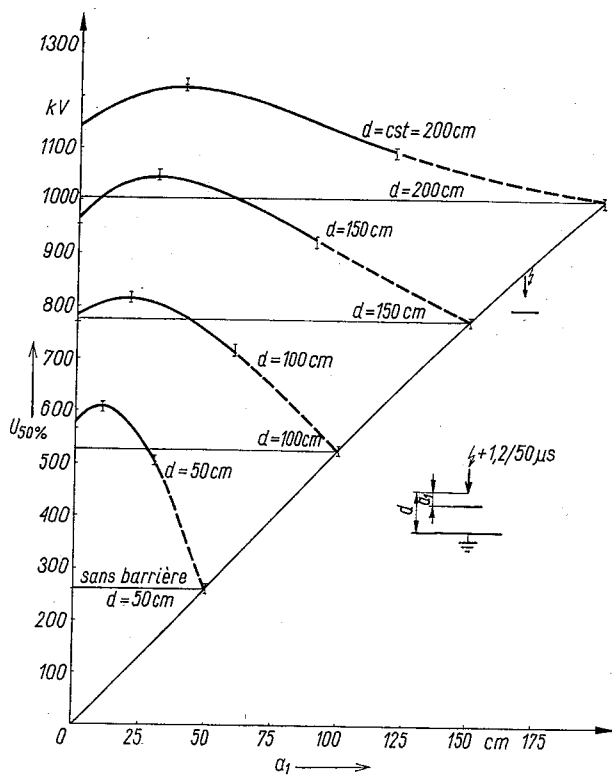


Fig. 2. Tension 50% de décharge disruptive de choc de foudre en fonctions de la position a_1 de la barrière

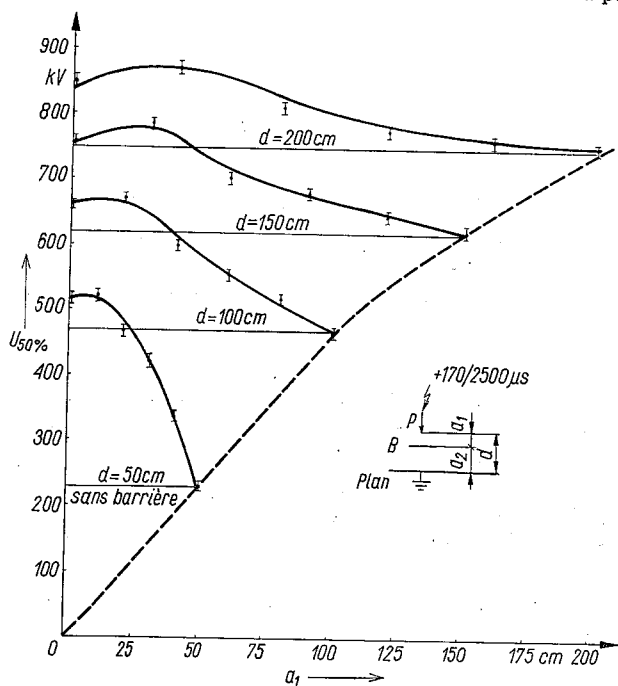


Fig. 3. Tension 50% de décharge disruptive de choc de manoeuvres en fonction de la position de la barrière

En effet quand la barrière isolante à surfaces propres (résistivité superficielle de l'ordre de $10^{11} \Omega$) s'approche de la pointe, la tension disruptive de l'air augmente considérablement et atteint une valeur maximale quand la barrière se trouve près de la pointe à environ 20% de la distance pointe-plan ($a_1 = 0,2d$). Aux tensions de choc de foudre (Fig. 2) et de choc de manoeuvres (Fig. 3) l'augmentation optimale de la tension 50% de décharge disruptive est de l'ordre de 130% pour la plus petite distance *pointe-plan* ($d = 50$ cm) et de l'ordre de 20% pour la plus grande distance ($d = 200$ cm). Pour le cas de la tension à fréquence industrielle (Fig. 4) l'augmentation optimale de la tension disruptive est de l'ordre de 135% pour la distance d égale à 40 cm et de l'ordre de 40% pour $d = 120$ cm.

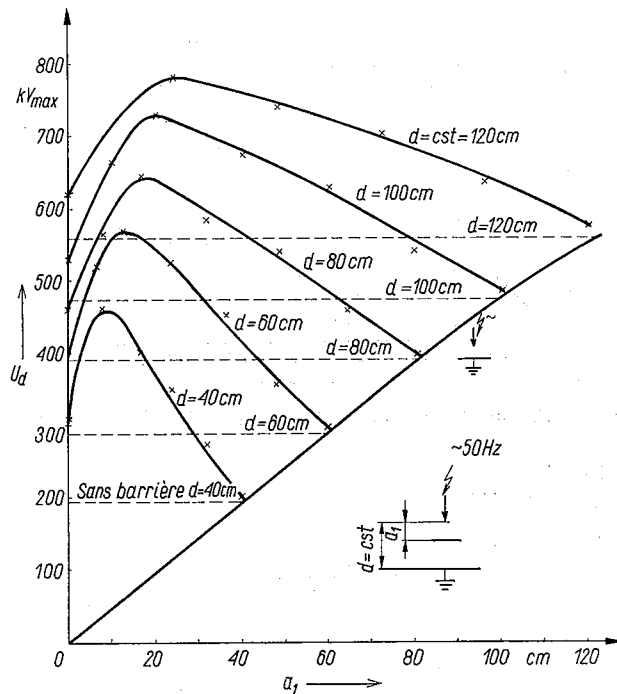


Fig. 4. Tension disruptive à fréquence industrielle en fonction de la position de la barrière

L'amélioration de la rigidité diélectrique des moyens intervalles d'air *pointe-plan* est due surtout au fait que la barrière forme un obstacle géométrique à la décharge disruptive directe. Avec les barrières utilisées (épaisseur égale à 6 mm), la décharge se développe de la pointe vers le bord de la barrière et de ce bord vers le plan mis à la terre. Sur la Fig. 5 sont présentées quelques photographies de décharges disruptives, aux chocs de foudre (Fig. 5a et 5b), aux chocs de manoeuvres (Fig. 5c) et à fréquence industrielle (Fig. 5d). Sur ces photographies, il est aussi facile de voir que la longueur du canal de la décharge est en grandeur très proche de la distance d' *pointe-bord de la barrière-plan* (Fig. 1).

A partir des observations photographiques et de l'analyse des résultats numériques il fut établi [23] que la tension disruptive correspondant à la longueur du canal de la décharge d' est approximativement égale à la tension disruptive de l'arrangement *pointe*

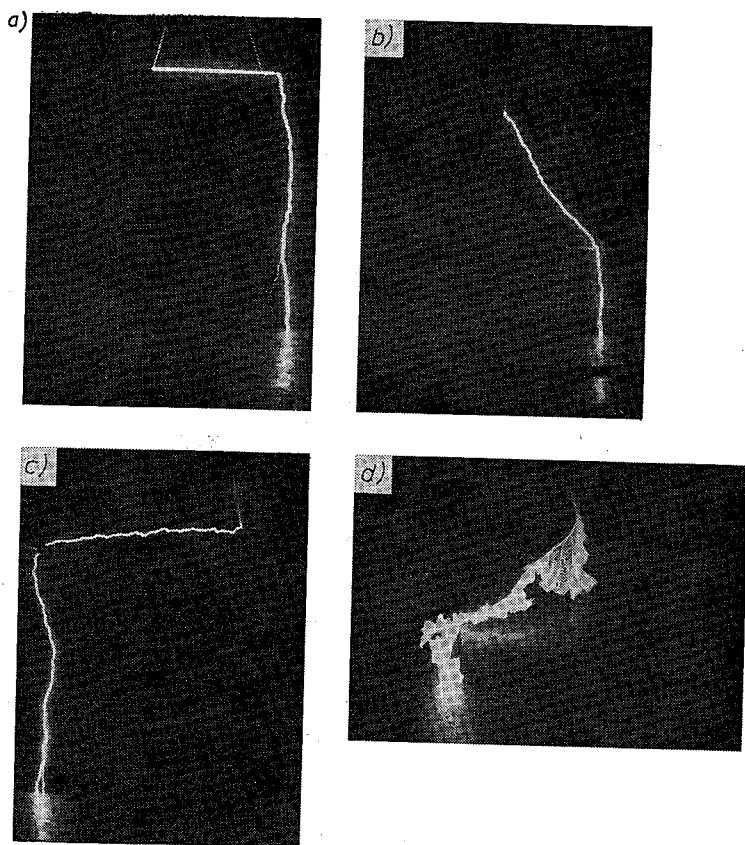


Fig. 5. Décharge disruptive

- a — tension $+1,2/50 \mu s$, $U = 1130 \text{ kV}$, $d = 200 \text{ cm}$, $a_1 = 0 \text{ cm}$
 b — tension $+1,2/50 \mu s$, $U = 920 \text{ kV}$, $d = 150 \text{ cm}$, $a_1 = 90 \text{ cm}$
 c — tension $+170/2500 \mu s$, $U = 620 \text{ kV}$, $d = 100 \text{ cm}$, $a_1 = 20 \text{ cm}$
 d — tension 50 Hz , $U = 570 \text{ kV}$, $d = 100 \text{ cm}$, $a_1 = 50 \text{ cm}$

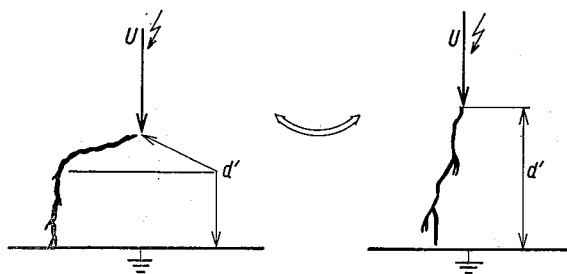


Fig. 6. Equivalence électro-géométrique des systèmes avec barrière et des systèmes sans barrière

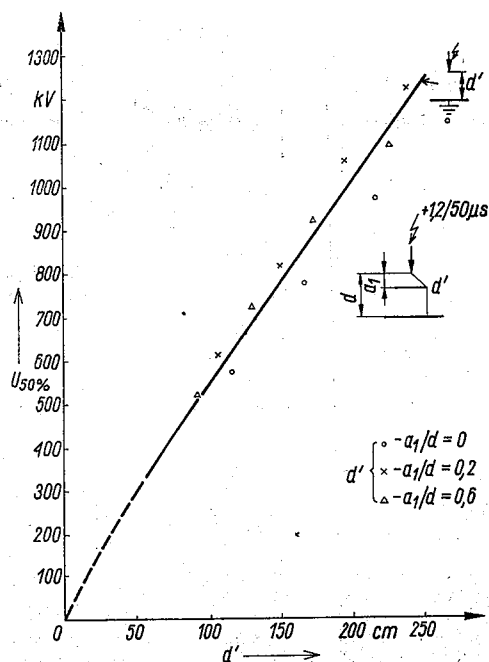


Fig. 7. Comparaison entre le système *pointe-barrière-plan* et le système *pointe-plan*
Tension $+1,2/50 \mu s$

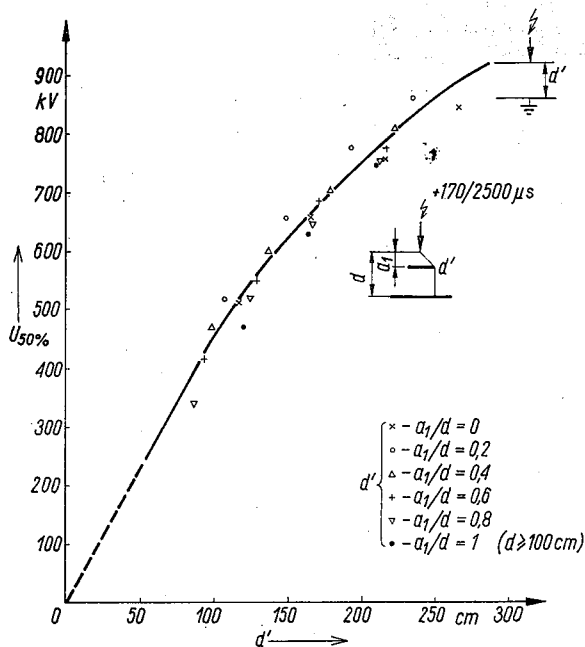


Fig. 8. Comparaison entre le système *pointe-plan-barrière* et le système *pointe-plan*
Tension $+170/2500 \mu s$

plan sans barrière ayant une distance entre électrodes égale à d' (Fig. 6). Autrement dit, un intervalle d pointe-barrière-plan est, du point de vue rigidité, équivalent à un intervalle d' pointe-plan sans barrière. La comparaison entre les intervalles d' pointe-barrière-plan et pointe-plan est présentée sur la Fig. 7 pour les chocs de foudre, sur la Fig. 8 pour les chocs de manoeuvres et sur la Fig. 9 pour la tension à fréquence industrielle 50 Hz. L'équivalence entre les deux systèmes reste généralement valable pour les positions de la barrière comprises entre $a_1/d = 0,2$ et $a_1/d = 0,6$. Elle n'est par contre plus valable dans certains cas quand la barrière est tout près de l'une des électrodes ($a_1/d = 0$ et $a_1/d = 0,8$). Dans

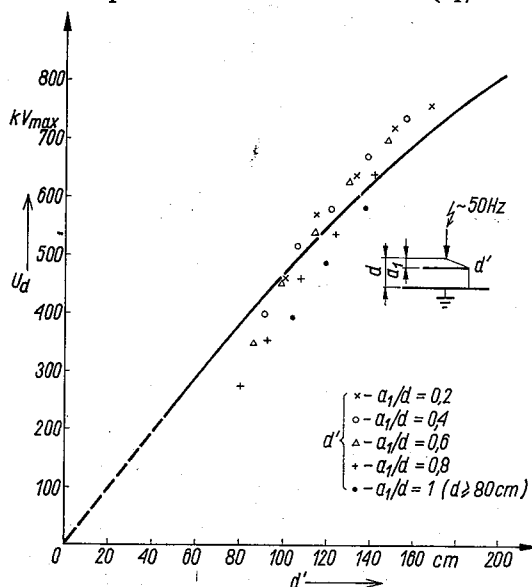


Fig. 9. Comparaison entre le système *pointe-barrière-plan* et le système *pointe-plan*
Tension à 50 Hz

ces cas précis, l'apparition des décharges superficielles glissantes sur la surface de la barrière en face de la pointe facilite l'amorçage des intervalles *pointe-barrière-plan* à des tensions beaucoup plus basses que celles correspondant à des intervalles *pointe-plan* d' sans barrière. La forme de ces décharges est présentée sur les Fig. 10 et Fig. 11 respectivement pour le cas des chocs de manoeuvres et celui des chocs de foudre (photographie de contact) et pour le cas de la tension à fréquence industrielle sur la Fig. 12 (photographie statique). Ces figures correspondent à la position de la barrière en contact avec la pointe (Fig. 11 et Fig. 12) ou tout près de la pointe (Fig. 10). Quand la barrière est proche du plan on a d'abord le développement de la décharge de façon directe de la pointe vers le milieu de la barrière, ensuite l'apparition de décharges glissantes contournant la surface en face de la pointe et finalement l'amorçage de l'intervalle d'air entre le bord de la barrière et le plan entraînant la rupture totale (Fig. 13).

En conclusion, en omettant les cas des ruptures avec présence de décharges glissantes ($a_1/d = 0$ et $a_1/d = 0,8$), l'augmentation de la tension disruptive des intervalles d'air *pointe-plan* provoquée par la barrière est conséquente à la déviation du canal de la décharge disruptive vers le bord de la barrière. Cette décharge a lieu approximativement pour une

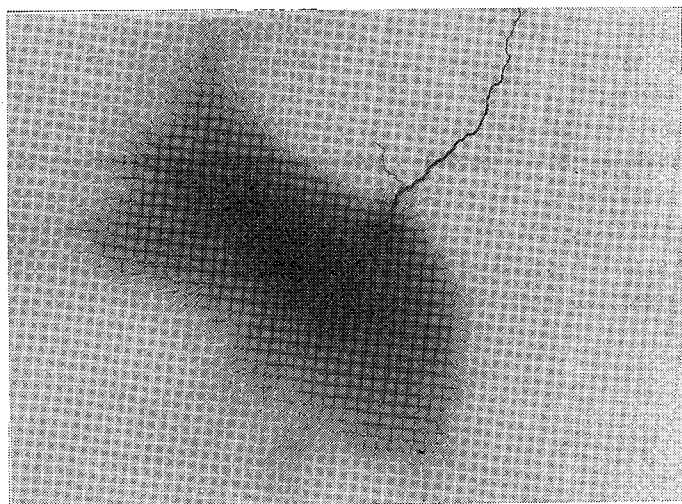


Fig. 10. Décharges superficielles (tension: choc de manoeuvre +170/2500 μ s)

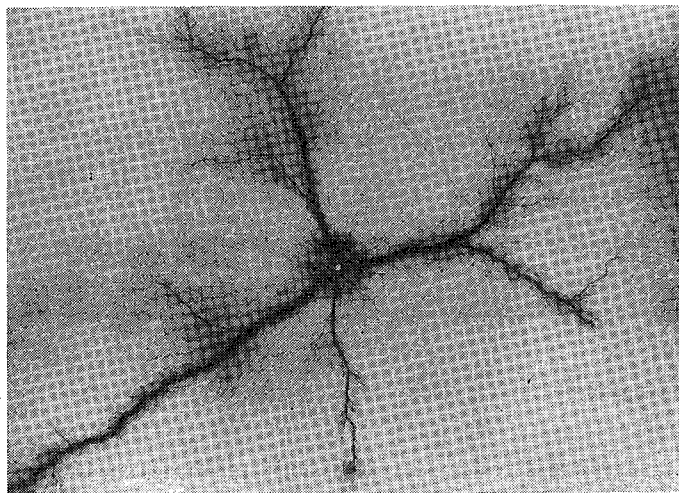


Fig. 11. Décharges superficielles (tension: choc de foudre +1,2/50 μ s)

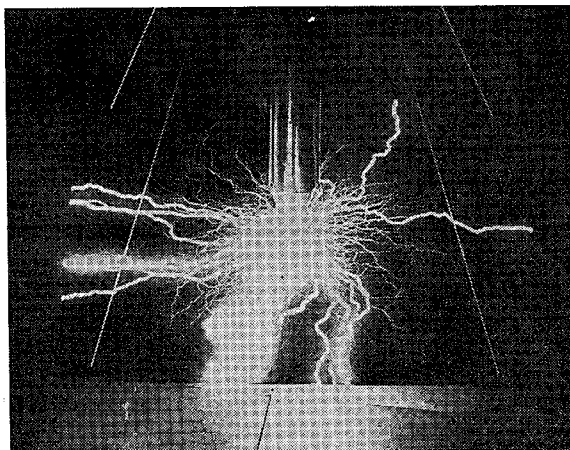


Fig. 12. Décharges superficielles (tension à 50 Hz)

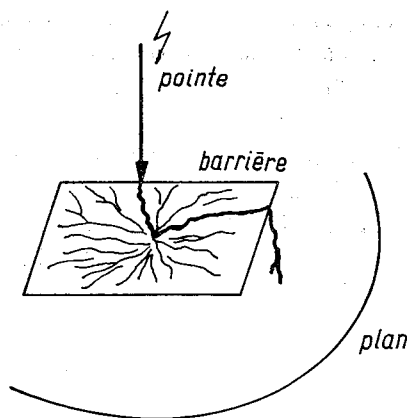
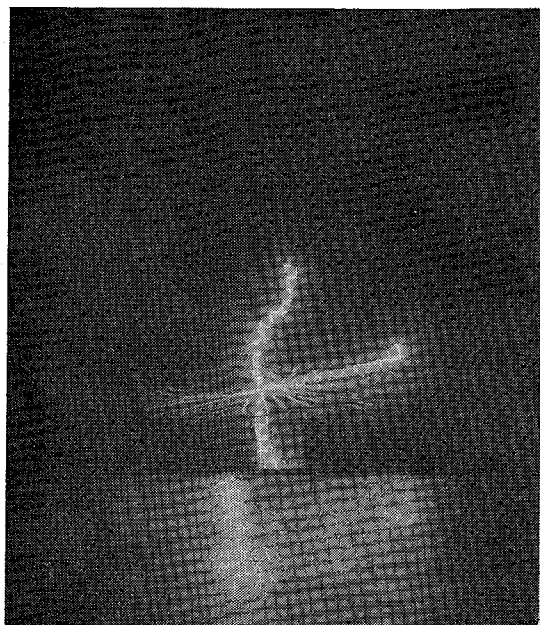


Fig. 13. Rupture en présence des décharges glissantes

Tension $+170/2500 \mu s$, $U = 380 \text{ kV}$, $d = 50 \text{ cm}$, $a_1 = 35 \text{ cm}$

même tension nécessaire à l'apparition d'une décharge de même longueur dans un intervalle *pointe-plan* sans barrière.

La déviation de la décharge vers le bord de la barrière n'est pas provoquée seulement par la présence de la barrière en tant qu'obstacle géométrique mais aussi et surtout par l'obstacle électrostatique que forme la charge électrique déposée sur la surface de la barrière du côté de la pointe. Ceci a été confirmé par l'étude des barrières trouées en leur milieu. En effet même avec un trou de 10 mm de diamètre en face de la pointe [23], la décharge disruptive passe par le bord de la barrière. Ceci toute fois quand la barrière n'est pas trop près de la pointe ($a_1 \ll 0,2 d$), cas où la décharge passe directement par le trou l'entraînant, de ce fait, aucune augmentation de la tension disruptive.

4.2. Barrière polluée

Dans le paragraphe précédent a été présentée l'importance de l'utilisation des barrières isolantes à surfaces propres dans les intervalles d'air *pointe-plan*. L'emploi des barrières en pratique doit aussi tenir compte des pollutions pouvant éventuellement exister. Pour cela une étude fut effectuée avec des barrières isolantes recouvertes de couches solides semi-conductrices sèches. Il n'est pas tenu compte ici de la nature de la pollution mais des conséquences qu'elle peut engendrer. De ce fait seule la conductivité superficielle des couches a été prise en considération. Les couches utilisées étaient composées d'un mélange de peinture semi-conductrice à base de graphite colloïdal avec de l'alcool isopropylique ($\text{iso-C}_3\text{H}_7\text{OH}$) et du bioxyde de silicium (SiO_2). Un choix adéquat des con-

centrations de chaque composant a permis d'avoir des couches ayant une conductivité superficielle, après séchage, comprise entre $0,01 \mu S$ et $50 \mu S$ [23]. L'étude a été limitée au cas le plus important c. à d. à la surface de la barrière en face de la pointe et à la tension de choc positif de manoeuvres. Les couches semi-conductrices ont été appliquées de façon uniforme sur la surface de la barrière et leur conductivité superficielle a été déterminée à l'aide d'une sonde. Cette méthode fut déjà utilisée dans d'autres recherches concernant les isolateurs pollués [24].

Avec une barrière à surface semi-conductrice, la tension disruptive des moyens intervalles d'air *pointe-plan* varie entre les valeurs obtenues avec la barrière isolante à surfaces propres et celles obtenues avec une barrière métallique de même forme (bords non arrondis). Un fait important (Fig. 14) est que la barrière polluée entraîne pratiquement

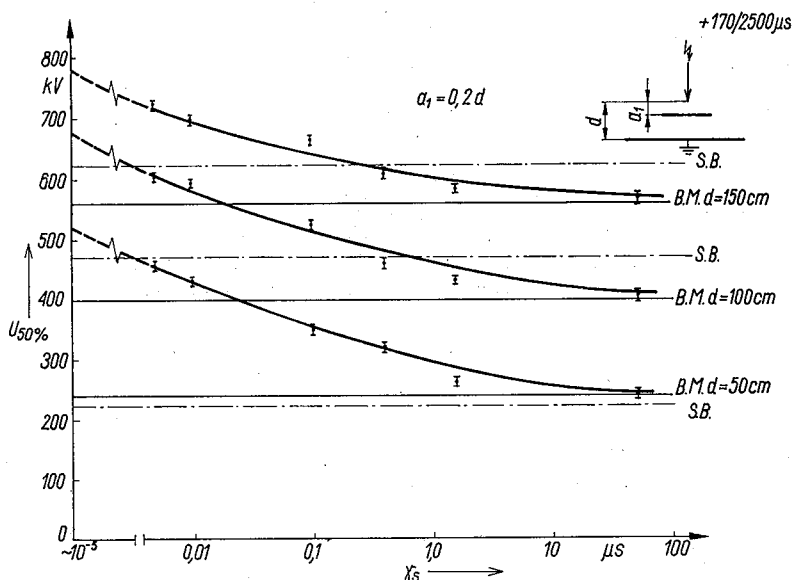


Fig. 14. Variation de la tension 50% de décharge disruptive de choc de manoeuvres $+170/2500 \mu s$ en fonction de la conductivité superficielle des couches semi-conductrices

les mêmes résultats que la barrière métallique dès que la conductivité superficielle des couches semi-conductrices dépasse $1,6 \mu S$. Ceci est aussi nettement remarqué par la forme de la décharge disruptive. Dans le cas des couches à conductivité superficielle inférieure à $0,4 \mu S$, cette décharge se développe de la même manière que dans le cas des barrières à surfaces propres, c. à d. qu'elle part de la pointe vers le bord de la barrière et ensuite vers le plan (Fig. 5). Tandis qu'à partir des conductivités de l'ordre de $1,6 \mu S$, la décharge se fait en deux étapes comme dans le cas de la barrière métallique. Il y a d'abord la décharge *pointe-milieu de la barrière* et ensuite la décharge *bord de la barrière-plan* (Fig. 15). Avec une barrière polluée ayant des bords à faible rayon de courbure, comme il en est le cas ici, après l'amorçage de l'intervalle *pointe-milieu de la barrière*, entraînant directement le contournement de la barrière, vu la faible d.d.p. sur la surface polluée, des décharges préliminaires (streamers) apparaissent aux bords de la barrière et facilitent la rupture

totale à des tensions relativement faibles. Les bords de la barrière jouent dans ce cas un rôle de pointe. Pour les distances *pointe-plan* supérieures ou égales à 100 cm, la tension disruptive devient même inférieure à celle du système sans barrière.

L'étude n'a pas été élargie aux cas des barrières ayant des bords à grand rayon de courbure, mais il est à rappeler que d'autres recherches [19] ont prouvé qu'une barrière métallique, dont les bords sont suffisamment arrondis, entraîne aussi une augmentation de la tension disruptive. De ce fait, il serait important d'apporter une amélioration à la forme des bords de la barrière isolante en cas d'utilisation dans des conditions de pollution. Dans les recherches futures il faudrait aussi étudier le cas de la barrière complètement

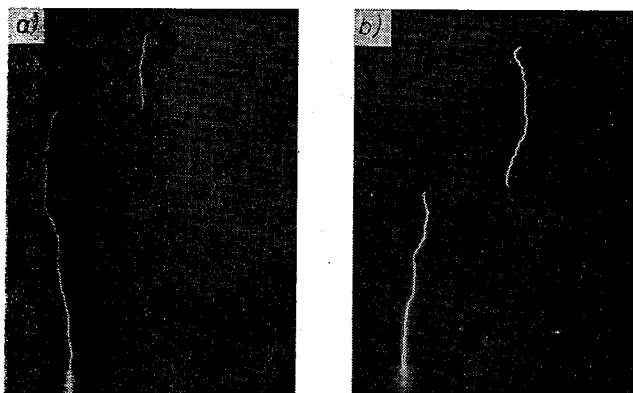


Fig. 15. Décharge disruptive en présence de barrière conductrice

a — tension $+1,2/50 \mu\text{s}$, $U = 890 \text{ kV}$, $d = 200 \text{ cm}$, $a_1 = 40 \text{ cm}$

b — tension $+170/2500 \mu\text{s}$, $U = 610 \text{ kV}$, $d = 200 \text{ cm}$, $a_1 = 80 \text{ cm}$

polluée, c. à d. dont les deux surfaces seraient recouvertes d'une couche polluante. Dans notre étude, nous avons effectué des essais avec une barrière isolante ayant sa surface en face du plan recouverte d'une couche métallique en argent colloïdal et sa surface en face de la pointe laissée propre. Ces essais ont confirmé que la surface en face du plan n'aurait pas un rôle important en présence de pollution. En effet, avec ce dernier type de barrière la tension disruptive reste supérieure à celle obtenue avec la barrière dont la surface en face de la pointe était recouverte d'une couche ayant une conductivité superficielle de $0,1 \mu\text{S}$.

5. CONCLUSIONS

1. Une barrière isolante, à surfaces propres (grande résistivité superficielle), entraîne une augmentation considérable de la tension disruptive des intervalles d'air *pointe-plan* quand elle est placée près de la pointe et surtout à environ 20% de la distance entre électrodes. Cette augmentation est causée surtout par l'allongement du canal de la décharge disruptive qui suit un chemin allant de la pointe vers le bord de la barrière et de ce bord vers le plan mis à la terre.

2. La tension disruptive des moyens intervalles *pointe-barrière-plan* est approximativement égale à celle des intervalles *pointe-plan* ayant un écart entre électrodes égal à la distance *pointe-bord de la barrière-plan*. Cette conclusion cesse d'être valable quand la barrière est tout près de l'une des électrodes.

3. La charge spatiale déposée sur la barrière provoque la déviation de la décharge vers le bord de la barrière. Cette charge apparaît quand les têtes de streamers touchent la barrière, entraînant l'apparition des décharges superficielles. Elle se manifeste aussi dans le cas où, vis-à-vis de la pointe, la barrière possède un trou de 10 mm de diamètre et la décharge passe par le bord de la barrière.

4. Quand la barrière isolante est polluée par une couche semiconductrice, uniformément répartie sur la surface du côté de la pointe, la tension disruptive des intervalles d'air *pointe-barrière-plan* devient comprise entre celles des intervalles avec une barrière complètement isolante et celle des intervalles avec une barrière conductrice (métallique). La barrière polluée commence à avoir la même influence que la barrière métallique dès que la conductivité superficielle des couches étrangères dépasse $1,6 \mu S$. Cette influence est caractérisée par une diminution de la tension disruptive par rapport à celle des intervalles sans barrière.

5. Dans les recherches futures, la conclusion 2. doit être complétée par l'étude des barrières ayant des largeurs plus grandes. De telles recherches permettront aussi une explication plus précise du mécanisme de rupture en présence des décharges superficielles glissantes.

Pour ce qui est de l'utilisation des barrières en pratique, il faudrait prendre en considération les cas des barrières complètement et aussi partiellement polluées. L'étude de la forme des bords de la barrière serait de première importance.

BIBLIOGRAPHIE

1. C. E. I., *Coordination de l'isolement*, publication 71—2, 1976.
2. L. Paris, A. Taschini, K. H. Schneider, K. H. Weck, *Distances d'isolement dans l'air, entre phase et terre ainsi qu'entre phases, dans les postes*, Electra, 29, 1973, p. 29.
3. L. Paris, *Influence of air gap characteristics on line-to-ground switching surge strength*, IEEE Transactions on PAS, Vol. PAS 86, Nr 8, 1967, p. 936.
4. G. Leroy, G. Gallet, R. Kosztaluk, I. L. Kromer, *Réseaux aériens à ultra haute tension: calcul des distances d'isolement*, RGE, numéro spécial, Juin 1974, p. 27.
5. R. Kosztaluk, *Izolacja sieci skrajnie wysokich napięć*. Prace I. En., z. 3, Warszawa 1975.
6. A. Roth, *Hochspannungstechnik*, Springer-Verlag, Wien 1927.
7. E. Marx, *Der elektrische Durchschlag von Luft im unhomogenen Felde*, Archiv für Elektrotechnik Bd 24, 1930, p. 61.
8. E. Marx, *Der Durchschlag der Luft im unhomogenen elektrischen Feld bei verschiedenen Spannungssarten*, ETZ, H. 33, 1930, p. 1161.
9. H. Roser, *Dünne Schirme im raumladungsbeschwerten Entladungsfeld zum Zwecke der Erhöhung der Durchschlagspannung und Begrenzung der Trägerströme in Luft*, Dissertation TH Braunschweig 1930.
10. H. Roser, *Schirme zur Erhöhung der Durchschlagspannung in Luft*, ETZ, H. 17, Bd 53, 1932, p. 411.
11. N. F. Wolochchenko, *Problèmes relatifs au mécanisme de l'effet de barrière* (en russe), Elektrichestvo, nr 6, 1947, p. 61.

12. M. P. Verma, *Durchschlagspannung und Durchschlagvorgang für die Anordnung Spitze-platte mit Schirm*, Dissertation, TU Dresden, 1961.
13. J. Pillig, *Luftisolierung mit Isolierstoffbarrieren und verkleidungen bei Wechselfspannung*, Dissertation, TU Dresden, 1968.
14. V. S. Komelkov, A. M. Lifchits, *Influence of barrier on the development of the electric discharge in long gaps*, Izv. Akad. Nauk. URSS, otdel T. N., nr 10, Moscou 1950, p. 1463.
15. O. Salka, H. Norinder, *Screens in long discharge gaps*, Arkiv för Fysik, Bd 6, nr 17, 1952, p. 151.
16. R. Finsterwalder, *Isolierschirm im inhomogenen Feld bei positiven Stoßspannungen unterschiedlicher Anstiegszeit*, Dissertaion, TU Stuttgart, 1973.
17. M. Darveniza, B. Holcombe, *The influence of a thin floating barrier on the switching impulse strength of a long air gap with a plastic sheet barrier*, Proc. IEEE Letters, Vol. 64, Nr 6, 1976, p. 1017.
18. M. Darveniza, B. Holcombe, *The switching impulse strength of a long air gap with a plastic sheet barrier*, Congrès Electrotechnique Mondial, rap. 104, Section 2, Moscou 1977.
19. A. Seigny, I. B. Jordan, S. St-Arnaud, *Floating barrier influence upon point-to-plane breakdown*, 3rd Int. Conf. on Gas Discharge, IEE Conf. publication 118, London 1974, p. 438.
20. C. E. I., *Technique des essais à haute tension*, publications 60-1 et 60-2, édit. 1973.
21. W. J. Dixon, A. M. Mood, *A method for obtaining and analysing sensitivity data*, Journal of ASA, Vol. 43, 1948, p. 109.
22. J. Kučera, *Confidence limits if 50% flashover voltage in Up-and-Down method*, Acta Technica ČSAV, Vol. 17, nr 1, 1972, p. 61.
23. A. Boubakeur, *Influence des barrières sur la tension de décharge disruptive des moyens intervalles d'air pointe-plan*, Thèse de doctorat, Politechniki Warszawaska, Varsovie 1979.
24. Z. Pohl, E. Sojda, *O niektórych zasadniczych zagadnieniach metodyki prób zabrudzeniowych*, Prace Naukowe IPEiE Politechniki Wrocławskiej, nr 4, Wrocław 1971, p. 109.

A. BOUBAKEUR

THE INFLUENCE OF BARRIERS ON DIELECTRIC STRENGTH OF MEAN AIR GAPS IN POINT-PLANE ARRANGEMENTS

Summary

The paper discusses the results of investigations of polluted and clean barriers influence on the breakdown voltage between point and plane electrodes. Voltage 1,2/50 μ s, 170/2500 μ s, 50 Hz was used. The distance between electrodes changes from 30 cm to 300 cm. The conditions of existence of flashover around the barrier without breakdown and without surface discharge are determined.

A. BOUBAKEUR

WPLYW PRZEGRÓD NA WYTRZYMAŁOŚĆ ELEKTRYCZNĄ POWIETRZA DLA ŚREDNICH ODSTĘPÓW UKŁADU „OSTRZE-PŁYTA”

Streszczenie

Artykuł omawia rezultaty badań dotyczących wpływu przegród czystych i zabrudzonych na napięcie przeskoku układu ostrze-płyta. Napięcie stosowane: +1,2/50 μ s, +170/2500 μ s, 50 Hz; odstęp elektrod: 0—300 cm. Określono warunki występowania przeskoku naokoło przegrody bez przebicia i bez wyładowań ślizgowych.

A. BOUBAKEUR

EINFLUß VON TRENNWÄNDEN AUF DIE ELEKTRISCHE LUFTFESTIGKEIT
BEI MITTLEREN ABSTÄNDEN DER ANORDNUNG „SPITZE-PLATTE“

Zusammenfassung

Man erwägt die Untersuchungsergebnisse bei sauberen und schmutzigen Trennwänden der Anordnung „Spitze-Platte“. Angewandte Spannungen: $+1,2/50 \mu\text{s}$, $+170/2500 \mu\text{s}$, 50 Hz. Elektrodenabstände: 30÷300 cm. Es wurden Bedingungen bestimmt, bei denen Überschläge rund um die Trennwand ohne Durchschlag und ohne Gleitentladungen stattfinden.

A. БОУБАКЭУР

ВЛИЯНИЕ БАРЬЕРОВ НА НАПРЯЖЕНИЕ ПЕРЕКРЫТИЯ ВОЗДУШНОГО ПРОМЕЖУТКА
ДЛЯ СРЕДНИХ РАССТОЯНИЙ ОСТРИЕ-ПЛОСКОСТЬ

Резюме

Обсуждены результаты исследований влияния чистых и загрязненных барьеров на напряжение перекрытия промежутка острие-плоскость. Применяемые напряжения: $+1,2/50 \mu\text{s}$, $+170/2500 \mu\text{s}$; 50 гц. Расстояние электродов: 30—300 см. Определены условия выступления перекрытия вокруг барьера при отсутствии скользящих разрядов и пробоя твердого материала барьера.

Zachowanie się kontaktronów gazowanych w środowisku ciekłego azotu

BOGDAN MIEDZIŃSKI (WROCŁAW)

Instytut Energoelektryki Politechniki Wrocławskiej

HORST KIELOCH (WROCŁAW)

Instytut Podstaw Elektrotechniki i Elektrotechnologii

Politechniki Wrocławskiej

TADEUSZ ŻOŁNIERCZYK (WROCŁAW)

Centrum Naukowo-Produkcyjne Podzespołów i Urządzeń Elektronicznych

UNITRA-DOLAM

Otrzymano 3.12.1979

W artykule przedstawiono wyniki badań właściwości kontaktronów wodorowych zwieranych w środowisku ciekłego azotu. Zbadano wpływ środowiska ciekłego azotu na wartość i charakter zmian rezystancji kontaktronu w zależności od wartości prądu obciążenia kontaktronu i czasu jego przepływu. Określono wartości przepływów zadziałania i odpadania kontaktronów. Pomierzono wytrzymałość elektryczną szczeliny międzystykowej w zależności od rodzaju materiału stykowego i ilości cykli łączeniowych wykonanych przez kontaktron. Zbadano specyfikę wyładowania elektrycznego kontaktronów przerywających prądy indukcyjne o wartościach powyżej 0,5 A w środowisku ciekłego azotu i określono zdolność oraz trwałość łączeniową tych kontaktronów. Przeprowadzono również, za pomocą elektronowego mikroskopu scanningowego, obserwacje erozji styków kontaktronów. Badania wykonano dla kontaktronów, których stycki wykonane były z takich stopów żelaznikowych, jak FeNi29Co17 i FeNi54. Materiałami stykowymi były wolfram, rod i złoto. Bańki kontaktronów wykonane były ze szkła typu SL54.1. Badania zachowania się kontaktronów gazowanych w środowisku ciekłego azotu wykonano w celu określenia przydatności kontaktronów, zarówno jako elementów mierzających jak i wykonawczych, do pracy w niskich temperaturach.

1. WSTĘP

Zastosowanie niskich temperatur w procesie wytwarzania, przesyłu i rozdziału energii elektrycznej wymaga między innymi łączników elektrycznych nadających się do pracy w tych temperaturach. Na świecie przeprowadzane są już pierwsze próby z krio-łącznikami użytych mocy przeznaczonymi zarówno dla krio-generatorów jak i generatorów konwencjonalnych [4]. Mając na uwadze potencjalne zapotrzebowanie krioteknik na stykowy łącznik małej mocy, wydaje się celowe rozpoczęcie również odpowiednich prac zmierzających do skonstruowania takiego łącznika. W artykule przedstawiono wyniki badań niektórych właściwości kontaktronów gazowanych w środowisku ciekłego azotu. Badania wykonano w celu określenia przydatności kontaktronów jako elementów mierzających i wykonawczych do pracy w niskich temperaturach w obwodach silnoprądowych. Przedmiotem badań

były standardowe kontaktrony zwierne wodorowe typu ZW-103 posiadające stycзки wykonane ze stopów typu FeNi29Co17 lub FeNi54. Ciśnienie wodoru w bańkach kontaktronów wykonanych ze szkła SL54,1 wynosiło w temperaturze 20°C około 500—600 Tr. Materiałami stykowymi były W, Rh oraz Au.

2. PRZEPŁYW ZADZIAŁANIA I ODPADANIA

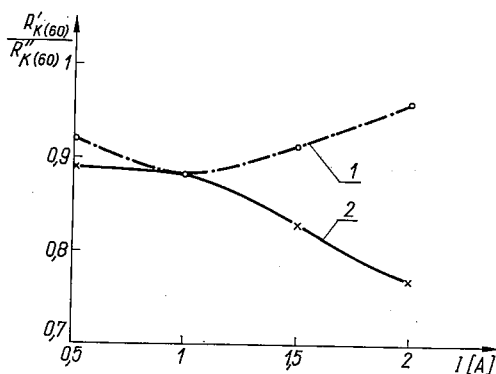
Możliwości wykorzystania kontaktronów do pracy w środowisku ciekłego azotu zależą przede wszystkim od wytrzymałości złącza szkło-metal oraz właściwości magnetycznych materiału styczek w niskiej temperaturze. Wytrzymałość złącza styczka-szkło decyduje w podstawowy sposób o mechanicznym uszkodzeniu się kontaktronu przy jego umieszczeniu w ciekłym azocie. Właściwości zaś magnetyczne styczek określają wartości przepływu zadziałania i odpadania kontaktronu. Firmy produkujące szkło na złącza szkło-metal gwarantują temperaturowy zakres pracy tych złączy bez wyjątków w granicach od -65°C do +250°C, podkreślając równocześnie możliwość rozszerzenia tego zakresu (od -270°C do +350°C) przez zastosowanie odpowiednich stopów. Również producenci stosowanych na stycзки materiałów magnetycznie miękkich typu Kovar dopuszczają możliwość pracy tych materiałów w temperaturach poniżej -150°C, nie ograniczając na przykład praktycznie wartości niskiej temperatury pracy materiału typu FeNi54 [3]. Dokładne zatem dopasowanie szkła i materiału styczek kontaktronu pod względem temperaturowej rozszerzalności, przy odpowiedniej technologii zatapiania zapewniającej staranne odprężenie złącza styczka-szkło, pozwoli na wykonanie kontaktronów przeznaczonych dla celów krioteknik. Badania kontaktronów typu ZW-103, przeznaczonych do pracy w temperaturze nie niższej od -55°C, wykazały niewielką procentowo liczebność uszkodzonych kontaktronów wskutek mikropełnięć szkła w złączu szkło-metal podczas gwałtownego umieszczania tych kontaktronów w ciekłym azocie. Nie zaobserwowano żadnych pęknięć bańki szklanej w obszarze zestyku kontaktronów mimo iż wydzielana w tym obszarze energia łuku podczas komutowania obwodów indukcyjnych była dość znaczna. Nie stwierdzono również w żadnym przypadku zmiany struktury (w martenzyt) materiału styczek co oznaczałoby nieodwracalną utratę właściwości magnetycznych tego materiału. Wartości przepływów zadziałania Θ_i i odpadania Θ_o badanych kontaktronów, umieszczonych w środowisku ciekłego azotu, niewiele różniły się od wartości pomierzonych w temperaturze otoczenia +20°C. Różnice te nie przekraczały $\pm 5\%$.

3. REZYSTANCJA KONTAKTRONU

Rezystancja kontaktronu R_k , będąca sumą rezystancji styczek oraz rezystancji przejścia R_p , odgrywa w obwodach silnoprądowych równie ważną rolę jak w obwodach o niewielkich wartościach prądu obciążenia. Wymagania odnośnie do rezystancji R_k , a zwłaszcza do podstawowego jej składnika — rezystancji przejścia, są przy większych wartościach napięć i prądów komutowanych nieco inne. Od materiału stykowego wymaga się bowiem przede wszystkim dużej odporności na wysokie temperatury z uwagi na zjawisko łuku elektrycznego, występujące w tych obwodach między otwierającymi się styczkami. Wartość

rezystancji R_p nie może być jednak zbyt duża, aby nie powodować nadmiernego nagrzewania się zestyku podczas przepływu prądu obciążenia, co działa stymulująco na łuk elektryczny podczas łączenia. Całkowita rezystancja kontaktronu określa graniczny prąd obciążenia powodujący samoczynne otwieranie się tego kontaktronu [6].

Przeprowadzone badania kontaktronów ZW-103 w środowisku ciekłego azotu wykazały, że R_k nowych kontaktronów jest o 10 do 60 procentów mniejsza od wartości pomierzonej w temperaturze 20°C. Największym zmniejszeniem wartości R_k charakteryzują się kontaktrony ze złożonym pokryciem styków. Wartość rezystancji kontaktronu nie zależy od czasu przepływu prądu (do wartości skutecznej 3 A), zależy natomiast (podobnie jak w temperaturze 20°C, [7, 8]) od wartości tego prądu, co dla nowego kontaktronu ilustruje rys. 1. Rezystancja $R'_{k(60)}$ (mierzona po czasie 60 s od chwili rozpoczęcia przepływu



Rys. 1. Zależność rezystancji $R'_{k(60)}$ nowego kontaktronu (styczki z FeNi29Co17) od skutecznej wartości stałego prądu obciążenia w środowisku ciekłego azotu

przepływ Θ cewki przekątnika równy 1,5-krotnemu przepływowi rozruchowemu Θ_r , $R'_{k(60)}$ — rezystancja nowego kontaktronu w 20°C, 1 — pokrycie Rh; 2 — pokrycie (W + Co) (Θ — przepływ sterujący)

prądu obciążenia) kontaktronów rodowanych rośnie ze wzrostem prądu obciążenia powyżej 1 A. Odwrotnie jest natomiast dla kontaktronów wolframowych (krzywa 2 rys. 1). Zmiany rezystancji $R'_{k(60)}$ powodowane są głównie nagrzewaniem się zestyku podczas przepływu prądu obciążenia. Średnia wartość temperatury ϑ_k w °K, cieplnego stanu ustalonego w zestyku kontaktronu ZW-103 podczas przepływu prądu obciążenia, określona zależnościami

$$\vartheta_k = \left(\frac{U_k^2}{9,6} \cdot 10^8 + \vartheta_o \right)^{1/2}, \quad (1)$$

gdzie:

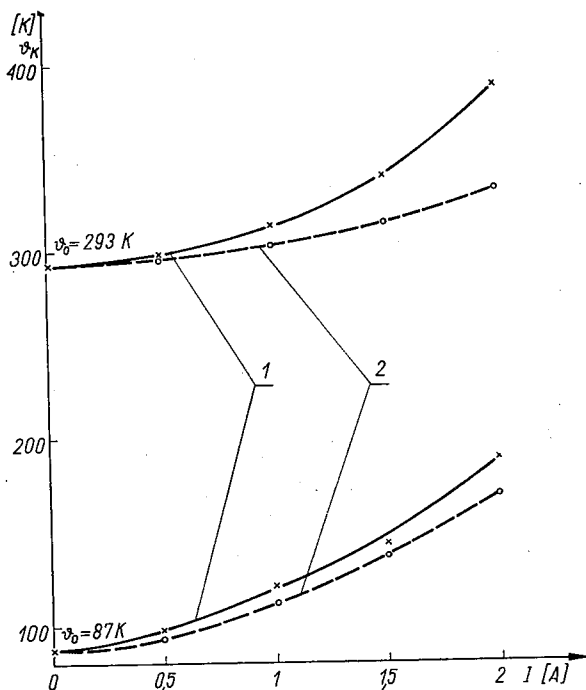
U_k — spadek napięcia na kontaktronie, w V,

ϑ_o — temperatura końców styków, w °K,

osiąga niewielkie wartości (rys. 2). Duże jednak gęstości prądów w rzeczywistych miejscach styku [6] mogą powodować lokalny wzrost temperatury do znacznych wartości.

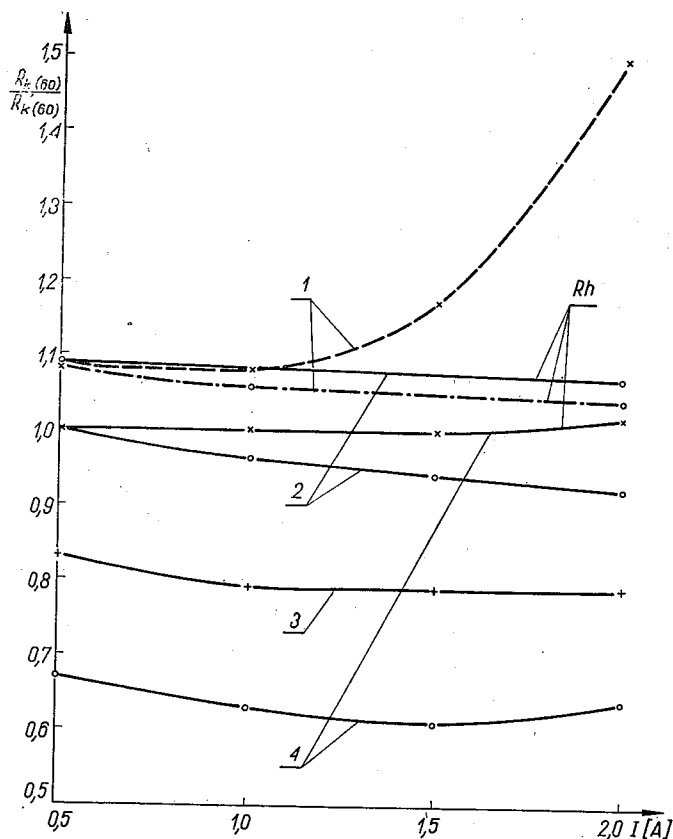
W trakcie badań stwierdzono, że w pewnych przypadkach rezystancja nowych kontaktronów gwałtownie wzrastała po umieszczeniu ich w ciekłym azocie — niezależnie od

rodzaju zastosowanego materiału stykowego. $R'_{k(60)}$ osiągała nawet 7-krotną wartość rezystancji $R''_{k(60)}$, pomierzonej w temperaturze 20°C. Wzrost rezystancji kontaktronu spowodowany jest wzrostem rezystancji przejścia wskutek prawdopodobnie zwiększonej fizycznej adsorpcji gazu w niskiej temperaturze oraz zmianą właściwości obecnych w zestyku zanieczyszczeń, głównie organicznych, naniesionych w procesie technologicznej obróbki kontaktronu lub znajdujących się w samym materiale stykowym. O obecności warstw adsorbowanych na powierzchni materiału stykowego takich kontaktronów świadczyć może, między innymi, szybkie ich uszkodzanie się podczas komutacji, o ile nie zasto-



Rys. 2. Zależność obliczeniowej średniej temperatury ϑ_k w zestyku kontaktronu, od wartości stałego prądu obciążenia w środowisku ciekłego azotu i temperaturze 293 K (Θ — przepływ sterujący)
 $\Theta = 1,5$ Θ_r ; 1 — pokrycie (W + Co); 2 — pokrycie Rh

sowano uprzednio kondycjonowania tych kontaktronów, na przykład komutowaniem prądu niewielkiej wartości [1]. Kontaktrony posiadające złoto jako materiał stykowy wykazywały dość stabilną wartość rezystancji R_k niezależnie od temperatury w jakiej wykonywano pomiary. Wartości natomiast R_k kontaktronów wolframowych oraz rodowych, komutowanych w środowisku ciekłego azotu, wzrastają po wyjęciu ich z tego środowiska. Rezystancja tych kontaktronów maleje jednak już po kilku cyklach łączeniowych ($I \leq 0,5$ A), wykonanych zarówno w temperaturze 20°C jak i w ciekłym azocie po ponownym umieszczeniu tam kontaktronu. Zmienność wartości rezystancji kontaktronu, łączącego w środowisku ciekłego azotu obwód indukcyjny, w funkcji liczby łączeń ilustruje rys. 3.



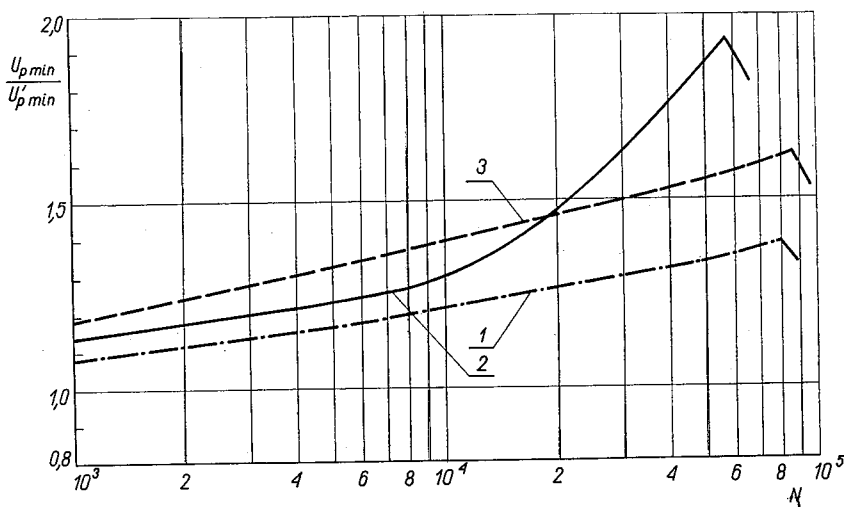
rys. 3. Zależność rezystancji $R_{k(60)}$ kontaktronu, w środowisku ciekłego azotu, od wartości stałego prądu obciążenia i liczby łączy (110 V —, 0,5 A, 6 ms) dla pokryw styków wykonanych z (W+Co) oraz Rh
 $R_{k(60)}$ — rezystancja kontaktronu nowego, $\Theta = 1,5\Theta_r$; 1 — kontaktron nowy, $\Theta = \Theta_r$; 2 — po 10 łączy, $\Theta = 1,5\Theta_r$; 3 — po 100 łączy, $\Theta_r = 1,5\Theta_r$; 4 — po 1000 łączy, $\Theta = 1,5\Theta_r$ (Θ_r — przepływ sterujący)

4. WYTRZYMAŁOŚĆ ELEKTRYCZNA

Wytrzymałość elektryczną kontaktronów określono dla napięcia przemiennego 50 Hz osuszając rezystor ograniczający połączony szeregowo z badanym kontaktronem. Średnie napięcie przeskoku U_{ps} (wartość skuteczna) szczeliny stykowej obliczano z 10 pomiarów serii. Określono też minimalną skuteczną wartość napięcia przeskoku U_{pmin} . Stwierdzono, że przeskok w badanych kontaktronach ZW-103 niezależnie od rodzaju zastosowanego materiału stykowego (W, Rh lub Au) występuje przy niższym napięciu niż w 20°C (wartość średnia jest niższa od 5 do 20 procent) [8]. Obniżenie się napięcia przeskoku szczeliny stykowej wynika ze zmniejszenia ciśnienia gazu ochronnego w bańce kontaktronu. Ciśnienie wodoru w kontaktronie, wynoszące średnio około 500 Tr w temperaturze $\vartheta = 293$ K, maleje w temperaturze $\vartheta = 87$ K do wartości p równej

$$p = p_0 \frac{\vartheta}{\vartheta_0} \cong 150 \text{ Tr.} \quad (2)$$

Dla średniej szczeliny δ_0 badanych kontaktronów, wynoszącej około 0,2 mm, iloczyn $p\delta_0$ zmaleje więc od 100 do 30 Trmm co dla wodoru na krzywej Paschena (określającej zależność napięcia przeskoku od iloczynu ciśnienia gazu i odległości między elektrodami) oznacza zmniejszenie wartości napięcia przeskoku. Napięcie przeskoku zależy również od liczby łączy wykonywanych przez kontaktron. W odróżnieniu jednak od pracy w temperaturze 20°C [8], napięcie to praktycznie cały czas wzrasta ze wzrostem liczby wykonywanych cykli łączeniowych; zaczyna maleć dopiero tuż przed ostatecznym uszkodzeniem się kontaktronu. Tę interesującą zależność dla U_{pmin} w funkcji liczby łączy ilustruje rys. 4 (U_{pmin} w badanych kontaktronach było niższe od wartości średniej U_{ps} o 2 do 20 procentów).



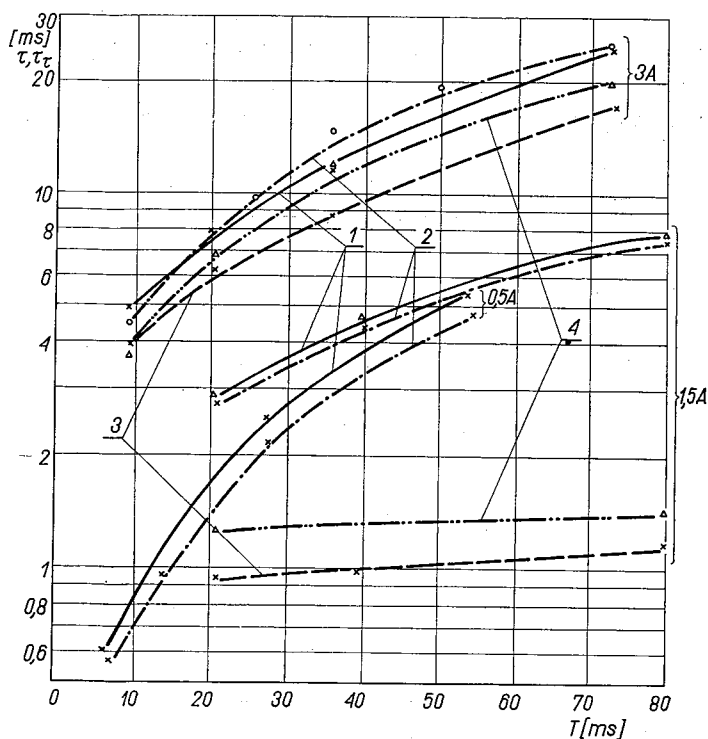
Rys. 4. Zależność minimalnego napięcia przeskoku U_{pmin} (wartość skuteczna) kontaktronu ZW-103 z pokryciem (W+Co) od liczby N łączy (1 A, 110 V, 30 ms) w środowisku ciekłego azotu
 U'_{pmin} — minimalne napięcie przeskoku nowego kontaktronu; 1 — $U'_{pmin} = 545$ V; 2 — $U'_{pmin} = 482$ V; 3 — $U'_{pmin} = 580$ V

Podczas badań stwierdzono, że w wielu przypadkach napięcie przeskoku kontaktronów gwałtownie wzrastało po ich umieszczeniu w środowisku ciekłego azotu, niezależnie o szczeliny δ_0 oraz powierzchni s_p pokrycia się stycek. Wzrost ten, dochodzący do 4-krotności wartości U_{ps} (w badanych kontaktronach nie przekraczał 2 kV wartości skutecznej) występował zawsze w kontaktronach ze stwierdzonymi mikropeknieniami szkła w złączu styeczka-szkło. Napięcie przeskoku w tego typu przypadkach malało z reguły do poziomu U_{ps} po kilku zainicjowanych przeskokach, w niektórych jednak kontaktronach trzeba było wykonywać cykle łączeniowe, aby efekt ten uzyskać. Znaczny wzrost wartości napięcia pierwszych przeskoków w kontaktronie, po jego umieszczeniu w środowisku ciekłego azotu, może być wywołany zanieczyszczeniami gazu ochronnego tego kontaktronu. Zanieczyszczenia bowiem takie jak N_2 , O_2 , SO_2 , CnH_{2n+2} , H_2O i inne wydzielane są, jak stwierdzono w [5], w dość znacznych ilościach ze stycek oraz szkła głównie podczas procesu hermetyzowania kontaktronu. Ponieważ adsorpcja, zwłaszcza fizyczna, wzrasta przy obniżeniu się temperatury, stąd przy dostarczonej dodatkowo energii podczas pomiaru napięcia przeskoku, wydzielane gazy mogą zostać zaadsorbowane w materiale

zestyku. Zaadsorbowana zaś warstwa gazu na powierzchni metalu zwiększa jego pracę wyjścia, powodując wzrost napięcia przeskoku kontaktronu [2]. Pewnym potwierdzeniem wzrostu napięcia przeskoku wskutek zwiększonej adsorpcji może być fakt powtarzalności wyników pomiarów tego napięcia w kontaktronach z pokryciami W, Rh lub Au przy każdorazowym wyjmowaniu ich i ponownym umieszczeniu w środowisku ciekłego azotu.

5. WŁAŚCIWOŚCI ŁĄCZENIOWE

Badania zdolności i trwałości łączeniowej kontaktronów w środowisku ciekłego azotu, pod kątem ich wykorzystania w obwodach silnoprądowych, wykonano analogicznie jak w [6, 8]. Stwierdzono, że w porównaniu z kontaktronami pracującymi w temperaturze 20°C , średni czas trwania τ wyładowania całkowitego, inicjowanego otwieraniem się styeczek, może być większy lub mniejszy. Zależy to od udziału czasu trwania wyładowania łukowego τ_r w całkowitym wyładowaniu [7]. Przy większych wartościach prądu komutowanego czas τ jest mniejszy z uwagi na znaczny udział czasu τ_r w wyładowaniach. Wyładowanie łukowe zaś w środowisku ciekłego azotu, z powodu głównie lepszych warunków chłodzenia kontaktronu, trwa znacznie krócej niż w temperaturze 20°C . Odwrotnie



rys. 5. Zależność średniego czasu τ trwania całkowitego wyładowania oraz wyładowania łukowego τ_r w kontaktronie (W+Co), od stałej czasowej T indukcyjnego obwodu napięcia stałego 24 V dla różnych skutecznych wartości prądu łączzonego

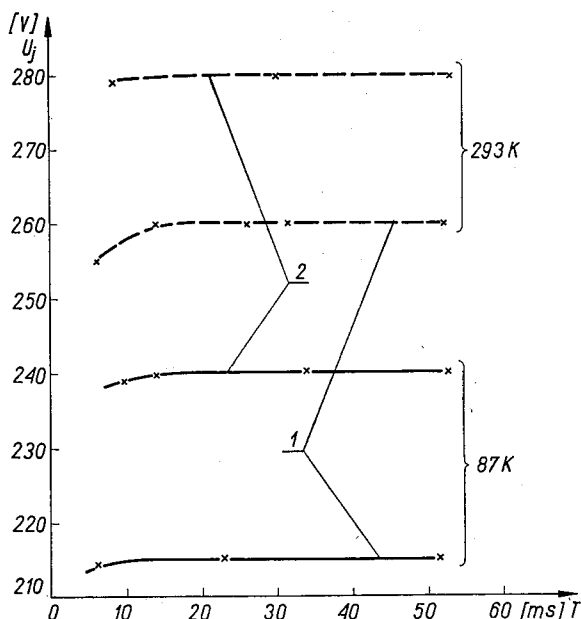
($\delta_0 = 0,2 \text{ mm}$, $S_p \approx 1,5 \text{ mm}^2$); 1) — τ w 87 K; 2) — τ w 293 K; 3) — τ_r w 87 K; 4) — τ_r w 293 K

natomiast sytuacja wygląda w przypadkach łączenia mniejszych prądów ($\leq 0,5$ A), gdzie dominującą rolę odgrywa wyładowanie jarzeniowe, bowiem czas trwania wyładowania jarzeniowego kontaktronu w środowisku ciekłego azotu wzrasta. Porównanie czasów wyładowania w temperaturze 20°C oraz temperaturze ciekłego azotu dla kontaktronu z pokryciem (W + Co) ilustruje rys. 5. Zwiększenie czasu trwania wyładowania jarzeniowego τ_j w kontaktronach umieszczonych w środowisku ciekłego azotu związane jest ze zmniejszeniem się wartości napięcia jarzenia U_j . Bowiem czas τ_j (w sekundach), na przykład wyładowania jarzeniowego normalnego, można wyrazić zależnością [9]

$$\tau_j = T \ln \frac{E}{U_j - E}, \quad (3)$$

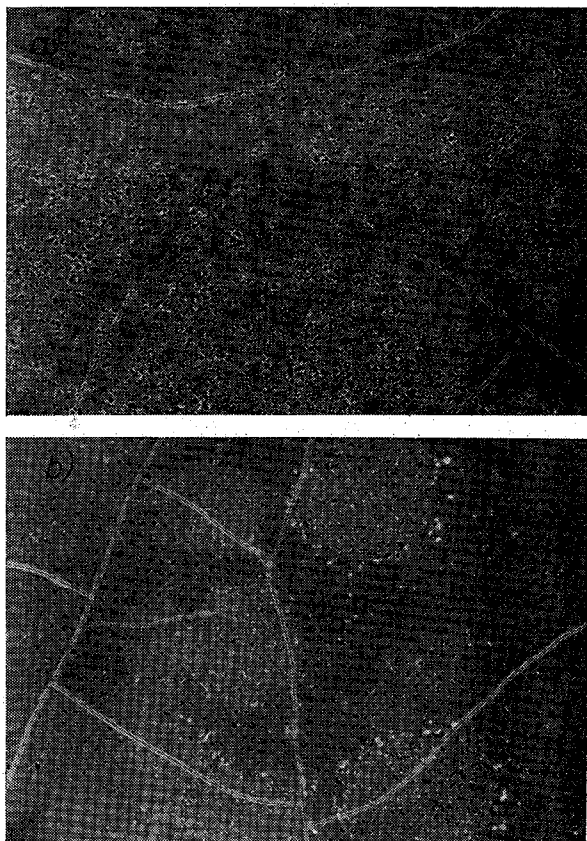
gdzie: R , T — odpowiednio, napięcie (w V) i stała czasowa (w s) indukcyjnego obwodu komutowanego.

Przeprowadzone badania wykazały, że wartości napięcia U_j (-197°C) są niższe o około 40 do 80 V od wartości pomierzonych w kontaktronach pracujących w temperaturze 20°C (rys. 6). Zależność U_j od wartości łączonego prądu świadczy o występowaniu wyładowań jarzeniowego nienormalnego, co wyjaśnia między innymi rozbieżności wyników pomiarów i wyników obliczeń czasu τ_j zgodnie z zależnością (3). Badania kontaktronów ZW-103 w środowisku ciekłego azotu, łączących w obwodach indukcyjnych prądy o wartościach skutecznych powyżej 0,5 A, wykazały, że najlepszym materiałem stykowym w tych warunkach jest wolfram. Potwierdziły to obserwacje mikroskopowe erozji stycek wykonane za pomocą elektronowego mikroskopu scanningowego typu JEOL JSM-35. W przypadkach kontaktronów rodowanych stwierdzono bowiem pęknięcia powierzchni materiału styko-



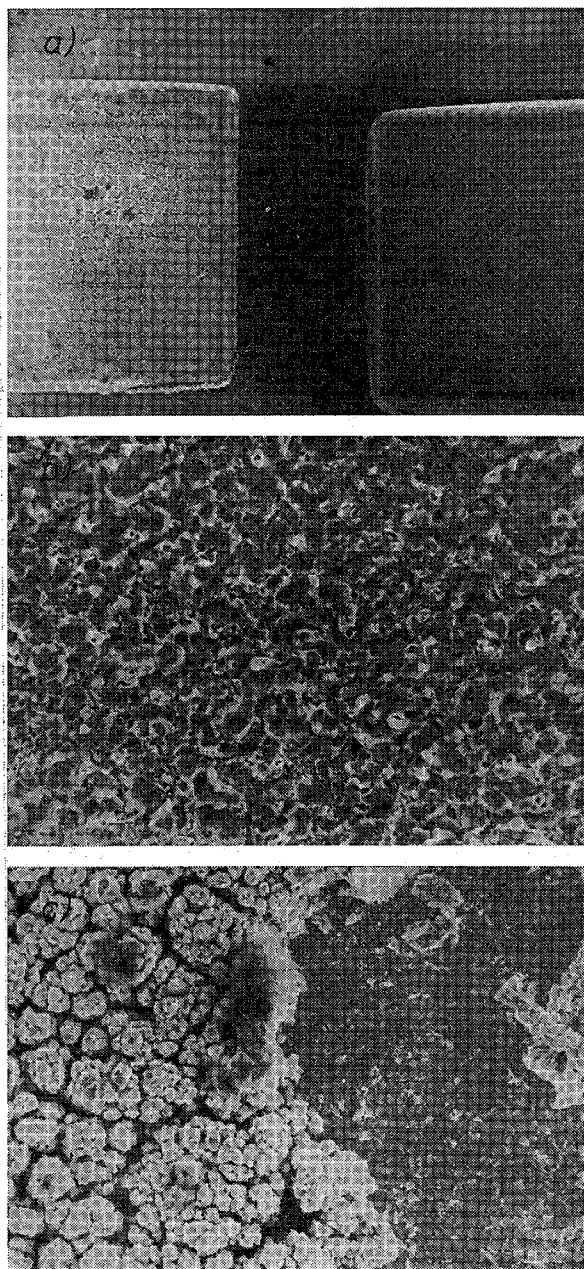
Rys. 6. Zależność wartości napięcia U_j w temperaturze 293 K oraz w ciekłym azocie, dla kontaktronów (W+Co), od stałej czasowej T indukcyjnego obwodu napięcia stałego 24 V ($\delta_0 \approx 0,23$ mm; $S_p \approx 3,9$ mm²), 1) — 0,5 A; 2) — 1 A

wego wskutek istnienia naprężeń wewnętrznych tego materiału. Pęknięcia te są przyczyną zwiększonej erozji stycek prowadzącej do szybkiego uszkodzania się kontaktronu (rys. 7). Na powierzchniach stycek rodowanych i złożonych, będących anodami, zaobserwowano ponadto miejscowe termiczne przegrzania objawiające się charakterystycznym tęczowym zabarwieniem metalu. Przegrzania te świadczą o istnieniu wysokich lokalnych temperatur w kontaktronie podczas komutacji i o śladowych zanieczyszczeniach tlenem. Widoczny



Rys. 7. Fragmenty powierzchni stykowej kontaktronu z pokryciem rodowym po wykonaniu 1000 łążeń (0,5 A, 110 V, 6 ms) w środowisku ciekłego azotu
a — katoda; b — anoda; powiększenie $\times 500$

był też niekiedy, na anodach kontaktronów z pokryciami wykonanymi z Rh lub Au, brązowy osad. Występowanie takiego osadu, i to zarówno w temperaturze 20°C jak i w temperaturze ciekłego azotu, związane jest prawdopodobnie z istnieniem organicznych zanieczyszczeń kontaktronu, stymulujących rozwój łuku elektrycznego podczas komutacji. Istotnie kontaktrony z brązowym osadem wykazywały znacznie większą erozję stycek. Największą erozję stycek, niezależnie od rodzaju zastosowanego materiału stykowego, wykazywały zawsze kontaktrony ze stwierdzonymi pęknięciami szkła. W takich przypadkach gwałtownemu zerodowaniu ulegała głównie cała powierzchnia zakładki katody, prowadząc do szybkiego uszkodzenia się kontaktronu, co dla kontaktronu z pokryciem rodowym ilustruje rys. 8.



Rys. 8. Fragmenty powierzchni styczek, pokrytych rodem, kontaktronu ZW-103 ze stwierdzonym mikro-
 pęknięciem w złączu styczka-szkło, po 5100 cyklach łączeniowych (1 A, 110 V, 30 ms) w środowisku ciekłego azotu
 a) styczki od wewnętrznej strony zakładki (katoda po lewej stronie) $\times 20$, b) fragment katody ($\times 500$); c) fragment anody ($\times 500$)

Obserwacje mikroskopowe charakteru i wielkości erozji styczek kontaktronów ZW-103 z pokryciami wykonanymi z W oraz Rh wykazały mniejsze ilości materiału stykowego przenoszonego podczas komutacji w temperaturze ciekłego azotu niż w temperaturze otoczenia 20°C. Zdolność i trwałość łączeniowa kontaktronów ZW-103 z wolframowym pokryciem styków również jest wyższa w środowisku ciekłego azotu niż w 20°C. Wyniki badań zdolności załączania i wyłączania tych kontaktronów przeprowadzonych, zgodnie z wymaganiami stawianymi w automatyce elektroenergetycznej, w obwodach napięcia stałego i przemiennego 220 V (analogicznie jak w [8]) przedstawiono w tablicach 1 i 2.

Trwałość łączeniowa kontaktronów ZW-103 wzrasta w środowisku ciekłego azotu minimum 2-krotnie w porównaniu z trwałością w temperaturze 20°C [8].

Tablica 1

Zdolność załączania kontaktronów wodorowych w środowisku ciekłego azotu

Typ kontaktronu	Warunki pomiarów	Największa skuteczna wartość prądu załączanego	Zdolność załączania kontaktronu (wartość skuteczna)
—	—	A	A
ZW-103/III	220 V $T = 6 \text{ ms}$	20,0	$\leq 5,0$
(W + Co)	220 V; 50 Hz $\cos \varphi = 0,7 \pm 0,1$	17,2	$\leq 4,3$

Tablica 2

Zdolność wyłączania kontaktronów wodorowych w środowisku ciekłego azotu

Typ kontaktronu	Warunki pomiarów	Graniczna zdolność wyłączania (wartość skuteczna)	Zdolność wyłączania kontaktronu (wartość skuteczna)
—	—	A	A
ZW-103/III	220 V $T = 40 \text{ ms} \pm 5 \text{ ms}$	3,0	$\leq 1,5$
(W + Co)	220 V; 50 Hz $\cos \varphi = 0,4 \pm 0,1$	7,0	$\leq 3,5$

6. WNIOSKI

— Kontaktrony wodorowe ZW-103/III mogą być z powodzeniem stosowane, nawet w obecnym ich wykonaniu, w środowisku ciekłego azotu jako elementy w silnoprądowych pomiarowych i sterujących obwodach, takich jak obwody automatyki elektroenergetycznej czy przemysłowej. Z uwagi jednak na występujące pęknięcia szkła (niewielkiej wprawdzie

liczby ok. 5% ogólnej populacji), zwiększenie niezawodności tych kontaktronów wymaga zwiększenia wytrzymałości termicznej głównie złącza styčka-szkło.

— Wartości przepływów zadziałania i odpadania kontaktronów (ze stycznymi wykonanymi z FeNi29Co17 oraz FeNi54) umieszczonych w środowisku ciekłego azotu nie różnią się więcej niż o ± 5 procentów od wartości pomierzonych w temperaturze 20°C.

— Rezystancja nowego kontaktronu w temperaturze ciekłego azotu jest mniejsza o 10 do 60 procentów od R_k w temperaturze 20°C. Wartość R_k maleje (najbardziej dla pokrycia wolframowego) w funkcji liczby łączy wykonanych przez kontaktron.

— Napięcie przeskoku (w porównaniu do 20°C) jest mniejsze o 5 do 20 procentów, z tym że pierwsze przeskoki mogą występować przy napięciu dochodzącym do 2 kV wartości skutecznej. Wartość napięcia przeskoku wzrasta (od 10 do 90 procentów w stosunku do napięcia nowego kontaktronu) w funkcji liczby łączy praktycznie aż do uszkodzenia się kontaktronu.

— Najlepszym materiałem stykowym kontaktronów ZW-103 umieszczonych w środowisku ciekłego azotu i łączących w obwodach indukcyjnych napięcia stałego i przemiennego prądu powyżej 0,5 A wartości skutecznej, jest wolfram.

— Zdolność załączania kontaktronów ZW-103 z wolframowym pokryciem styków nieznacznie wzrasta, natomiast zdolność wyłączania jest średnio większa o około 30 procentów w porównaniu do wartości w temperaturze 20°C. Trwałość łączeniowa tych kontaktronów w środowisku ciekłego azotu wzrasta minimum 2-krotnie.

LITERATURA

1. J. A. Augis, *Influence of adsorbed layers on the duration of short arc on make*, VII Congres International sur les Phenomenes de Contact Electrique, Paris 1974.
2. S. Dushman, *Scientific Found of Vacuum Technique*, New York 1962.
3. *Firmenblatt GOO2*, Vacuumschmelze CMBH Hanau.
4. R. V. Harrowell, *New superconducting switch field-circuit version*, Proc. IEE, 1972, nr 2.
5. M. Kubo, Y. Hayashi, *Analysis of released gases from reed blade and glass when sealed and their effect on contact resistance*, Electr. Cont. 20-th Annual Holm Sem. Illinois Inst. Technol, 1974, October, 29—31.
6. B. Miedziński, J. Pytel, *Właściwości komutacyjne kontaktronów próżniowych*, PAK, 1975, nr 10.
7. B. Miedziński, J. Rydosz, *Wpływ obwodu komutowanego na wyladowanie w kontaktronie gazowym*, PAK, 1976, nr 1.
8. J. Pytel, B. Miedziński, *Właściwości komutacyjne kontaktronów typu ZW-103/III i ZA-103/III*, Przegląd Elektrotechniczny, 1975, nr 8/9.
9. M. Takahashi, *Metal Transfer, Erosion and Welding of Reed Switch Contacts*, Elektrische Kontakte, VDE-Verlag, 1970.

B. MIEDZIŃSKI, H. KIELOCH, T. ŻOŁNIERCZYK

BEHAVIOUR OF GAS FILLED REEDS IN LIQUID NITROGEN ENVIRONMENT

Summary

The article presents the results of the study of the properties of hydrogen filled reeds in a liquid nitrogen environment. The influence of the liquid nitrogen environment on the value and nature of reed resistance

changes according to the value of the reed load current and to the time of its passage was studied. Passage values necessary for the engagement and disengagement of the reed were defined. The electric strength of the gap between contacts was measured according to the type of contact material and to the number of switching cycles performed by the reed. The specific nature was studied of the reed electric discharge, interrupting inductive current over 0.5 A in a liquid nitrogen environment, as well as the switching ability and durability of the reeds. An electron scanning microscope was employed to observe the erosion of reed contacts. The investigation was performed for reeds whose contact points were made of such ferronickel alloys as FeNi29Co17 and FeNi54. The contacts were made of tungsten, radium and gold. The reed bubbles were made of type SL 54.1 glass. The study was performed in order to determine the usefulness of the reeds, both as measuring and switching elements in low temperature power circuits.

B. MIEDZIŃSKI, H. KIELOCH, T. ŻOŁNIERCZYK

COMPOTEMENT DES ILS À GAZ DANS LE MILIEU D'AZOTE LIQUIDE

Résumé

L'article présente les résultats des essais des particulatirés des ILS (Interrupteur à lames scellés) à hydrogène, munis de contacts ouverts au repos et fonctionnant dans le milieu d'azote liquide. On a examiné l'influence du milieu en question sur les variations de la résistance de ILS en fonction de la valeur du courant de charge et de la durée de son écoulement. On a déterminé les valeurs des forces magnétomotrices de la fermeture et de l'ouverture des ILS. On a mesuré la rigidité diélectrique de la fente entre les contacts en fonction du matériau de contact et la quantité d'opérations effectuées par un ILS. On a examiné le caractère de la décharge électrique de ILS en interrompant, dans le milieu d'azote liquide, les courants inductifs dépassant 0,5 A et on a déterminé la capacité de coupure et la durée de service des ILS. De plus à l'aide d'un microscope scanning, on examine l'érosion des contacts de ILS. Cet examen a été effectué pour les ILS munis de contacts en alliage ferronickel tels que FeNi29Co17 et FeNi54.1. La présente étude, concernant le comportement des ILS à gaz dans le milieu d'azote liquide, a été effectuée afin de déterminer l'utilité des ILS comme éléments de mesure et de commande, fonctionnant aux basses températures.

B. MIEDZIŃSKI, H. KIELOCH, T. ŻOŁNIERCZYK

EIGENSCHAFTEN DER SCHUTZROHRKONTAKTE IM MEDIUM FLÜSSIGEN STICKSTOFFES

Zusammenfassung

Im vorliegenden Beitrag wurden die Prüfergebnisse bezüglich der Eigenschaften von Schutzrohrkontakten im Medium flüssigen Stickstoffes dargestellt. Es wurde der Einfluß des Stickstoffmediums auf den Wert und Charakter der Widerstandsveränderungen der Schutzrohrkontakte in Abhängigkeit vom Laststrom und der Durchflußzeit untersucht. Außerdem wurden Zahlenwerte der Amperewindungen bei Ein- und Ausschaltung und die Werte des kommutierten Stromes bestimmt. Es wurde auch die Überschlagnspannung des Zwischenspaltes in Abhängigkeit vom Kontaktenstoff und der Zyklusschaltzahl gemessen. Auch wurden Untersuchungen der Bogen-Entladungseigenschaften bei Stromunterbrechung von über 0,5 A im Induktionskreis durchgeführt. Beobachtungen der Kontaktflächen wurden mit Hilfe eines elektronischen Scanning-Mikroskops für Kontaktstücke aus Legierungen FeNi29Co17 und FeNi54 durchgeführt. Wolfram, Rhodium und Gold wurden als Kontaktwerkstoffe angewandt. Glasbirnen der Schutzrohrkontakte sind aus SL 54.1-Glass. Der Zweck der Untersuchungen beruhte auf der Bestimmung der Anwendbarkeit von Schutzrohrkontakten als Schalt- und Messungselemente in niedrigen Temperaturen.

Б. МЕДЗИНСКИ, Х. КЕЛЬОХ, Т. ЖОЛНЕРЧИК

ПОВЕДЕНИЕ ГАЗОВЫХ ГЕРКОНОВ В СРЕДЕ ЖИДКОГО АЗОТА

Резюме

В статье представлены результаты исследований свойств водородных нормальнооткрытых герконов в среде жидкого азота, проведенных с целью определения пригодности герконов для работы в качестве измерительных и исполнительных элементов в сильноточковых цепях работающих в низких температурах. Исследовано влияние среды жидкого азота на значение и характер изменений резистанса геркона в зависимости от величины тока нагрузки и времени его протекания. Определены значения ампервитков замыкания и размыкания контактронов. Измерена диэлектрическая прочность межконтактного зазора в зависимости от материала контактов и количества циклов работы геркона. Исследованы особенности электрического разряда герконов прерывающих индукционные токи порядка свыше 0,5 а в среде жидкого азота и определена жизнеспособность герконов в этих условиях. При использовании анализирующего электронного микроскопа исследовано разрушение поверхности контактов геркона. Исследованы контакты на пружинах из таких сплавов как FeNi29Co17 и FeNi54. Контактными материалами были вольфрам, радий и золото. Балоны герконов были изготовлены из стекла типа SL54.1.

Metody poprawiania stabilności algorytmów numerycznego rozwiązywania sztywnych układów równań różniczkowo-algebraicznych

PRZEMYSŁAW DYMARSKI (WARSZAWA)

Instytut Telekomunikacji Politechniki Warszawskiej

Otrzymano 6.12.1979

W pracy bada się warunki stabilności metody różnicowej Gearsa o zmiennym rzędzie i kroku całkowania, zastosowanej do rozwiązywania układów równań algebraiczno-różniczkowych, wynikających np. w analizie układów elektronicznych. Szczególnie uwzględnia się wpływ zastosowanego algorytmu wyboru kroku i rzędu na stabilność rozwiązania. Najpierw zbadano warunki stabilności metody różnicowej o stałym rzędzie i zmiennym kroku całkowania, a następnie metody różnicowej o zmiennym rzędzie i zmiennym kroku całkowania. Wyznaczono amplitudę i częstość rozwiązania pasożytniczego oraz określono dla jakich układów mogą wystąpić kłopoty ze stabilnością. Zwrócono uwagę na wpływ niestabilności na obniżenie efektywności metody różnicowej Gearsa. Podano pewne metody polepszania stabilności, które porównano z dotychczas stosowanymi.

1. WSTĘP

Efektywność programów analizy w dziedzinie czasu nieliniowych układów elektronicznych zależy w głównej mierze od zastosowanego algorytmu rozwiązywania równań różniczkowych opisujących te układy. Do najefektywniejszych algorytmów rozwiązywania równań różniczkowych zwyczajnych zalicza się obecnie metody różnicowe Gearsa o zmiennym kroku całkowania i zmiennym rzędzie [1], [8], [9], [10]. Potwierdzeniem tego jest ich zastosowanie w uniwersalnych programach analizy układów elektronicznych, takich jak ASTAP [12], ECAPII [14] i NAP2 [4].

Badania porównawcze [6] potwierdziły wysoką efektywność metod różnicowych Gearsa w zastosowaniu do rozwiązywania szerokiej klasy układów równań różniczkowych zwyczajnych z wyjątkiem tych układów, których wartości własne zbliżały się do osi urojonej. Przyczyną pogorszenia się właściwości algorytmu przy analizie tego rodzaju układów jest to, że metody różnicowe Gearsa rzędu większego od 2 nie są A -stabilne [7]. Najprostszym sposobem ominięcia tych trudności jest ograniczenie się do metod rzędu pierwszego i drugiego [12], [15]. Wiąże się to jednak z obniżeniem efektywności algorytmu (wzrost czasu obliczeń). Lepszym rozwiązaniem jest wykrywanie niestabilności metody różnicowej i obniżanie rzędu tylko wtedy, gdy jest to konieczne. W [2] opisano skuteczny algorytm detekcji niestabilności. Ma on jednak tę wadę, że nadaje się do współpracy

wyłącznie z mało efektywną wersją metody różnicowej Geara, wymagającą utrzymywania odcinkami stałej długości kroku całkowania i stałego rzędu.

W niniejszej pracy dokonano analizy zachowania się metody Geara o zmiennym kroku i rzędzie w warunkach niestabilności. Opracowano również kilka algorytmów detekcji niestabilności nadających się do współpracy z bardziej efektywną wersją metody Geara [1]. Opracowane algorytmy porównano pod względem skuteczności i efektywności.

2. STOSOWANE WARIANTY METOD RÓŻNICOWYCH GEARA DO ROZWIĄZYWANIA SZTYWNYCH UKŁADÓW RÓWNAŃ ALGEBRAICZNO-RÓŻNICZKOWYCH

2.1. Ogólna postać i właściwości analizowanych układów

Dowolny układ elektroniczny o elementach skupionych można opisać za pomocą układu $N+W$ równań różniczkowo-algebraicznych

$$g[\dot{x}(t), x(t), a(t), t] = 0, \quad (2.1)$$

gdzie:

$x(t) = [x_1(t), \dots, x_N(t)]^T$ — wektor zmiennych różniczkowych (np. napięcia na pojemnościach i prądy w indukcyjnościach),

$a(t) = [a_1(t), \dots, a_W(t)]^T$ — wektor zmiennych algebraicznych (inne napięcia i prądy).

Analiza układu elektronicznego w dziedzinie czasu polega na otrzymaniu za pomocą

metody numerycznej przybliżonego rozwiązania $\begin{bmatrix} x(t_n) \\ a(t_n) \end{bmatrix}$, $n = 0, 1, \dots$, układu (2.1),

w przedziale całkowania $t \in \langle t_0, t_{max} \rangle$, spełniającego warunek początkowy $x(t_0) = x_0 = [x_{01}, \dots, x_{0N}]^T$. Od otrzymanego rozwiązania wymaga się ponadto, aby nie odbiegało

ono zbyt od rozwiązania dokładnego $\begin{bmatrix} x^*(t) \\ a^*(t) \end{bmatrix}$. Ściślej, wymaga się aby $\hat{n} t_n \in \langle t_0, t_{max} \rangle$,

$|e_i(t_n)| = |x_i^*(t_n) - x_i(t_n)| < e_{max}$, $i = 1, \dots, N$, gdzie $e(t_n) = x^*(t_n) - x(t_n)$ jest błędem globalnym w chwili t_n , a e_{max} jest ograniczeniem narzuconym na błąd globalny. Dla uproszczenia zrezygnowano z analogicznych ograniczeń dla zmiennych algebraicznych zakładając, że spełnienie ww. nierówności dla zmiennych różniczkowych gwarantuje wystarczającą

dokładność całego rozwiązania $\begin{bmatrix} x(t_n) \\ a(t_n) \end{bmatrix}$.

Właściwości zagadnienia różniczkowego (2.1) zależą od rodzaju analizowanego układu. W większości przypadków zagadnienie to cechuje tzw. sztywność, polegająca na tym że rozwiązanie $x(t)$ posiada składowe wolno i szybko zmieniające się w czasie. Innym słowy, macierz współczynników przybliżenia liniowego zagadnienia (2.1) posiada wartości własne różniące się między sobą o wiele rzędów wielkości.

2.2. Metody różnicowe Geara

Metodę różnicową Geara w wersji umożliwiającej zmianę długości kroku całkowania na każdym etapie rozwiązania ([1], [4], [8]) można opisać następującym wzorem

$$\dot{x}(t_{n+1}) = (1/H_1) \sum_{i=0}^k \alpha_i x(t_{n+1-i}), \quad (2.2)$$

gdzie $H_1 = t_{n+1} - t_n$ jest krokiem całkowania w bieżącym, $n+1$ -szym etapie rozwiązania (tzn. przy przejściu z t_n do t_{n+1}). Podobnie $H_i = t_{n+2-i} - t_{n+1-i}$, $i = 2, 3, \dots$ są krokami całkowania w poprzednich etapach rozwiązania. Współczynniki zmiany kroku całkowania są określone następująco: $r_1 = H_1/H_2$, $r_2 = H_2/H_3$, $\dots$, $r_{k-1} = H_{k-1}/H_k$, $r_k = H_k/H_{k+1}$.

Współczynniki metody różnicowej rzędu k są funkcjami współczynników zmiany kroku całkowania ($\alpha_i = \alpha_i(r_1, \dots, r_{k-1})$, $i = 0, 1, \dots, k$) i są wyznaczone w ten sposób, aby wzór (2.2) był spełniony dla wielomianów stopnia $\leq k$.

Tablica 1

Współczynniki metod różnicowych i stałe błędu dla metod Geara o stałym kroku całkowania

k	1	2	3	4	5	6
a_0	1	$\frac{3}{2}$	$\frac{11}{6}$	$\frac{25}{12}$	$\frac{137}{60}$	$\frac{147}{60}$
a_1	-1	-2	-3	-4	-5	-6
a_2		$\frac{1}{2}$	$\frac{3}{2}$	3	5	$\frac{15}{2}$
a_3			$-\frac{1}{3}$	$-\frac{4}{3}$	$-\frac{10}{3}$	$-\frac{20}{3}$
a_4				$\frac{1}{4}$	$\frac{5}{4}$	$\frac{15}{4}$
a_5					$-\frac{1}{5}$	$-\frac{6}{5}$
a_6						$\frac{1}{6}$
C_k	$\frac{1}{2}$	$\frac{2}{9}$	$\frac{3}{22}$	$\frac{12}{152}$	$\frac{10}{137}$	$\frac{60}{1029}$

W tabl. 1 podano wartości współczynników metod różnicowych Geara o stałym kroku całkowania ($a_i = \alpha_i(1, \dots, 1)$).

W każdym etapie rozwiązania wektor $x(t)$ obarczony jest błędem lokalnym metody różnicowej, którego wielkość w $n+1$ -szym etapie można oszacować następującym wzorem

$$E(t_{n+1}) = -\gamma_k(r_1, \dots, r_{k-1})(H_1)^{k+1}x^{k+1}(\xi), \quad \xi \in \langle t_{n+1-k}, t_{n+1} \rangle. \quad (2.3)$$

Dla metod o stałym kroku całkowania współczynnik γ_k przechodzi w stałą błędu C_k ($C_k = \gamma_k(1, \dots, 1)$, patrz tablica 1).

2.3. Implementacja metod różnicowych Geara w programach analizy układów elektronicznych w dziedzinie czasu

Rozwiązanie zadania postawionego w punkcie 2.1 przebiega etapami, z których każdy polega na znalezieniu rozwiązania układu (2.1) w pewnej chwili t_n . Przeanalizujemy $n+1$ -szy etap rozwiązania: mając dane rozwiązanie przybliżone $\begin{bmatrix} x(t_i) \\ a(t_i) \end{bmatrix}$, $i = 0, 1, \dots, n$, poszukuje się rozwiązania w chwili t_{n+1} . Etap ten można podzielić na następujące pod-etapy.

1° *Znalezienie pierwszego przybliżenia rozwiązania.* Najczęściej jako pierwszego przybliżenia używa się wektora $\begin{bmatrix} \bar{x}^{(k+1)}(t_{n+1}) \\ \bar{a}^{(k+1)}(t_{n+1}) \end{bmatrix}$ otrzymanego w wyniku predykcji wielomianem k -tego stopnia przechodzącym przez punkty

$$\begin{bmatrix} x(t_n) \\ a(t_n) \end{bmatrix}, \begin{bmatrix} x(t_{n-1}) \\ a(t_{n-1}) \end{bmatrix}, \dots, \begin{bmatrix} x(t_{n-k}) \\ a(t_{n-k}) \end{bmatrix}.$$

2° *Rozwiązanie układu równań algebraicznych.* Przez podstawienie (2.2) do (2.1) otrzymuje się następujący układ równań algebraicznych

$$g[(\alpha_0/H_1)x(t_{n+1}) + \bar{S}_k, x(t_{n+1}), a(t_{n+1}), t_{n+1}] = 0, \quad (2.4)$$

gdzie

$$\bar{S}_k = (1/H_1) \sum_{i=1}^k \alpha_i x(t_{n+1-i}).$$

Układ równań (2.4) rozwiązuje się ze względu na $\begin{bmatrix} x(t_{n+1}) \\ a(t_{n+1}) \end{bmatrix}$ metodą Newtona-Raphsona

$$\begin{aligned} [g_x(\alpha_0/H_1) + g_x, g_a] \begin{bmatrix} x^{j+1}(t_{n+1}) - x^j(t_{n+1}) \\ a^{j+1}(t_{n+1}) - a^j(t_{n+1}) \end{bmatrix} = \\ = -g[(\alpha_0/H_1)x^j(t_{n+1}) + \bar{S}_k, x^j(t_{n+1}), a^j(t_{n+1}), t_{n+1}]. \end{aligned} \quad (2.5)$$

Macierz $J = [(\alpha_0/H_1)g_x + g_x, g_a]$ jest jacobianem układu (2.4) wyliczonym w punkcie $\begin{bmatrix} x^j(t_{n+1}) \\ a^j(t_{n+1}) \end{bmatrix}$.

Iterację rozpoczyna się z punktu

$$\begin{bmatrix} x^0(t_{n+1}) \\ a^0(t_{n+1}) \end{bmatrix} = \begin{bmatrix} \bar{x}^{(k+1)}(t_{n+1}) \\ \bar{a}^{(k+1)}(t_{n+1}) \end{bmatrix}.$$

Iteracje (2.5) wykonuje się aż do uzyskania wystarczającej zbieżności:

$$\begin{aligned} |x_i^{j+1}(t_{n+1}) - x_i^j(t_{n+1})| &< \varepsilon, \quad i = 1, \dots, N, \\ |a_i^{j+1}(t_{n+1}) - a_i^j(t_{n+1})| &< \varepsilon, \quad i = 1, \dots, W. \end{aligned}$$

Współczynnik ε wybiera się tak, aby błąd rozwiązania układu (2.4) ze względu na $x(t_{n+1})$ nie przewyższył błędu lokalnego metody różnicowej w bieżącym etapie rozwiąza-

nia. W praktyce wystarczy, aby ε był 10 razy mniejszy od dopuszczalnego błędu lokalnego w bieżącym etapie rozwiązania.

Dzięki stosowaniu przyrostowych modeli elementów nieliniowych i wykorzystaniu (2.2) do przekształcenia równań elementów reaktancyjnych, macierz J można otrzymać wprost z analizowanej sieci [11], [12]. Najbardziej czasochłonne jest rozwiązanie równania liniowego (2.5), mimo stosowania dla większych układów techniki macierzy rzadkich. Dlatego też stosuje się modyfikację metody Newtona-Raphsona, polegającą na wyznaczeniu jacobianu i dokonaniu jego dekompozycji na macierz trójkątną górną U i dolną L tylko raz w każdym etapie rozwiązania, w punkcie $\begin{bmatrix} x^0(t_{n+1}) \\ a^0(t_{n+1}) \end{bmatrix}$. Dopiero gdy nie można uzyskać szybkiej zbieżności, przechodzi się do klasycznej metody Newtona-Raphsona, wymagającej obliczenia jacobianu w każdej iteracji. W przypadku układów quasi-liniowych można, stosując odcinkami stały krok H , obliczać J raz na kilka etapów rozwiązania [10]. W przypadku układów silnie nieliniowych przydatność takiego algorytmu jest niewielka. Dla większości układów konieczne jest obliczenie jacobianu przynajmniej raz w każdym etapie rozwiązania.

3° *Kontrola błędu.* Błąd lokalny metody różnicowej, powstały w bieżącym etapie rozwiązania, szacuje się ze wzoru (2.3) zastępując $x^{(k+1)}(\xi)$ odpowiednim ilorazem różnicowym, lub dokonując porównania wektorów $\bar{x}^{(k+1)}(t_{n+1})$ z $x(t_{n+1})$ [1], [8]. Następnie sprawdza się, czy składowe błędu lokalnego $E(t_{n+1}) = [E_1(t_{n+1}), \dots, E_N(t_{n+1})]^T$ nie przekraczają dopuszczalnej wartości błędu lokalnego w $n+1$ -szym etapie rozwiązania $E_{\max}(t_{n+1})$. Gdy błąd jest zbyt duży, otrzymane rozwiązanie ignoruje się i ponawia się obliczenia ze zredukowaną długością kroku całkowania H_1 .

W praktyce stosuje się następujące rodzaje kontroli błędu:

a) kontrola błędu względnego [15]

$$\frac{|E_i(t_{n+1})|}{|x_i(t_{n+1})|} \leq E_{\max}(t_{n+1}), \quad (2.6)$$

b) kontrola błędu bezwzględnego [1], [8]

$$|E_i(t_{n+1})| \leq E_{\max}(t_{n+1}), \quad (2.7)$$

c) mieszana kontrola błędu [4], [10]

$$\frac{|E_i(t_{n+1})|}{\max_{j=1, \dots, n+1} |x_j(t_j)|} \leq E_{\max}(t_{n+1}). \quad (2.8)$$

Z trzech ww. wariantów najbardziej racjonalny wydaje się drugi, gdyż stwarza pewną możliwość kontroli błędu globalnego $e(t_{n+1}) = x^*(t_{n+1}) - x(t_{n+1})$. Jeżeli metoda różnicowa zachowuje się stabilnie, tzn. składowe $E_i(t_{n+1})$, $i = 1, \dots, N$, błędu wprowadzonego w $n+1$ -szym etapie nie są wzmacniane w następnych etapach, to przyjmując $E_{\max}(t_{n+1}) = \frac{e_{\max}}{t_{\max} - t_0} H_1$ mamy gwarancję, że składowe błędu globalnego $e(t)$ nie przekroczą wartości $e_{\max}$ dla $t \in \langle t_0, t_{\max} \rangle$.

Z tej przyczyny z dalszej części pracy będzie brany pod uwagę taki właśnie rodzaj kontroli błędu (jest to tzw. kontrola błędu lokalnego przypadającego na jednostkę czasu).

4° Wybór optymalnego rzędu k' i optymalnej długości kroku całkowania H'_1 dla następnego etapu rozwiązania. W obliczeniach wykorzystuje się zależność (2.3) zakładając, że $|x^{(k+1)}(\xi)|$ nie zmienia się znacznie w $n+1$ -szym i $n+2$ -gim etapie rozwiązania. Dla danego k wylicza się $H_1^{(k)}$ takie, aby otrzymać błąd lokalny równy

$$\delta E_{\max}(t_{n+2}) = \gamma_k \left[\frac{H_1^{(k)}}{H_1}, r_1, \dots, r_{k-2} \right] [H_1^{(k)}]^{k+1} |x_1^{(k+1)}(\xi)|. \quad (2.9)$$

Współczynnik bezpieczeństwa δ przyjmuje się rzędu 0,5+0,9, aby zmniejszyć prawdopodobieństwo naruszenia warunku (2.7) w $n+2$ -gim etapie rozwiązania. Obliczenia (2.9) wykonuje się dla tej składowej i wektora $x(t)$, która zachowywała się najbardziej czynnie w sprawdzaniu nierówności (2.7) (maksimum $E_i(t_{n+1})$). Za $E_{\max}(t_{n+2})$ we wzorze (2.9) podstawia się wartość $\frac{e_{\max}}{t_{\max}-t_0} H_1^{(k)}$.

Jeżeli warunek (2.7) nie był spełniony i otrzymane rozwiązanie zostało zignorowane, to długość kroku dla $n+1$ -go etapu rozwiązania wylicza się ze wzoru

$$\delta E_{\max}(t_{n+1}) = \gamma_k \left[\frac{H_1^{(k)}}{H_2}, r_2, \dots, r_{k-1} \right] [H_1^{(k)}]^{k+1} |x_1^{(k+1)}(\xi)|, \quad (2.10)$$

gdzie

$$E_{\max}(t_{n+1}) = \frac{e_{\max}}{t_{\max}-t_0} H_1^{(k)}.$$

Równania (2.9) i (2.10) można rozwiązywać metodami iteracyjnymi, jak proponuje się w [1]. Można też, zakładając, że krok nie zmienia się znacznie z etapu na etap, podstawiać $\gamma_k = C_k$; wówczas wyliczenie $H_1^{(k)}$ jest trywialne [8].

Obliczenia przeprowadza się dla kilku wartości k (najczęściej dla $k := k$, $k := k-1$ i $k := k+1$, gdzie k stojące po prawej stronie jest rzędem metody różnicowej w $n+1$ -szym etapie rozwiązania) i wybiera się takie k' , dla którego uzyskuje się największą wartość $H_1^{(k')}$. Tym kończy się $n+1$ -szy etap obliczeń. Następny, $n+2$ -gi etap wykonuje się z rzędem k' i krokiem $H'_1 = H_1^{(k')}$.

Wybierając rząd zapewniający użycie maksymalnej długości kroku całkowania, uzyskuje się maksymalną efektywność metody Geara o zmiennym kroku i zmiennym rzędzie. Można wyjaśnić to w następujący sposób. Ilość operacji arytmetycznych, a co za tym idzie czas obliczeń, zależy głównie od ilości obliczeń jakobianu i wartości funkcji g . Ponieważ przyjęliśmy, że jakobian musi być obliczony i zdekomponowany na macierze L i U co najmniej raz w każdym etapie rozwiązania, a liczba iteracji (2.5) nie zależy silnie od długości kroku całkowania, można przyjąć za miarę efektywności algorytmu liczbę etapów (tzw. liczbę kroków) w przedziale $t \in \langle t_0, t_{\max} \rangle$.

Efektywność algorytmu zależy od szybkości jego adaptacji do rozwiązywanego zagadnienia, tzn. od możliwości szybkich zmian rzędu i kroku całkowania. W dotychczasowych implementacjach metod różnicowych Geara można wyróżnić następujące warianty wyboru kroku całkowania i rzędu.

1° Warianty wyboru kroku całkowania

a) Stały krok $H_1 = H_2 = \dots = H$. W zastosowaniu do rozwiązywania układów sztywnych, wariant ten jest używany tylko wtedy, gdy nie zależy nam na dokładnym odwzoro-

waniu składowych szybko zmieniających się w czasie, a tylko składowych wolnozmiennych.

b) *Możliwość zmiany długości kroku raz na $k+1$ etapów rozwiązania.* Wariant ten stosuje się w tzw. metodzie Nordsiecka-Geara [9], [10], gdzie utrzymywanie odcinkami stałego kroku całkowania i rzędu jest koniecznością.

c) *Możliwość zmian długości kroku w każdym etapie rozwiązania* [1], [4], [8]. Wariant c) jest bardziej efektywny niż b), gdyż zapewnia szybszą adaptację długości kroku do zmian rozwiązania. Z tej przyczyny metoda Nordsiecka-Geara jest rzadko używana w programach analizy układów elektronicznych (np. w programach ASTAP i NAP-2 nie korzysta się z ww. metody, tylko z pewnych modyfikacji algorytmu przedstawionego w [8]).

2° Warianty wyboru rzędu

a) *Stały rząd $k = k_{FIX} = \text{const.}$, k_{FIX} przyjmuje wartości od 1 do 6 (wyższych rzędów nie używa się ze względu na stabilność).* Wariant ten jest rzadko używany we współczesnych programach analizy układów elektronicznych. Efektywność programu zależy od wybranej wartości k_{FIX} . W procesie rozwiązywania układów liniowych efektywność rośnie wraz ze wzrostem wartości rzędu. Tendencja ta występuje tym silniej, im ostrzejsze są wymagania dotyczące dokładności rozwiązania (tzn. im mniejsze e_{max}) [16]. Podobna tendencja występuje też dla dużej klasy układów nieliniowych, chociaż przewaga metod o wyższych rzędach może nie być tak duża jak dla układów liniowych [4], [8].

b) *Rząd zmienny w przedziale od 1 do k_{max} , $k_{max} = 2, 3, 4, 5$ lub 6 — z możliwością zmiany raz na $k+1$ etapów rozwiązania* [8]. Wariant ten zapewnia większą efektywność algorytmu niż użycie stałego rzędu. Gdy w układzie występuje stan ustalony ($x(t) \simeq \text{const.}$), metody różnicowe niskich rzędów zapewniają użycie dłuższego kroku całkowania, niż metody wyższych rzędów. W [4] rozwiązywano równanie oscylatora Van der Pola posługując się metodą różnicową stałego rzędu ($k_{FIX} = 6$), a następnie metodą o zmiennym rzędzie k w przedziale od 1 do k_{max} , $k_{max} = 6$ — uzyskując efektywność większą o 20%. Dla większości analizowanych układów im większe k_{max} , tym większa efektywność algorytmu potwierdzają to częściowo eksperymenty z rozwiązywaniem równań wtórника emitowanego, opisane w [8]).

c) *Rząd zmienny w przedziale od 1 do k_{max} , $k_{max} = 2, 3, 4, 5$ lub 6 z możliwością zmiany w każdym etapie rozwiązania* [1], [4]. Ten wariant zapewnia najszybszą adaptację kroku rzędu metody numerycznej do zmian rozwiązania — jest więc najbardziej efektywny. Potwierdzają to badania porównawcze dotyczące rozwiązywania równań liniowych [1] uzyskano o 20% większą efektywność w porównaniu z wariantem b) i nieliniowych [4] rozwiązując równania oscylatora Van der Pola uzyskano dwa razy większą efektywność w porównaniu z wariantem b). Ponadto przeprowadzono badania efektywności algorytmu funkcji k_{max} [4]. Uzyskane wyniki potwierdzają celowość stosowania jak największej wartości k_{max} , nawet dla układów nieliniowych, zwłaszcza gdy wymagana jest duża dokładność obliczeń.

Wersja algorytmu Geara opublikowana w [1] umożliwia zmianę rzędu o 1 w każdym etapie rozwiązania, natomiast wersja opublikowana w [4] umożliwia zmianę w pełnym zakresie od 1 do k_{max} . Badania porównawcze obu wersji algorytmu wykazały znikomą przewagę tej drugiej (zmiany rzędu o więcej niż 1 zdarzają się bardzo rzadko).

W podsumowaniu można stwierdzić, że nieomal wszystkie przeprowadzone dotychczas badania (pewnym wyjątkiem jest [13], gdzie zakwestionowano opłacalność stosowania rzędu wyższego niż 2) wskazują, że metoda różnicowa Geara jest tym efektywniejsza, im większa jest jej zdolność do adaptowania się do zagadnienia. Najbardziej efektywna jest implementacja dopuszczająca użycie zmiennego rzędu k w przedziale od 1 do 6 i zmiennego kroku całkowania z możliwością zmiany zarówno rzędu jak i kroku w każdym etapie rozwiązywania. Taka też implementacja będzie dalej brana pod uwagę.

3. STABILNOŚĆ METOD RÓŻNICOWYCH GEARA

3.1. Metody badania stabilności

W procesie rozwiązywania zagadnienia (2.1) ważne jest, aby błąd lokalny wprowadzony w danym etapie obliczeń nie powiększał się w następnych etapach. Takie zachowanie się metody różnicowej nazywa się stabilnością bezwzględną [9].

Definicja. Metoda różnicowa jest bezwzględnie stabilna dla danego zagadnienia (2.1) i przy danej sekwencji kroków całkowania (tzn. przy danych $t_1, t_2, \dots \in \langle t_0, t_{\max} \rangle$) gdy zaburzenie $\Delta x(t_{n+1})$ wprowadzone w $n+1$ -szym etapie rozwiązania nie powiększa się w następnych etapach rozwiązania.

Użycie metody różnicowej nie posiadającej tej własności prowadzi najczęściej do odejsścia rozwiązania przybliżonego od rozwiązania dokładnego zagadnienia (2.1). W tych warunkach nie jest też możliwa kontrola błędu globalnego (mimo spełnienia warunków (2.6)–(2.8) nie można nic powiedzieć o zachowaniu się błędu globalnego $e(t_n) = x^*(t_n) - x(t_n)$). Zdefiniowana powyżej stabilność bezwzględna zależy nie tylko od samej metody różnicowej, ale również od rozwiązywanego zagadnienia. Dla pewnej klasy zagadnień (problemy niestateczne¹⁾) nie można wymagać od metody różnicowej stabilności bezwzględnej. Dla tego typu problemów wymaga się, aby błąd globalny spowodowany zaburzeniem w n -tym etapie rozwiązania nie rósł szybciej niż rozwiązanie dokładne (tzw. stabilność względna).

Stabilności bezwzględnej należałoby więc wymagać od metody różnicowej zastosowanej do rozwiązania każdego zagadnienia statecznego. W praktyce sprawdzenie, czy dana metoda różnicowa posiada właściwość stabilności bezwzględnej dla każdego zagadnienia statecznego, napotyka na trudności. Wobec tego właściwość tę sprawdza się dla konkretnego równania różniczkowego, tzw. równania testującego [7]

$$\dot{x} = \lambda x, x(t_0) = x_0, \operatorname{Re} \lambda < 0, x \text{ — skalar.} \quad (3.1)$$

Taka metoda badania stabilności jest powszechnie stosowana. Udowodniwszy, że dana metoda różnicowa jest bezwzględnie stabilna w zastosowaniu do rozwiązania zagadnienia (3.1) nie mamy gwarancji, że jest ona stabilna dla dowolnego statecznego zagadnienia opisanego wzorem (2.1). Badanie stabilności opisaną wyżej metodą ma jednak sens z następujących względów:

¹⁾ Problem jest niestateczny, gdy rozwiązanie dokładne $x^*(t)$ i rozwiązanie zaburzone w chwili t_n $x'(t_i) = x^*(t_i) + \Delta x(t_i)$ różnią się w pewnej chwili $t_n > t_i$ o więcej niż $\Delta x(t_i)$, tzn. $|x'(t_n) - x^*(t_n)| > |\Delta x(t_i)|$.

- a) jeżeli dana metoda różnicowa jest niestabilna w zastosowaniu do rozwiązania zagadnienia (3.1), to nie będzie się ona również nadawała do rozwiązywania szerokiej klasy zagadnień nieliniowych;
- b) jeżeli metoda różnicowa jest stabilna w zastosowaniu do zagadnienia (3.1), to jest również stabilna w procesie rozwiązywania zagadnień liniowych o stałych współczynnikach

$$\dot{x} = Ax + B, \quad x(t_0) = x_0, \quad (3.2)$$

gdzie wartości własne macierzy A : $\lambda_1, \dots, \lambda_N$ mają ujemną część rzeczywistą;

- c) jeżeli metoda różnicowa jest stabilna dla (3.1), to można przypuszczać, że będzie lokalnie stabilna dla tych zagadnień (2.1), których przybliżenie liniowe jest stateczne i zmienia się niewiele na przestrzeni kilku kolejnych etapów rozwiązania [11].

Zbadajmy stabilność metody różnicowej Geara w zastosowaniu do rozwiązania zagadnienia testującego (3.1). Rozwiązanie przybliżone $x(t_n)$ równania (3.1) musi spełniać następującą równość ((2.2), (3.1))

$$\frac{1}{H_1} \sum_{i=0}^k \alpha_i x(t_{n+1-i}) = \lambda x(t_{n+1}), \quad (3.3)$$

stąd

$$x(t_{n+1}) = \frac{1}{H_1 \lambda - \alpha_0} \sum_{i=1}^k \alpha_i x(t_{n+1-i}).$$

Wprowadźmy zaburzenie $\Delta x(t_{l+1})$ w $l+1$ -szym etapie rozwiązania (może to być np. błąd lokalny $E(t_{l+1})$). Rezultatem będzie rozwiązanie zaburzone $x'(t_n)$. Różnica tych dwóch rozwiązań $e(t_n)$, czyli błąd globalny spowodowany zaburzeniem w $l+1$ -szym etapie, będzie spełniać następujące równanie błędu globalnego

$$e(t_{n+1}) = \begin{cases} 0, & n = 0, \dots, l-1, \\ \Delta x(t_{l+1}), & n = l, \\ \frac{1}{H_1 \lambda - \alpha_0} \sum_{i=1}^k \alpha_i e(t_{n+1-i}), & n > l. \end{cases} \quad (3.4)$$

Należy tu przypomnieć, że $\alpha_i = \alpha_i(r_1, \dots, r_{k-1})$, a współczynniki $H_1, k, r_1, \dots, r_{k-1}$ zmieniają się z etapu na etap. Równanie błędu globalnego (3.4) jest więc liniowym równaniem różnicowym o zmiennych współczynnikach, o rozwiązaniu którego w ogólnym przypadku nie można nic powiedzieć.

3.2. Stabilność metod różnicowych Geara o stałym kroku całkowania i stałym rzędzie

Nieomal wszystkie publikacje dotyczące zagadnienia stabilności metod różnicowych odnoszą się do stałego kroku całkowania [7], [9]. W tym wypadku równanie błędu glo-

balnego (3.4) jest równaniem różnicowym liniowym o stałych współczynnikach, którego rozwiązanie można zapisać w postaci

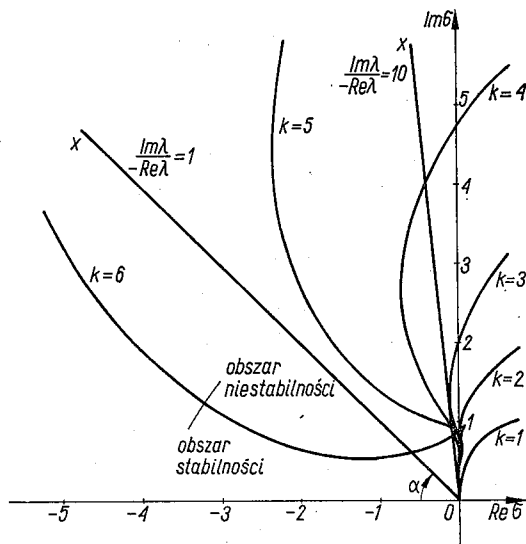
$$e(t_n) = c'_1(z_1)^n + \dots + c'_k(z_k)^n, \quad (3.5)$$

gdzie $z_1, \dots, z_k$ są pierwiastkami (zakładamy, że jednokrotnymi) równania

$$\sigma z^k - \sum_{i=0}^k a_i z^{k-i} = 0. \quad (3.6)$$

Wyrażenie stojące po lewej stronie jest wielomianem charakterystycznym metody różnicowej, $a_i = \alpha_i(1, \dots, 1)$, $i = 0, \dots, k$, są współczynnikami metody różnicowej, a $\sigma = H_1 \lambda$. Aby metoda była bezwzględnie stabilna, wszystkie pierwiastki muszą leżeć w kole jednostkowym, a pierwiastki leżące na okręgu o promieniu jednostkowym muszą być jednokrotne.

Na rys. 1 przedstawiono obszar stabilności bezwzględnej metod Geara, tzn. zbiór wartości $\sigma = H_1 \lambda$, dla których pierwiastki (3.6) leżą w kole jednostkowym.



Rys. 1. Obszar stabilności metody różnicowej Geara k -tego rzędu (wg [4])

Aby zapewnić stabilność dla każdego λ i H_1 , należałoby żądać, aby obszar stabilności wypełniał całą lewą półpłaszczyznę (tzw. stabilność typu A). Ten warunek spełnia tylko metoda Geara rzędu 1 i 2, zgodnie z twierdzeniem Dahlquista, że nie ma A-stabilnych metod różnicowych rzędu większego niż 2 [7].

Gear [9] argumentuje, że stabilność typu A jest zbyt ostrym wymaganiem, gdyż przy użyciu metody różnicowej o stałym kroku nie pracuje się przy σ znajdującym się w obszarze położonym blisko osi urojonej i oddalonym od osi $Re \sigma$ o więcej niż $0,5 \div 1$ ze względu na małą dokładność otrzymanych w ten sposób rozwiązań (dla λ leżącego blisko osi urojonej rozwiązanie jest oscylacyjne — aby uzyskać co najmniej kilka punktów na okresie drgań, trzeba wziąć dostatecznie małe H). Metody różnicowe posiadające własność sta

stabilności bezwzględnej dla σ leżącego w lewej półpłaszczyźnie z wyjątkiem ww. obszarów nazywają się metodami „sztywno-stabilnymi” (*stiffly-stable*). Tę własność posiadają metody Geara rzędu 3 ÷ 6. Metody sztywno-stabilne pozwalają na uzyskanie stabilnych rozwiązań układów sztywnych zarówno przy użyciu małego kroku całkowania (σ w otoczeniu początku układu współrzędnych) jak i bardzo dużego kroku całkowania ($\sigma \rightarrow \infty$, $\text{Re } \sigma < 0$).

3.3. Stabilność metod różnicowych Geara o stałym rzędzie i zmiennym kroku całkowania

Gdy krok całkowania nie jest stały, badanie rozwiązania równania różnicowego (3.4) napotyka na duże trudności. Jedyne dla $k = 1$, gdzie $\alpha_i = \text{const}$, równanie (3.4) daje się sprowadzić do postaci

$$e(t_{n+1}) = \frac{1}{1 - H_1 \lambda} e(t_n). \quad (3.7)$$

Ponieważ dla każdego $\sigma = H_1 \lambda$, $\text{Re } \sigma < 0$, spełniona jest zależność $|e(t_{n+1})| < |e(t_n)|$, można stwierdzić, że metoda pierwszego rzędu o zmiennym kroku całkowania jest bezwzględnie stabilna.

Stabilność metody różnicowej rzędu 2 i 3 bada się w [5]. Uzyskane wyniki są dość pesymistyczne. Stwierdzono, że metoda 2 rzędu nie jest w ogólnym przypadku A -stabilna. Znalezione pewne ograniczenia na parametr σ i współczynnik zmiany kroku całkowania r_1 , przy spełnieniu których metoda jest stabilna. Wyniki te mają jednak małą przydatność praktyczną. Wiąże się to z tym, że w [5] poczyniono szereg założeń upraszczających, a mianowicie:

1) Założono, że w procesie rozwiązywania (3.1) krok całkowania zmienia się dowolnie z etapu na etap (nakłada się ograniczenie wyłącznie na współczynniki zmiany kroku całkowania r_i). W praktyce krok jest wybierany wg pewnego algorytmu, np. opisanego wzorem (2.9). Należałoby więc badać stabilność metody różnicowej w powiązaniu z konkretnym algorytmem wyboru kroku całkowania, a nie metody różnicowej „jako takiej”.

2) Założono, że σ może dowolnie zmieniać się z etapu na etap (nakłada się ograniczenie wyłącznie na obszar, w którym σ ma pozostawać). Ignoruje się w ten sposób fakt, że σ zmienia się proporcjonalnie do długości kroku całkowania ($\sigma_{n+1} = r_1 \sigma_n$), a współczynniki zmiany kroku całkowania w kolejnych etapach są ze sobą związane (i -ty współczynnik w n -tym etapie przechodzi w $i+1$ -szy współczynnik w $n+1$ -szym etapie).

Rezygnując z ww. założeń upraszczających, można stwierdzić stabilność metod różnicowych Geara jedynie w dwóch przypadkach, bardzo jednak istotnych dla układów sztywnych.

1) $|\sigma| = H_1 |\lambda| \gg 1$. Przypadek ten zachodzi wtedy, gdy szybkozmienne stany przejściowe w układzie wygasły, a na algorytm wyboru kroku całkowania oddziałują jedynie składowe wolnozmienne rozwiązania, związane z małymi wartościami własnymi. Wówczas dla wartości własnych dużych, odpowiadających składowym szybkozmiennym rozwiązaniom, σ może przyjmować duże wartości. Jeżeli współczynniki zmiany kroku

całkowania utrzymują się w rozsądnych granicach, to współczynnikami metody różnicowej $\alpha_i(r_1, \dots, r_{k-1})$ są ograniczone co do wartości bezwzględnej pewną liczbą M . Wówczas dla odpowiednio dużych wartości σ nierówność $\sum_{i=1}^k \left| \frac{\alpha_i}{\sigma - \alpha_0} \right| < 1$ jest spełniona i słuszne jest następujące oszacowanie:

$$|e(t_{n+1})| \leq \sum_{i=1}^k \left| \frac{\alpha_i}{\sigma - \alpha_0} \right| |e(t_{n+1-i})| \leq \max_{i=1, \dots, k} |e(t_{n+1-i})| \sum_{i=1}^k \left| \frac{\alpha_i}{\sigma - \alpha_0} \right| < \max_{i=1, \dots, k} |e(t_{n+1-i})|.$$

Oznacza to, że metoda różnicowa Geara jest bezwzględnie stabilna, niezależnie od tego w jaki sposób zmienia się krok całkowania (pod warunkiem, że pozostanie wystarczająco duży).

2) $|\sigma| = H_1 |\lambda| \ll 1$. Przypadek ten zachodzi wtedy, gdy na algorytm wyboru kroku całkowania oddziałują składowe szybkozmiennne rozwiązania. Wówczas dla małych wartości własnych odpowiadających składowym wolnozmiennym rozwiązaniom ww. warunek jest spełniony. Załóżmy dodatkowo, że krok zmienia się z etapu na etap w postępie geometrycznym (założenie to nie jest bardzo odległe od rzeczywistości, gdyż w procesie rozwiązania (3.1) występuje tendencja do ciągłego wzrostu kroku całkowania — z niewielkimi fluktuacjami w przypadku rozwiązań oscylacyjnych), czyli $r_1 = r_2 = \dots = r = \text{const}$. W tych warunkach (3.4) jest równaniem różnicowym liniowym o stałych współczynnikach i do analizy stabilności można użyć tego samego aparatu matematycznego co w przypadku metod o stałym kroku. Badania dla metod rzędu 2 [15] i wyższych [4] wykazały, że pierwiastki $\sum_{i=1}^k \alpha_i(r) z^{k-i}$ pozostają w obrębie koła jednostkowego pod warunkiem, że $r < r_{\max}$. Odpowiednie wartości $r_{\max}$ podano w tabl. 2. Ponieważ dla metod

Tablica 2

Maksymalne wartości współczynnika zmiany kroku całkowania zapewniające stabilność metody rzędu k dla $|\sigma| \ll 1$ (dla $k = 3 \div 6$ podano wartości przybliżone)

k	1	2	3	4	5	6
$r_{\max}$	∞	$1 + \sqrt{2}$	1,55	1,25	1,1	1,02

rzędu $1 \div 6$ $r_{\max} > 1$, to można zapewnić stały wzrost kroku całkowania H_1 w warunkach stabilności, ograniczając współczynnik wzrostu kroku całkowania wartością $r_{\max}$. Stały wzrost H_1 pociąga za sobą wzrost σ i po pewnej liczbie etapów warunek $|\sigma| \ll 1$ nie będzie mógł być spełniony. O stabilności metody nie będzie można wówczas nic powiedzieć aż do momentu, gdy H_1 osiągnie na tyle dużą wartość, że stabilność będzie znów zagwarantowana (przypadek 1). W obszarze o niepewnej stabilności punkt σ może pozostawać

(w warunkach stałego wzrostu kroku całkowania) przez skończoną liczbę etapów — nie należy się więc obawiać znacznego wzrostu błędu globalnego. Znaczna akumulacja błędów może powstać jedynie wówczas, gdy σ będzie znajdować się w obszarze niestabilności przez wiele etapów, a więc tylko wtedy, gdy obszary niestabilności istnieją już przy $r = 1$, tzn. dla metod nie charakteryzujących się stabilnością typu A. Tak więc dla metod rzędu 1 i 2 można nie obawiać się trudności wynikających z niestabilności [15] — mogą one jednak wystąpić dla metod rzędu $3 \div 6$.

Rozważmy, jak zachowuje się metoda różnicowa Geara o stałym rzędzie i zmiennym kroku całkowania w procesie rozwiązywania (3.1), w sytuacji gdy występują kłopoty ze stabilnością. Punkt $\sigma = \lambda H_1$ przesuwają się po półprostej (na rys. 1 półprosta OX), której nachylenie zależy od wartości własnej λ ($\operatorname{tg} \alpha = \frac{\operatorname{Im} \lambda}{-\operatorname{Re} \lambda}$). Załóżmy, że półprosta OX przecina granicę obszaru niestabilności dla $H = \text{const}$. Niech metoda różnicowa współpracuje z algorytmem wyboru kroku całkowania, opisanym wzorem (2.9). Wartość kroku całkowania dla następnego etapu H'_1 będzie dobierana automatycznie w zależności od wartości $k+1$ pochodnej rozwiązania

$$\delta \frac{e_{\max}}{t_{\max} - t_0} H'_1 = \gamma_k \left[\frac{H'_1}{H_1}, r_1, \dots, r_{k-2} \right] (H'_1)^{k+1} |x^{k+1}(\xi)|. \quad (3.8)$$

Po podstawieniu za $|x^{(k+1)}(\xi)|$ wartości bezwzględnej $k+1$ pochodnej rozwiązania dokładnego $x^*(t) = x_0 e^{\lambda t}$, tzn. $|x^{(k+1)}(t)| = x_0 |\lambda|^{k+1} e^{\operatorname{Re} \lambda t}$ otrzymuje się

$$H'_1 = \sqrt[k]{\frac{\delta e_{\max}}{(t_{\max} - t_0) \gamma_k x_0 |\lambda|^{k+1} e^{\operatorname{Re} \lambda t}}}. \quad (3.9)$$

Z (3.9) wynika, że w miarę upływu czasu t krok całkowania H_1 rośnie i w pewnym momencie punkt $\sigma = H_1 \lambda$ osiągnie wartość, przy której metoda różnicowa przestanie być stabilna. Wówczas, wskutek wzmacniania wprowadzanych błędów lokalnych, błąd globalny zacznie wzrastać i wkrótce zdominuje rozwiązanie dokładne. Otrzymane w ten sposób rozwiązanie pasożytnicze będzie oddziaływać na algorytm kontroli kroku całkowania tak, że krok przestanie rosnąć i punkt σ przesunie się w kierunku początku układu współrzędnych. Ostatecznie σ ustali się w punkcie leżącym na granicy obszaru niestabilności. Ponieważ odpowiada to stałemu krokowi całkowania, do badania zachowania się metody różnicowej można użyć aparatu matematycznego stosowanego do opisu zachowania się metod różnicowych o stałym kroku całkowania.

I tak, rozwiązanie można przedstawić w postaci (3.5), przy czym $z_1, \dots, z_k$ są pierwiastkami równania (3.6) dla σ leżącego na granicy obszaru stabilności dla $H = \text{const}$. Niech $z_i = e^{\xi_i} = e^{\operatorname{Re} \xi_i} e^{j \operatorname{Im} \xi_i}$ będzie tym pierwiastkiem, który jest odpowiedzialny za niestabilność, tzn. dla σ znajdującego się na granicy obszaru niestabilności $|z_i| = e^{\operatorname{Re} \xi_i} = 1$, a dla σ leżącego wewnątrz obszaru niestabilności $|z_i| > 1$. Po wygaśnięciu stanów niestabilnych rozwiązanie równania (3.1) przyjmie następującą postać

$$x(t_n) = e(t_n) = c'_i (z_i)^n = c'_i e^{\operatorname{Re} \xi_i n} e^{j \operatorname{Im} \xi_i n}. \quad (3.10)$$

Gdy σ leży na granicy obszaru niestabilności, to $\operatorname{Re} \xi_i = 0$. Odpowiada to drganiom niegasnącym (o ile $\operatorname{Im} \xi_i \neq 0$). Aby określić częstość tych drgań, obliczono wartości pierwiastków równania (3.6) dla kilku wartości σ leżących na granicy obszaru niestabilności.

Wartości pierwiastka z_l odpowiadającego za niestabilność podano w tabl. 3. Na podstawie wyników przedstawionych w tabl. 3 można stwierdzić, że dla σ leżących na granicy obszaru niestabilności w punktach niezbyt odległych od początku układu współrzędnych spełniona jest w przybliżeniu następująca równość

$$\text{Im} \xi_l \approx \text{Im} \sigma = H \text{Im} \lambda. \quad (3.11)$$

Oznacza to, że częstość drgań pasożytniczych jest w przybliżeniu równa częstości drgań własnych układu (rozwiązanie dokładne wynosi $x^*(t_n) = x_0 e^{\text{Re} \lambda H n} e^{j \text{Im} \lambda H n}$, co należy porównać z (3.10)). Wyjątkiem jest metoda różnicowa rzędu 6, gdzie częstość drgań pasożytniczych przekracza nawet kilkakrotnie częstość drgań własnych układu.

Tablica 3

Wartości pierwiastka z_l odpowiedzialnego za niestabilność metody różnicowej Geara rzędu k

k	σ		$z_l = \exp \xi_l$	
	$\text{Re} \sigma$	$\text{Im} \sigma$	$ z_l = \exp \text{Re} \xi_l$	$\text{Im} \xi_l$
3	-0,0014	0,28	1	0,28
	-0,083	1,30	1	1,047
4	-0,002	0,434	1	0,436
	-0,548	3,539	1	1,72
5	-0,0037	0,751	1	0,785
	-0,041	0,854	1	0,942
	-0,812	1,259	1	1,396
	-1,742	2,266	1	1,640
6	-0,0029	0,840	0,9997	1,2
	-0,0606	0,793	1	1,25
	-1,08	0,51	1	1,52
	-4,92	3,07	1	1,9

Z punktu widzenia dokładności otrzymanego rozwiązania ważne jest, aby amplituda drgań pasożytniczych nie przekraczała dopuszczalnej wartości błędu globalnego $e_{\max}$. Niech A_k będzie amplitudą tych drgań uzyskaną przy użyciu metody różnicowej k -tego rzędu. Do oszacowania tej wielkości posłużymy się wzorem (3.8), podstawiając $\gamma_k = C_k$ (mamy do czynienia ze stałym krokiem całkowania H)

$$x(t_n) = A_k \exp(j \text{Im} \xi_l n) = A_k \exp(j \text{Im} \xi_l H n / H) = A_k \exp\left(j \frac{\text{Im} \xi_l}{H} t\right) \Big|_{t=Hn}$$

$$\text{oraz } |x^{(k+1)}(t)| = A_k' (|\text{Im} \xi_l| / H)^{k+1}.$$

Równanie (3.8) można więc przepisać w następującej postaci

$$\delta \frac{e_{\max}}{t_{\max} - t_0} H = C_k H^{k+1} A_k \frac{|\text{Im} \xi_l|^{k+1}}{H^{k+1}}. \quad (3.12)$$

po wprowadzeniu współczynnika p , oznaczającego liczbę pełnych cykli drgań na odcinku $\in \langle t_0, t_{max} \rangle$

$$t_{max} - t_0 = p T_l = p \frac{2\pi H}{|\operatorname{Im} \xi_l|}$$

otrzymamy następujące oszacowanie amplitudy drgań pasożytniczych

$$A_k = \frac{\delta e_{max}}{2\pi p C_k |\operatorname{Im} \xi_l|^k} \quad (3.13)$$

A_k jest proporcjonalne do dopuszczalnej wartości błędu globalnego e_{max} . Współczynnik proporcjonalności jest tym większy, im mniejszy jest argument pierwiastka odpowiedzialnego za niestabilność obliczonego w punkcie $\sigma = H\lambda$ na granicy obszaru niestabilności, tzn. im bliżej osi $\operatorname{Re} \sigma$ znajduje się punkt wejścia w obszar niestabilności (patrz 3.11). W związku z tym największa amplituda drgań pasożytniczych wystąpi przy analizie układów o wartości własnej leżącej bardzo blisko osi urojonej (patrz rys. 1). Podstawiając do wzoru (3.13) bardzo niekorzystne wartości parametrów: $p = 10$, $\delta = 0,55$ oraz $\operatorname{Im} \xi_l$ obliczone do bardzo mało tłumionych drgań $\left(\frac{\operatorname{Im} \lambda}{-\operatorname{Re} \lambda} = 200 \right.$,

patrz tabl. 3) otrzymamy współczynniki proporcjonalności dla metod rzędu 3, 4, 5, 6 odpowiednio 2,9, 3,0, 0,4, 0,05. Wynika stąd, że dla bardzo niewielkiej klasy układów tylko dla metod rzędu 3 i 4 istnieje możliwość, że amplituda drgań pasożytniczych prześciera wartość e_{max} .

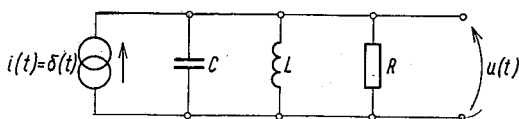
Uzyskane powyżej wyniki sprawdzono posługując się najprostszym układem elektrycznym, przy analizie którego mogą wystąpić kłopoty ze stabilnością.

Przykład

Poszukiwana jest odpowiedź impulsowa tłumionego obwodu rezonansowego jak na rys. 2. Wyżej wymieniony układ elektroniczny można opisać następującym układem równań różniczkowych liniowych o stałych współczynnikach:

$$\dot{x} = Ax, \quad x(t) = \begin{bmatrix} x_1(t) \\ x_2(t) \end{bmatrix} = \begin{bmatrix} u(t) \\ \dot{u}(t) \end{bmatrix}, \quad A = \begin{bmatrix} 0 & 1 \\ -\frac{1}{LC} & -\frac{1}{RC} \end{bmatrix},$$

$$x(0) = \begin{bmatrix} \frac{1}{C} \\ -\frac{1}{RC^2} \end{bmatrix} \quad (3.14)$$

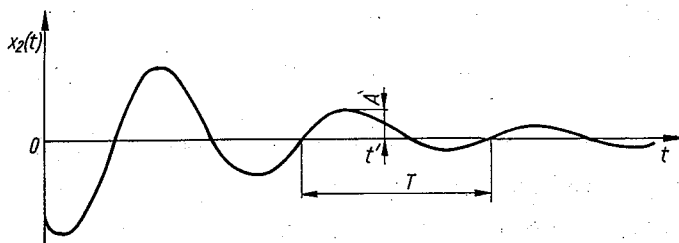


Rys. 2. Układ testujący

Wartości własne macierzy A wynoszą $\lambda_{1,2} = \alpha \pm j\beta$, gdzie $\alpha = -\frac{1}{2RC}$, $\beta = \sqrt{\frac{1}{LC}}$.

W obliczeniach będzie stosowany następujący układ jednostek: μs , MHz , nF , mA , mH , V , mA . Na błąd globalny rozwiązania $e(t) = \begin{bmatrix} e_1(t) \\ e_2(t) \end{bmatrix}$ będzie nałożone ograniczenie $|e_1(t)| < e_{\max}$, $|e_2(t)| < e_{\max}$. Wartość dopuszczalna błędu globalnego $e_{\max}$ będzie dobrana tak, aby $\frac{e_{\max}}{t_{\max} - t_0} = \text{const} = 0,002$ (np. dla $t_{\max} = 50$, $t_0 = 0$, $e_{\max} = 0,1$).

Równanie (3.14) będzie rozwiązywane dla dużego stosunku $\frac{\beta}{-\alpha} = \frac{|\text{Im}\lambda|}{|\text{Re}\lambda|}$, gdyż wtedy mogą wystąpić problemy ze stabilnością metod różnicowych Geara. Rozwiązanie równania (3.14) jest wtedy oscylacyjne, o charakterze takim jak na rys. 3 (przedstawiony przykładowy przebieg drugiej składowej wektora stanu $x_2(t)$).



Rys. 3. Przykładowy przebieg rozwiązania układu (3.14)

A — amplituda rozwiązania odniesiona do chwili t' , T — okres drgań odniesiony do chwili t'

W występujących w dalszej części tekstu tablicach 4 i 5 porównano parametry (A , T rozwiązania dokładnego $x^*(t)$ równania (3.14) z parametrami rozwiązania $x(t)$ uzyskanego za pomocą metody różnicowej (dla większej przejrzystości wzięto pod uwagę tylko drugą składową wektora stanu). Przez n oznaczono ilość etapów rozwiązania od momentu $t_0 = 0$ do bieżącej chwili t .

Rozwiązania przybliżone równania (3.14) uzyskano za pomocą algorytmu PBDF, który jest adaptacją algorytmu PBDF opisanego w [1] z następującymi niewielkimi zmianami:

- w celu polepszenia stabilności szybkość wzrostu kroku całkowania ograniczono do wartości $r_{\max}$ (tabl. 2);
- krok całkowania dla następnego etapu jest wyliczany ze wzoru (2.9) lub (2.10), a współczynnik bezpieczeństwa δ ma wartość 0,5; w [1] $\delta = 0,57$, lecz do następnego etapu bierze się wartość kroku całkowania zmniejszoną w stałym stosunku $w = 0,88$ w porównaniu z wielkością uzyskaną ze wzorów (2.9) lub (2.10).

Algorytm opisany w [1] wzięto pod uwagę z tego względu, że jest to (obok [4]) jedno z najbardziej efektywnych implementacji metod różnicowych Geara, umożliwiającą zmianę kroku całkowania i rzędu (w zakresie $k-1 \div k+1$) w każdym etapie rozwiązania (badania porównawcze przedstawione w [1] wykazały jej większą efektywność od algorytmu opisanego w [8]).

W tabl. 4 przedstawiono parametry rozwiązania dokładnego $x^*(t)$ równania (3.14) dla $\beta/-\alpha = 10$ i rozwiązania przybliżonego $x(t)$ uzyskanego za pomocą algorytmu PBDF1 o stałym rzędzie $k = \text{const} = 4$. Rozwiązanie przybliżone charakteryzuje się występowaniem drgań niegasnących. Częstość tych drgań jest w przybliżeniu równa częstości drgań własnych układu, co potwierdza słuszność wzoru (3.11). Amplituda tych drgań jest niemal 10-krotnie mniejsza od dopuszczalnej wartości błędu globalnego ($e_{\max} = 0,002 \cdot 50 = 0,1$), co potwierdza dokonane uprzednio oszacowanie. Punkt σ ustalił się w pobliżu granicy obszaru niestabilności dla metody Gears 4 rzędu (rys. 1).

Tablica 4

Parametry rozwiązania równania (3.14) ($\beta/-\alpha = 10$, $t_0 = 0$, $t_{\max} = 50$) za pomocą algorytmu PBDF1 (rzęd stały $k = 4$)

t [μs]	n	amplituda $x_2^*(t)$	amplituda $x_2(t)$	okres $x_2^*(t)$	okres $x_2(t)$	$\text{Im } \sigma$
0	—	2,36	2,36	2,35	2,46	0,29
10	101	0,19	0,19	—,,—	2,14	0,38
20	149	$1,1 \cdot 10^{-2}$	$1,8 \cdot 10^{-2}$	—,,—	2,19	0,74
30	182	$8,8 \cdot 10^{-4}$	$1,1 \cdot 10^{-2}$	—,,—	2,8	0,73
40	214	$5,1 \cdot 10^{-5}$	$1,1 \cdot 10^{-2}$	—,,—	2,14	0,93
50	245	$6,2 \cdot 10^{-6}$	$1,1 \cdot 10^{-2}$	—,,—	2,0	0,89

W tabl. 5 przedstawiono wyniki uzyskane dla metody 6 rzędu przy silniej tłumionym obwodzie rezonansowym ($\beta/-\alpha = 1$). Ponieważ półprosta $\beta/-\alpha = 1$ przecina granicę obszaru niestabilności dla $k = 6$ (patrz rys. 1), powstają podobne efekty jak przy użyciu metody 4 rzędu. Nowym zjawiskiem jest występowanie drgań pasożytniczych o częstości kilkukrotnie większej od częstości drgań własnych układu. Potwierdza to uzyskany uprzednio wynik, że dla metody 6 rzędu $\text{Im } \xi_i > H \text{Im } \lambda$. Drgania pasożytnicze, których istnienie potwierdzają tabl. 4 i 5, wynikają z braku A -stabilności metod różnicowych Gears 3÷6 i są zjawiskiem niekorzystnym z następujących przyczyn:

• wywołują złudne wrażenie występowania drgań niegasnących w tłumionych układach, dla metody 6 rzędu wprowadzają znaczne zniekształcenia częstotliwościowe, powodują ustalenie się σ w punkcie na granicy obszaru niestabilności, uniemożliwiając dalsze zwiększenie kroku całkowania.

Najbardziej istotna jest przyczyna 3°. W przypadku układów sztywnych, gdy oprócz szybkozmiennych składowych oscylacyjnych występują wolnozmiennne składowe rozwiązania, efekt ten może mieć katastrofalne skutki: czas obliczeń może wydłużyć się o wiele rządów wielkości. Uniemożliwia to stosowanie metod różnicowych Gears o stałym rzędzie $k = 3, 4, 5$ lub 6 do analizy liniowych układów sztywnych, które posiadają choć jedną parę dużych wartości własnych leżących w pobliżu osi urojonej. Podobnych efektów można oczekiwać przy analizie sztywnych układów nieliniowych, których przybliżenie liniowe posiada wartości własne leżące w pobliżu osi urojonej.

Tablica

Parametry rozwiązania równania (3.14) ($\beta/\alpha = 1, t_0 = 0, t_{max} = 50$) za pomocą algorytmu PBDF1 (rzędz stały $k = 6$)

t [μs]	n	amplituda $x_2^*(t)$	amplituda $x_2(t)$	okres $x_2^*(t)$	okres $x_2(t)$	$\text{Im } \sigma$
0	—	0,133	0,133	6,28	6,2	0,14
10	89	$2,5 \cdot 10^{-4}$	$1,0 \cdot 10^{-3}$	—,,—	3,1	0,37
20	112	$2,0 \cdot 10^{-8}$	—,,—	—,,—	2,6	0,42
30	133	$1,5 \cdot 10^{-12}$	$1,4 \cdot 10^{-3}$	—,,—	1,9	0,51
40	154	$5,5 \cdot 10^{-18}$	$1,6 \cdot 10^{-3}$	—,,—	2,0	0,52
50	175	$4,5 \cdot 10^{-22}$	$1,2 \cdot 10^{-3}$	—,,—	2,6	0,51

3.4. Stabilność algorytmu o zmiennym rzędzie i zmiennym kroku całkowania

W poprzednim podrozdziale wyjaśniono, że niestabilne zachowanie się metody różnicowej o zmiennym kroku całkowania i stałym rzędzie $k = 3, 4, 5$ lub 6 polega na występowaniu drgań pasożytniczych o amplitudzie A_k (3.13) i częstości $\text{Im } \xi_i/H$.

W obecnym przypadku rząd metody jest wybierany automatycznie. Dla każdego etapu wybiera się ten rząd k , który zapewnia maksymalną długość kroku całkowania $H_1^{(k)}$ (wg oszacowania 2.9 i 2.10). W związku z tym efekty związane z niestabilnością algorytmu mogą wystąpić tylko wtedy, gdy wybrana zostanie metoda różnicowa odpowiednio wysokiego rzędu. Gdy rząd się obniży tak, że nie wystąpi przecięcie półprostej OX (rys. 1) z granicą obszaru niestabilności dla $H = \text{const}$, drgania pasożytnicze nie wystąpią. W tej sytuacji będziemy mówili, że metoda różnicowa zachowuje się stabilnie lub że nie wystąpiły efekty związane z niestabilnością, zdając sobie sprawę z tego, że musi to oznaczać stabilność bezwzględnej w sensie definicji z punktu 3.1.

W związku z powyższym celowe jest określenie, w jakich warunkach występuje tendencja do obniżania rzędu. Załóżmy, że algorytm korzysta z metody różnicowej k -tego rzędu. Niech wartość kroku całkowania wynosi w $H^{(k)}$ i nie zmienia się znacznie z etapem na etap (takie warunki występują przy generacji drgań pasożytniczych). Wprowadzenie współczynnika w pozwoli rozszerzyć nasze rozważania na tę klasę algorytmów wyboru rzędu i kroku całkowania, w których długość kroku całkowania dla następnego etapu jest wyliczana ze wzorów (2.9) i (2.10), a następnie jest zmniejszana w stałym stosunku $w < 1$ [1]. Analiza warunków stabilności algorytmu Geara wyposażonego w taką własną kontrolę kroku i rzędu zostanie przeprowadzona w p. 4.4. Obecnie założymy, że $w =$

Oznaczmy przez $z_i = \exp(\xi_i)$ dominujący pierwiastek (tzn. pierwiastek o największym module) wielomianu (3.6) dla $\sigma = \lambda w H^{(k)}$. Dla σ leżącego w otoczeniu początku układu współrzędnych pierwiastek dominujący jest jednocześnie pierwiastkiem głównym, tzn. $\xi_i = \xi_i(w H^{(k)} \lambda) \simeq w H^{(k)} \lambda$. Załóżmy, że rozwiązanie przybliżone równania (3.1) je

zależne wyłącznie od pierwiastka dominującego (to założenie jest w pełni słuszne dla σ leżącego na granicy obszaru niestabilności, gdy występują drgania pasożytnicze), tzn.

$$x(t_n) = \exp(\xi_l n) = \exp\left(\frac{\xi_l}{wH^{(k)}} wH^{(k)} n\right) = \exp\left(\frac{\xi_l}{wH^{(k)}} t\right)\bigg|_{t=t_n}. \quad (3.15)$$

Krok całkowania $H^{(k)}$ jest wyliczany z etapu na etap ze wzorów (2.9) i (2.10), musi być więc spełniona następująca równość

$$\frac{\delta e_{\max}}{t_{\max} - t_0} H^{(k)} = C_k [H^{(k)}]^{k+1} |x^{(k+1)}(t)|, \quad (3.16)$$

gdzie na podstawie (3.15)

$$|x^{(k+1)}(t)| = \frac{|\xi_l|^{k+1}}{w^{k+1} [H^{(k)}]^{k+1}} \exp\left(\frac{\operatorname{Re} \xi_l}{wH^{(k)}} t\right).$$

Stąd otrzymuje się

$$H^{(k)} = \frac{t_{\max} - t_0}{\delta e_{\max}} C_k \frac{|\xi_l|^{k+1}}{w^{k+1}} \exp\left(\frac{\operatorname{Re} \xi_l}{wH^{(k)}} t\right). \quad (3.17)$$

Jednocześnie sprawdza się, czy metoda różnicowa rzędu $q \neq k$ pozwoli na użycie większego kroku całkowania $H^{(q)}$; $H^{(q)}$ wylicza się ze wzoru analogicznego do (3.16)

$$\frac{\delta e_{\max}}{t_{\max} - t_0} H^{(q)} = C_q [H^{(q)}]^{q+1} |x^{(q+1)}(t)|, \quad (3.18)$$

gdzie

$$|x^{(q+1)}(t)| = \frac{|\xi_l|^{q+1}}{w^{q+1} [H^{(k)}]^{q+1}} \exp\left(\frac{\operatorname{Re} \xi_l}{wH^{(k)}} t\right).$$

Oznaczając przez h_q iloraz $H^{(q)}/H^{(k)}$ otrzymuje się

$$H^{(q)} = \frac{t_{\max} - t_0}{\delta e_{\max}} C_q [h_q]^{q+1} \frac{|\xi_l|^{q+1}}{w^{q+1}} \exp\left(\frac{\operatorname{Re} \xi_l}{wH^{(k)}} t\right). \quad (3.19)$$

Dzieląc (3.19) przez (3.17) otrzymuje się

$$h_q = \frac{C_q}{C_k} |\xi_l|^{q-k} w^{k-q} [h_q]^{q+1}; \quad (3.20)$$

stąd

$$[h_q]^q = \frac{C_k}{C_q} |\xi_l|^{k-q} w^{q-k}. \quad (3.21)$$

W algorytmie opisanym w [1], w każdym etapie istnieje możliwość zmiany rzędu tylko o 1, tzn. $q = k \pm 1$. Nas interesuje, w jakich warunkach dojdzie do obniżenia rzędu. Podstawiając $q = k-1$ otrzymamy

$$[h_{k-1}]^{k-1} = \frac{C_k}{C_{k-1}} \frac{|\xi_l|}{w}. \quad (3.22)$$

Gdy $h_{k-1} > 1$, algorytm obniży rząd o 1. Na podstawie (3.22) stanie się to wtedy, gdy

$$|\xi_l| > \frac{C_{k-1}}{C_k} w. \quad (3.23)$$

Do obniżenia rzędu dojdzie więc wtedy, gdy pierwiastek dominujący wielomianu (3.6) $z_l = \exp(\xi_l)$ spełni nierówność (3.23). Będzie to miało miejsce dla σ leżącego w obszarze oddalonym od początku układu współrzędnych, gdyż dla $|\sigma| \ll 1$, $\xi_l \simeq wH^{(k)}\lambda = \sigma$. W celu określenia obszarów, w których spełnione jest (3.23) znajdziemy wpierw zbiór punktów, dla których $|\xi_l| = \frac{C_{k-1}}{C_k}w$. Powyższa równość jest spełniona dla

$$\xi_l = w \frac{C_{k-1}}{C_k} e^{j\varphi} = w \frac{C_{k-1}}{C_k} (\cos \varphi + j \sin \varphi), \quad \varphi \in \langle 0, 2\pi \rangle, \quad (3.24)$$

a więc dla

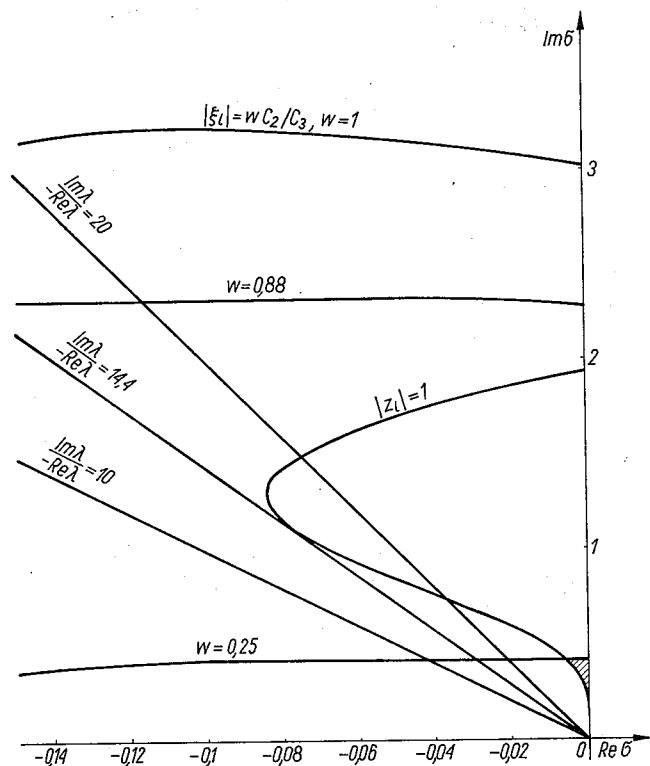
$$z_l = \exp(\xi_l) = \exp\left(\frac{wC_{k-1}}{C_k} \cos \varphi\right) \exp\left(j \frac{wC_{k-1}}{C_k} \sin \varphi\right) \quad \varphi \in \langle 0, 2\pi \rangle. \quad (3.25)$$

Z równania (3.6) można wyliczyć σ

$$\sigma = z^{-k} \sum_{i=0}^k a_i z^{k-i} = \sum_{i=0}^k a_i z^{-i}. \quad (3.26)$$

Podstawiając (3.25) do (3.26) otrzymamy szukany zbiór punktów. Tą samą metodą otrzymuje się granicę obszaru stabilności, podstawiając do (3.26) $z_l = \exp(j\psi)$, $\psi \in \langle 0, 2\pi \rangle$.

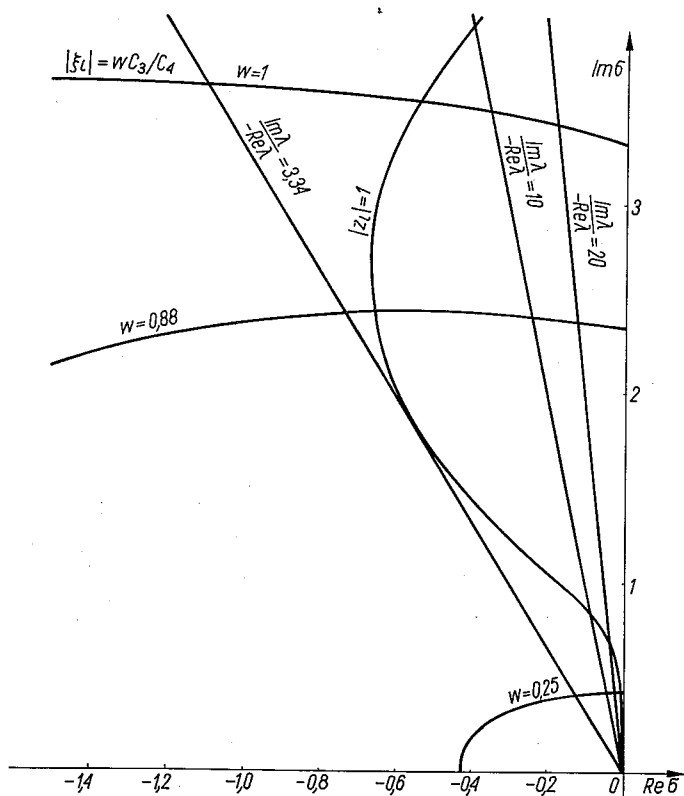
Na rys. 4. przedstawiono wyniki uzyskane dla $k = 3$. Trajektorja $|\xi_l| = C_2/C_3$ dzieli obszar rysunku na dwa podobszary. W podobszary leżącym pod krzywą $|\xi_l| = C_2/C_3$



Rys. 4. Granica obszaru niestabilności ($|z_l| = 1$) i trajektorie $|\xi_l| = w C_2/C_3$, $w = 1, 0.88, 0.25$ dla metody różnicowej 3 rzędu

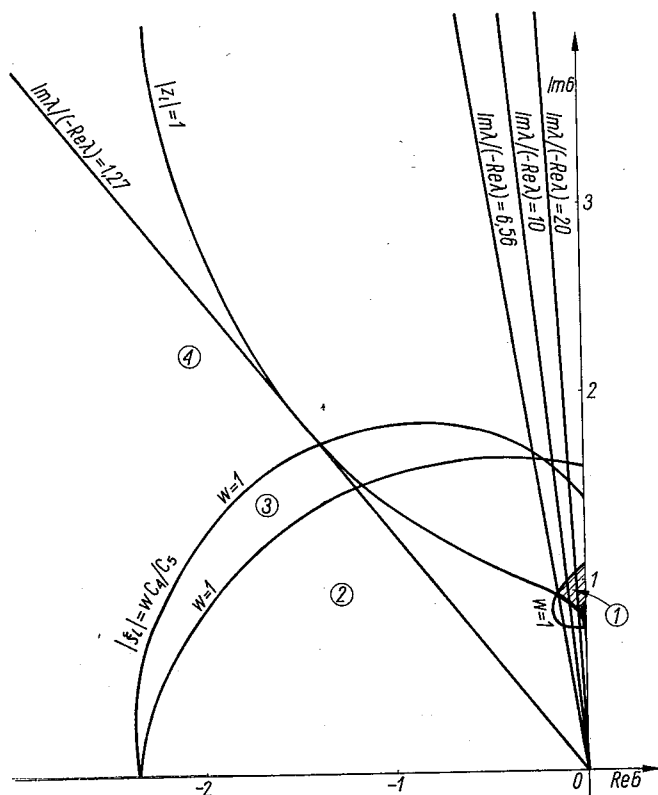
wszystkie pierwiastki wielomianu (3.6) spełniają warunek $|\xi_i| < C_2/C_3$. Oznacza to, że nie występuje tam tendencja do obniżania rzędu z $k = 3$ na $q = 2$. W związku z tym, w każdym punkcie σ leżącym na przecięciu się półprostej OX o nachyleniu $\frac{\text{Im } \lambda}{-\text{Re } \lambda}$ z granicą obszaru niestabilności mogą powstać drgania pasożytnicze. Efektów związanych z niestabilnością można się spodziewać dla $\text{Im } \lambda / -\text{Re } \lambda > 14,4$.

Podobna sytuacja występuje dla $k = 4$ (patrz rys. 5). Tu również wejście w obszar niestabilności odbywa się w warunkach, gdzie wszystkie pierwiastki spełniają nierówność $|\xi_i| < C_3/C_4$. Efektów związanych z niestabilnością (drgań pasożytniczych) można się spodziewać dla $\frac{\text{Im } \lambda}{-\text{Re } \lambda} > 3,34$.



rys. 5. Granica obszaru niestabilności ($|z_i| = 1$) i trajektorie $|\xi_i| = w C_3/C_4$, $w = 1, 0,88, 0,25$ dla metody różnicowej 4 rzędu

Zupełnie inna sytuacja występuje dla $k = 5$ (patrz rys. 6). Trajektorja $|\xi_i| = C_4/C_5$ przecina kilka podobszarów, w obrębie których podano liczbę pierwiastków spełniających nierówność (3.23). Dla nas najbardziej interesująca jest sytuacja w otoczeniu granicy obszaru niestabilności. Tę część granicy obszaru niestabilności, w której pierwiastek dominujący (o module $|z_i| = 1$) nie spełnia nierówności (3.23) oznaczono grubą linią. W pozostałej części granicy obszaru niestabilności pierwiastek dominujący spełnia (3.23)



Rys. 6. Granica obszaru niestabilności ($|z_i| = 1$) i trajektoria $|\xi_i| = w C_4 / C_5$, $w = 1$, dla metody 5 rzędu cyframi oznaczono liczbę pierwiastków, dla których $|\xi_i| > C_4 / C_5$

a więc nastąpi obniżenie rzędu i drgania pasożytnicze nie wystąpią (jeżeli drgania wystąpią to będzie za nie odpowiedzialna metoda różnicowa rzędu 4 lub 3, a nie 5).

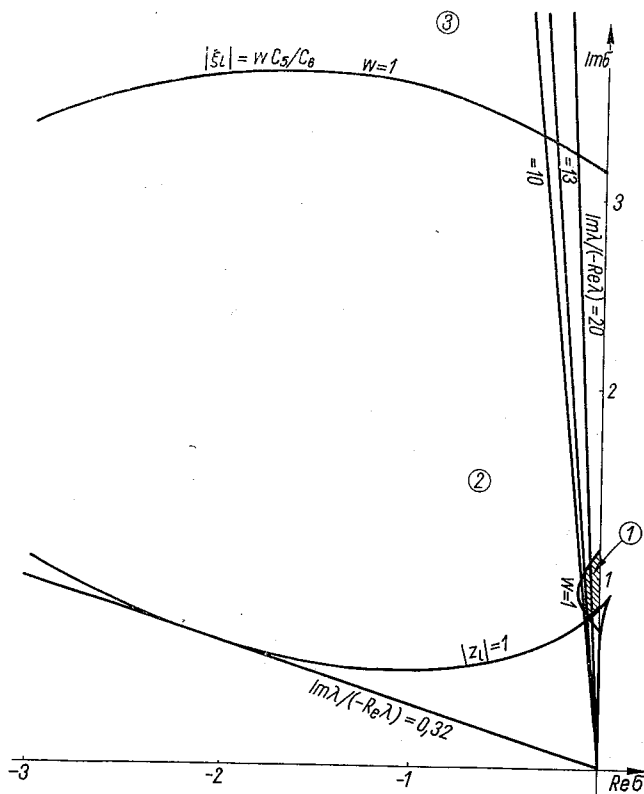
Metoda 5 rzędu może być przyczyną wystąpienia drgań pasożytniczych tylko dla $\text{Im } \lambda / -\text{Re } \lambda > 6,56$. Dla porównania, metoda stałego rzędu $k = 5$ może dawać efekty drgań pasożytniczych dla $\text{Im } \lambda / -\text{Re } \lambda > 1,27$. Obecnie efektywny obszar niestabilności został zredukowany do obszaru zakresowanego, w którym pierwiastek dominujący o module $|z_i| > 1$ nie spełnia nierówności (3.23). W pozostałej części obszaru niestabilności nastąpi obniżenie rzędu do $k = 4$. Ponieważ metoda rzędu 4 może dawać efekty niestabilności dla $\text{Im } \lambda / -\text{Re } \lambda > 3,34$, to samo ograniczenie obowiązuje i dla metod o zmiennym rzędzie w zakresie od 1 do 5.

Podobna sytuacja występuje dla metody rzędu 6 (rys. 7). Drgania pasożytnicze spowodowane przez metodę 6 rzędu mogą powstać wyłącznie przy $\text{Im } \lambda / -\text{Re } \lambda > 13$ (dla porównania, przy stałym rzędzie $k = 6$ dla $\text{Im } \lambda / -\text{Re } \lambda > 0,32$). Efektywny obszar niestabilności redukuje się do części zakreskowej, w pozostałej części obszaru niestabilności nastąpi obniżenie rzędu. Ostatecznie dla metody o zmiennym rzędzie w zakresie od 1 do 6 decydująca będzie stabilność metody 4 rzędu ($\text{Im } \lambda / -\text{Re } \lambda < 3,34$).

Reasumując można stwierdzić, że algorytm o zmiennym kroku całkowania i zmienny rzędzie w zakresie od 1 do $k_{\max}$ posiada dla $k_{\max} = 2, 3$ i 4 te same granice stabilności

to analogiczne metody o stałym rzędzie. Dla $k_{\max} = 5$ i 6 algorytm jest znacznie stabilniejszy od analogicznych metod stałego rzędu.

Ilustracją potwierdzającą słuszność ww. wniosków może być eksperyment przedstawiony w tabl. 6, poz. 2. Układ z rys. 2 był dla $\beta/-\alpha = 10$ rozwiązywany metodą Geara o zmiennym rzędzie k w zakresie od 1 do 6 i zmiennym kroku całkowania. Ponieważ dla $\beta/-\alpha = 10$ następuje przecięcie półprostej OX z granicą obszaru niestabilności dla metod rzędu 4 i 5 (patrz rys. 5 i 6) algorytm zachował się niestabilnie (wystąpiły drgania pasyżnicze, punkt σ zatrzymał się w pobliżu granicy obszaru niestabilności).



rys. 7. Granica obszaru niestabilności ($|z_i| = 1$) i trajektoria $|\xi_i| = w C_5/C_6$, $w = 1$, dla metody 6 rzędu cyframi oznaczono liczbę pierwiastków, dla których $|\xi_i| > C_5/C_6$

W tabl. 6, poz. 1 przedstawiono wyniki analogicznego eksperymentu przeprowadzonego z bardziej tłumionym układem ($\beta/-\alpha = 1$). Ponieważ $\beta/-\alpha < 3,34$, algorytm zachował się stabilnie (krok uległ zwiększeniu, nie wystąpiły drgania pasyżnicze). Dla równania (tabl. 5), w tych samych warunkach metoda stałego rzędu $k = 6$ zachowała się niestabilnie.

Należy tutaj podkreślić, że stwierdzenie, że algorytm o zmiennym rzędzie w zakresie 1 do 6 i zmiennym kroku całkowania zachowuje się stabilnie dla $\text{Im } \lambda / -\text{Re } \lambda < 3,34$ oznacza, że przy rozwiązywaniu rzeczywistych układów o wartościach własnych spełniających warunek $\text{Im } \lambda / -\text{Re } \lambda > 3,34$ zawsze wystąpią efekty związane z niestabilnością.

Wiąże się to z tym, że dla rzeczywistych układów (w odróżnieniu od (3.1)) oszacowania wielkości $|x^{(k+1)}(t)|$ i $|x^{(q+1)}(t)|$ dane wzorami (3.16) i (3.18) nie są ściśle (ww. oszacowania dają kres górny modułu pochodnej). W związku z tym zarówno krok całkowania jak i rząd mogą ulegać pewnym fluktuacjom, nawet w warunkach występowania drgań pasożytniczych. Co za tym idzie, istnieje pewne prawdopodobieństwo, że w warunkach gdy półprosta OX przecina granicę obszaru niestabilności nie wnikając zbyt głęboko w jego

Tablica

Parametry rozwiązania równania (3.14) za pomocą różnych wersji algorytmu PBDF1

Poz.	Wersja algorytmu	β/α	t_{max}	Liczba etapów n	$\text{Im } \sigma$ w ostatnim etapie	Wartości rzędu k w kilku ostatnich etapach
1	PBDF1: zmienny rząd, $k_{max} = 6$	1	50	38	49	1
2	jw.	10	50	202	1,1	3÷5
3	odcinkami stały rząd i krok całkowania	10	50	209	2,8	2
4	jw.	20	100	440	0,8	4÷6
5	zmienny rząd, $k_{max} = 2$	10	50	588	55	1
6	det. niestabilności met. badania normy estymatora błędu lokalnego	20	100	442	52	1
7	jw.	10	50	254	16	1
8	det. niestabilności oparta na zmodyfik. algorytmie badania normy estymatora błędu lokalnego	10	50	205	58	1
9	jw.	20	100	307	113	1
10	zredukowany krok całkowania, $w = 0,88$	10	50	139	63	1
11	j.w.	20	100	430	0,8	3÷6
12	zredukowany krok całkowania, $w = 0,25$	10	50	244	16,1	1
13	jw.	20	200	571	416	1
			100	567	1	2
14	det. niestabilności met. wykrywania oscylacji	10	50	166	240	1
15	jw.	20	100	331	176	1

wnętrze, drgania pasożytnicze nie wystąpią. Prawdopodobieństwo to zwiększa się, gdy stosuje się metody o rzędzie i kroku całkowania odcinkami stałych (np. [10]), gdzie po każdej zmianie kroku całkowania i rzędu wielkości te są niezmiennie przez $k+1$ etapów.

W tabl. 6 poz. 3 przedstawiono wynik rozwiązywania układu (3.14) metodą identyczną jak poprzednio, ale z możliwością zmiany kroku całkowania i rzędu co $k+1$ etapów. W tym przypadku mimo, że $\beta/-\alpha = 10$, efekty związane z niestabilnością nie wystąpiły (porównaj tabl. 6, poz. 2).

Stosowanie algorytmu z możliwością zmiany rzędu i kroku całkowania co pewną liczbę etapów nie jest jednak uniwersalnym lekarstwem i nie rozwiązuje problemu niestabilności. Poza tym prowadzi to do obniżenia efektywności algorytmu (punkt 2.3). Przykładem niech będzie użycie ww. metody do rozwiązywania mniej tłumionego układu ($\beta/-\alpha = 20$, patrz tabl. 6, poz. 4), gdzie wystąpiły efekty związane z niestabilnością.

4. SPOSOBY POPRAWIANIA STABILNOŚCI METODY RÓŻNICOWEJ GEARA O ZMIENNYM RZĘDZIE I ZMIENNYM KROKU CAŁKOWANIA

4.1. Obniżenie maksymalnego rzędu metody różnicowej do wartości $k_{max} = 2$

Z rozważań przedstawionych w punkcie 3.4 wynika, że metoda różnicowa Geara o zmiennym kroku całkowania i zmiennym rzędzie w zakresie od 1 do k_{max} , $k_{max} = 3, 4, 5$ lub 6, nie nadaje się do rozwiązywania układów sztywnych o wartościach własnych leżących blisko osi urojonej (występują pasożytnicze drgania niegasnące, uniemożliwiające zwiększenie kroku całkowania). Najprostszym sposobem wyeliminowania tych efektów jest ograniczenie się do metod różnicowych 1 i 2 rzędu poprzez zmniejszenie wartości k_{max} do dwóch. W tej sytuacji drgania pasożytnicze nie wystąpią. Z rozważań przedstawionych w punkcie 3.3 wynika, że metoda różnicowa rzędu pierwszego jest bezwzględnie stabilna dla dowolnych zmian kroku całkowania, a metoda 2 rzędu, przy ograniczeniu współczynnika wzrostu kroku całkowania do $1 + \sqrt{2}$, może dać jedynie ograniczony wzrost błędu globalnego (patrz też [15]). Sposobu tego użyli twórcy programu ASTAP [12]. Jest on niezawodny, natomiast prowadzi niekiedy do znacznego obniżenia efektywności programu ([4], [8], [16]).

W tabl. 6, poz. 5 przedstawiono wyniki eksperymentu, w trakcie którego rozwiązano układ (3.14) za pomocą metody różnicowej o maksymalnym rzędzie równym 2. Z tablicy wynika, że nie wystąpiły efekty związane z niestabilnością (drgania pasożytnicze), ale liczba etapów rozwiązania (kroków) uległa znacznemu zwiększeniu (do 588) w porównaniu z algorytmem o maksymalnym rzędzie równym 6 (202 kroki, patrz tabl. 6, poz. 2). Okazuje się, że ww. sposób poprawiania stabilności jest niekiedy bardzo kosztowny.

4.2. Detekcja niestabilności metodą badania normy estymatora błędu lokalnego

Z rozważań przedstawionych w punkcie 3.4 wynika, że stabilność algorytmu o zmiennym rzędzie i kroku całkowania zależy zarówno od użytych metod różnicowych, jak i od

metody wyboru kroku i rzędu. Począwszy od niniejszego punktu będą rozpatrywane takie sposoby poprawiania stabilności, które polegają na modyfikacji algorytmu wyboru kroku całkowania i rzędu. Algorytm wyboru kroku całkowania i rzędu opisany wzorami (2.9) i (2.10) jest algorytmem lokalnie optymalnym (zostaje wybrany taki rząd, który zapewnia maksymalną długość kroku całkowania). Jest więc prawdopodobne, że pewne modyfikacje tego algorytmu będą prowadziły do obniżenia efektywności programu.

W niniejszym punkcie będzie scharakteryzowana metoda opublikowana w [2]. Jej celem jest uzyskanie algorytmu o zmiennym kroku całkowania i zmiennym rzędzie w przedziale od 1 do 6, w którym rząd będzie obniżony w przypadku, gdy punkt $\sigma = H_1^{(k)}\lambda$ znajdzie się wewnątrz obszaru niestabilności metody rzędu k . W związku z tym konieczne jest opracowanie testu sprawdzającego, czy σ znajduje się wewnątrz obszaru niestabilności.

Założmy, podobnie jak w punkcie 3.4, że rozwiązanie przybliżone równania (3.1) zależy wyłącznie od pierwiastka dominującego $z_l = e^{\xi_l}$ i wynosi (3.15)

$$x(t_n) = \exp(\xi_l n) = \exp\left(\frac{\xi_l}{H^{(k)}} t\right) \Big|_{t=t_n}.$$

Niech s_k będzie minimalną wartością $|\xi_l|$ dla σ leżącego na tej części granicy obszaru niestabilności, gdzie może nastąpić wejście punktu σ w obszar niestabilności (przecięcie się granicy obszaru niestabilności z półprostą OX). Ponieważ tak określone s_k ma dla $k = 3$ i 4 wartość zero (dla $H^{(k)}|\lambda| \ll 1$, $H^{(k)}|\lambda| \approx |\xi_l|$), przyjmijmy, że wejście w obszar niestabilności może odbywać się w warunkach, gdy $\text{Im } \lambda / -\text{Re } \lambda < \nu$, $\nu \approx 200$ (w [2] przyjęto $\nu = 20$). Odpowiednie wartości s_k podano w tabl. 7. Niestabilność może mieć miejsce tylko wówczas, gdy przy użyciu metody różnicowej rzędu k , $|\xi_l| > s_k$, czyli spełniony jest łatwy do sprawdzenia warunek

$$\frac{|E^{(k)}(t_{n+1})|}{C_k} > \frac{|E^{(k-1)}(t_{n+1})|}{C_{k-1}} s_k, \quad (4.1)$$

gdzie $E^{(k)}(t_{n+1})$ jest błędem lokalnym uzyskanym przy użyciu metody rzędu k , a $E^{(k-1)}(t_{n+1})$ jest estymatorem błędu lokalnego, jaki byłby otrzymany, gdyby przy tej samej długości kroku całkowania $H^{(k)}$ użyć metody rzędu $k-1$.

Tablica 7

Minimalna wartość argumentu pierwiastka odpowiedzialnego za występowanie drgań pasożytniczych przy użyciu metody rzędu k

k	3	4	5	6
s_k	0,28	0,44	0,78	1,2

Istotnie, dla stałego kroku całkowania $H^{(k)}$ mamy

$$\begin{aligned} |E^{(k)}(t_{n+1})| &\simeq C_k [H^{(k)}]^{k+1} |x^{(k+1)}(t_{n+1})| = C_k [H^{(k)}]^{k+1} \left[\frac{|\xi_l|}{H^{(k)}} \right]^{k+1} \exp\left(\frac{\text{Re } \xi_l}{H^{(k)}} t_{n+1}\right) = \\ &= C_k |\xi_l|^{k+1} \exp\left(\frac{\text{Re } \xi_l}{H^{(k)}} t_{n+1}\right) \end{aligned} \quad (4.2)$$

oraz

$$|E^{(k-1)}(t_{n+1})| \simeq C_{k-1} [H^{(k)}]^k \left[\frac{|\xi_l|}{H^{(k)}} \right]^k \exp \left(\frac{\operatorname{Re} \xi_l}{H^{(k)}} t_{n+1} \right) = C_{k-1} |\xi_l|^k \exp \left(\frac{\operatorname{Re} \xi_l}{H^{(k)}} t_{n+1} \right). \quad (4.3)$$

Dzieląc stronami (4.2) przez (4.3) otrzymuje się

$$\frac{|E^{(k)}(t_{n+1})|}{|E^{(k-1)}(t_{n+1})|} = \frac{C_k}{C_{k-1}} |\xi_l|. \quad (4.4)$$

Wobec tego dla σ leżącego wewnątrz obszaru niestabilności, gdzie $|\xi_l| > s_k$, warunek 4.1) musi być spełniony.

W procesie rozwiązywania wielowymiarowych układów rzeczywistych, błąd lokalny jest wektorem

$$E^{(m)}(t_{n+1}) = [E_1^{(m)}(t_{n+1}), \dots, E_N^{(m)}(t_{n+1})]^T, \quad m = k, \quad k-1.$$

Składowe wektora błędu lokalnego odpowiedzialne za niestabilność mogą przyjmować do modułu dowolne wielkości nie większe od wielkości określonych wzorami (4.2) (4.3). W związku z tym nie jest możliwe bezpośrednie sprawdzenie warunku (4.1) dla każdej składowej wektora $E^{(m)}(t_{n+1})$ — trzeba znaleźć normę wektora $E^{(m)}(t_{n+1})$, która pozostawałaby stała mimo oscylacyjnych zmian składowych tego wektora. W [2] proponuje się normę

$$||E^{(m)}(t_{n+1})||^2 = \sum_{i=1}^N \left[\frac{|E_i^{(m)}(t_{n+1})|}{\max_{j=1, \dots, n+1} |x_i(t_j)|} \right]^2. \quad (4.5)$$

Norma ta pozostaje stała, mimo oscylacyjnych zmian składowych wektora błędu, jedynie w procesie rozwiązywania autonomicznych układów liniowych. W przypadku układów nieautonomicznych, gdy pewne składowe wektora $x(t)$ mają duży podkład składowej stałej, wartość wyrażenia (4.5) może podlegać silnym wahaniom na przestrzeni kilku kolejnych etapów rozwiązania.

Rozwiązanie określa się jako niestabilne, gdy warunek (4.1) z normą (4.5) jest spełniony oraz przy stałym kroku $H^{(k)}$ i rzędzie k w ciągu k kolejnych etapów norma estymatora błędu lokalnego ma tendencję do wzrostu

$$||E^{(k)}(t_{n+1})|| > ||E^{(k)}(t_n)|| > \dots > ||E^{(k)}(t_{n-k+1})||. \quad (4.6)$$

W praktyce dopuszcza się możliwość wystąpienia 1 ÷ 2 wyjątków od reguły (4.6) na przestrzeni k kolejnych etapów rozwiązania.

W przypadku wykrycia niestabilności (4.1 i 4.6 spełnione) obniża się rząd do wartości q ($1 \leq q \leq k-1$), gdzie q jest rzędem, dla którego wyrażenie $||E^{(q)}(t_{n+1})||/C_q$ osiąga minimum.

Aby sprawdzić skuteczność opisanej metody detekcji niestabilności, zastosowano ją w algorytmie o zmiennym kroku i rzędzie opisanym w [1]. Wymagało to przeróbki algorytmu PBDF1 na algorytm o kroku i rzędzie odcinkami stałych oraz zmiany reakcji na wykrycie niestabilności (w przypadku wykrycia niestabilności obniża się rząd do wartości $q-1$, gdyż algorytm PBDF1 nie stwarza możliwości zmiany rzędu o więcej niż 1 z etapu na etap). W tabl. 6, poz. 6 i 7 przedstawiono wyniki eksperymentu. Algorytm detekcji

niestabilności działa skutecznie (rozwiązanie zachowuje się stabilnie: po wygaśnięciu stanów nieustalonych, krok całkowania ulega zwiększeniu). Pewną jego niekorzystną cechą było kilkakrotne zadziałanie w początkowej fazie rozwiązania, gdy punkt σ leżał w otoczeniu początku układu współrzędnych, daleko od granicy obszaru niestabilności („fałszywy alarm”).

Otrzymane rezultaty można porównać z wynikami otrzymanymi za pomocą algorytmu PBDF1 o odcinkami stałym rzędzie i kroku (tabl. 6, poz. 4 dla $\beta/-\alpha = 20$ i tabl. 6, poz. 3 dla $\beta/-\alpha = 10$). W przypadku bardziej tłumionego układu ($\beta/-\alpha = 10$) algorytm detekcji niestabilności był właściwie zbędny, gdyż odpowiedni algorytm bez detekcji niestabilności zachował się stabilnie (patrz tabl. 6, poz. 3).

Reasumując, opisana w [2] metoda detekcji niestabilności jest skuteczna, posiada jednak następujące mankamenty:

- 1° działa dopiero po wejściu punktu $\sigma = H^{(k)}\lambda$ w obszar niestabilności (dopiero w tych warunkach (4.6) może być spełnione), co może doprowadzić do wzrostu błędu globalnego,
- 2° współpracuje z algorytmem o kroku i rzędzie odcinkami stałych; jak wynika z rozważań przedstawionych w części 2, ograniczenie to może doprowadzić do znacznego obniżenia efektywności algorytmu, zwłaszcza w procesie rozwiązywania równań nieliniowych ([1], [4]),
- 3° wymaga znalezienia właściwej normy estymatora błędu lokalnego; w przypadku złej dobranej normy wzrasta prawdopodobieństwo wystąpienia „fałszywego alarmu”.

4.3. Zmodyfikowany algorytm badania normy estymatora błędu lokalnego

Najbardziej istotną wadą algorytmu opisanego w punkcie 4.2 jest konieczność utrzymywania odcinkami stałego kroku i rzędu. Obecnie będzie rozpatrywana możliwość adaptacji ww. algorytmu do metody o kroku i rzędzie zmieniających się z etapu na etap. Niestabilne zachowanie się tego rodzaju metody polega na tym, że punkt $\sigma = H^{(k)}\lambda$ utrzymuje się na granicy obszaru niestabilności. Co za tym idzie, nie należy oczekiwać stałego wzrostu estymatora błędu lokalnego i sprawdzanie warunku (4.6) nie ma sensu. Pozostaje sprawdzanie warunku (4.1), co ma sens zważywszy na to, że w warunkach niestabilności krok całkowania nie ulega dużym zmianom z etapu na etap. Rozwiązanie będzie więc określone jako niestabilne, gdy warunek (4.1) będzie spełniony. Zmianie ulegnie również reakcja na wykrycie niestabilności. O ile w poprzednim przypadku obniżony rząd metody $q = k-1$ był utrzymywany przez co najmniej k następnych etapów o tyle w obecnych warunkach test przeprowadzany jest w każdym etapie i obniżony rząd dotyczy wyłącznie następnego etapu.

Tak modyfikowany algorytm został zastosowany do rozwiązywania układu (3.14). Wyniki zamieszczono w tabl. 6, poz. 8 i 9. Algorytm jest skuteczny (efekty związane z niestabilnością nie wystąpiły, porównaj tabl. 6, poz. 2). Największym jego mankamentem jest duża liczba „fałszywych alarmów”.

4.4. Algorytm oparty na redukcji kroku całkowania

Rozważmy obecnie inną modyfikację algorytmu wyboru rzędu i kroku całkowania. Niech krok całkowania dla następnego etapu będzie wyliczany ze wzorów (2.9) i (2.10), z tym że wyliczona długość kroku będzie następnie zmniejszana w stałym stosunku $w < 1$.

Modyfikacja taka była zasugerowana w [1] z myślą o zmniejszeniu prawdopodobieństwa przekroczenia dopuszczalnej wartości błędu lokalnego w następnym etapie rozwiązania (w [1] $w = 0,88$). W toku eksperymentów z układem opisanym równaniami (3.14) przeprowadzanych przez autora okazało się, że modyfikacja ta ma wpływ na stabilność algorytmu. Ściślej, wraz ze zmniejszaniem się współczynnika w rośnie tendencja do obniżania rzędu, co powoduje poprawę stabilności algorytmu. Uzasadnienie takiego zachowania się algorytmu stanowi wzór (3.23). Gdy wartość współczynnika w ulega zmniejszeniu, wówczas już przy mniejszych wartościach $|\xi_l|$, a więc dla σ leżącego bliżej początku układu współrzędnych, następuje obniżenie się rzędu. To stwierdzenie można uściślić, wyznaczając zbiór punktów na płaszczyźnie σ , dla których $|\xi_l| = wC_{k-1}/C_k$, $w < 1$.

Na rys. 4 przedstawiono wyniki dla $k = 3$ i $w = 0,88$ (wartość zalecana w [1]). Trajektoria $|\xi_l| = 0,88C_2/C_3$ dzieli płaszczyznę rysunku na dwa obszary. W obszarze leżącym pod wymienioną krzywą wszystkie pierwiastki spełniają warunek $|\xi_l| < wC_{k-1}/C_k$ i nie występuje tendencja do obniżania się rzędu. Ponieważ wejście w obszar niestabilności może nastąpić poniżej tej krzywej, warunki stabilności metody 3 rzędu pozostaną niezmiennione. Dla $w = 0,25$ trajektoria $|\xi_l| = 0,25C_2/C_3$ przechodzi znacznie niżej i obszar niestabilności redukuje się do małej zakresowanej części.

Podobnie dla $k = 4$ i $w = 0,88$ (rys. 5) wejście w obszar niestabilności może się odbyć tylko w sytuacji, gdy $|\xi_l| < wC_3/C_4$ i warunki stabilności nie powinny się zmienić. Należy jednak zauważyć, że zmalał efektywny obszar niestabilności, w związku z czym wzrosło prawdopodobieństwo jego „przebicia”. Dla $w = 0,25$ wejście w obszar niestabilności następuje w sytuacji, gdy pierwiastek dominujący spełnia nierówność (3.23), w związku z czym metoda 4 rzędu powinna być stabilna nawet dla dużego stosunku $\text{Im } \lambda / -\text{Re } \lambda$.

Dla $k = 5$ i $w = 0,88$ obszar, w którym pierwiastek dominujący o module $|z_l| \geq 1$ nie spełnia warunku (3.23), zredukował się jeszcze bardziej (patrz rys. 6 — obszar zaczerntoniony). Dla $w = 0,25$ wejście w obszar niestabilności może nastąpić wyłącznie w sytuacji, gdy pierwiastek dominujący spełnia (3.23), w związku z czym drgania pasożytnicze nie mogą wystąpić.

Dla $k = 6$ już przy $w = 0,88$ wejście w obszar niestabilności może nastąpić wyłącznie w sytuacji, gdy (3.23) jest spełnione. Co za tym idzie, metoda różnicowa 6 rzędu wchodząca w skład algorytmu o zmiennym rzędzie i kroku całkowania nie może być przyczyną powstawania drgań pasożytniczych.

Z przedstawionych rozważań wynika, że dla $w = 0,88$ warunki stabilności algorytmu o zmiennym rzędzie i kroku całkowania nie powinny ulec większym zmianom, przy czym najbardziej aktywną rolę powinna nadal odgrywać stabilność metody różnicowej 4 rzędu. Dla $w = 0,25$ algorytm powinien być stabilny (obszar niestabilności dla metody 3 rzędu występuje przy $\beta / -\alpha > 80$, jednak z uwagi na znikome pole tego obszaru, powinien on być łatwy do „przebicia”).

Zaprezentowany tu algorytm sprawdzono w praktyce, posługując się układem (3.14).

W tabl. 6, poz. 10 i 11 przedstawiono rezultaty uzyskane przy $w = 0,88$ odpowiednio dla $\beta/-\alpha = 10$ i $\beta/-\alpha = 20$. W przypadku bardziej tłumionego układu (tabl. 6, poz. 10), mimo że wejście w obszar niestabilności dla $k = 4$ następuje w sytuacji gdy warunek (3.23) nie jest spełniony, nastąpiło „przebiecie” obszaru niestabilności i niegasnące drgania pasożytnicze nie wystąpiły. Jest to wynikiem zmniejszenia się obszaru niestabilności w porównaniu z obszarem dla $w = 1$ (porównaj z tabl. 6, poz. 2).

W przypadku mniej tłumionego układu (tabl. 6, poz. 11) wystąpiły efekty związane z niestabilnością (drgania pasożytnicze). W związku z tym potwierdza się przypuszczenie, że wartość $w = 0,88$ jest zbyt duża i nie gwarantuje stabilności algorytmu.

W tabl. 6, poz. 12 i 13 przedstawiono wyniki dla $w = 0,25$. Tu w obu przypadkach osiągnięto stabilne zachowanie się algorytmu, zapłacono za to jednak znacznym zwiększeniem liczby kroków w przedziale $t \in \langle t_0, t_{max} \rangle$.

4.5. Detekcja niestabilności metodą wykrywania oscylacyjnej składowej błędu lokalnego

Jak wynika z rozważań przeprowadzonych w punkcie 3.3, efekty związane z niestabilnością metod o zmiennym kroku całkowania polegają na wystąpieniu oscylacji pasożytniczych dających się opisać następującym wzorem

$$\begin{aligned} x(t_n) &= A_k \exp(\xi_l n) = A_k \exp\left(\frac{\xi_l}{H^{(k)}} H^{(k)} n\right) = A_k \exp\left(\frac{\xi_l}{H^{(k)}} t\right) \Big|_{t=t_n} = \\ &= A_k \exp\left(j \frac{\text{Im } \xi_l}{H^{(k)}} t\right) \Big|_{t=t_n}, \end{aligned} \quad (4.7)$$

gdzie A_k można oszacować ze wzoru (3.13), a $z_l = \exp(\xi_l) = \exp(j \text{Im } \xi_l)$ jest pierwiastkiem dominującym wielomianu (3.6) dla σ leżącego w miejscu przecięcia półprostej $\text{Im } \lambda / -\text{Re } \lambda$ z granicą obszaru niestabilności metody rzędu k .

Dla rzeczywistych układów oscylacje pasożytnicze będą miały następującą postać (zakładamy, że i -ta składowa wektora $x(t)$ jest odpowiedzialna za niestabilność)

$$x_i(t_n) = A_k \cos\left(\frac{\text{Im } \xi_l}{H^{(k)}} t_n + \varphi\right). \quad (4.8)$$

Ich częstość $\frac{\text{Im } \xi_l}{H^{(k)}}$ musi być większa od $\frac{s_k}{H^{(k)}}$ (patrz tabl. 7), przy założeniu, że wejście

w obszar niestabilności może się odbywać dla $\frac{\text{Im } \lambda}{-\text{Re } \lambda} < 200$. Tę samą częstość drgań posiada również i -ta składowa estymatora błędu lokalnego $E_i(t)$, gdyż jest proporcjonalna do $k+1$ pochodnej rozwiązania $x_i(t)$ (wzór 2.3).

W związku z tym, rozwiązanie określa się jako niestabilne, gdy częstość drgań estymatora błędu $E_i(t)$ dla metody rzędu k jest większa niż $s_k/H^{(k)}$.

Dla wykrywania niestabilności można zastosować pomysł Krogha [3], który mając do czynienia z zupełnie inną klasą metod rozwiązywania równań różniczkowych badań, czy rozwiązanie użyteczne nie jest zdominowane przez składową o wyższej częstości (po-

chodzącą od innego niż główny pierwiastka równania charakterystycznego), zliczając liczbę przejść przez zero estymatora błędu. Ideę tę, po odpowiedniej modyfikacji, można zastosować do wykrywania niestabilności algorytmu opartego na metodach Geara, gdzie najczęściej pierwiastek główny ($\xi_l \simeq H^{(k)}\lambda$) jest odpowiedzialny za niestabilność.

Okres drgań pasożytniczych jest równy $T_l = \frac{2\pi H^{(k)}}{\text{Im } \xi_l}$ i nie przekracza wielkości $\frac{2\pi H^{(k)}}{s_k}$. Wobec tego w warunkach niestabilności na jeden okres drgań przypada $\frac{T_l}{H^{(k)}} = \frac{2\pi}{\text{Im } \xi_l} \leq \frac{2\pi}{s_k}$ etapów rozwiązania. Na półokres drgań, czyli na odcinek pomiędzy dwoma kolejnymi przejściami przez zero estymatora błędu lokalnego, przypada $\frac{\pi}{\text{Im } \xi_l} \leq \frac{\pi}{s_k}$ etapów (patrz tabl. 8).

Na tej zasadzie oparty jest następujący test na wykrywanie niestabilności algorytmu o zmiennym kroku i rzędzie w zakresie od 1 do $k_{\max}$. Rozwiązanie określa się jako niestabilne, jeżeli pomiędzy v przejściami przez zero estymatora błędu lokalnego $E_i(t)$ nie

Tablica 8

Parametry algorytmu detekcji niestabilności metodą wykrywania oscylacji pasożytniczych

k	$\frac{\pi}{s_k}$	d_k	u_k
3	11,2	12	8
4	7,14	8	12
5	4,03	5	15
6	2,62	4	25

zdarzy się, że w ciągu $d_{k_{\max}}$ kolejnych etapów estymator błędu ma ten sam znak. W przypadku układów wielowymiarowych test przeprowadza się tylko w stosunku do tych składowych estymatora błędu $E(t)$, które mają na tyle dużą wartość, że czynnie oddziałują na proces wyboru kroku całkowania.

W przypadku wykrycia niestabilności zapamiętuje się długość kroku całkowania $H_{(k_{\max})}$, przy którym nastąpiło wykrycie niestabilności i obniża się maksymalny rząd $k_{\max}$ o 1. Za $d_{k_{\max}}$ przyjęto liczbę całkowitą większą od maksymalnej liczby etapów przypadających na półokres drgań pasożytniczych: $d_{k_{\max}} > \pi/s_{k_{\max}}$ (tabl. 8).

Im większe $d_{k_{\max}}$, tym większe prawdopodobieństwo mylnego wyłączenia z algorytmu metod o wyższych rzędach. Gdy $d_{k_{\max}}$ ma za małą wartość, algorytm może być nieskuteczny (nie wykryje niestabilności).

Wartość parametru v przyjęto równą 4. Im większe v , tym większe prawdopodobieństwo pojawienia się sekwencji $d_{k_{\max}}$ etapów z jednakowym znakiem estymatora błędu

i tym dłuższe działanie algorytmu testującego niestabilność. Im mniejsze v , tym łatwiejsze mylne wyłączenie z algorytmu metod różnicowych o wyższych rzędach wskutek przypadkowych zmian znaku estymatora błędu. Te przypadkowe zmiany znaku estymatora błędu mogą być spowodowane zmianami rzędu metody różnicowej.

Gdy niestabilność zostanie wykryta i maksymalny rząd k_{max} zmniejszony do odpowiedniej wartości, krok całkowania ulegnie bez przeszkód zwiększeniu do wartości koniecznej dla śledzenia wolnozmiennych składowych rozwiązania, odpowiadających mniejszym wartościom własnym. Celowe jest wtedy zwolnienie ograniczeń nałożonych na k_{max} , aby algorytm odzyskał swą pełną efektywność. Gdy krok całkowania przekroczy wielkość $u_{k_{max}+1} H_{(k_{max}+1)}$, gdzie $H_{(k_{max}+1)}$ jest zapamiętaną długością kroku całkowania, przy której wykryto niestabilność metody o maksymalnym rzędzie $k_{max}+1$, zwiększa się k_{max} o 1. Współczynniki $u_{k_{max}}$ dobiera się tak, aby punkt σ nie znalazł się w obszarze niestabilności metody różnicowej rzędu k_{max} . Praktycznie $u_{k_{max}}$ jest stosunkiem maksymalnej wartości $|\sigma|$, przy której wystąpić może niestabilność metody rzędu k_{max} do minimalnej wartości $|\sigma|$, przy której może wystąpić niestabilność (tabl. 8).

Zaproponowany tu algorytm został sprawdzony na układzie (3.14). Wyniki przedstawiono w tabl. 6, poz. 14, 15. Algorytm działa skutecznie, zdarzają się jednak przypadki mylnego wyłączenia metod o wyższych rzędach („fałszywy alarm”). Prawdopodobieństwo „fałszywego alarmu” jest duże zwłaszcza w początkowej fazie rozwiązania, gdy szybko zmienia się rząd, a co za tym idzie, często zmienia się znak estymatora błędu lokalnego.

4.6. Porównanie metod poprawiania stabilności

Z przedstawionych w punktach 4.1÷4.5 metod poprawiania stabilności wszystkie okazały się skuteczne. Dzięki użyciu tych metod, wyeliminowano szkodliwe zjawisko, jakim jest występowanie drgań pasożytniczych, uniemożliwiających zwiększenie kroku całkowania. Każda z wymienionych metod wpływa w różny sposób na efektywność algorytmu rozwiązywania układów równań algebraiczno-różniczkowych. Jak wyjaśniono w punkcie 2.3, miarą efektywności algorytmu jest ilość etapów rozwiązania w zadanym przedziale całkowania $\langle t_0, t_{max} \rangle$.

W tabl. 6, poz. 5÷9, 12÷15 porównano wymienione metody posługując się układem (3.14) w wersji z większym ($\beta/-\alpha = 10$) i mniejszym ($\beta/-\alpha = 20$) tłumieniem. Z powyższego zestawienia wynika, że dla bardziej tłumionego układu najefektywniejszy okazał się algorytm oparty na wykrywaniu oscylacji pasożytniczych (poz. 14), a dla mniej tłumionego układu — zmodyfikowany algorytm badania normy estymatora błędu lokalnego (poz. 9).

Nieźle wyniki daje w obu przypadkach algorytm opisany w [2]. Należy tu jednak podkreślić, że współpracuje on wyłącznie z algorytmem o rzędzie i kroku całkowania odcinkami stałych. Pociąga to za sobą obniżenie efektywności, zwłaszcza w procesie rozwiązywania układów nieliniowych [4]. Wobec powyższego algorytm ten należy uznać za zbyt kosztowny.

Algorytmy oparte na obniżeniu rzędu k_{max} do 2 i redukcji kroku całkowania, mimo że współpracują z metodą Geara o rzędzie i kroku mogącym zmieniać się w każdym etapie, okazały się zbyt kosztowne.

Z dwóch pozostałych metod detekcji niestabilności trudno jest wybrać zdecydowanie lepszą. Na niekorzyść algorytmu opisanego w punkcie 4.3 przemawia fakt, że nie można znaleźć takiej normy wektora błędu lokalnego, której wartość zmieniałaby się w dostatecznie małych granicach przy oscylacyjnym zachowaniu się składowych tego wektora. Norma (4.5) ma w zasadzie sens wyłącznie dla autonomicznych układów liniowych.

Algorytm oparty na detekcji oscylacji pasożytniczych nie wymaga stosowania normy, wymaga jednak większej ilości obliczeń numerycznych i doprowadza niekiedy do mylnego odrzucenia metod o wyższych rzędach wskutek przypadkowych zmian znaku estymatora błędu lokalnego. W związku z tym dokonano modyfikacji ww. algorytmu detekcji niestabilności, poddając testowi na wykrywanie oscylacji pasożytniczych nie estymator błędu lokalnego $E(t)$, a estymator pierwszej pochodnej rozwiązania $\dot{x}(t)$. Dzięki temu wyeliminowano wpływ zmian rzędu metody różnicowej na działanie algorytmu detekcji niestabilności i bardzo znacznie zmniejszono prawdopodobieństwo wystąpienia „fałszywego alarmu”.

LITERATURA

1. W. M. G. Bokhoven, *Linear implicit differentiation formulas of variable step and order*, IEEE Trans. CAS-22, Febr. 1975.
2. S. Skelboe, *The control of order and steplength for backward differentiation methods*, BIT 17, 1977, ss. 91—107.
3. F. T. Krogh, *A test for instability in the numerical solution of ordinary differential equations*, J. ACM 14/1967, ss. 351—354.
4. T. Rübner-Petersen, *An efficient algorithm using backward time-scaled differences for solving stiff differential-algebraic systems*, Report of the Institute of Circuit Theory and Telecommunication, Technical University of Denmark, September 1973.
5. R. K. Brayton, C. C. Conley, *Some results on the stability and instability of the backward differentiation methods with nonuniform time steps*, Topics in Numerical Analysis, Proc. of the Royal Irish Academy Conf. on Numerical Analysis, Dublin, 1972.
6. W. H. Enright, T. E. Hull, B. Lindberg, *Comparing numerical methods for stiff systems of O. D. Es.*, BIT 15/1975, ss. 10—48.
7. G. G. Dahlquist, *A special stability problem for linear multistep methods*, BIT 3, 1963.
8. R. K. Brayton, F. G. Gustavson, C. D. Hachtel, *A new efficient algorithm for solving differential-algebraic systems using implicit backward differentiation formulas*, Proc. of the IEEE, January 1972.
9. C. W. Gear, *The automatic integration of stiff ordinary differential equations*, Proc. of IFIP Congress 1968, North Holland Publishing Company, Amsterdam 1969.
10. C. W. Gear, *Algorithm 407, DIFSUB for solution of ordinary differential equations*, Comm ACM 14, ss. 185—190.
11. L. O. Chua, P. M. Lin, *Computer-aided analysis of electronic circuits*, Prentice-Hall Inc., Englewood Cliffs, New Jersey 1975.
12. W. T. Weeks, A. J. Jimenez, G. W. Mahoney, D. Mehta, H. Qassemzadeh, T. R. Scott, *Algorithms for ASTAP — a Network Analysis Program*, IEEE Trans. CT-20, Nr 6.
13. L. W. Nagel, *SPICE-2: A computer program to simulate semiconductor circuits*, Memorandum № ERL-M520, University of California, May 1975.
14. F. H. Branin, *ECAP-2, A new electronic circuit analysis program*, IEEE Trans. SC-6, № 4, August 1971.
15. P. Dymarski, *Numeryczne rozwiązywanie równań stanu nieliniowego układów elektronicznych*

przy użyciu stabilnego algorytmu adaptacyjnego, Referat zesz. 38, Instytut Telekomunikacji Politechniki Warszawskiej, 1976.

16. P. Dymarski, *Metody poprawiania stabilności algorytmu numerycznego rozwiązywania równań nieliniowych dynamicznych układów elektronicznych*, Referat zesz. 77, Instytut Telekomunikacji Politechniki Warszawskiej, 1979.

P. DYMARSKI

METHODS FOR IMPROVING THE STABILITY OF AN ALGORITHM FOR SOLUTION OF STIFF DIFFERENTIAL-ALGEBRAIC SYSTEMS

Summary

Stability conditions of the backward differentiation (Gear's) methods of a variable step and order are investigated. These methods are used extensively for the solution of differential-algebraic systems resulting e.g. from transient analysis of electronic circuits. It is shown how the step and order control algorithm affects the stability of a solution. The stability conditions of a variable step constant order and a variable step variable order methods are consecutively determined. The amplitude and frequency of a parasitic solution are computed and it is suggested which systems may yield unstable solutions. It is shown how instability affects the efficiency of the backward differentiation methods. Some methods for improving the stability are given and compared with the other approaches.

P. DYMARSKI

MÉTHODES D'AMÉLIORATION DE LA STABILITÉ D'UN ALGORITHME DESTINÉ À RÉSOUDRE LES SYSTÈMES D'ÉQUATIONS ALGÈBRIQUES DIFFÉRENTIELLES

Résumé

On analyse les conditions de stabilité de la méthode multi-pas de Gear à ordre et pas variables. Cette méthode est utilisée souvent pour résoudre les systèmes d'équations algébriques différentielles, par exemple dans l'analyse des circuits électriques. On étudie l'influence d'un algorithme de détermination d'ordre et de pas sur la stabilité de la solution. On analyse les conditions de stabilité de la méthode de Gear à pas variable et ordre constant et variable. On calcule l'amplitude et fréquence erronée et on détermine la classe des systèmes à solution instable et faible efficacité de calcul. On a donné quelques méthodes permettant d'améliorer la stabilité et on les a comparées avec d'autres techniques.

P. DYMARSKI

METHODEN FÜR DIE STABILITÄTSVERBESSERUNG DES VERFAHRENS ZUR LÖSUNG STEIFER SYSTEME VON ALGEBRAISCH-DIFFERENTIALGLEICHUNGEN

Zusammenfassung

Man untersuchte die Stabilitätsbedingungen des Gear-Verfahrens mit variabler Ordnung und Schrittweite. Dieses Verfahren ist oft für die Lösung von Systemen der Algebraisch-Differentialgleichungen angewandt worden, z.B. bei der Analyse elektrischer Netzwerke. Es wird nachgewiesen, daß die Stabilität der Lösung der Steuerung von der Konsistenzordnung und Schrittweite abhängt. Zuerst untersuchte man die Stabilitätsbedingungen des Verfahrens mit konstanter Ordnung. Man berechnete Amplitude und Frequenz der Fehlerlösung und bestimmte, bei Analyse welcher Systeme Stabilitätsschwierigkeiten auftreten können. Es wurde darauf hingewiesen, daß Stabilitätsschwierigkeiten zur Effektivitätssenkung des Verfahrens beitragen. Es werden manche Methoden für die Erreichung besserer Stabilität angeführt und diese mit den bisher gebräuchlichen Methoden verglichen.

П. ДЫМАРСКИ

МЕТОДЫ УЛУЧШЕНИЯ УСТОЙЧИВОСТИ АЛГОРИТМА РЕШЕНИЯ ЖЁСТКИХ СИСТЕМ АЛГЕБРО-ДИФФЕРЕНЦИАЛЬНЫХ УРАВНЕНИЙ

Резюме

Излагаются условия устойчивости многошагового метода Гира с изменяющимся шагом и порядком. Этот метод часто применяется для решения систем алгебро-дифференциальных уравнений, на пример в программах анализа электронных схем. Рассмотрено влияние алгоритма выбора шага и порядка на устойчивости решения. Во первых излагаются условия устойчивости метода Гира с постоянным порядком и изменяющимся шагом, а потом с переменным шагом и порядком. Вычислена величина и частота паразитного решения и показано, при анализе каких систем может появиться такое решение. Показано, что неустойчивость метода ухудшает эффективность программы. Предлагаются некоторые методы улучшения устойчивости, которые сравниваются с методами применяемыми до сих пор.

Analiza, projektowanie i badania eksperymentalne wzmacniacza klasy E w optymalnych warunkach pracy

MARIAN KAZIMIERCZUK (WARSZAWA)

Instytut Radioelektroniki Politechniki Warszawskiej

Otrzymano 10.1.1980

W pracy przedstawiono najważniejsze wyniki analizy teoretycznej, sposób projektowania oraz wyniki badań eksperymentalnych wzmacniacza mocy wielkiej częstotliwości klasy E z równoległym kondensatorem pracującego w warunkach optymalnych. Na podstawie uproszczonego modelu wzmacniacza wyznaczono niektóre parametry układu. Następnie oszacowano straty mocy w kolektorze i sprawność kolektorową, a także moc wejściową i wzmocnienie mocy wzmacniacza. Przedstawiono cztery typy obwodów kolektorowych posiadających właściwości dopasowujące oraz wzory opisujące wartości elementów tych obwodów. Zamieszczono również obszerne wyniki badań eksperymentalnych.

1. WSTĘP

Specyficzną właściwością wzmacniacza mocy wielkiej częstotliwości klasy E jest duża sprawność kolektorowa [1—11], która w praktycznych układach może wynosić nawet ponad 96%. Warunki osiągnięcia dużej sprawności kolektorowej tego wzmacniacza omówiono w pracy [1]. W pracy [2] wykazano, że największą moc wyjściową można uzyskać przy kącie przepływu prądu kolektora równym 180° . Na podstawie wyników obu tych prac określono optymalne warunki pracy wzmacniacza klasy E . Przypadek ten ma duże znaczenie praktyczne i dlatego został szczegółowo omówiony w niniejszej pracy.

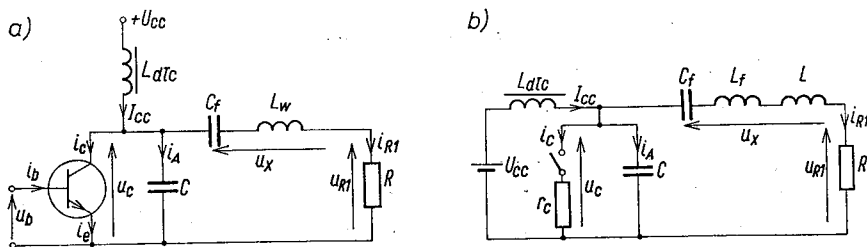
2. ANALIZA WZMACNIACZA KLASY E

2.1. Układ i model wzmacniacza

Schemat podstawowego układu tranzystorowego wzmacniacza mocy wielkiej częstotliwości klasy E przedstawiono na rys. 1a, a uproszczony model tego wzmacniacza na rys. 1b. Wzmacniacz ten składa się z tranzystora pracującego jako klucz, obwodu kolektorowego C , C_f , L_w , R i dławika L_{dlc} . Kolektor tranzystora jest zasilany napięciem stałym U_{cc} . Przez dławik L_{dlc} płynie w przybliżeniu prąd stały, równy prądowi zasilania wzmacniacza I_{cc} . Obwód kolektorowy składa się z kondensatora C włączonego równolegle z tranzystorem oraz szeregowego obwodu rezonansowego C_f , L_w , R . W szerego-

wym obwodzie rezonansowym płynie w przybliżeniu prąd sinusoidalny i_{R1} . Stąd na obciążeniu wzmacniacza R występuje w przybliżeniu sinusoidalne napięcie u_{R1} . W optymalnych warunkach pracy układu szeregowy obwód rezonansowy jest rozstrojony i ma charakter indukcyjny. Z tego względu wypadkową indukcyjność L_w w modelu wzmacniacza rozdzielono na dwie składowe L_f i L , w taki sposób, że elementy L_f , C_f są dostrojone do częstotliwości podstawowej napięcia sterującego wzmacniacz $f = \omega/2\pi$, czyli

$$\omega = \frac{1}{\sqrt{L_f C_f}}. \quad (1)$$



Rys. 1. Wzmacniacz klasy E:

a) schemat podstawowego układu wzmacniacza klasy E, b) model wzmacniacza klasy E

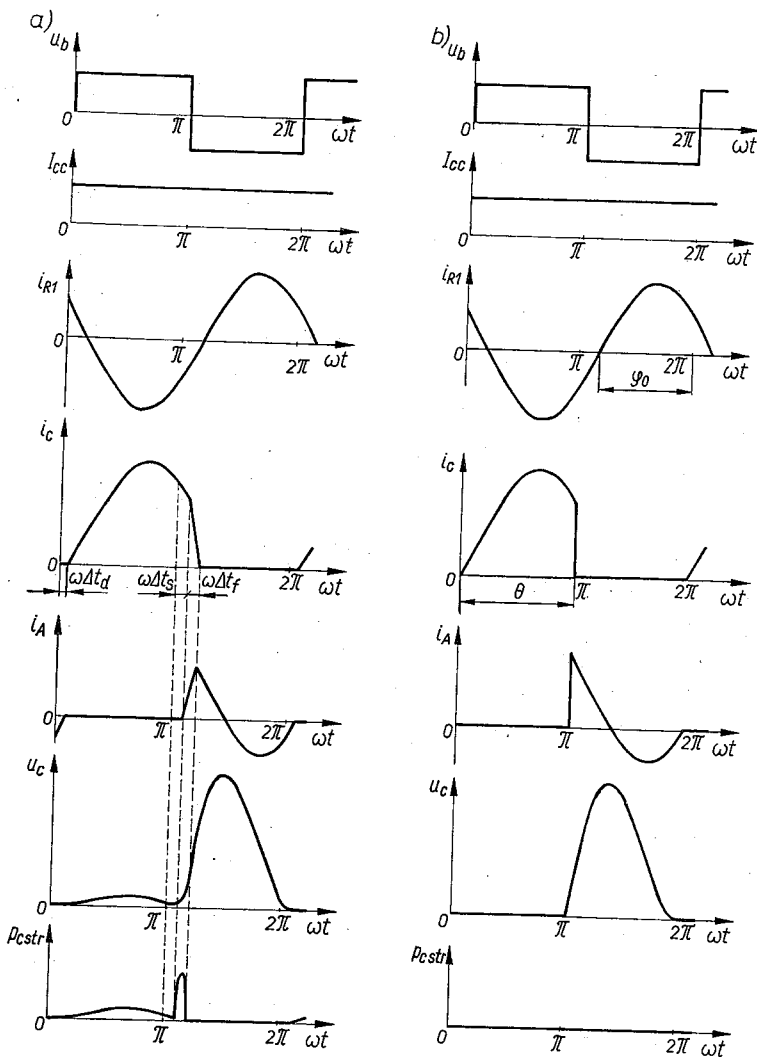
Ze względu na tłumienie harmonicznych w obciążeniu wzmacniacza R powinna być dostatecznie duża dobroć wypadkowa szeregowego obwodu rezonansowego

$$Q_w = \frac{\omega L_w}{R} = \frac{\omega(L + L_f)}{R}. \quad (2)$$

Rzeczywiste przebiegi chwilowe prądów, napięć i mocy strat w kolektorze we wzmacniaczu klasy E pracującym w warunkach optymalnych pokazano na rys. 2a, a przebiegi idealizowane na rys. 2b. W czasie włączenia tranzystora napięcie kolektora u_c jest bliskie zera, a różnica prądów I_{cc} oraz i_{R1} płynie w przybliżeniu tylko przez tranzystor. W czasie wyłączenia tranzystora różnica prądów I_{cc} oraz i_{R1} płynie przez kondensator C wywołując na nim spadek napięcia, który jest jednocześnie napięciem kolektora u_c . W czasie opóźnienia Δt_d tranzystor znajduje się w stanie odcięcia, w czasie magazynowania Δt_s — w stanie nasycenia, a w czasie opadania prądu kolektora Δt_f — w stanie aktywnym.

Zasadniczą zaletą pracy tranzystora jako klucza są małe straty mocy w kolektorze. W stanie odcięcia tranzystora straty te wynoszą zero, ponieważ prąd kolektora nie płynie, a w stanie nasycenia tranzystora straty te są małe, ponieważ napięcie kolektora przyjmuje małe wartości chwilowe (rys. 2a). Napięcie to z założenia jest mniejsze niż przy pracy tranzystora jako źródła prądowego, występującej na przykład we wzmacniaczu klasy C. Natomiast istotną wadą pracy tranzystora jako klucza są duże straty mocy w kolektorze związane z jego przełączaniem, które w ogólnym przypadku odbywa się poprzez stan aktywny [12].

We wzmacniaczu klasy E można wyeliminować moc strat w kolektorze związaną z włączaniem tranzystora. W tym celu powinny być spełnione dwa warunki, a mianowicie napięcie kolektora oraz pochodna tego napięcia w chwili włączania tranzystora (w chwili otwarcia złącza baza-emiter) powinny wynosić zero. Jeżeli przyjąć, że włączanie tranzysto-



rys. 2. Przebiegi chwilowe prądów, napięć i mocy strat w kolektorze we wzmacniaczu klasy E pracującym w warunkach optymalnych
a) przebiegi rzeczywiste, b) przebiegi idealizowane

następuje w chwili odpowiadającej kątowi fazowemu $\omega t = 2\pi$ (rys. 2b), to warunki można przedstawić w następującej postaci [1]:

$$u_c(2\pi) = 0, \tag{3}$$

$$\left. \frac{du_c(\omega t)}{d(\omega t)} \right|_{\omega t = 2\pi} = 0. \tag{4}$$

Jeżeli są spełnione powyższe warunki, to napięcie kolektora u_c oraz prąd kolektora i_c chwili włączania tranzystora wynoszą zero, przy włączaniu tranzystora nie występują bkie zmiany prądu i napięcia kolektora, bowiem napięcie kolektora u_c łagodnie opada zera przed włączeniem tranzystora, a prąd kolektora i_c łagodnie narasta od zera po

włączeniu tranzystora, a ponadto przebiegi prądu i napięcia kolektora nie przyjmują wartości ujemnych (rys. 2). Stąd przejście tranzystora ze stanu odcięcia do normalnego stanu nasycenia odbywa się natychmiastowo, z pominięciem normalnego stanu aktywnego oraz stanów inwersyjnych. Warunki (3) i (4) można spełnić poprzez odpowiedni dobór wartości elementów obwodu kolektorowego.

Moc strat w kolektorze we wzmacniaczu klasy E wynika tylko z niezerowego napięcia kolektora w stanie nasycenia tranzystora oraz z czasu opadania prądu kolektora Δt_f przy wyłączaniu tranzystora (rys. 2a).

Największą moc wyjściową wzmacniacza można osiągnąć przy kącie przepływu kolektora $\Theta = 180^\circ$ [2]. W tym przypadku, pomijając czas Δt_f , tranzystor przez pół okresu znajduje się w stanie odcięcia, a przez pozostałe pół okresu w stanie nasycenia. Ponieważ zwykle czas magazynowania Δt_s jest większy od czasu opóźnienia Δt_d , stąd w praktycznym układzie może być wymagany współczynnik wypełnienia prostokątnego napięcia sterującego mniejszy od 0,5.

Zgodnie z powyższymi rozważaniami wzmacniacz klasy E pracuje w warunkach optymalnych, gdy są spełnione warunki (3) i (4) (maksimum sprawności kolektorowej) oraz kąt przepływu prądu kolektora $\Theta = 180^\circ$ (maksimum mocy wyjściowej).

Analizę wzmacniacza przeprowadzono na podstawie modelu pokazanego na rys. 1b przy następujących założeniach upraszczających:

- 1) indukcyjność dławika L_{dlc} jest na tyle duża, że prąd płynący przez ten dławik jest stały i równy prądowi zasilania wzmacniacza I_{cc} ,
- 2) dobroć wypadkowa szeregowego obwodu rezonansowego Q_w jest na tyle duża, że prąd i_{R1} płynący w tym obwodzie można uważać za sinusoidalny,
- 3) pojemność wyjściowa tranzystora nie zależy od napięcia kolektora i wraz z równoległą pojemnością zastępczą dławika oraz pojemnością montażu wchodzi w skład pojemności C ,
- 4) elementy obwodu kolektorowego są idealne, tzn. pracują przy częstotliwości znacznie mniejszej od częstotliwości pierwszego rezonansu własnego i są bezstratne (zastępczą rezystancję szeregową strat cewki L_w można ewentualnie włączyć do rezystancji obciążenia wzmacniacza R).
- 5) napięcie kolektora w stanie nasycenia tranzystora u_c równa się zero, czyli rezystancja klucza r_c w modelu wzmacniacza jest nieskończenie mała,
- 6) szybkość przełączania klucza jest nieskończenie duża, czyli czasy Δt_d , Δt_s i Δt_f wynoszą zero.

Przebiegi prądów i napięć w modelu wzmacniacza przy powyższych założeniach pokazano na rys. 2b. Konsekwencją założenia 4) jest pominięcie mocy strat w obwodzie kolektorowym, a założeń 5) i 6) pominięcie mocy strat w kolektorze związanej z niezerowym napięciem kolektora w stanie nasycenia tranzystora oraz z czasem opadania prądu kolektora Δt_f przy wyłączaniu tranzystora.

2.2. Wyniki analizy wzmacniacza

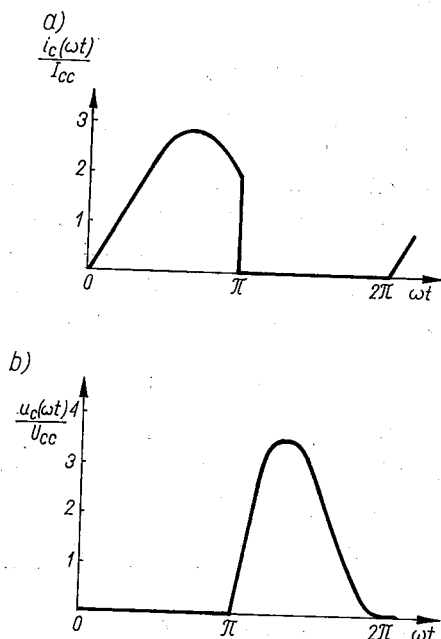
W niniejszym rozdziale przedstawiono najistotniejsze wyniki analizy wzmacniacza klasy E przeprowadzonej przy założeniach upraszczających zamieszczonych w rozdz. 2

Przebiegi prądu kolektora i_c oraz napięcia kolektora u_c wyrażają się wzorami (Dodatek 1):

$$\frac{i_c(\omega t)}{I_{cc}} = \begin{cases} \frac{\pi}{2} \sin \omega t - \cos \omega t + 1 & \text{dla } 0 < \omega t \leq \pi, \\ 0 & \text{dla } \pi < \omega t \leq 2\pi, \end{cases} \quad (5)$$

$$\frac{u_c(\omega t)}{U_{cc}} = \begin{cases} 0 & \text{dla } 0 < \omega t \leq \pi, \\ \pi \left(\omega t - \frac{3\pi}{2} - \frac{\pi}{2} \cos \omega t - \sin \omega t \right) & \text{dla } \pi < \omega t < 2\pi. \end{cases} \quad (6)$$

Przebiegi te pokazano na rys. 3.



Rys. 3. Idealizowane przebiegi prądu kolektora i_c oraz napięcia kolektora u_c we wzmacniaczu klasy E pracującym w warunkach optymalnych:

a) przebieg prądu kolektora i_c , b) przebieg napięcia kolektora u_c

Amplituda U_{R1} napięcia występującego na obciążeniu wzmacniacza R oraz amplituda I_{R1} i faza początkowa φ_0 prądu płynącego w szeregowym obwodzie rezonansowym wynosi (Dodatki 1 i 2):

$$U_{R1} = \frac{4}{\sqrt{\pi^2 + 4}} U_{cc} \cong 1,074 U_{cc}, \quad (7)$$

$$I_{R1} = \frac{\sqrt{\pi^2 + 4}}{2} I_{cc} \cong 1,862 I_{cc}, \quad (8)$$

$$\varphi_0 = \pi - \arctg \frac{2}{\pi} \cong 2,57468 \text{ rad} \cong 147,518^\circ. \quad (9)$$

Ze wzoru (7) wynika, że we wzmacniaczu klasy E możliwa jest modulacja amplitudy przez zmianę napięcia zasilania U_{cc} .

Maksymalne wartości chwilowe prądu kolektora i_{cmax} oraz napięcia kolektora u_{cmax} wyrażają się wzorami:

$$i_{cmax} = \left(\frac{\sqrt{\pi^2 + 4}}{2} + 1 \right) I_{cc} \cong 2,8621 I_{cc}, \quad (10)$$

$$u_{cmax} = 2\pi(\pi - \varphi_o) U_{cc} \cong 3,562 U_{cc}. \quad (11)$$

W celu spełnienia warunków (3) i (4) wartości elementów obwodu kolektorowego L , C , R przy kącie przepływu prądu kolektora Θ równym 180° powinny spełniać następujące związki (Dodatek 2):

$$\omega CR = \frac{8}{\pi(\pi^2 + 4)} \cong 0,1836, \quad (12)$$

$$\frac{\omega L}{R} = \frac{\pi}{16} (\pi^2 - 4) \cong 1,1525. \quad (13)$$

Wartości elementów L_f i C_f można wyznaczyć na podstawie wzorów (1) i (2), zakładając dobroć wypadkową szeregowego obwodu rezonansowego Q_w .

Stosunek $\omega L/R$ opisany wzorem (13) jest miarą rozstrojenia szeregowego obwodu rezonansowego. Wskutek tego rozstrojenia między napięciem u_{R1} i składową podstawową napięcia kolektora u_{c1} występuje przesunięcie fazowe ψ_o . Przesunięcie to wynosi (Dodatek 2)

$$\psi_o = \arctg \left[\frac{\pi(\pi^2 - 4)}{16} \right] \cong 49,052^\circ. \quad (14)$$

Przy modulacji amplitudy istotna jest rezystancja tranzystora dla prądu stałego R_m . Rezystancja ta wyraża się wzorem (Dodatek 2)

$$R_m = \frac{U_{cc}}{I_{cc}} = \frac{\pi^2 + 4}{8} R \cong 1,7337 R. \quad (15)$$

Amplituda I_{c1} i faza początkowa γ_1 składowej podstawowej prądu kolektora oraz amplituda U_{c1} i faza początkowa ϑ_1 składowej podstawowej napięcia kolektora wynoszą (Dodatek 3):

$$I_{c1} = I_{cc} \sqrt{\left(\frac{\pi}{4} + \frac{2}{\pi} \right)^2 + \frac{1}{4}} \cong 1,5074 I_{cc}, \quad (16)$$

$$\gamma_1 = -\arctg \left(\frac{2\pi}{\pi^2 + 8} \right) \cong -19,372^\circ, \quad (17)$$

$$U_{c1} = U_{cc} \sqrt{\frac{16}{\pi^2 + 4} + \frac{\pi^2(\pi^2 - 4)^2}{16(\pi^2 + 4)}} \cong 1,6389 U_{cc}, \quad (18)$$

$$\vartheta_1 = \varphi_o + \psi_o \cong 196,571^\circ. \quad (19)$$

Związki między amplitudami składowych podstawowych i maksymalnymi wartościami chwilowymi prądu i napięcia kolektora są następujące:

$$I_{c1} \cong 0,5267 i_{cmax}, \quad (20)$$

$$U_{c1} \cong 0,4601 u_{cmax}. \quad (21)$$