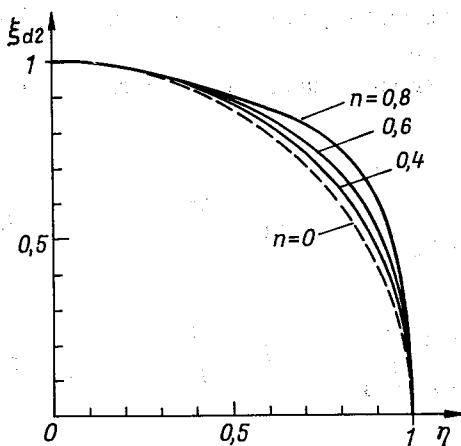
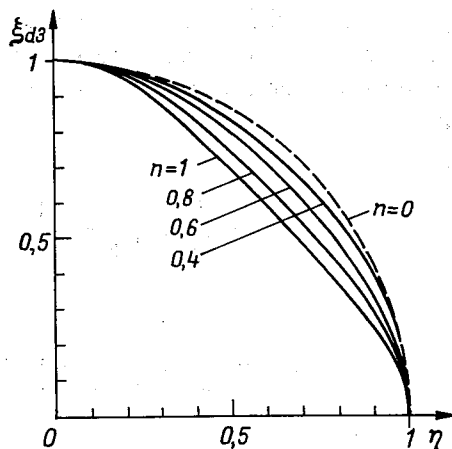


Rys. 13. Wykres zależności $\xi_{d2}(\eta)$ Rys. 14. Wykres zależności $\xi_{d3}(\eta)$

gdzie: n — parametr.
Wykresy zależności

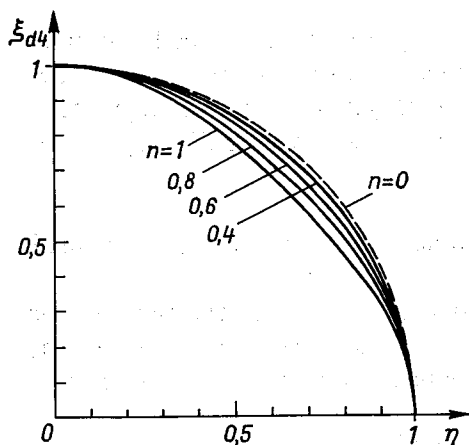
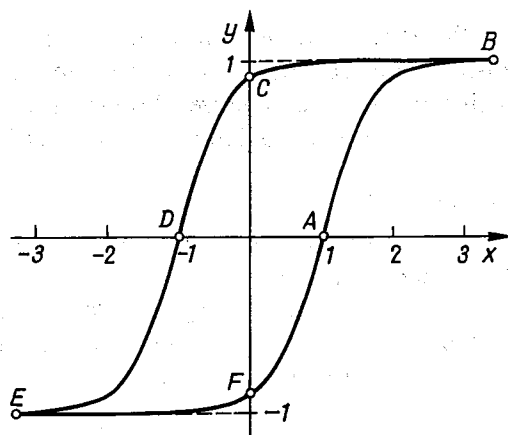
$$\xi_{d3} = \frac{\sqrt{1-\eta^2}}{1+n^2\eta^2} \quad (48)$$

oraz

$$\xi_{d4} = \sqrt{\frac{1-\eta^2}{1+n^2\eta^2}} \quad (49)$$

są dla kilku wartości parametru n przedstawione na rysunkach 14 oraz 15.

Przez porównanie wykresu zależności $\xi_d(\eta)$ dla aproksymowanej pętli (rys. 11) z wykresami $\xi_{d1}(\eta)$ oraz $\xi_{d2}(\eta)$ (rys. 12 i 13) lub z wykresami $\xi_{d3}(\eta)$ oraz $\xi_{d4}(\eta)$ (rys. 14

Rys. 15. Wykres zależności $\xi_{d4}(\eta)$ 

Rys. 16. wykres pętli histerezy do przykładu 1

i 15) można ustalić zarówno wartość współczynnika korekcyjnego, jaki należy zastosować, oraz potrzebną wartość parametru n .

Po wprowadzeniu współczynnika korekcyjnego funkcja aproksymująca przybiera postać:

dla ξ_{k1}

$$x = -\frac{1}{2k} \ln \frac{1-y}{1+y} \pm \left[1 + n^2 \left(\frac{y}{y_m} \right)^2 \right] \cdot \sqrt{1 - \left(\frac{y}{y_m} \right)^2} \quad (50)$$

dla ξ_{k2}

$$x = -\frac{1}{2k} \ln \frac{1-y}{1+y} \pm \left[1 + n^2 \left(\frac{y}{y_m} \right)^4 \right] \cdot \sqrt{1 - \left(\frac{y}{y_m} \right)^2} \quad (51)$$

dla ξ_{k3}

$$x = -\frac{1}{2k} \ln \frac{1-y}{1+y} \pm \frac{\sqrt{1 - \left(\frac{y}{y_m} \right)^2}}{1 + n^2 \left(\frac{y}{y_m} \right)^2} \quad (52)$$

dla ξ_{k4}

$$x = -\frac{1}{2k} \ln \frac{1-y}{1+y} \pm \sqrt{\frac{1 - \left(\frac{y}{y_m} \right)^2}{1 + n^2 \left(\frac{y}{y_m} \right)^2}} \quad (53)$$

gdzie: $-1 < y < 1$.

4. PRZYKŁADY

Przedstawione w poprzednim rozdziale sposoby aproksymacji pętli histerezy zostaną zilustrowane dwoma przykładami, z których pierwszy dotyczy aproksymacji pętli o praktycznie stałej szerokości i o wierzchołku praktycznie leżącym na asymptocie $Y = Y_s = \text{const}$, a drugi aproksymacji pętli o wierzchołku leżącym poniżej asymptoty.

Przykład 1

Przeprowadzić aproksymację eksperymentalnie uzyskanej normowanej pętli, przedstawionej na rys. 16.

Aproksymowana pętla jest pętlą o stałej szerokości, dlatego rozwiązania zadania szukamy w postaci wyrażonej równaniami (28) i (29). Jak z rys. 16 wynika, mamy $y_R = 0,9$. Wartość parametru k wyznaczamy posługując się wzorem (25). Otrzymujemy

$$k = -\frac{1}{2} \ln \frac{1-0,9}{1+0,9} = 1,4722 \quad (54)$$

Równanie lewej gałęzi ma postać

$$F_l(x) = \frac{2}{1 + \exp[-2,9444(x+1)]} - 1, \quad (55)$$

a prawej gałęzi

$$F_p(x) = \frac{2}{1 + \exp[-2,9444(x-1)]} - 1. \quad (56)$$

Obliczone wg. (55) i (56) wartości rzędnych są praktycznie takie same jak wartości rzędnych zmierzone na wykresie, tzn. odchyłki nie przekraczają 1% wartości rzędnej dla stanu nasycenia.

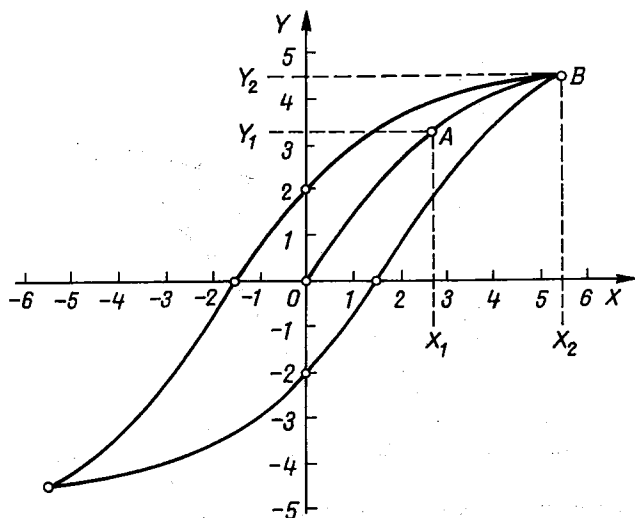
Dla $x = 3$ obliczona rzędna prawej gałęzi ma wartość $F_p(3) = 0,99971$, a lewej gałęzi $F_l(3) = 0,99998$. Różnica rzędnych jest więc równa 0,00017, tzn. można przyjąć, że wykresy funkcji $F_p(x)$ oraz $F_l(x)$ dla $x > 3$ oraz $x < -3$ się pokrywają.

Przykład 2

Dla uzyskanej w wyniku eksperymentu pętli histerezy, której wykres jest przedstawiony na rys. 17, należy wyznaczyć funkcje aproksymujące gałęzie pętli.

Z wykresu pętli histerezy wynika, że jej wierzchołek B leży poniżej asymptoty $Y_s = \text{const}$, a szerokość pętli maleje w miarę zbliżania się do wierzchołka. Dlatego aproksymację zaczynamy od wykreślenia łuku $\widehat{OAB}$ krzywej średniej pętli leżącego w pierwszej ćwiartce i odczytania wartości współrzędnych punktów A i B. Są one równe

$$\begin{aligned} X_1 &= 2,75 & Y_1 &= 3,3 \\ X_2 &= 5,5 & Y_2 &= 4,5. \end{aligned}$$



Rys. 17. Wykres pętli histerezy do przykładu 2

Następnie wg. wzoru (21) obliczamy wartość X_s

$$Y_s = Y_1 \sqrt{\frac{Y_2}{2Y_1 - Y_2}} = 3,3 \sqrt{\frac{4,5}{2 \cdot 3,3 - 4,5}} = 4,83 \quad (57)$$

oraz wg. wzorów (23) i (24) obliczamy współrzędne punktów A i B leżących na normowanej krzywej średniej. Mamy

$$x_1 = \frac{X_1}{X_c} = \frac{2,75}{1,5} = 1,833 \quad (58)$$

$$x_2 = \frac{X_2}{X_c} = \frac{5,5}{1,5} = 3,667 \quad (59)$$

$$y_1 = \frac{Y_1}{Y_s} = \frac{3,3}{4,83} = 0,683 \quad (60)$$

$$y_2 = \frac{Y_2}{Y_s} = \frac{4,5}{4,83} = 0,932 \quad (61)$$

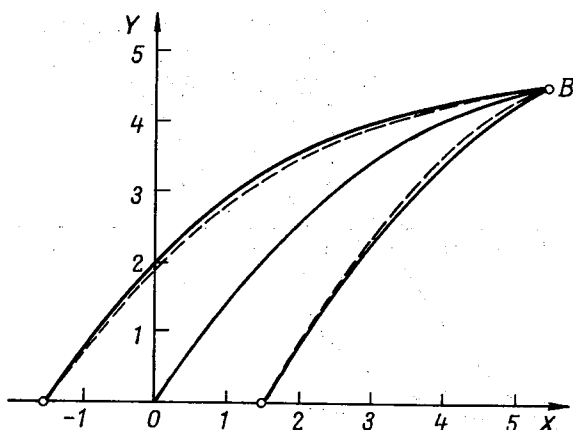
gdzie $X_c = 1,5$. Następnie wg. wzoru (26) obliczamy wartości współczynnika k

$$k = -\frac{1}{2x_1} \ln \frac{1-y_1}{1+y_1} = -\frac{1}{2 \cdot 1,833} \ln \frac{1-0,683}{1+0,683} = 0,455 \quad (62)$$

Uwzględniając, że $y_2 = y_m$, otrzymujemy wg. (41) równanie gałęzi pętli histerezy w postaci

$$x = -\frac{1}{2k} \ln \frac{1-y}{1+y} \pm \sqrt{1 - \left(\frac{y}{y_m}\right)^2} = -\frac{1}{2 \cdot 0,456} \ln \frac{1-y}{1+y} \pm \sqrt{1 - \left(\frac{y}{0,932}\right)^2}, \quad (63)$$

$-1 < y < 1.$



Rys. 18. Wykres górnej połówki aproksymowanej pętli histerezy (linia ciągła) i funkcji aproksymującej (64) (linia przerywana)

Dla kolejnych wartości y obliczamy wg. (63) odpowiadające im wartości zmiennej x . Następnie obliczamy wartości zmiennych X oraz Y wg. zależności

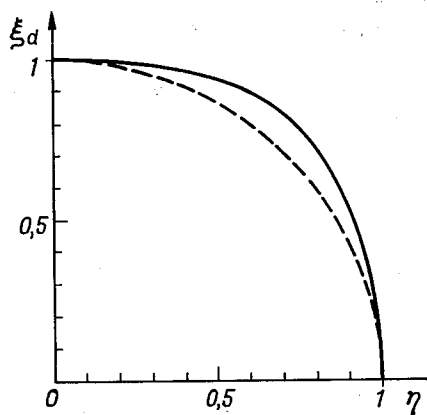
$$X = 1,5 \cdot x \quad (64)$$

oraz

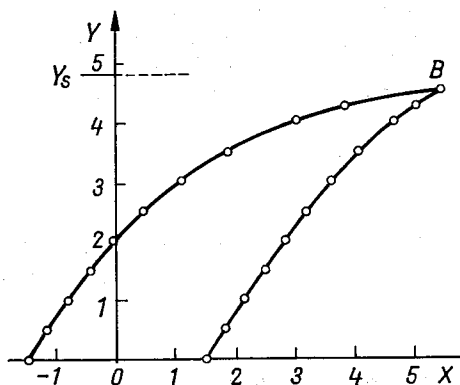
$$Y = 4,83 \cdot y \quad (65)$$

Wyniki obliczeń nanosimy na wykres przedstawiony na rys. 18, na którym linią ciągłą zaznaczono zadaną, aproksymowaną pętlę histerezy, a linią przerywaną wykres funkcji aproksymującej (63). Jak z wykresu wynika, funkcja (63) dobrze przybliża zadaną pętlę dla $0 < Y < 1$ oraz $4 < Y < 4,5$. W pozostałym przedziale moduł bezwzględny błędu przybliżenia $|\Delta Y|$ nie przekracza 2,5% maksymalnej wartości Y . W celu uzyskania lepszego przybliżenia wykonujemy wykres zależności $\xi_d(\eta)$. Jest on przedstawiony na rys. 19, na którym linią przerywaną wykreślono łuk koła o promieniu $r = 1$. Jak widać, wykres zależności $\xi_d(\eta)$ leży powyżej łuku koła. Porównanie tego wykresu z wykresami przedstawionymi na rys. 12 i 13 prowadzi do wniosku, że znacznie lepsze przybliżenie uzyska się przy zastosowaniu współczynnika korekcyjnego ξ_{k1} z parametrem $n \approx 0,5$. Poprawiona funkcja aproksymująca ma postać

$$x = -\frac{1}{2 \cdot 0,456} \ln \frac{1-y}{1+y} \pm \left[1 + 0,5^2 \left(\frac{y}{0,932} \right)^2 \right] \cdot \sqrt{1 - \left(\frac{y}{0,932} \right)^2} \quad (66)$$



Rys. 19. Wykres zależności $\xi_d(\eta)$ dla aproksymowanej pętli histerezy



Rys. 20. Wykres górnej połówki aproksymowanej pętli histerezy — kółkami zaznaczono wyniki obliczeń wg. wzoru (67)

Na rys. 20 przedstawiony jest wykres górnej połówki aproksymowanej pętli histerezy, na którym kółkami zaznaczono wyniki obliczeń wg. wzoru (66). Z porównania rysunków (18) oraz (20) wynika, że w rezultacie wprowadzenia współczynnika korekcyjnego $\xi_{k1} = 1 + 0,5^2 \eta^2$ uzyskano wyraźnie lepsze przybliżenie gałęzi pętli histerezy niż przy posługiwaniu się wzorem (63) nie zawierającym tego współczynnika.

WNIOSKI

Proponowany sposób aproksymacji histerezy ma następujące zalety:

1. funkcja aproksymująca jest opisana prostym wzorem, łatwym do zaprogramowania przy posługiwaniu się komputerem,
2. funkcja aproksymująca jest ciągła i posiada ciągłe pochodne,
3. funkcje aproksymujące opisują gałęzie pętli histerezy w całym ich przebiegu,
4. funkcje aproksymujące nadają się do przybliżenia gałęzi pętli histerezy zarówno w przypadku pętli o stałej szerokości, jak i pętli o szerokości malejącej w miarę zbliżania się do wierzchołka.

BIBLIOGRAFIA

1. *Wielka Encyklopedia Powszechna*, t. 4, Warszawa, PWN, 1964
2. *Meyers Neues Lexikon*, t. 6, Leipzig: Bibliogr. Institut, 1973
3. *Bolszaja Sowietskaja Enciklopedia*, t. 6, Moskwa, 1971
4. *Encyklopedia Techniki. Materiałoznawstwo*, Warszawa: WNT, 1969
5. *Encyklopedia Fizyki*, t. 1, Warszawa: PWN, 1972
6. G. Klaus: *Wörterbuch der Kybernetik*, Dietz Verlag, Berlin 1968
7. H. W. Katz: *Współczesne elementy magnetyczne i dielektryczne*, WNT, Warszawa 1963
8. B. Konorski: *Podstawy elektrotechniki*, t. 1, Warszawa: PWN, 1965
9. M. Dąbrowski: *Pola i obwody magnetyczne maszyn elektrycznych*, Warszawa: WNT, 1971
10. E. Philippow: *Nichtlineare Elektrotechnik*, Akadem. Verlagsges Geest u. Portig K. G., Leipzig 1963 (tłum. ros. 1968)
11. R. Kurdziel: *Obwody magnetyczne*, Warszawa: PWN, 1962
12. S. E. Winogradow, E. M. Nicenko: *Sposob approksimacii osnovnoj kriwoj namagnicziwanija*, Elektricestwo 1974, Nr 12 ss. 68—70
13. L. J. Weiss, T. Wysocki: *Przybliżenia wykładnicze przebiegów i rozkładów o charakterze skokowym lub impulsowym*, Rozprawy Elektrotechniczne 1986, t. 32, z. 2, ss. 395—410

L. J. WEISS, T. WYSOCKI

APPROXIMATION OF STATIC HYSTERESIS LOOP BY MEANS OF EXPONENTIAL FUNCTIONS

Summary

The paper deals with a method suitable for approximation of static hysteresis mean curve and loop branches by means of exponential functions. The properties of approximating function and of correcting function for achieve a better approximation, are discussed. The ways for determination of approximating function parameters for two cases: 1) when hysteresis loop shows saturation and loop width is almost constant and 2) when hysteresis loop shows no saturation and the loop width diminishes by loop peak approaching, are given. Two numerical examples are also given.

On the constitutive theory of the distributed parameter modelling for the phenomental control of the multi-component electric fields

WACŁAW NIEMIEC

Instytut Elektroniki, Politechnika Śląska w Gliwicach

Received 1989.10.12

Authorized 1989.12.28

Deduced in Part I of this article an original constitutive distributed parameter mathematical model — in the form of the system of the partial differential constitutive state equations of the single component plasma, has been solved analitically from the point of view of the control influences on the yield and quality aspects of the appearance and existence of this kind of the plasma, by its physical phenomena. The results of above considerations comparative to the classic distributed parameter control theory have been discussed. Consequently to this approach an original structure of the phenomental distributed parameter control for the single component plasma is suggested and analysed.

PART I

The partial differential constitutive state equations describing the multi-component electric fields

1. INTRODUCTION

The world literature on the application of electricity with the transformations of its state variables covers a very wide scope of processes based on processing of the multi-component electric fields. The following domains of sciences can serve as characteristic examples:

- electrochemistry [1—3], [8], — telecommunication [4—7], [8],
- electrometallurgy and metal thermal processing [8—9],
- biophysics [10], — manufacturing processes of materials [11],
- electric energy production [12—14], and many others whose list is very long, indeed. Problems of distributed parameter modelling of these real processes have not yet been exactly determined and solved, especially, in cases taking into account

applications of physical phenomena. The problem is to pin pertinent processing and phenomenological elements of the physical phenomena decisive for the yield and quality of final products or for the final effects of industrial processes. For realizing this task it is necessary to introduce a complete form describing the charge, energy and momentum as the most effective modelling method for multi-component electric fields. It is necessary to stress here the great importance of the physical features of multi-component electric fields on the products or the effects of real processes based on the application of electricity. It is important to pay attention to the procedure of this modelling method of considered processes:

- the determination of the processing and phenomenologic conditions with respect to the yield and quality of these processes,
- the measuring conditions of coefficients of chosen physical phenomena,
- the balance forms of processing and phenomenologic aspects for the charge, energy and momentum of multi-component electric fields,
- the construction of partial differential constitutive state equations and partial differential equations of continuity of every state variable as a new form of the constitutive invariance for these real processes.

2. A LOCALLY DISTRIBUTED PARAMETER METHOD OF THE MODELLING OF MULTI-COMPONENT ELECTRIC FIELDS

In my work [16] I have considered a way of modelling of locally distributed parameters of a single component electric field. This modelling is connected with consequently introduced:

- point $Z(x, y, z, t)$ of transformation of the single-component electric field,
- nearest surroundings represented by a locally selected volume-time element $\Omega_{(a,b,c)t}$ of the physical conditions of this transformation, with surface $F_{(a,b,c)t}$ which is externally oriented by the normal external surface orientation vector $n^{(z)}$ with pertinent circulation decisive for the yield and quality questions of the transformation of this kind of electric fields,
- determination of the physical phenomena worthy of acknowledgment as having importance for the yield and quality of the transformation of these electric fields,
- physical conditions inside the whole working space volume-time of the electric fields Ω_{Rt} of the apparatuses,
- control boundary conditions on the surface E_{Rt} of Ω_{Rt} is respect to the circulation of the normal external surface orientation vector $n_R^{(z)}$ relative to the circulation of $n^{(z)}$. This general procedure is very well applicable to the modelling by the space-time distributed mass/charge, energy and momentum balance condition for every real process of phenomenologically continuous media [17—19]. In the frames of this paper we shall focus our attention on “i” — component electric fields where $i = l/l$ fulfilling “l” transformations where $l = l/l/L$ which create a special kind of incompressible continuous media with space and time memories [20—22]. For the

particular analysis of multi-component electric fields, it is necessary to suggest the following theorem:

THEOREM 1. Let exist points $Z_i(x_i, y_i, z_i, t)$ of every "i" electric field undergoing the "l" transformation where $i = 1/I$ and $l = 1/L$. It is possible to introduce an equivalent point $Z(x, y, z, t)$ for which the summations $\sum_{i=1}^I \sum_{l=1}^L$ of the effects of these electric fields are valid.

PROOF. For every point $Z_i(x_i, y_i, z_i, t)$ exist:

- the scalar function of the activity of the transformation of the "i" component with "l" transformation R_{il} , and/or
- the vector function of the activity of the transformation of of the "i" component with "l" transformation S_{il} , as the characteristic magnitudes of the multi-component electric fields. There exist resulting function having the form

$$R^s = \sum_{i=1}^I \sum_{l=1}^L R_{il} \quad (2.1)$$

and/or

$$S^s = \sum_{i=1}^I \sum_{l=1}^L S_{il} \quad (2.2)$$

where:

R_{il} —represent all scalar state variables characteristic for multi-component electric fields (charge, voltage),

S_{il} —represent all vector state variables characteristic for multi-component electric fields (field vector velocity),

consequently for $i = 1/I$ and $l = 1/L$. The equivalent point $Z(x, y, z, t)$ is determined by the complete forms of the functions of the scalar and vector coordinates of the transformation activities of all scalar and vector state variables of multi-component electric fields, so that

$$Z(x, y, z, t) \xleftrightarrow{\text{def}} \begin{bmatrix} R^s \\ S^s \end{bmatrix}. \quad (2.3)$$

The thesis of the Theorem 1 leads to the assertion:

THEOREM 2. If for every point $Z_i(x_i, y_i, z_i, t)$ exists the reciprocal variable point $Q_i^R(\xi_i^R, \eta_i^R, \zeta_i^R, \tau)$ where $i = 1/I$, pertinent to the functions of activity of transformations R_{il} —for the scalar case and/or S_{il} —for the vector case, where $l = 1/L$, it is possible to determine the equivalent variable point $Q^R(\xi^R, \eta^R, \zeta^R, \tau)$.

PROOF. Case 1. We can prove that the existence of the equivalent variable point

$$Q^R(\xi^R, \eta^R, \zeta^R, \tau) \xleftrightarrow{\text{def}} \begin{bmatrix} \tilde{R}^s \\ \tilde{S}^s \end{bmatrix}, \quad (2.4)$$

where:

$$\tilde{R}^s = \sum_{i=1}^I \sum_{l=1}^L \tilde{R}_{il} \quad (2.5)$$

and/or

$$\tilde{S}^s = \sum_{i=1}^I \sum_{l=1}^L \tilde{S}_{il} \quad (2.6)$$

For formulae (2.5) and/or (2.6):

$$\tilde{R}_{il} \equiv R_{il} \quad (2.7)$$

in the points:

$$Q_i^R(\xi_i^R, \eta_i^R, \zeta_i^R, \tau) \quad Z_i(x_i, y_i, z_i, t)$$

and/or

$$\tilde{S}_{il} \equiv_{\text{in points}} S_{il} \quad (2.8)$$

$$Q_i^R(\xi_i^R, \eta_i^R, \zeta_i^R, \tau) \quad Z_i(x_i, y_i, z_i, t)$$

Case2. Theorem 2 can straightfully be proved on the basis of the reciprocity principle for isotropic and anisotropic heterogeneous media with space and time memories, so that the variable point $Q^R(\xi^R, \eta^R, \zeta^R, \tau)$ is the reciprocal constant image of the constant coordinates point $Z(x, y, z, t)$ [21–22], [16–19].

In accordance with the presented ideas we have

DEFINITION 1. The locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ represents the nearest surroundings with the dimensions $(a, b, c)t$ of the equivalent point $Z(x, y, z, t)$ of all transformations of multi-component electric fields.

COROLLARY 1. The dimensions $(a, b, c)t$ describe cuboidal dimensions “ $x = a$, $y = b$, $z = c$ ” for time “ t ”, of the cuboidal volume element.

As the extension of the problems defined above we can reflect

THEOREM 3. The locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ is phenomenologically closed, so the influences of every physical phenomenon can be treated in a separate way.

PROOF. The proof of this theorem is based on the summations of influences of the physical phenomena of the potential fields and rotational fields on the mass/charge, energy and momentum coordinates of the complete state vector of multi-component electric fields [23], [22].

Hence we can introduce consequently to the presented approach:

DEFINITION 2. The locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ has regular surface $F_{(a,b,c)t}$ externally oriented by the normal external surface orientation vector $n^{(z)}$.

Now we should define the influences of the surface conditions on the mass/charge, energy and momentum balances.

DEFINITION 3. The circulation of the normal external surface orientation vector $n^{(z)}$ has a substantial significance for the signs of summations of the influences of phenomena of the potential fields and the rotational fields, in the mass/charge, energy and momentum balances.

Definition 3 leads to the discussion:

COROLLARY 2. The effectivenesses of the influences of physical phenomena on the mass/charge, energy and momentum balances are connected to the signs: (+) for potential fields, (–) for rotational fields,

or

(–) for potential fields, (+) for rotational fields, depending on the circulation of the vector $n^{(z)}$ [24], [23], [22]. This question exists for a very wide scope of real processes not including the composite processes.

For the complicated processes we can consider

COROLLARY 3. For minimum two composite processes the following possibilities exist:

1°. The first process is/of the same kind as the second one, so that the I^p -process and II^p -process are bound with the signs:

I^p, II^p ; (+) for the potential fields and (–) for the rotational fields. As an example of this class of processes are these for which the second is the continuation of the first process or on the contrary.

2°. The composite processes are competitive in the working space volume-time of the pertinent apparatus, with the signs for the processes:

I^p ; (+) for the potential fields and (–) for the rotational fields,
and

II^p ; (–) for the potential fields and (+) for the rotational fields.

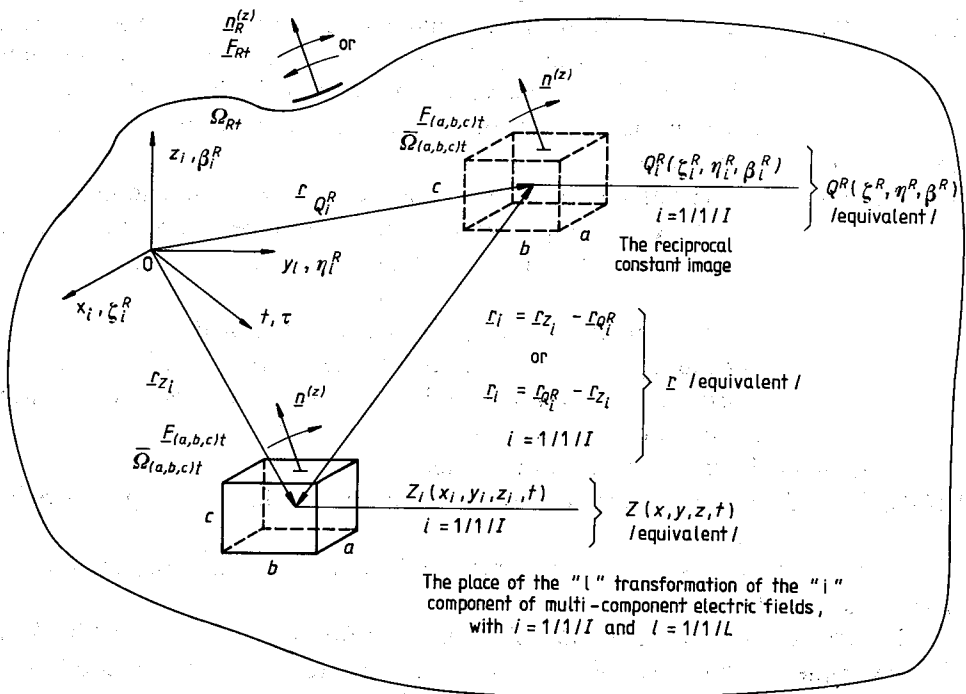
The presented above reflections can be extended to the processes having more then two composite processes under the condition that their mutual dependences are determined by physical phenomena.

The theoretical considerations are completed by the following assumptions for multicomponent electric fields:

ASSUMPTION 1. Multi-component electric fields are determined as continuous media, having a reciprocal constant image and space and time memories pertinent for all state variables [20], [1], [12–16].

ASSUMPTION 2. Assumption 1 assures the importance of the constitutive invariance for continuous media with space and time memories in the form of the continuity partial differential equations of every state variable of multicomponent electric fields [20], [16], [22], [12–15].

ASSUMPTION 3. The actual known physical phenomena of the multi-component electric fields are used for the construction of the constitutive distributed parameter mathematical model of these electric fields [1], [12–16].



Part I. Fig. 1. Interpretation of the locally distributed parameters for the multi-component electric fields. $Z_i(x_i, y_i, z_i, t)$ — point of the activity function for the "i" component and "l" transformation of multi-component electric fields,

$Z(x, y, z, t)$ — equivalent point for all components and all transformations of multi-component electric fields,

$\bar{\Omega}(a, b, c)t$ — locally selected volume-time element (phenomenologically closed — every phenomenon can be treated in a separate way) around point $Z(x, y, z, t)$,

$\vec{F}_{(a, b, c)t}$ — the externally oriented surface of the locally selected volume-time element $\bar{\Omega}(a, b, c)t$,

$\vec{n}^{(z)}$ — the normal external surface orientation vector for the surface $\vec{F}_{(a, b, c)t}$;

$(a, b, c)t$ — the dimensions of the locally selected volume-time element $\bar{\Omega}(a, b, c)t$,

$Q_i^R(\xi_i^R, \eta_i^R, \zeta_i^R, \tau)$ — the picture of the point of the activity function for the "i" component and "l" transformation from the point of view of the space-time memories for multi-component electric fields,

$Q^R(\xi^R, \eta^R, \zeta^R, \tau)$ — the picture of the equivalent point for the point $Z(x, y, z, t)$;

$\vec{r}_{Zi}$ — the vector lead radius for the point $Z_i(x_i, y_i, z_i, t)$,

$\vec{r}_{Q_i^R}$ — the vector lead radius for the point $Q_i^R(\xi_i^R, \eta_i^R, \zeta_i^R, \tau)$,

$\vec{r}_i$ — the reciprocal vector radius between points $Z_i(x_i, y_i, z_i, t)$ and $Q^R(\xi_i^R, \eta_i^R, \zeta_i^R, \tau)$,

$\vec{r}$ — the equivalent reciprocal vector radius between points $Z(x, y, z, t)$ and $Q^R(\xi^R, \eta^R, \zeta^R, \tau)$,

Ω_{Rt} — the working space volume-time of multi-component electric fields,

$\vec{F}_{Rt}$ — the externally oriented surface of the working space volume-time of electric fields Ω_{Rt} ;

$\vec{n}_R^{(n)}$ — the normal external surface orientation vector of the surface $\vec{F}_{Rt}$ (the circulation of the vector $\vec{n}_R^{(n)}$ is the same as for the vector $\vec{n}^{(z)}$ or is in opposition to the vector $\vec{n}^{(z)}$ — pertinent to the control problems by the use of boundary conditions on the surface $\vec{F}_{Rt}$)

ASSUMPTION 4. The physical phenomena mentioned in Assumption 3 belong to potential fields and rotational fields, and their influences can be treated in a separate way depending on the properties of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ whose dimensions $(a,b,c)t$ are determined by the locally distributed physical conditions of the equivalent source at the point $Z(x,y,z,t)$ [17—19].

ASSUMPTION 5. The locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ has its reciprocal constant image inside the working space volume-time of the electric apparatus [22], [17—19].

ASSUMPTION 6. For all considerations of this paper the dynamic and static forms of the reciprocity principle for isotropic and anisotropic heterogeneous media with space and time memories are valid [21].

The geometric interpretation of the above ideas is presented in figure 1.

The considerations of this chapter are consequently used as a basis for the development of a mathematical model with space-time distributed parameters of multi-component electric fields.

3. THE PRESENTATION OF KNOWN PHYSICAL ASPECTS OF THE CHARGE, ENERGY AND MOMENTUM BALANCES OF MULTI-COMPONENT ELECTRIC FIELDS

Let us consider a moving coordinate system, related to the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$ (phenomenologically closed — every phenomenon can be treated in a separate way) with dimensions $x = a, y = b, z = c$ for the time “ t ”, in the condition of a constant coordinate system being connected with multi-component electric fields. In addition let us assume the field vector velocity “ v ” of the volume-time element $\bar{\Omega}_{(x,y,z)t}$ from the aspect of constant coordinates of the working space volume-time of electric apparatuses Ω_{Rt} . Moreover, let “ i ” represent the “ i ” transformation of the “ i ” component of multi-component electric fields. With these elements of a new approach at hand, Table 1 displays the physical conditions of multi-component electric fields [1], [12—16].

For the construction of the partial differential constitutive state equations the following remarks are essential:

REMARK 1. In Table 1, it is easy to see that the charge convection flux for every “ i ” component of multi-component electric fields contains the field vector velocity “ v_{ki} ” although the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$ has a different field vector velocity “ v ”. This difference results from the mutual influence of the effects of the electric charge of every “ i ” component itself and of other components of the multi-component electric fields, especially for multi-component electrically active continuous media [1—16]. By a separate treatment independent of the

Table 1.

The presentation of the physical conditions of multi-component electric fields

Balances Phenomena	Charge	Energy	Momentum
	1	2	3
Diffusion	$j_{ei} = D_{ei} \text{grad } q_i$ $\left[\frac{A}{m^2} \right]$ and $q_i = q_{ki} + q_{wi}$ $\left[\frac{C}{m^3} \right]$ for $i = 1/I$	$I_d = \sum_{i=1}^I U_{ei} \iiint_{\Omega_{(x,y,z)}} \text{div} v / D_{ei} \cdot$ $\cdot \text{grad } q_i / d\Omega_{(x,y,z)}$ $[W]$	
Conductivity	$j_{ei} = \kappa_i \text{grad } U_i$ $\left[\frac{A}{m^2} \right]$ with $i = 1/I$	$I_c = \sum_{i=1}^I U_{ei} \iiint_{\Omega_{(x,y,z)}} \text{div} v / \kappa_i \cdot \text{grad } U_i / d\Omega_{(x,y,z)}$ $[W]$	
Electrostatics		$I_s = \frac{1}{2} p \sum_{i=1}^I U_{ei} \iiint_{\Omega_{(x,y,z)}} \text{div} v / \epsilon_i \cdot \text{grad } U_i / d\Omega_{(x,y,z)}$ for $i = 1/I$, see Appendix 1 $[W]$	$N_p = \sum_{i=1}^I q_i \text{grad } U_i \cdot d\Omega_{(x,y,z)}$ and $i = 1/I$. $[N]$
Convection and excitation	$j_{ki} = q_{ki} v_{ki}$ $\left[\frac{A}{m^2} \right]$ and $j_{wi} = q_{wi} v$ $\left[\frac{A}{m^2} \right]$ for $i = 1/I$	$i_{ei} = q_i U_i v$ $\left[\frac{W}{m^2} \right]$ $i = 1/I$, by $q_i = q_{ki} + q_{wi}$	$N_{ki} = -f_i \iiint_{\Omega_{(x,y,z)}} [v \text{div}(q_i v) + (q_i v \text{grad} v)] \cdot$ $\cdot d\Omega_{(x,y,z)}$ $i = 1/I$ $[N]$

	1	2	3
Conductor properties			$N_v = \iiint_{\Omega_{(x,y,z)t}} [\delta \nabla^2 v + \delta \text{grad} \cdot \text{div} v] d\Omega_{(x,y,z)}$ [N]
Source	$J_Q = \pm \sum_{i=1}^I \sum_{t=1}^L \iiint_{\Omega_{(x,y,z)t}} J_{it}(t) \cdot d\Omega_{(x,y,z)}$ $\cdot d\Omega_{(x,y,z)} =$ $= \pm \sum_{i=1}^I \sum_{t=1}^L \iiint_{\Omega_{(x,y,z)t}} [J_{wit}(t) + J_{pit}(t)] d\Omega_{(x,y,z)}$ [A]	$I_Q = \pm \sum_{i=1}^I \sum_{t=1}^L \iiint_{\Omega_{(x,y,z)t}} P_{it}(t) \cdot d\Omega_{(x,y,z)}$ [W]	$N_Q = \pm \sum_{i=1}^I \sum_{t=1}^L \iiint_{\Omega_{(x,y,z)t}} T_{it}(t) \cdot d\Omega_{(x,y,z)}$ [N]
Increase	$\frac{\partial}{\partial t} \iiint_{\Omega_{(x,y,z)t}} \varrho d\Omega_{(x,y,z)}$ [A]	$\frac{\partial}{\partial t} \iiint_{\Omega_{(x,y,z)t}} (\varrho U) d\Omega_{(x,y,z)}$ [W]	$\frac{\partial}{\partial t} \iiint_{\Omega_{(x,y,z)t}} (\varrho v) d\Omega_{(x,y,z)}$ [N]

source, of component "i" by its movement this problem has a peculiar importance for multi-component electric fields.

REMARK 2. The excitation flux $j_{wi} = q_{wi} v$ depends strictly on the field vector velocity "v" of the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$ for $i = 1/I$. The movement of $\bar{\Omega}_{(x,y,z)t}$ implies the excitation of "i" charge for $i = 1/I$ depending on the properties of the "i" component of the multi-component electric fields [1], [12—16].

REMARK 3. We consider the physical phenomena related to the complete form of the "i" component electric charge $q_i = q_{ki} + q_{wi}$ for $i = 1/I$ and the charge balance curves for the complete form of the multi-component electric charge $q = \sum_{i=1}^I q_i$. On the other hand, the source of the multi-component electric charge is based on its complete form. The existence of the source for every "i" component of the charge leads to the complete form of the "i" component convection:

$$j_i = j_{ki} + j_{wi} = q_{ki} v_{ki} + q_{wi} v. \quad (3.1)$$

From the definition of "i" component source charge

$$q_i = q_{ki} + q_{wi} \quad (3.2)$$

we have the complete form of the source ordered "i" component convection

$$j_i = q_i v \quad \text{for} \quad v_{ki} = v \quad \text{and} \quad i = 1/I \quad (3.3)$$

Eq. (3.3) represents the "i" component charge orientation to the "i" component charge source [1], [12—16].

Table 1 contains the charge, energy and momentum physical elements treated separately for the distributed parameter balances for the state variables of multi-component electric fields.

4. THE CONSTRUCTION OF THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATIONS DESCRIBING THE PHYSICAL CONDITIONS OF MULTI-COMPONENT ELECTRIC FIELDS

4.1. THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATION OF THE CHARGE BALANCE OF MULTI-COMPONENT ELECTRIC FIELDS

Using the charge coordinate of Table 1 with Remark 3 and the surface conditions of $F_{(x,y,z)t}$ of the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$, the dynamical balance equation for the "i" component electric charge is written as:

$$\frac{\partial}{\partial t} \iiint_{\bar{\Omega}_{(x,y,z)t}} q_i^{\text{d}\Omega}(x,y,z) = \oint_F [D_{ei} \text{grad } q_i - q_i v + \gamma_i \text{grad } U_i] dF_{(x,y,z)} \pm \sum_{l=1}^L \iiint_{\bar{\Omega}_{(x,y,z)t}} J_{il}(t) d\Omega_{(x,y,z)}. \quad (4.1.1)$$

Applying the Gauss-Ostrogradski law to eq. (4.1.1) [22—23], [16—19]

$$\frac{\partial q_i}{\partial t} = \text{grad } D_{ei} \text{ grad } q_i + D_{ei} \nabla^2 q_i - \mathbf{v} \text{ grad } q_i - q_i \text{ div } \mathbf{v} + \text{grad } \gamma_i \text{ grad } U_i + \\ + \gamma_i \nabla^2 U_i \pm \sum_{l=1}^L J_{il}(t). \quad (4.1.2)$$

Subsequently to eq. (4.1.2), the definition of the following derivative is utilized [22], [16—19], with the result

$$\frac{Dq_i}{Dt} = D_{ei} \nabla^2 q_i - q_i \text{ div } \mathbf{v} \pm \sum_{l=1}^L J_{il}(t) + \gamma \nabla^2 U_i. \quad (4.1.3)$$

From the condition for the following derivative $\frac{Dq_i}{Dt} = 0$ [22], [16—19], the partial differential equations of the “ i ” component charge continuity can be obtained:

— for $D_{ei} = D_{ei}(x, y, z, t)$ and $\gamma_i = \gamma_i(x, y, z, t)$

$$\frac{\partial q_i}{\partial t} = \text{grad } D_{ei} \text{ grad } q_i - \mathbf{v} \text{ grad } q_i + \text{grad } \gamma_i \text{ grad } U_i \quad (4.1.4)$$

— for $D_{ei} \neq D_{ei}(x, y, z, t)$ and $\gamma_i \neq \gamma_i(x, y, z, t)$

$$\frac{\partial q_i}{\partial t} = -\mathbf{v} \text{ grad } q_i. \quad (4.1.5)$$

For the realization of the programme of the complete description of multi-component electric fields it is necessary to introduce the following theorems:

THEOREM 4. Let $\mathbf{j}_{di} = D_{ei} \text{ grad } q_i$ be the “ i ” component diffusion transport of “ i ” component electric charge and $i = 1/I$ of the multi-component electric fields. Then the complete diffusion vector $\mathbf{j}_d = D_e \text{ grad } q$, where $q = \sum_{i=1}^I q_i$ and $D_e = \sum_{i=1}^I D_{ei}$.

PROOF. We can introduce the definition of the “ i ” point charges and $i = 1/I$ which fulfil relation [1], [12—15]

$$q = \sum_{i=1}^I q_i. \quad (4.1.6)$$

On the other hand, the whole form of the diffusion charge transport for this case is the vector summation

$$\mathbf{j}_d = \sum_{i=1}^I D_{ei} \text{ grad } q_i \quad (4.1.7)$$

now introducing for $q_i \geq 0$ and $D_{ei} \geq 0$ packet summation criterions [27]:

$$D_e \text{ grad } q = \sum_{i=1}^I D_{ei} \text{ grad } q_i. \quad (4.1.8)$$

From definition (4.1.6) appears [23]

$$D_e = \sum_{i=1}^I D_{ei} \quad (4.1.9)$$

As a consequence of accepting this point of view, we can suggest:

THEOREM 5. If $j_{ci} = \gamma_i \text{grad } U_i$ is the "i" component form of the conductivity phenomenon and $i = 1/I$, there exists a complete form of the conductivity phenomenon $j_c = \gamma \text{grad } U$ for which

$$U = \sum_{i=1}^I u_i \text{ and } \gamma = \sum_{i=1}^I \gamma_i.$$

PROOF. The proof is based on the introduction of the voltage " U_i " of the "i" component of multi-component electric fields, as a point voltage, so that [1], [12—15]

$$U = \sum_{i=1}^I U_i. \quad (4.1.10)$$

Introducing the complete form of the conductivity phenomenon

$$j_c = \sum_{i=1}^I \gamma_i \text{grad } U_i \quad (4.1.11)$$

and comparing elements of the equation of packet summation [27]

$$(U_i \geq 0 \text{ and } \gamma_i \geq 0) \quad \gamma \text{grad } U = \sum_{i=1}^I \gamma_i \text{grad } U \quad (4.1.12)$$

with respect to the eq. (4.1.10), one can obtain [23]

$$\gamma = \sum_{i=1}^I \gamma_i. \quad (4.1.13)$$

Making use of the above considerations the resultant form of the charge balance of the multi-component electric fields is:

$$\frac{Dq}{Dt} = D_e \nabla^2 q - q \text{div } v + \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) + \gamma \nabla^2 U \quad (4.1.14)$$

with the partial differential equations of the charge continuity from the condition

$$\frac{Dq}{Dt} = 0 \quad [16], [22], [17—19]:$$

— for $D_e = D_e(x, y, z, t)$ and $\gamma = \gamma(x, y, z, t)$

$$\frac{\partial q}{\partial t} = \text{grad } D_e \text{ grad } q - v \text{ grad } q + \text{grad } \gamma \text{ grad } U \quad (4.1.15)$$

— for $D_e \neq D_e(x, y, z, t)$ and $\gamma \neq \gamma(x, y, z, t)$

$$\frac{\partial q}{\partial t} = -v \text{ grad } q \quad (4.1.16)$$

COROLLARY 3. The presented approach may be interpreted when the volume-time element $\bar{\Omega}_{(x,y,z)t}$ driven with the field vector velocity “ $\mathbf{v}$ ” contains complete electric charge $q = \sum_{i=1}^I q_i$. We can consider this problem separately for every “ i ” component of the multi-component electric fields but for the realizing this task we must define:

- $\bar{\Omega}_{(x,y,z)t}^i$ for every “ q_i ” and $i = 1/I$,
- field vector velocity “ $\mathbf{v}_i$ ” for every “ q_i ” and $i = 1/I$ with $\mathbf{v} = \sum_{i=1}^I \mathbf{v}_i$, ($\mathbf{v}_i \geq 0$) packet summation [27] under the condition that all “ i ” components exist independently to each other and without relation to one another. From this point of view the volume-time element $\bar{\Omega}_{(x,y,z)t}$ must contain all volume-time elements $\bar{\Omega}_{(x,y,z)t}^i$ for $i = 1/I$, treated separately.

4.2. THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATIONS OF THE ENERGY BALANCE OF MULTI-COMPONENT ELECTRIC FIELDS

We use the “energy” elements of the Table 1, to deduce adequate energy balance of multi-component electric fields, in the form:

$$\begin{aligned} \frac{\partial}{\partial t} \iiint_{\bar{\Omega}_{(x,y,z)t}} (qU) d\Omega_{(x,y,z)} &= \frac{1}{2} p \sum_{i=1}^J U_{\rho_i} \iiint_{\bar{\Omega}_{(x,y,z)t}} \operatorname{div}(\varepsilon_i \operatorname{grad} U_i) d\Omega_{(x,y,z)} - \\ &- \sum_{i=1}^I \oint_{F_{(x,y,z)t}} (q_i U_i \mathbf{v}) dF_{(x,y,z)} + \sum_{i=1}^I U_{d_i} \iiint_{\bar{\Omega}_{(x,y,z)t}} \operatorname{div}(D_{ei} \operatorname{grad} q_i) d\Omega_{(x,y,z)} + \\ &+ \sum_{i=1}^I U_{ci} \iiint_{\bar{\Omega}_{(x,y,z)t}} \operatorname{div}(\gamma_i \operatorname{grad} U_i) d\Omega_{(x,y,z)} \pm \sum_{i=1}^I \sum_{l=1}^L \iiint_{\bar{\Omega}_{(x,y,z)t}} P_{il}(t) d\Omega_{(x,y,z)}. \end{aligned} \quad (4.2.1)$$

Applying the Gauss-Ostrogradski law [23], [16—19], to the above equation, it is possible to obtain

$$\begin{aligned} \frac{\partial}{\partial t} (qU) &= \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \operatorname{div}(\varepsilon_i \operatorname{grad} U_i) - \sum_{i=1}^I \operatorname{div}(q_i U_i \mathbf{v}) + \sum_{i=1}^I U_{d_i} \operatorname{div}(D_{ei} \operatorname{grad} q_i) + \\ &+ \sum_{i=1}^I U_{ci} \operatorname{div}(\gamma_i \operatorname{grad} U_i) \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t) \end{aligned} \quad (4.2.2)$$

Hence, from the partial differential equation (4.2.2) one gets

$$\begin{aligned} U \frac{\partial q}{\partial t} + q \frac{\partial U}{\partial t} &= \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \operatorname{div}(\varepsilon_i \operatorname{grad} U_i) - \sum_{i=1}^I U_i \mathbf{v} \operatorname{grad} U_i - \sum_{i=1}^I q_i \mathbf{v} \operatorname{grad} U_i - \\ &- \sum_{i=1}^I q_i U_i \operatorname{div} \mathbf{v} + \sum_{i=1}^I U_{d_i} \operatorname{div}(D_{ei} \operatorname{grad} q_i) + \sum_{i=1}^I U_{ci} \operatorname{div} \gamma_i \operatorname{grad} U_i \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t). \end{aligned} \quad (4.2.3)$$

Marking use of the partial differential equation of the continuity of the "i" component electric charge (4.1.5)

$$\frac{\partial q_i}{\partial t} = -\mathbf{v} \operatorname{grad} q_i \quad (4.2.4)$$

modified to the continuity of the whole electric charge and presented in the energy form

$$U \frac{\partial q}{\partial t} = - \sum_{i=1}^I U_i \mathbf{v} \operatorname{grad} q_i \quad (4.2.5)$$

to the eq. (4.2.3), the result is:

$$\begin{aligned} q \frac{\partial U}{\partial t} = & \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \operatorname{div}(\epsilon_i \operatorname{grad} U_i) - \sum_{i=1}^I q_i \mathbf{v} \operatorname{grad} U_i - \sum_{i=1}^I q_i U_i \operatorname{div} \mathbf{v} + \\ & + \sum_{i=1}^I U_{c_i} \operatorname{div}(\gamma_i \operatorname{grad} U_i) \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t). \end{aligned} \quad (4.2.6)$$

Now, we introduce the whole derivative of the element "(qU)"

— see Appendix 2, consequently to:

— the time operation;

$$q \frac{\partial U}{\partial t} = q \frac{\partial U}{\partial t} + \sum_{i=1}^I U_{di} \frac{\partial q_i}{\partial t} \quad (4.2.7)$$

where

$$U_{di} = \left. \frac{\partial(qU)}{\partial q_i} \right|_{U, q, \neq i} \quad (4.2.8)$$

and consequently,

— the space operation,

$$q \operatorname{grad} U = q \operatorname{grad} U + \sum_{i=1}^I U_{di} \operatorname{grad} q_i \quad (4.2.9)$$

by " U_{di} " defined with the application of the eq. (4.2.8).

The problems concerning the derivational formulae (4.2.7), (4.2.8) and (4.2.9) are presented in detail in Appendix 2.

Both sides of eq. (4.2.9) can now be multiplied by the field vector velocity " $\mathbf{v}$ ",

$$\mathbf{v} q \operatorname{grad} U = \mathbf{v} q \operatorname{grad} U + \mathbf{v} \sum_{i=1}^I U_{di} \operatorname{grad} q_i \quad (4.2.10)$$

and rewritten in the form

$$\mathbf{v} \sum_{i=1}^I q_i \operatorname{grad} U_i = \mathbf{v} \sum_{i=1}^I q_i \operatorname{grad} U_i + \mathbf{v} \sum_{i=1}^I U_{di} \operatorname{grad} q_i \quad (4.2.11)$$

The application of eq. (4.2.7) and eq. (4.2.11) to eq. (4.2.6), gives:

$$q \frac{\partial U}{\partial t} + \sum_{i=1}^I U_{di} \frac{\partial q_i}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \operatorname{div}(\epsilon_i \operatorname{grad} U_i) - \mathbf{v} \sum_{i=1}^I q_i \operatorname{grad} U_i -$$

$$\begin{aligned}
& -\nu \sum_{i=1}^I U_{di} \operatorname{grad} q_i - \sum_{i=1}^I q_i U_i \operatorname{div} \nu + \sum_{i=1}^I U_{di} \operatorname{grad} D_{ei} \operatorname{grad} q_i + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i + \\
& + \sum_{i=1}^I U_{ci} \operatorname{grad} \gamma_i \operatorname{grad} U_i + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t).
\end{aligned} \quad (4.2.12)$$

For continuing the analysis of eq. (4.2.12) the following cases are possible:

Case 1. Let us consider the partial differential equation of the continuity of the “ i ” component electric charge in the form (4.1.4)

$$\frac{\partial q_i}{\partial t} = \operatorname{grad} D_{ei} \operatorname{grad} q_i - \nu \operatorname{grad} q_i + \operatorname{grad} \gamma_i \operatorname{grad} U_i. \quad (4.2.13)$$

Multiplying both sides by the voltage of the diffusion phenomenon “ U_{di} ” and modifying it to the form

$$\sum_{i=1}^I U_{di} \frac{\partial q_i}{\partial t} = \sum_{i=1}^I U_{di} \operatorname{grad} D_{ei} \operatorname{grad} q_i - \nu \sum_{i=1}^I U_{di} \operatorname{grad} q_i + \sum_{i=1}^I U_{di} \operatorname{grad} \gamma_i \operatorname{grad} U_i. \quad (4.2.14)$$

Now we can subtract eq. (4.2.14) from eq. (4.2.12) with the result:

$$\begin{aligned}
q \frac{\partial U}{\partial t} &= \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \operatorname{div} (\varepsilon_i \operatorname{grad} U_i) - \nu \sum_{i=1}^I q_i \operatorname{grad} U_i - \sum_{i=1}^I q_i U_i \operatorname{div} \nu + \\
&+ \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i + \sum_{i=1}^I (U_{ci} - U_{di}) \operatorname{grad} \gamma_i \operatorname{grad} U_i + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t).
\end{aligned} \quad (4.2.15)$$

For eq. (4.2.15) it is necessary to prove for every “ i ” component of multi-component electric fields the relation between voltages “ U_{di} ” and “ U_{ci} ”, according to the theorem:

THEOREM 6. If the multi-component electric fields possess “ i ” component electric charges “ q_i ” related to their voltages “ U_i ” where $i = 1/I$, and if the diffusion electric charge bases on “ q_i ” consequently with the conductivity phenomenon based on the voltage “ U_i ”, the pertinent voltages “ U_{di} ” — for diffusion phenomenon and “ U_{ci} ” for the conductivity phenomenon fulfil the equation

$$U_{ci} = U_{di} \quad (\text{or} \quad |U_{di}| = |U_{ci}|) \quad (4.2.16)$$

and consequently

$$\sum_{i=1}^I (U_{ci} - U_{di}) \operatorname{grad} \gamma_i \operatorname{grad} U_i = 0. \quad (4.2.17)$$

PROOF. Let be present the whole derivative for the element (qU) of the multi-component electric fields — Appendix 2:

— for the time operation,

$$q \frac{\partial U}{\partial t} = q \frac{\partial U}{\partial t} + \sum_{i=1}^I (U_{ci} - U_{di}) \frac{\partial q_i}{\partial t} \quad (4.2.18)$$

— for the space operation,

$$\sum_{i=1}^I q_i \text{grad } U_i = \sum_{i=1}^I q_i \text{grad } U_i + \sum_{i=1}^I (U_{ci} - U_{di}) \text{grad } q_i. \quad (4.2.19)$$

Eqs. (4.2.18) and (4.2.19) are fulfilled from the condition

$$U_{ci} = U_{di}. \quad (4.2.20)$$

We note that in order to apply the conditions for the voltages “ U_{di} ” and “ U_{ci} ” from the Theorem 6 the result is:

$$\begin{aligned} \varrho \frac{\partial U}{\partial t} = & \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \text{div}(\varepsilon_i \text{grad } U_i) - \nu \sum_{i=1}^I q_i \text{grad } U_i - \sum_{i=1}^I q_i U_i \text{div } \nu + \\ & + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t). \end{aligned} \quad (4.2.21)$$

Case 2. Another way to obtain eq. (4.2.21) is the modification of eq. (4.2.13) to the form

$$\varrho \frac{\partial q_i}{\partial t} = -\nu \text{grad } q_i + \text{grad } \gamma_i \text{grad } U_i \quad (4.2.22)$$

by eliminating the term “ $\text{grad } D_{ei} \text{grad } q_i$ ” from eq. (4.2.15). Making then use of eq. (4.2.18) and eq. (4.2.19) we obtain:

$$\begin{aligned} \varrho \frac{\partial U}{\partial t} + \sum_{i=1}^I (U_{ci} - U_{di}) \frac{\partial q_i}{\partial t} = & \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \text{div}(\varepsilon_i \text{grad } U_i) - \\ & - \nu \sum_{i=1}^I q_i \text{grad } U_i - \nu \sum_{i=1}^I (U_{ci} - U_{di}) \text{grad } q_i \sum_{i=1}^I q_i U_i \text{div } \nu + \\ & + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i + \sum_{i=1}^I (U_{ci} - U_{di}) (\text{grad } \gamma_i \text{grad } U_i + \\ & + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t)). \end{aligned} \quad (4.2.23)$$

Subsequently for this case we can multiply both sides of eq. (4.2.22) by “ $(U_{ci} - U_{di})$ ” to read:

$$\begin{aligned} \sum_{i=1}^I (U_{ci} - U_{di}) \frac{\partial q_i}{\partial t} = & -\nu \sum_{i=1}^I (U_{ci} - U_{di}) \text{grad } q_i + \\ & + \sum_{i=1}^I (U_{ci} - U_{di}) \text{grad } \gamma_i \text{grad } U_i. \end{aligned} \quad (4.2.24)$$

Subtracting now eq. (4.2.24) from eq. (4.2.23) we can easily get eq. (4.2.21). The following derivative form of eq. (4.2.21) is:

$$\varrho \frac{DU}{Dt} = \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \varepsilon_i \nabla^2 U_i - \varrho U \text{div } \nu + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \pm$$

$$\pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t) + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 U_i. \quad (4.2.25)$$

For the partial differential constitutive state equation (4.2.25) exist the partial differential equations of continuity, from the condition $\varrho \frac{DU}{Dt} = 0$ [16—19], [22]

— for $D_{ei} = D_{ei}(x, y, z, t)$, $\gamma_i = \gamma_i(x, y, z, t)$ and $\varepsilon_i = \varepsilon_i(x, y, z, t)$

$$\varrho \frac{\partial U}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \text{grad} \varrho_i \text{grad} U_i - \nu \sum_{i=1}^I \varrho_i \text{grad} U_i \left[+ \sum_{i=1}^I U_{di} \text{grad} D_{ei} \cdot \text{grad} \varrho_i + \sum_{i=1}^I U_{ci} \text{grad} \gamma_i \text{grad} U_i \right]^* \quad (4.2.26)$$

[*] — included term only for the complete presentation of the “ i ” component energy continuity in the fluxes of the conductivity and diffusion for multi-component electric fields.

— for $D_{ei} \neq D_{ei}(x, y, z, t)$, $\gamma_i \neq \gamma_i(x, y, z, t)$ and $\varepsilon_i \neq \varepsilon_i(x, y, z, t)$,

$$\frac{\partial U}{\partial t} = -\nu \text{grad} U. \quad (4.2.27)$$

Case 3. Although the two above cases of the derivation of the partial differential equation (4.2.25) are equivalent and give the same result, it is necessary to present one more case due to the existence of the relation:

$$|U_{\rho_i}| = |U_{di}| = |U_{ci}| \quad \text{for } i = 1/I, \quad (4.2.28)$$

which results from the fact that every “ i ” component electric charge “ q_i ” equals unity [1], [12—16], Appendix 2.

COROLLARY 4. The considerations show the energy influences of every “ i ” component physical phenomena of the multi-component electric fields on the resulting total electric energy of these fields.

4.3. THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATION OF THE MOMENTUM BALANCE OF MULTI-COMPONENT ELECTRIC FIELDS

In order to continue the previous considerations we need to use the “momentum” elements of Table 1, so that:

$$\begin{aligned} \frac{\partial}{\partial t} \iiint_{\Omega_{(x,y,z)t}} (f \varrho \nu) d\Omega_{(x,y,z)} &= \sum_{i=1}^I \iiint_{\Omega_{(x,y,z)t}} \varrho_i \text{grad} U_i d\Omega_{(x,y,z)} - \sum_{i=1}^I f_i \iiint_{\Omega_{(x,y,z)t}} [\nu \text{div}(\varrho_i \nu) + \\ &+ (\varrho_i \nu \text{grad}) \nu] d\Omega_{(x,y,z)} + \iiint_{\Omega_{(x,y,z)t}} [\hat{a} \nabla^2 \nu + \hat{b} \text{grad} \text{div} \nu] d\Omega_{(x,y,z)} \pm \end{aligned}$$

$$\pm \sum_{i=1}^I \sum_{l=1}^L \iiint_{\Omega_{(x,y,z)}} T_{il}(t) d\Omega_{(x,y,z)} \quad (4.3.1)$$

After simple transformations with the application of the rules of summation of a series we get:

$$\begin{aligned} f\mathbf{v} \frac{\partial \varrho}{\partial t} + f\varrho \frac{\partial \mathbf{v}}{\partial t} &= \Delta \nabla^2 \mathbf{v} - f\mathbf{v} \operatorname{div}(\varrho \mathbf{v}) - f(\varrho \mathbf{v} \operatorname{grad}) \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \pm \\ &\pm \sum_{i=1}^I \sum_{l=1}^L (t) + \varrho \operatorname{grad} U \end{aligned} \quad (4.3.2)$$

and hence

$$\begin{aligned} f\mathbf{v} \frac{\partial \varrho}{\partial t} + f\varrho \frac{\partial \mathbf{v}}{\partial t} &= \Delta \nabla^2 \mathbf{v} - f\mathbf{v}(\mathbf{v} \operatorname{grad} \varrho) - f\mathbf{v} \varrho \operatorname{div} \mathbf{v} - f(\varrho \mathbf{v} \operatorname{grad}) \mathbf{v} + \\ &+ \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \pm \sum_{i=1}^I \sum_{l=1}^L T_{il}(t) + \varrho \operatorname{grad} U. \end{aligned} \quad (4.3.3)$$

Now the continuity equation of the total electric charge of multi component electric fields is taken in the form (4.1.16) multiplied on both sides by “ $f\mathbf{v}$ ” to the form

$$f\mathbf{v} \frac{\partial \varrho}{\partial t} = -f\mathbf{v}(\mathbf{v} \operatorname{grad} \varrho). \quad (4.3.4)$$

After subtraction of eq. (4.3.4) from eq. (4.3.3) we obtain

$$f\varrho \frac{\partial \mathbf{v}}{\partial t} = \Delta \nabla^2 \mathbf{v} - f\varrho \mathbf{v} \operatorname{div} \mathbf{v} - f(\varrho \mathbf{v} \operatorname{grad}) \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \pm \sum_{i=1}^I \sum_{l=1}^L T_{il}(t) + \varrho \operatorname{grad} U \quad (4.3.5)$$

Eq. (4.3.5) can be presented in the form of the following derivative

$$f\varrho \frac{D\mathbf{v}}{Dt} = \Delta \nabla^2 \mathbf{v} - f\varrho \mathbf{v} \operatorname{div} \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \pm \sum_{i=1}^I \sum_{l=1}^L T_{il}(t) + \varrho \operatorname{grad} U. \quad (4.3.6)$$

From the condition $f\varrho \frac{D\mathbf{v}}{Dt} = 0$ [16—19], [22] the partial differential equation of continuity for the momentum of the multi component electric fields is:

$$\frac{\partial \mathbf{v}}{\partial t} = -(\mathbf{v} \operatorname{grad}) \mathbf{v}. \quad (4.3.7)$$

COROLLARY 5. For the considerations of this chapter and of the whole paper the contents of the Corollary 3 concerning the possibility of the existence of the “ i ” component field vector velocity “ $\mathbf{v}_i$ ”, is valid. In consequence of this point of view the following is fulfilled:

$$\mathbf{v} = \sum_{i=1}^I \mathbf{v}_i \quad (4.3.8)$$

5. THE ANALYSIS OF THE PROPERTIES OF THE RESULTANT FORMS OF THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATIONS AND THEIR PARTIAL DIFFERENTIAL EQUATIONS OF CONTINUITY FOR MULTI-COMPONENT ELECTRIC FIELDS

The partial differential constitutive state equation of the "i" component electric charge takes the form

$$\frac{Dq_i}{Dt} = D_{ei}\nabla^2 q_i - q_i \operatorname{div} \mathbf{v} \pm \sum_{l=1}^L J_{il}(t) + \gamma_i \nabla^2 U_i \quad (1)$$

and leads to a system of partial differential constitutive state equations of the whole charge, energy and momentum balances of the multi-component electric fields:

$$\frac{Dq}{Dt} = D_e \nabla^2 q - q \operatorname{div} \mathbf{v} \pm \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) + \gamma \nabla^2 U \quad (2)$$

$$\begin{aligned} q \frac{DU}{Dt} = & \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \varepsilon_i \nabla^2 U_i - q U \operatorname{div} \mathbf{v} + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \pm \\ & \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t) + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i \end{aligned} \quad (3)$$

$$f q \frac{D\mathbf{v}}{Dt} = \hat{a} \nabla^2 \mathbf{v} - f q \mathbf{v} \operatorname{div} \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \pm \sum_{i=1}^I \sum_{l=1}^L T_{il}(t) + q \operatorname{grad} U \quad (4)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} + \mathbf{v} \operatorname{grad} - (\text{sum of adequate gradient operations}) [22], [16-19].$$

The existence of eq. (A.1) assures "i" component partial differential equations of continuity of the "i" component electric charge of multi-component electric fields:
— for $D_{ei} \neq D_{ej}(x, y, z, t)$, $\gamma_i \neq \gamma_j(x, y, z, t)$ and $\varepsilon_i \neq \varepsilon_j(x, y, z, t)$ with $D_e(x, y, z, t)$, $\gamma \neq \gamma(x, y, z, t)$

$$\frac{\partial q_i}{\partial t} = -\mathbf{v} \operatorname{grad} q_i \quad (1')$$

$$\frac{\partial q}{\partial t} = -\mathbf{v} \operatorname{grad} q \quad (2')$$

$$\text{B) } \frac{\partial U}{\partial t} = -\mathbf{v} \operatorname{grad} U \quad (3')$$

$$\frac{\partial \mathbf{v}}{\partial t} = -(\mathbf{v} \operatorname{grad}) \mathbf{v} \quad (4')$$

— for $D_{ei} = D_{ej}(x, y, z, t)$, $\gamma_i = \gamma_j(x, y, z, t)$ and $\varepsilon_i = \varepsilon_j(x, y, z, t)$ with $D_e = D_e(x, y, z, t)$, $\gamma = \gamma(x, y, z, t)$

$$\frac{\partial q_i}{\partial t} = \text{grad } D_{ei} \text{ grad } q_i - v \text{ grad } q_i \text{ grad } \gamma_i \text{ grad } U_i \quad (1'')$$

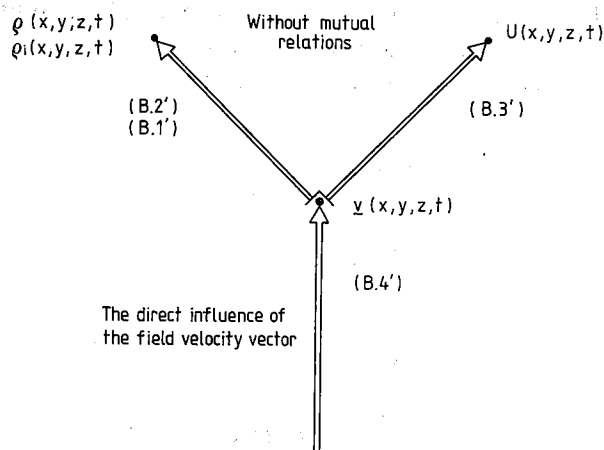
$$\frac{\partial q}{\partial t} = \text{grad } D_e \text{ grad } q - v \text{ grad } q + \text{grad } \gamma \text{ grad } U \quad (2'')$$

$$C) \quad q \frac{\partial U}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \text{ grad } \varepsilon_i \text{ grad } U_i - q v \text{ grad } U \left[+ \sum_{i=1}^I U_{di} \text{ grad } D_{ei} \cdot \right. \\ \left. \cdot \text{ grad } q_i + \sum_{i=1}^I U_{ci} \text{ grad } \gamma_i \text{ grad } U_i \right]^* \quad (3'')$$

$$\frac{\partial v}{\partial t} = -(v \text{ grad}) v \quad (4'')$$

[*] — included term only for the complete presentation of the “i” component energy continuity in the fluxes of conductivity and diffusion for the multi-component electric fields.

The classification of the set of the partial differential constitutive state equations, in the form of the following derivatives A), allows to catch its properties as for a QUASILINEAR set of partial differential equations, because all highest differential operations of this system $(\nabla^2 q_i)^{n=1}$, $(\nabla^2 q)^{n=1}$, $(\nabla^2 U_i)^{n=1}$, $(\nabla^2 U)^{n=1}$, $(\nabla^2 v)^{n=1}$, $(\text{grad div } v)^{n=1}$, have an identical power $n = 1$ [25], [22], [16—19]. The system A) is INVARIANT with respect to the physical coefficients “ D_{ei} , D_e , γ_i , γ , ε_i , $\hat{a}$, $\hat{b}$ ” and, from this point of view, the constancy of these coefficients for the properties of the system A) is proved by the partial differential equations of continuity B) or C). The interpretation of the constitutive invariance by the system B) is shown in figure 2.



Part I. Fig. 2. Interpretation of the constitutive invariance for multi-component electric fields by the system of partial differential equations of continuity B).

(B. 1') — has the form $(-v \text{ grad } q_i)$, (B. 2') — is defined as $(-v \text{ grad } q)$, (B. 3') — bases on $(-v \text{ grad } U)$, (B. 4') — assumes $[(-v \text{ grad}) v]$

6. CONCLUSIONS

The considerations of this part of the paper are an original extension of my theory suggested for distributed parameter modelling of a single component electric field [16]. The derived set of partial differential constitutive state equations, in the forms of the following derivatives A), is a novel in the world literature of phenomenological distributed parameter modelling of multi-component electric fields. The sets of partial differential equations of continuity B) or C) assure a constitutive invariance of the system A) with respect to continuity and space-time memories of all state variables of multi-component electric fields and the space-time constancy to the ground state variables of physical coefficients for the phenomena in multi-component electric fields. All real cases of multi-component electric fields have been analysed in this part of the paper by the application of the physical phenomena for the construction of the system A) adequate to the systems B) or C). Part II of this paper is an analytical solution of system A) on the basis of the application of:

- a separation of the system A) for the partial differential equations of the potential fields (diffusion, conductivity and penetrability) and consequently in the forms of the following derivatives of the rotational fields (field vector velocity), with the possibility of an analysis of the (individual solution of every phenomenological partial differential equation owing to the absence of all other ones) the idealization of the electric field — coefficients of the other phenomenological equations equal zero [1], [12—16],

- the application of the properties of the phenomenological Green functions with respect to the initial and boundary conditions,

- the splice forms of the phenomenological Green functions with their source functions.

The complete form of the solution of system A) is presented in the Part II from the point of view of the distributed parameter phenomenological control theory [26], [16—19]. For the original approach to distributed parameter modelling of the multicomponent electric fields deduced in this part of the paper a special importance should be ascribed to:

- the questions of the definition of the following derivative $\frac{D}{Dt}$ and of the defini-

tion of the partial differential equations of continuity from the condition $\frac{D}{Dt} = 0$, for every state variable of continuous media with the space and time memories analyzed by B. Średniawa in [22],

- the reciprocal constant image considerations according to the reciprocity principle for the isotropic and anisotropic heterogeneous media with space and time memories, proved by W. Nowacki [21],

- the selection of state variables for distributed parameter phenomenological modelling with regard to the yield and quality control aspects of the products of real

processes by the classic distributed parameter control, for the different physical phenomena, e.g. according to S. Węgrzyn [26].

APPENDIX 1

The conditions of the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$ with externally oriented surface $F_{(x,y,z)t}$ by the normal external surface orientation vector $\mathbf{n}^{(z)}$ of adequate circulation, imply the existence of the electrostatic energy flux

$$I_s = \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \oint_{F_{(x,y,z)}} \mathbf{D}_i d\mathbf{F}_{(x,y,z)} = \frac{1}{2} p \sum_{i=1}^I U_{\rho_i} \iiint_{\bar{\Omega}_{(x,y,z)t}} \operatorname{div}(\varepsilon_i \operatorname{grad} U_i) d\Omega_{(x,y,z)} [\text{W}] \quad (1)$$

APPENDIX 2

Let be defined the function

$$\varrho U = \Psi(U, \varrho_1, \dots, \varrho_i, \dots, \varrho_I) \quad (1)$$

and let, for every "i" component electric charge " ϱ_i " for $i = 1/I$, be defined the voltage " U_{di} ", so that

$$U_{di} = \left. \frac{\partial(\varrho U)}{\partial \varrho_i} \right|_{U, \varrho_j \neq i} \quad (2)$$

The whole derivative of the whole electric voltage is [16—19]:

— for the time coordinate

$$\varrho \frac{\partial U}{\partial t} = \varrho \frac{\partial U}{\partial t} + \sum_{i=1}^I U_{di} \frac{\partial \varrho_i}{\partial t}, \quad (3)$$

— for the space coordinates

$$\varrho \operatorname{grad} U = \varrho \operatorname{grad} U + \sum_{i=1}^I U_{di} \operatorname{grad} \varrho_i \quad (4)$$

For the constructive application of equation (4), this equation is multiplied on both sides by the field vector velocity " $\mathbf{v}$ " and modified to the form

$$\mathbf{v} \sum_{i=1}^I \varrho_i \operatorname{grad} U_i = \mathbf{v} \sum_{i=1}^I \varrho_i \operatorname{grad} U_i + \mathbf{v} \sum_{i=1}^I U_{di} \operatorname{grad} \varrho_i. \quad (5)$$

While operations

$$\left\{ \begin{array}{l} \gamma_i \operatorname{grad} U_i \\ p \varepsilon_i \operatorname{grad} U_i \end{array} \right\} \equiv D_{ei} \operatorname{grad} \varrho_i \quad \left[\frac{\text{A}}{\text{m}^2} \right] \quad (6)$$

are exactly related to the same "i" component electric charge " ϱ_i ", the voltages " U_{ρ_i} " and " U_{ci} " are defined in the same way as in eq. (2), and the equality

$$|U_{di}| = |U_{\rho_i}| = |U_{ci}| \quad \text{for } i = 1/I \quad (7)$$

is fulfilled owing to the fact that the electric charge for every "i" component of multi-component electric fields equals unity [1], [12—16].

APPENDIX 3

In my paper devoted to distributed parameter modelling of a single-component electric field [16], are established the material coefficients of media applied to the multi-component electric fields processing "a" and "b", by the formulae:

— for the conductor in fluid phase,

$$\hat{a} = \eta_w \quad \text{and} \quad \hat{b} = \frac{\eta_w}{3} \quad (1)$$

— for the conductor in crystal phase,

$$\hat{a} = \frac{h}{p} \quad \text{and} \quad \hat{b} = \frac{h\bar{k}}{p(\bar{k}-2)}. \quad (2)$$

The coefficients "a" and "b" show the conditions of the physical properties of the charge motion in different materials.

REMARK 4.

The procedure of identification of the coefficients for the physical phenomenological partial differential equations can be found in the book of S. Węgrzyn [26].

NOTATIONS

q_{ki} —density of the "i" component electric charge in the flux of convection	$\frac{C}{m^3},$
q_{wi} —density of the "i" component electric charge of the excitation phenomenon	$\frac{C}{m^3},$
q_i —total density of the "i" component electric charge $q_i = q_{ki} + q_{wi}$	$\frac{C}{m^3},$
v —the external generated field vector velocity of the electric charge in the form of ions or electrons for all components of multi-component electric fields	$\frac{m}{s},$
D_{ei} —coefficient of diffusion of the "i" component electric charge	$\frac{m^2}{s},$

γ_i —coefficient of the “ i ” component electric conductivity	$\frac{A}{Vm'}$,
ε_i —coefficient of the “ i ” component electric permeability	$\frac{C}{Vm'}$,
U —the total voltage of multi-component electric fields from a external source	V ,
U_{di} —the voltage of multi-component electric fields for the “ i ” component shift current and epidermis phenomenon	V ,
U_{ci} —the voltage of multi-component electric fields for the “ i ” component electric conductivity current	V ,
U_{ρ_i} —the voltage of the electric fields of the “ i ” component electric charge with density “ q_i ” $\frac{C}{m^3}$	V ,
D_i —the vector of the “ i ” component induction of the multi—component electric fields	$\frac{C}{m^3}$,
f_i —coefficient of the mass of the “ i ” component electric charge for multi-component electric fields	$\frac{kg}{C}$,
Ψ —the function describing the energy of multi-component electric fields	
$\hat{a}, \hat{b}$ —coefficients of material properties of the conductor of the electric current	$\frac{Ns}{m^2}$,
η_w —coefficient of the dynamic viscosity	$\frac{Ns}{m^2}$,
h —the crystal pression module	$\frac{N}{m^2}$,
$\bar{k}$ —the Poisson constant of the crystal	
p —coefficient of the time transformation of the multi-component electric fields	$\frac{1}{s}$,
$J_{ii}(t)$ —the external source of the “ i ” component electric charge undergoing to “ P ” processing, J_{wil} —for the excitation current and J_{pil} —for the source current	$\frac{A}{m^3}$,
$P_{ii}(t)$ —the external source of the “ i ” component electric energy undergoing to “ P ” processing	$\frac{W}{m^3}$,

$T_{it}(t)$ —the external source of the “ i ” component electric momentum undergoing to “ l ” processing	$\frac{N}{m^3}$,
$\bar{\Omega}_{(x,y,z)t}$ —the locally selected volume-time element (phenomenon logically closed — every phenomenon can be treated in a separate way) around point $Z(x, y, z, t)$ of the processing of multi-component electric fields	m^3-s ,
Ω_{Rt} —the working space volume-time of multi-component electric fields	m^3-s ,
$(a, b, c)t$ —the dimensions of the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$, by: $x = a, y = b, z = c$	m, m, m, s ,
(x, y, z, t) —the constant space-time coordinates of the working space volume-time of multi-component electric fields	m, m, m, s ,
$F_{(a,b,c)t}$ —the externally oriented surface of the locally selected volume—time element $\bar{\Omega}_{(a,b,c)t}/F_{(x,y,z)t}$ for $x = a, y = b, z = c$	m^2-s ,
$n^{(z)}$ —the normal external surface orientation vector of the surface $F_{(a,b,c)t}$	
F_{Rt} —the externally oriented surface of the working space volume-time of multi-component electric fields Ω_{Rt}	m^2-s ,
$n_R^{(z)}$ —the normal external surface orientation vector of the surface F_{Rt}	
t —time	s .

REFERENCES

1. A. Piekara: *Elektryczność, materia i promieniowanie*. Warszawa: PWN, 1986
2. C. Nantell: *Elektrochemia przemysłowa*. Warszawa: WNT, 1965
3. *Struktura materii, przewodnik encyklopedyczny*. Warszawa: PWN, 1980
4. M. James: *Przyszłość telekomunikacji*. Warszawa: PWN, 1975
5. T. Masewicz: *Telewizja dla praktyków*. Warszawa: WKŁ, 1982
6. T. Masewicz: *Radiotechnika dla praktyków*. Warszawa: WKŁ, 1980
7. Proceedings of IEEE Global Telecommunication Conference “GLOBECOM’85”, December 2—5, 1985, New Orleans, Louisiana, USA
8. *Encyklopedia fizyki*, t. 1 i 3, Warszawa: PWN, 1973
9. *Encyklopedia techniki*, t. “Metalurgia”, Katowice: Wydawnictwo Śląsk, 1978
10. W. Leyko: *Biofizyka dla biologów*. Warszawa: PWN, 1987
11. *Encyklopedia fizyki współczesnej*. Warszawa: PWN, 1983
12. J. Jackson: *Elektrodynamika klasyczna*. Warszawa: PWN, 1982
13. R. Sikora: *Teoria pola elektromagnetycznego*. Warszawa: WNT, 1985
14. T. Cholewicki: *Elektrotechnika teoretyczna*, t. 1 i 2, Warszawa: WNT, 1971
15. L. Landau, E. Lifszyc: *Krótki kurs fizyki teoretycznej*. Tom 1. Mechanika. Elektrodynamika. Warszawa: PWN, 1986
16. W. Niemiec: *On the constitutive distributed parameter modelling of the electric field. Part I: The construction of partial differential constitutive state equations of the electric field*. *Rozprawy elektrotechniczne*, 1986, 32, z. 4, ss. 1095—1108

17. W. Niemiec: *Mathematical model for distributed parameter control of multi-component chemical processes in fluid and solid phases*. 2nd Joint ASCE/ASME Mechanics Conference, June 23—26, 1985, University of New Mexico, Albuquerque, NM, USA, Session FWA. 2
18. W. Niemiec: *On the constitutive theory of modelling and information for distributed parameter control of the continuous mass crystallization process: Part I: Partial differential constitutive state equations describing phenomena of the continuous mass crystallization process*. Third International Conference on Liquid Metal Engineering and Technology in Energy Production, 9—13 April, 1984, Oxford, England, No. 191
19. W. Niemiec: *The constitutive theory of modelling and information for phenomenal distributed parameter control of multi-component chemical processes in gas, fluid and solid phase. Part I. The constitutive distributed parameter model of multi-component chemical processes in gas, fluid and solid phase*. 7th Miami International Conference on Alternative Energy Sources, 9—11 December, Miami Beach, Florida, USA, Hydrocarbons/Energy Transfer, pp. 579—588
20. A. Morgan: *On the construction of constitutive equations for continuous media*. *Archiwum Mechaniki Stosowanej*, vol. 17, 1965, No. 1, pp. 145—174
21. W. Nowacki: *Dynamiczne zagadnienia termosprężystości*. Warszawa: PWN, 1966
22. B. Średniawa: *Hydrodynamika i teoria sprężystości*. Warszawa: PWN, 1977
23. T. Trajdos: *Matematyka dla inżynierów*. Warszawa: PWN, 1974
24. W. Smirnow: *Matematyka wyższa. t. 2*, Warszawa: PWN, 1960
25. C. Coulson, A. Jeffrey: *Fale modele matematyczne*. Warszawa: WNT, 1982
26. S. Węgrzyn: *Podstawy automatyki*. Warszawa: PWN, 1974
27. L. Schwartz: *Metody matematyczne w fizyce*. Warszawa: PWN, 1984

W. NIEMIEC

O KONSTYTUTYWNEJ TEORII ROZŁOŻONEGO PARAMETRYCZNIE MODELOWANIA DLA ZJAWISKOWEGO STEROWANIA WIELOSKŁADNIKOWYCH PÓL ELEKTRYCZNYCH

CZĘŚĆ I

RÓWNANIA RÓŻNICZKOWE CZĄSTKOWE KONSTYTUTYWNE STANU
OPISUJĄCE WIELOSKŁADNIKOWE POLA ELEKTRYCZNE

Streszczenie

Na bazie elementów źródłowych i zjawiskowych wieloskładnikowych pól elektrycznych odpowiadających bilansom: ładunków elektrycznych, energii elektrycznej i pędu ładunków elektrycznych wyprowadzono oryginalny, konstytutywny, rozłożony parametrycznie matematyczny model tych pól elektrycznych w przestrzeni roboczej i czasie dla aparatów elektrycznych. Ujmuje on tak wydajnościowe jak i jakościowe aspekty przetwarzania wieloskładnikowych pól elektrycznych, odpowiadającym poszczególnym zmiennym stanu dla każdego składnika i całości wieloskładnikowego pola elektrycznego. Wprowadzono bilansowe równania różniczkowe cząstkowe konstytutywne stanu: ładunków elektrycznych, energii elektrycznej i pędu ładunków elektrycznych dla odpowiednich zmiennych stanu każdego składnika i całości wieloskładnikowego pola elektrycznego. Wykazano inwariancję konstytutywną wyprowadzonego układu równań różniczkowych cząstkowych konstytutywnych stanu w stosunku do jego fizycznych współczynników dla udowodnienia ich stałości do następnych rozważań. W tym celu wyprowadzono i przeanalizowano równania różniczkowe cząstkowe ciągłości: ładunków elektrycznych, energii elektrycznej i pędu ładunków elektrycznych dla zmiennych i stałych współczynników fizycznych jako grupy transformacji ciągłości odpowiednich zmiennych stanu.

On the constitutive theory of the distributed parameter modelling for the phenomenal control of the multi-component electric fields

WACŁAW NIEMIEC

Instytut Elektroniki, Politechnika Śląska w Gliwicach

Received 1989. 10. 12

Authorized 1989. 12. 28

This part of the paper is devoted to the analytical solution of the constitutive mathematical model with distributed parameters of multi-component electric fields (developed in Part I [1]), from the point of view of the application of this solution to control the yield and quality of the very wide scope of the real processes. For the realization of this task the following ideas are of great importance:

- the possibility of separation of the mathematical model for influences of single physical phenomena of the multicomponent electric fields [4—5], [10—12],
- the properties of the elements of phenomenological solutions; the phenomenological Green functions and source functions [4—6], [10—12]
- the concept of summation of the effects of the activities of physical phenomena of potential fields and those of rotational fields [4—6], [10—12],
- the possibility of realization of the phenomenological distributed parameter control by the use of the physical phenomena of multi-component electric fields [10—12], [17—19].

The presented above elements have a substantial significance for the considerations covered in the frames of this paper.

PART II.

Physical yield and quality control problems for the processing of the multi-component electric fields

1. THE DISCUSSION OF THE CONSTITUTIVE MATHEMATICAL MODEL OF THE DISTRIBUTED PARAMETERS OF THE MULTI-COMPONENT ELECTRIC FIELDS

We shall consider the influences of separate physical phenomena from the system of partial differential constitutive state equations, in the form of the following derivatives of the multi-component electric fields signed A) in the Part I [1]. The system A) in Part I is invariant with respect to its physical phenomenological coefficients, according to the sets of the partial differential equations of continuity B) or C). The partial differential constitutive state equations of multi-component charge, energy and momentum balances A) are quasilinear in the sense of the theory of the partial differential equations [2], and can be treated as having, in addition, constant coefficients on the basis of the introduction of the system B) or C) in [1]. As an introduction point we need to quote the complete form of the partial differential constitutive state equations A) for multi-component electric fields:

— for the "i" component electric charge

$$\frac{Dq_i}{Dt} = D_{ei} \nabla^2 q_i - q_i \operatorname{div} \mathbf{v} \pm \sum_{l=1}^L J_{il}(t) + \gamma_i \nabla^2 U_i \quad (1)$$

— for the total charge, energy and momentum

$$\frac{Dq}{Dt} = D_e \nabla^2 q - q \operatorname{div} \mathbf{v} \pm \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) + \gamma_i \nabla^2 U \quad (2)$$

$$\begin{aligned} (I) \quad q \frac{DU}{Dt} = & \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i - q U \operatorname{div} \mathbf{v} + \sum_{i=1}^I U_{ei} \gamma_i \nabla^2 U_i \pm \\ & \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t) + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i \end{aligned} \quad (3)$$

$$f q \frac{D\mathbf{v}}{Dt} = \hat{a} \nabla^2 \mathbf{v} - f q \mathbf{v} \operatorname{div} \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \pm \sum_{i=1}^I \sum_{l=1}^L T_{il}(t) + q \operatorname{grad} U, \quad (4)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} + \mathbf{v} \operatorname{grad} \quad \text{— (sum of adequate gradient operations) [1], [4—5], [8], [10—12].}$$

The invariance of the system (I) is involved with the movement of the volume-time element $\bar{\Omega}_{(a,b,c)t}$ with the field vector velocity "v" and its reciprocal constant image, according to the reciprocity principle for isotropic and anisotropic heterogeneous media with space and time memories [3—4]. From this point of view, the coefficients D_{ei} , D_e , γ_i , γ , ε_i , $\hat{a}$, $\hat{b}$ are assumed to be constant in the subsequent considerations. Another question is to stress quasilinear properties of the system (I) as all its highest differential operations and their functions are of the same first order " $n = 1$ ", so that $(\nabla^2 \varrho_i)^{n=1}$, $(\nabla^2 U)^{n=1}$, $(\nabla^2 \varrho)^{n=1}$, $(\nabla^2 U)^{n=1}$, $(\nabla^2 v)^{n=1}$, $(\text{grad div } v)^{n=1}$ [2], [1], [4—5], [10—12]. From the properties of the system (I) appears the possibility of its solution in an analytical way. Two cases of the analytical solution of the system (I) are considered:

A — the existence of the initial conditions,

B — the existence of the initial and boundary conditions,
for the multi-component electric fields.

The system (I) contains all phenomenological and transformation elements of multi-component electric fields, constructed in the constitutive form [13—16].

2. THE SOLUTION OF THE SYSTEM OF THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATIONS DESCRIBING MULTI-COMPONENT ELECTRIC FIELDS

A. The existence of the initial conditions for multi-component electric fields.

A. The dynamical characteristics of the working space-time of the equipment for multi-component electric fields.

2.1. THE SOURCE FUNCTIONS OF MULTI-COMPONENT ELECTRIC FIELDS

For the equivalent point $Z(x, y, z, t)$ of the processing of the multi-component electric fields (defined in the Part I [1]), the system (I) is modified to the form

$$\frac{d\varrho_i}{dt} = \pm \sum_{l=1}^L J_{il}(t) \quad (1)$$

$$\frac{d\varrho_i}{dt} = \pm \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) \quad (2)$$

$$(II) \quad \varrho \frac{dU}{dt} = \pm \sum_{i=1}^I \sum_{l=1}^L P_{il}(t) \quad (3)$$

$$f\varrho \frac{dv}{dt} = \pm \sum_{i=1}^I \sum_{l=1}^L T_{il}(t). \quad (4)$$

The integration of the system (II) done in the time interval $[0, t)$ with the results of the pertinent differential equations:

$$q_i^z(t) = q_{0i} \pm \int_0^t \sum_{l=1}^L J_{il}(t) dt \quad (1)$$

$$q^z(t) = q_0 \pm \int_0^t \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) dt \quad (2)$$

(III)

$$U^z(t) = U_0 \pm \int_0^t \sum_{i=1}^I \sum_{l=1}^L \frac{P_{il}(t)}{|q_i^z(t)|} dt \quad (3)$$

$$v^z(t) = v_0 \pm \int_0^t \sum_{i=1}^I \sum_{l=1}^L \frac{T_{il}(t)}{|q_i^z(t)|} dt. \quad (4)$$

To include the effects of the integrations of eq. (III.3) and eq. (III.4) the modulus of eq. (III.1) is used to consider (4) the possibility of inverse transformations of the multi-component electric fields at every point $Z(x, y, z, t)$.

2.2. THE ANALYSIS OF THE HOMOGENEOUS FORM OF THE SYSTEM (I)

In accordance to the chapter A.1., the homogeneous form of the system of the partial differential constitutive state equations, in the forms of the following derivatives of the multipartial electric fields (I) is

$$\frac{Dq_i}{Dt} = D_{ei} \nabla^2 q_i - q_i \operatorname{div} v + \gamma_i \nabla^2 U_i \quad (1)$$

$$\frac{Dq}{Dt} = D_e \nabla^2 q - q \operatorname{div} v + \gamma \nabla^2 U \quad (2)$$

$$\begin{aligned} (IV) \quad q \frac{DU}{Dt} = & \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i - q U \operatorname{div} v + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i + \\ & + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i \end{aligned} \quad (3)$$

$$f_q \frac{Dv}{Dt} = \hat{a} \nabla^2 v - f_q v \operatorname{div} v + \hat{b} \operatorname{grad} \operatorname{div} v + q \operatorname{grad} U \quad (4)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} + v \operatorname{grad} \quad \text{— (sum of adequate gradient operations) [1], [4–5], [10–12].}$$

Now we should consider the mentioned previously constitutive invariance for the homogeneous form of the system (I), having the form (IV). This results from the existence of the partial differential equations of continuity for every state variable of multi-component electric fields, from the general condition $\frac{D}{Dt} = 0$ [4–5], [7],

[10—12], [8]. Two cases of the constitutive invariance for multi-component electric fields have been considered

— for $D_{ei} \neq D_{ei}(x, y, z, t)$, $\gamma_i \neq \gamma_i(x, y, z, t)$ and $\varepsilon_i \neq \varepsilon_i(x, y, z, t)$, with $D_e \neq D_e(x, y, z, t)$, $\gamma \neq \gamma(x, y, z, t)$

$$\frac{\partial \varrho_i}{\partial t} = -\mathbf{v} \text{grad} \varrho_i \quad (1')$$

$$\frac{\partial \varrho}{\partial t} = -\mathbf{v} \text{grad} \varrho \quad (2')$$

$$(V) \quad \frac{\partial U}{\partial t} = -\mathbf{v} \text{grad} U \quad (3')$$

$$\frac{\partial \mathbf{v}}{\partial t} = -(\mathbf{v} \text{grad}) \mathbf{v} \quad (4')$$

as the system B) in [1].

— for $D_{ei} = D_{ei}(x, y, z, t)$, $\gamma_i = \gamma_i(x, y, z, t)$ and $\varepsilon_i = \varepsilon_i(x, y, z, t)$, with $D_e = D_e(x, y, z, t)$, $\gamma = \gamma(x, y, z, t)$

$$\frac{\partial \varrho_i}{\partial t} = \text{grad} D_{ei} \text{grad} \varrho_i - \mathbf{v} \text{grad} \varrho_i + \text{grad} \gamma_i \text{grad} U_i \quad (1'')$$

$$\frac{\partial \varrho}{\partial t} = \text{grad} D_e \text{grad} \varrho - \mathbf{v} \text{grad} \varrho + \text{grad} \gamma \text{grad} U \quad (2'')$$

$$(VI) \quad \varrho \frac{\partial U}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \text{grad} \varepsilon_i \text{grad} U_i - \varrho \mathbf{v} \text{grad} U \left[+ \sum_{i=1}^I U_{di} \text{grad} D_{ei} \cdot \right. \\ \left. \cdot \text{grad} \varrho_i + \sum_{i=1}^I U_{ci} \text{grad} \kappa_i \text{grad} U_i \right]^* \quad (3'')$$

$$\frac{\partial \mathbf{v}}{\partial t} = -(\mathbf{v} \text{grad}) \mathbf{v} \quad (4'')$$

being representation of the system C) in [1].

In the system (VI) the part of eq. (VI.3'') inside brackets $[\]^*$ is included only for the complete presentation of the "i" component energy continuity in the fluxes of conductivity and diffusion for the multi-component electric fields.

The invariance of system (IV) with regard to the systems (V) and (VI) assures the constancy of the physical coefficients of system (IV). This constancy is of great importance for the classification of single phenomenological partial differential equations separated from system (IV). For the realization of this separation we need to introduce for system (IV) the definition of rotational fields [4—6], [8—12]: exits: A — the vector potential of the flow field and $\mathbf{v} = \text{rot} A$ with $\text{div} \text{rot} A \equiv 0$.

In consequence of this point of view, we can obtain:

$$\frac{Dq_i}{Dt} = D_{ei}\nabla^2 q_i + \gamma_i \nabla^2 U_i \quad (1)$$

$$(VII) \quad \frac{Dq}{Dt} = D_e \nabla^2 q + \gamma \nabla^2 U \quad (2)$$

$$e \frac{DU}{Dt} = \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i \quad (3)$$

$$f \frac{Dv}{Dt} = \text{grad} U. \quad (4)$$

For eq. (VII.4) the mathematical rule [4—6], [8—12] is used:

$$\text{rot rot} v = \text{grad div} v - \nabla^2 v \quad (5)$$

and the system (VII) takes its **potential field from**

$$\frac{\partial q_i}{\partial t} = D_{ei} \nabla^2 q_i + \gamma_i \nabla^2 U_i \quad (1)$$

$$(VIII) \quad \frac{\partial q}{\partial t} = D_e \nabla^2 q + \gamma \nabla^2 U \quad (2)$$

$$e \frac{\partial U}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i \quad (3)$$

$$f \frac{\partial v}{\partial t} = \text{grad} U. \quad (4)$$

According to this approach we have the set of partial differential equations, in the form of the following derivatives of **the rotational fields**

$$\frac{Dq_i}{Dt} = -q_i \text{div} v \quad (1)$$

$$(IX) \quad \frac{Dq}{Dt} = -q \text{div} v \quad (2)$$

$$e \frac{DU}{Dt} = -q U \text{div} v \quad (3)$$

$$f e \frac{Dv}{Dt} = a \nabla^2 v - f q v \text{div} v + \hat{b} \text{grad div} v \quad (4)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} + v \text{grad}, \quad [4—5], [8], 10—12].$$

3. THE SEPARATION OF THE PHENOMENOLOGICAL PARTIAL DIFFERENTIAL EQUATIONS FROM THE SYSTEM OF THE PARTIAL DIFFERENTIAL EQUATIONS OF THE POTENTIAL FIELDS (VIII).

As a first method of the separation of a single phenomenological partial differential equations from the system (VIII) the concept of **idealization of the media** can be introduced [5], [4], [9—16]. The principle of this approach is the successive consideration of the system (VIII) in the forms:

— without conductivity and penetrability phenomena (the physical coefficients $\gamma_i = 0$, $\gamma = 0$, $\varepsilon_i = 0$ independently to their constancy or variation with the space-time coordinates for the constitutive invariance) to obtain

$$\frac{\partial q_i}{\partial t} = D_{ei} \nabla^2 q_i \quad (1)$$

$$\frac{\partial q}{\partial t} = D_e \nabla^2 q \quad (2)$$

(VIII¹)

$$\rho \frac{\partial U}{\partial t} = \sum_{i=1}^I U_{di} D_{ei} \nabla^2 q_i \quad (3)$$

$$f \frac{\partial v}{\partial t} = \text{grad} U, \quad (4)$$

— without diffusion and conductivity phenomena (for the physical coefficients $D_{ei} = 0$, $D_e = 0$, $\gamma_i = 0$, $\gamma = 0$ independent to the constant or variable form to their space-time coordinates by the constitutive invariance) to assure

$$\frac{\partial q_i}{\partial t} = 0 \quad (1)$$

$$\frac{\partial q}{\partial t} = 0 \quad (2)$$

(VIII²)

$$\rho \frac{\partial U}{\partial t} = \frac{1}{2} \rho \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i \quad (3)$$

$$f \frac{\partial v}{\partial t} = \text{grad} U \quad (4)$$

— without diffusion and penetrability phenomena (the physical coefficients $D_{ei} = 0$, $D_e = 0$, $\varepsilon_i = 0$ independently to their constancy or variation with the space-time coordinates for the constitutive invariance) with the result

$$\frac{\partial q_i}{\partial t} = \gamma_i \nabla^2 U_i \quad (1)$$

$$(VIII^3) \quad \frac{\partial \varrho}{\partial t} = \gamma \nabla^2 U \quad (2)$$

$$\varrho \frac{\partial U}{\partial t} = \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \quad (3)$$

$$f \frac{\partial v}{\partial t} = \text{grad} U. \quad (4)$$

Another way of to obtain the systems (VIII¹), (VIII²) and (VIII³) is involved with the application of the theories deduced in Part I [1]. Let eq. (VIII.1) be multiplied on both sides by the voltage " U_{di} " and presented in the form

$$\sum_{i=1}^I U_{di} \frac{\partial \varrho_i}{\partial t} = \sum_{i=1}^I U_{di} D_{ei} \nabla^2 \varrho_i + \sum_{i=1}^I U_{di} \gamma_i \nabla^2 U_i. \quad (5)$$

According to the Appendix 2 in Part I [1] we have the form of the energy partial differential constitutive state equation

$$\varrho \frac{\partial U}{\partial t} + \sum_{i=1}^I U_{di} \frac{\partial \varrho_i}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i + \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 \varrho_i. \quad (6)$$

Now we need to subtract eq. (VIII.5) from eq. (VIII.6) with the result

$$\varrho \frac{\partial U}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i + \sum_{i=1}^I (U_{ci} - U_{di}) \gamma_i \nabla^2 U_i. \quad (7)$$

From the condition (Case 1) proved in the Part I [1]

$$U_{ci} = U_{di} \quad (8)$$

eq. (VIII.7) assumes the form

$$\varrho \frac{\partial U}{\partial t} = \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i \quad (9)$$

consequently to the system (VIII²). After the elimination of the partial differential equation of the permeability we have

$$\varrho \frac{\partial U}{\partial t} + \sum_{i=1}^I U_{di} \frac{\partial \varrho_i}{\partial t} = \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i + \sum_{i=1}^I U_{di} D_{ei} \nabla^2 \varrho_i. \quad (10)$$

Eq. (10) can be presented as the summation of the partial differential equations

$$\varrho \frac{\partial U}{\partial t} = \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i \quad (\text{for system (VIII}^3)) \quad (11)$$

and

$$\sum_{i=1}^I U_{di} \frac{\partial \varrho_i}{\partial t} = \sum_{i=1}^I U_{di} D_{ei} \nabla^2 \varrho_i \quad (\text{for system (VIII}^1)). \quad (12)$$

According to the constitutive invariance of the set of partial differential constitutive state equations (IV), with the partial differential equations of the continuity of the system (V) or the system (VI) this system is invariant in respect to its physical coefficients (SYSTEM (IV) INVARIANT (with respect to coefficients D_{ei} , D_e , γ_i , γ , ε_i , $\hat{a}$, $\hat{b}$).

From the above point of view, these coefficients are assumed to be constant in the below chapters.

4. THE ANALYTICAL SOLUTIONS OF THE OF THE SYSTEMS (VIII) AND (IX)

4.1. THE PARTIAL DIFFERENTIAL EQUATIONS OF THE POTENTIAL FIELDS

(variable point $Q^R(\xi^R, \eta^R, \varepsilon^R \tau) \equiv Q(\xi, \eta, \varepsilon, \tau)$ [1].

4.1.1. THE DIFFUSION "i" COMPONENT CHARGE TRANSPORT FOR THE MULTI-COMPONENT ELECTRIC FIELDS

The diffusion "i" component partial differential equation of its charge transport can be rewritten as

$$\frac{\partial \varrho_i}{\partial t} - D_{ei} \nabla^2 \varrho_i = 0. \quad (1)$$

For the solution of eq. (1), we introduce the relation

$$\varrho(x, y, z, t) = C(x, y, z) \bar{\varrho}_i(t). \quad (2)$$

Eq. (2) is valid under the condition of the assertion:

THEOREM 1. If the "i" component electric charge, with $i = 1/I$ for the diffusion transport can be defined in relation to the dimensions "a, b, c" of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ formula (2) reaches a great importance.

PROOF. Let be introduced an other from of the "i" component charge density

$$\varrho_i(x, y, z, t) = C_i(x, y, z) \bar{\varrho}_i(t) \quad (3)$$

Eq. (3) delivers the consequent from of the eq. (1).

$$\frac{\partial^2 C_i}{\partial x^2} + \frac{\partial^2 C_i}{\partial y^2} + \frac{\partial^2 C_i}{\partial z^2} + \lambda_i C_i = 0 \quad (4)$$

$$\frac{d\bar{\varrho}_i}{dt} = -\lambda_i D_{ei} \bar{\varrho}_i \quad (5)$$

The use of the proof of superposition of the potential fields and the rotational fields with the reciprocal constant image imposes the existence of an adequate volume-time element $\bar{\Omega}_{(a,b,c)t}^i$ with the boundary conditions on its boundary equalling zero [5—6], [8], [10—12].

$$\begin{aligned} C_i(0, y, 0) &= 0 & C_i(a_i, y, 0) &= 0 \\ C_i(x, 0, 0) &= 0 & \text{and} & C_i(x, b_i, 0) = 0 \\ C_i(0, 0, z) &= 0 & C_i(0, y, c_i) &= 0. \end{aligned} \quad (6)$$

For the presentation of the function

$$C_i(x, y, z) = X_i(x) Y_i(y) Z_i(z) \quad (7)$$

eq. (4) assumes the system form

$$\begin{aligned} X_i'' + \Theta_i X_i &= 0 & X_i(0) &= 0 & X_i(a_i) &= 0 \\ Y_i'' + \mu_i Y_i &= 0 & Y_i(0) &= 0 & \text{and} & Y_i(b_i) = 0 \\ Z_i'' + \gamma_i Z_i &= 0 & Z_i(0) &= 0 & Z_i(c_i) &= 0. \end{aligned} \quad (8)$$

The solutions of system (8) have the forms

$$\begin{aligned} X_n(x) &= \sin \frac{n\pi x}{a_i}; & \Theta_i &= \left(\frac{n\pi}{a_i} \right)^2 \\ Y_m(y) &= \sin \frac{m\pi y}{b_i}; & \mu_i &= \left(\frac{m\pi}{b_i} \right)^2 \\ Z_k(z) &= \sin \frac{k\pi z}{c_i}; & \gamma_i &= \left(\frac{k\pi}{c_i} \right)^2 \end{aligned} \quad (9)$$

and the eigenvalues have the form

$$\lambda_{i,n,m,k} = \left(\frac{n\pi}{a_i} \right)^2 + \left(\frac{m\pi}{b_i} \right)^2 + \left(\frac{k\pi}{c_i} \right)^2. \quad (10)$$

From the general assumption of the existence of **only a single volume-time element** $\bar{\Omega}_{(a,b,c)t}$, **being valid for all components of the multi-component electric charges**, it arises

$$a_i = a, \quad b_i = b, \quad \text{and} \quad c_i = c \quad (11)$$

so that

$$C_i(x, y, z) = C(x, y, z). \quad (12)$$

For Theorem 1, it is necessary to lead the following assertion: **COROLLARY 1.** While Theorem 1 is valid for all components "i" and $i = 1/I$ of the state quantities nad all physical phenomena of multi-component electric fields, the dimensions $(a, b, c)t$ of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ (phenomenologically closed — every phenomenon can be treated in a separate way), the eigenvalues and eigenfunctions are based on these dimensions and have the same formulae for all physical phenomena of these electric fields.

Theorem 1, in conformity with eq. (2) and eq. (3), gives their interpretations as follows

$$\frac{\partial^2 C}{\partial x^2} + \frac{\partial^2 C}{\partial y^2} + \frac{\partial^2 C}{\partial z^2} + \lambda C = 0 \quad (13)$$

$$\frac{d\bar{q}_i}{dt} = -\lambda D_{ei}\bar{q}_i. \quad (14)$$

It is conceivable that, subsequently, the following mathematical steps are covered by Theorem 1;

— the boundary conditions on the boundary of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ are equal to zero [5—6], [8], [10—12]

$$\begin{aligned} C(0, y, 0) &= 0 & C(a, y, 0) &= 0 \\ C(x, 0, 0) &= 0 & \text{and} & C(x, b, 0) = 0 \\ C(0, 0, z) &= 0 & C(0, y, c) &= 0 \end{aligned} \quad (15)$$

— the modification of the function $C(x, y, z)$ to the form

$$C(x, y, z) = X(x)Y(y)Z(z) \quad (16)$$

which allows for the presentation of eq. (13) in another way

$$\begin{aligned} X'' + \Theta X &= 0 & X(0) &= 0 & X(a) &= 0 \\ Y'' + \mu Y &= 0 & Y(0) &= 0 & \text{and} & Y(b) = 0 \\ Z'' + \gamma Z &= 0 & Z(0) &= 0 & Z(c) &= 0. \end{aligned} \quad (17)$$

This approach causes that the solutions of the above system of differential equations have the forms

$$\begin{aligned} X_n(x) &= \sin \frac{n\Pi x}{a}; & \Theta &= \left(\frac{n\Pi}{a}\right)^2 \\ Y_m(y) &= \sin \frac{m\Pi y}{b}; & \mu &= \left(\frac{m\Pi}{b}\right)^2 \\ Z_k(z) &= \sin \frac{k\Pi z}{c}; & \gamma &= \left(\frac{k\Pi}{c}\right)^2. \end{aligned} \quad (18)$$

For the above solutions the eigenvalues have the presentation

$$\lambda_{n,m,k} = \left(\frac{n\Pi}{a}\right)^2 + \left(\frac{m\Pi}{b}\right)^2 + \left(\frac{k\Pi}{c}\right)^2 \quad (19)$$

and consequently the eigenfunctions satisfy the equation

$$C_{n,m,k}(x, y, z) = A_{n,m,k} \sin \frac{n\Pi x}{a} \sin \frac{m\Pi y}{b} \sin \frac{k\Pi z}{c}. \quad (20)$$

It is possible to perceive that in eq. (20) $A_{n,m,k}$ can be considered from the point of view of the dimensions "a, b, c" being characteristic for the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$, by the application of the norm N_e of the eigenfunctions

$$N_e = \int_0^a \int_0^b \int_0^c C_{n,m,k}^2(x, y, z) dx dy dz = A_{n,m,k}^2 \int_0^a \sin^2 \frac{n\pi x}{a} dx \int_0^b \sin^2 \frac{m\pi y}{b} dy \cdot \int_0^c \sin^2 \frac{k\pi z}{c} dz = 1 \quad (21)$$

with the result of the integration

$$A_{n,m,k} = \sqrt{\frac{8}{abc}}. \quad (22)$$

In consequence of this approach one gets eigenfunctions related to the point $Z(x, y, z, t)$ of the volume-time element $\bar{\Omega}_{(a,b,c)t}$

$$C_{n,m,k}(x, y, z) = \sqrt{\frac{8}{abc}} \sin \frac{n\pi x}{a} \sin \frac{m\pi y}{b} \sin \frac{k\pi z}{c}. \quad (23)$$

For the variable point $Q(\xi, \eta, \beta, \tau)$ the eigenfunctions are written as

$$C_{n,m,k}(\xi, \eta, \beta) = \sqrt{\frac{8}{abc}} \sin \frac{n\pi \xi}{a} \sin \frac{m\pi \eta}{b} \sin \frac{k\pi \beta}{c}. \quad (24)$$

The integration of eq. (14) leads to the formula

$$\bar{q}_i(t) = D_{n,m,k}^i e^{-\lambda_{n,m,k} D_{ei} t} \quad (25)$$

The integration constant for eq. (25) is

$$D_{n,m,k}^i = \iiint_{\Omega_R} q_{0i} \sqrt{\frac{8}{abc}} \sin \frac{n\pi \xi}{a} \sin \frac{m\pi \eta}{b} \sin \frac{k\pi \beta}{c} d\xi d\eta d\beta. \quad (26)$$

From the integration over the time interval $[0, t]$ (of eq. (II.1)), the source function for the "i" component electric charge takes the form

$$q_i^z(t) = q_{0i} \pm \int_0^t \sum_{l=1}^L J_{il}(t) dt. \quad (27)$$

The resultant form of the "i" component diffusion charge transport, from the variable point $Q(\xi, \eta, \beta, \tau)$ to the constant coordinates point $Z(x, y, z, t)$ can be written as

$$q_i(Z, t) = \int_0^t \iiint_{\Omega_R} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^i e^{-\lambda_{n,m,k} D_{ei}(t-\tau)} \frac{8}{abc} \left(\sin \frac{n\pi x}{a} \sin \frac{n\pi \xi}{a} \right) \cdot \left(\sin \frac{m\pi y}{b} \sin \frac{m\pi \eta}{b} \right) \left(\sin \frac{k\pi z}{c} \sin \frac{k\pi \beta}{c} \right) \right\} \left[q_{0i} \pm \int_0^{t-\tau} \sum_{l=1}^L J_{il}(t) dt \right] d\xi d\eta d\beta d\tau \quad (28)$$

4.1.2. THE DIFFUSION TRANSPORT OF THE TOTAL ELECTRIC CHARGE OF THE MULTI-COMPONENT ELECTRIC FIELDS

Let us consider eq. (VIII¹.2), describing our question

$$\frac{\partial \varrho}{\partial t} - D_e \nabla^2 \varrho = 0 \quad (1)$$

For the solution of this problem an important role is played by the considerations of the chapter A.4.1.1. In terms of the total electric charge of multi-component electric fields $\varrho(x, y, z, t) = C(x, y, z) \bar{\varrho}(t)$, eq. (1) appears in the form

$$\frac{\partial^2 C}{\partial x^2} + \frac{\partial^2 C}{\partial y^2} + \frac{\partial^2 C}{\partial z^2} + \lambda C = 0 \quad (2)$$

$$\frac{d\bar{\varrho}}{dt} = -\lambda D_e \bar{\varrho}. \quad (3)$$

After the same considerations for eq. (2) as in the chapter A.4.1.1., by this way we obtain the eigenvalues in the form of

$$\lambda_{n,m,k} = \left(\frac{n\pi}{a} \right)^2 + \left(\frac{m\pi}{b} \right)^2 + \left(\frac{k\pi}{c} \right)^2 \quad (4)$$

and the eigenfunctions for the point $Z(x, y, z, t)$ are

$$C_{n,m,k}(x, y, z) = \sqrt{\frac{8}{abc}} \sin \frac{n\pi x}{a} \sin \frac{m\pi y}{b} \sin \frac{k\pi z}{c}. \quad (5)$$

The eigenfunctions related to the variable point $Q(\xi, \eta, \beta, \tau)$ lead to the formula

$$C_{n,m,k}(\xi, \eta, \beta) = \sqrt{\frac{8}{abc}} \sin \frac{n\pi \xi}{a} \sin \frac{m\pi \eta}{b} \sin \frac{k\pi \beta}{c}. \quad (6)$$

The solution of eq. (3) after its time integration is given by

$$\bar{\varrho}(t) = D_{n,m,k} e^{-\lambda_{n,m,k} D_e t}, \quad (7)$$

where the integration constant is

$$D_{n,m,k} = \iiint_{\Omega_r} \varrho_0 \sqrt{\frac{8}{abc}} \sin \frac{n\pi \xi}{a} \sin \frac{m\pi \eta}{b} \sin \frac{k\pi \beta}{c} d\xi d\eta d\beta. \quad (8)$$

The source function for the total electric charge is to obtain from the integration of eq. (II.2) in the time interval $[0, t)$ the form (III.2) which can be written as

$$\varrho^z(t) = \varrho_0 \pm \int_0^t \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) dt. \quad (9)$$

Making use of the results of the chapter A.4.1.1. and the above formulae for the solution of the diffusion transport of the total electric charge, we have

$$\varrho_i(Z, t) = \int_0^t \iiint_{\Omega_R} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k} e^{-\lambda_{n,m,k} D_{ei}(t-\tau)} \frac{8}{abc} \left(\sin \frac{n\pi x}{a} \sin \frac{n\pi \xi}{a} \right) \cdot \left(\sin \frac{m\pi y}{b} \sin \frac{m\pi \eta}{b} \right) \left(\sin \frac{k\pi z}{c} \sin \frac{k\pi \beta}{c} \right) \right\} \left[\varrho_0 \pm \int_0^{\tau} \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) dt \right] d\xi d\eta d\beta d\tau. \quad (10)$$

4.1.3. THE VOLTAGE OF THE MULTI-COMPONENT ELECTRIC FIELDS IN THE "i" COMPONENT DIFFUSION CHARGE TRANSPORT

Let us consider eq. (VIII¹.3) having the form

$$\varrho \frac{\partial U}{\partial t} = \sum_{i=1}^I U_{di} D_{ei} \nabla^2 \varrho_i. \quad (1)$$

Now we can use eq. (VIII¹.1) of the "i" component diffusion electric charge transport

$$\frac{\partial \varrho_i}{\partial t} = D_{ei} \nabla^2 \varrho_i. \quad (2)$$

It is possible to perceive that from eq. (1) and (2) one can obtain

$$\varrho \frac{\partial U}{\partial t} = \sum_{i=1}^I U_{di} \frac{\partial \varrho_i}{\partial t}. \quad (3)$$

For the solution of the problem described by eq. (3) we need to stress that $\varrho = \sum_{i=1}^I \varrho_i$.

This point of view enables to obtain the solution of the problem of eq. (3) with respect to the pertinent "i" component electric charge source function in the form

$$U(Z, t) = U_0(Z, t) + \sum_{i=1}^I U_{di} \ln \left| \int_0^t \iiint_{\Omega_R} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^i e^{-\lambda_{n,m,k} D_{ei}(t-\tau)} \cdot \frac{8}{abc} \left(\sin \frac{n\pi x}{a} \sin \frac{n\pi \xi}{a} \right) \left(\sin \frac{m\pi y}{b} \sin \frac{m\pi \eta}{b} \right) \left(\sin \frac{k\pi z}{c} \sin \frac{k\pi \beta}{c} \right) \right\} \cdot \left[\varrho_0 \pm \int_0^{\tau} \sum_{l=1}^L J_{il}(t) dt \right] d\xi d\eta d\beta d\tau \right| - \ln |\varrho_{0i}|. \quad (4)$$

The integration constant $D_{n,m,k}^i$ is defined in the chapter A.4.1.1.

4.1.4. THE VOLTAGE DISTRIBUTION FROM THE ELECTRIC PERMEABILITY OF THE "i" COMPONENT ELECTRIC VOLTAGES

Let us take into consideration eq. (VIII².3) whose form is

$$\varrho \frac{\partial U}{\partial t} - \frac{1}{2} p \sum_{i=1}^I U_{\rho i} \varepsilon_i \nabla^2 U_i = 0. \quad (1)$$

The solution of eq. (1) must be in consonance with Theorem 1 and Corollary 1, and we need remember that $\varrho = \sum_{i=1}^I \varrho_i$ and $U = \sum_{i=1}^I U_i$. Introducing now the "i" component voltage function $U_i(x, y, z, t) = W(x, y, z) \bar{U}_i(t)$, eq. (1) takes the form

$$\frac{\partial^2 W}{\partial x^2} + \frac{\partial^2 W}{\partial y^2} + \frac{\partial^2 W}{\partial z^2} + \lambda W = 0 \quad (2)$$

$$\frac{d\bar{U}_i}{dt} = -\lambda_2 \frac{p U_{\rho i} \varepsilon_i}{\varrho^2(t)} \bar{U}_i. \quad (3)$$

On the basis of the considerations of the chapter A.4.1.1. we can see that the eigenvalues are obtained as:

$$\lambda_{n,m,k} = \left(\frac{n\Pi}{a} \right)^2 + \left(\frac{m\Pi}{b} \right)^2 + \left(\frac{k\Pi}{c} \right)^2. \quad (4)$$

It is evident, therefore, that the eigenfunctions for the point $Z(x, y, z, t)$ are given by

$$W_{n,m,k}(x, y, z) = \sqrt{\frac{8}{abc}} \sin \frac{n\Pi x}{a} \sin \frac{m\Pi y}{b} \sin \frac{k\Pi z}{c} \quad (5)$$

and reciprocally, in the variable point $Q(\xi, \eta, \beta, \tau)$ the eigenfunctions fulfil the equation

$$W_{n,m,k}(\xi, \eta, \beta) = \sqrt{\frac{8}{abc}} \sin \frac{n\Pi \xi}{a} \sin \frac{m\Pi \eta}{b} \sin \frac{k\Pi \beta}{c}. \quad (6)$$

Eq. (3) can be transformed to

$$\frac{d\bar{U}_i}{\bar{U}_i} = -\lambda_{n,m,k} \frac{p U_{\rho i} \varepsilon_i}{2} \frac{dt}{\varrho_{0i} \pm \int_0^t \sum_{l=1}^L J_{il}(t) dt} \quad (7)$$

and the integration of eq. (7) leads to

$$\bar{U}_i(t) = F_{n,m,k}^i e^{\frac{p U_{\rho i} \varepsilon_i}{2} [K_i(t-\tau)]}, \quad (8)$$

where the function

$$[K_i(t)] = \int \frac{dt}{\varrho_{0i} \pm \sum_{l=1}^L J_{il}(t) dt}.$$

For eq. (8) the integration constant can be written as

$$F_{n,m,k}^i = \iiint_{\Omega_R} U_{0i} \sqrt{\frac{8}{abc}} \sin \frac{n\pi\xi}{a} \sin \frac{m\pi\eta}{b} \sin \frac{k\pi\beta}{c} d\xi d\eta d\beta. \quad (9)$$

The source function for this problem is to obtain from eq. (II.3) after integration in the time interval $[0, t]$ to eq. (III.3), and is

$$U_i^Z(t) = U_{0i} \pm \int_0^t \sum_{l=1}^L \frac{P_{il}(t)}{|q_i^Z(t)|} dt. \quad (10)$$

In accordance with the above considerations the voltage distribution by the penetrability phenomenon from the variable point $Q(\xi, \eta, \beta, \tau)$ to the point $Z(x, y, z, t)$ has the resultant form

$$U(Z, t) = \sum_{i=1}^I \int_0^t \iiint_{\Omega_R} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} F_{n,m,k}^i e^{-\lambda_{n,m,k} \frac{P_{0i} \varepsilon_i}{2} [K_i(t-\tau)]} \frac{8}{abc} \cdot \left(\sin \frac{n\pi x}{a} \sin \frac{n\pi \xi}{a} \right) \left(\sin \frac{m\pi y}{b} \sin \frac{m\pi \eta}{b} \right) \left(\sin \frac{k\pi z}{c} \sin \frac{k\pi \beta}{c} \right) \right\} \cdot \left[\varrho_{0i} \pm \int_0^{\tau} \sum_{l=1}^L \frac{P_{il}(t)}{|q_i^Z(t)|} dt \right] d\xi d\eta d\beta d\tau. \quad (11)$$

4.1.5. THE DISTRIBUTION OF THE "i" COMPONENT VOLTAGE OF THE CONDUCTIVITY PHENOMENON ON THE TOTAL ELECTRIC ENERGY

We use for our considerations eq. (VIII³.3), having the form

$$\varrho \frac{\partial U}{\partial t} - \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i = 0. \quad (1)$$

The solution method of eq. (1) coincides with the solution of the chapter A.4.1.4. and confirms the importance of Theorem 1 and Corollary 1, at $\varrho = \sum_{i=1}^I \varrho_i$ and $U = \sum_{i=1}^I U_i$.

We can present eq. (1), by the introduction of the "i" component electric voltage $U_i(x, y, z, t) = W(x, y, z) U_i(t)$ defined as having space and time elements;

$$\frac{\partial^2 W}{\partial x^2} + \frac{\partial^2 W}{\partial y^2} + \frac{\partial^2 W}{\partial z^2} + \lambda W = 0 \quad (2)$$

$$\frac{dU_i}{dt} = -\lambda_2 \frac{U_{ci}\gamma_i}{\varrho^z(t)} U_i. \quad (3)$$

After the same considerations as in chapter A. §.4.1.4. we receive the eigenfunctions in the form:

$$\lambda_{n,m,k} = \left(\frac{n\pi}{a}\right)^2 + \left(\frac{m\pi}{b}\right)^2 + \left(\frac{k\pi}{c}\right)^2. \quad (4)$$

The eigenfunctions in the point $Z(x, y, z, t)$ are

$$W_{n,m,k}(x, y, z) = \sqrt{\frac{8}{abc}} \sin \frac{n\pi x}{a} \sin \frac{m\pi y}{b} \sin \frac{k\pi z}{c} \quad (5)$$

and pertinent for the variable point $Q(\xi, \eta, \beta, \tau)$

$$W_{n,m,k}(\xi, \eta, \beta) = \sqrt{\frac{8}{abc}} \sin \frac{n\pi \xi}{a} \sin \frac{m\pi \eta}{b} \sin \frac{k\pi \beta}{c}. \quad (6)$$

Eq. (3) can be modified as follows

$$\frac{dU_i}{U_i} = -\lambda_{n,m,k} U_{ci} \gamma_i \frac{dt}{\varrho_{0i} \pm \int_0^t \sum_{l=1}^L J_{il}(t) dt}. \quad (7)$$

and integrated to the form

$$\bar{U}_i(t) = F_{n,m,k}^i e^{-\lambda_{n,m,k} U_{ci} \gamma_i [K_i(t)]}, \quad (8)$$

where

$$[K_i(t)] = \int \frac{dt}{\varrho_{0i} \pm \int_0^t \sum_{l=1}^L J_{il}(t) dt}.$$

For eq. (8) the integration constant is

$$F_{n,m,k}^i = \iiint_{\Omega_R} U_{0i} \sqrt{\frac{8}{abc}} \sin \frac{n\pi \xi}{a} \sin \frac{m\pi \eta}{b} \sin \frac{k\pi \beta}{c} d\xi d\eta d\beta. \quad (9)$$

From the integration of eq. (II.3) in the time interval $[0, t)$ to the form of eq. (III.3) we have

$$U_i^z(t) = U_{0i} \pm \int_0^t \sum_{l=1}^L \frac{P_{il}(t)}{|\varrho^z(t)|} dt. \quad (10)$$

In accordance with the above considerations the resultant solution of this problem can be written in the form

$$U(Z, t) = \sum_{i=1}^I \int_0^t \iiint_{\Omega_R} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} F_{n,m,k}^i e^{-\lambda_{n,m,k} U_{ci} \gamma_i [K_i(t-\tau)]} \frac{8}{abc} \cdot \left(\sin \frac{n\pi x}{a} \sin \frac{n\pi \xi}{a} \right) \left(\sin \frac{m\pi y}{b} \sin \frac{m\pi \eta}{b} \right) \left(\sin \frac{k\pi z}{c} \sin \frac{k\pi \beta}{c} \right) \right\} \cdot \left[U_{oi} \pm \int_0^{\tau} \sum_{l=1}^L \frac{P_{il}(t)}{|q^z(t)|} dt \right] d\xi d\eta d\beta d\tau. \quad (11)$$

4.1.6. THE DISTRIBUTION OF THE "P" COMPONENT ELECTRIC CHARGE OF THE CONDUCTIVITY PHENOMENON

Let be taken into consideration eq. (VIII³.1) having the form:

$$\frac{\partial q_i}{\partial t} = \gamma_i \nabla^2 U_i \quad (1)$$

and additionally eq. (VIII³.3) as

$$q \frac{\partial U}{\partial t} = \sum_{i=1}^I U_{ci} \gamma \nabla^2 U_i. \quad (2)$$

Introducing now relations $q = \sum_{i=1}^I q_i$ and $U = \sum_{i=1}^I U_i$ and using them with the combination of eq. (1) with eq. (2) one gets

$$\frac{\partial q_i}{\partial t} = \frac{i}{U_{ci}} \rho_i \frac{\partial U_i}{\partial t}. \quad (3)$$

The solution of the problem presented by eq. (3) takes the form:

$$q_i(Z, t) = q_{oi}(Z, t) e^{\left\langle \frac{i}{U_{ci}} \left\langle \int_0^t \iiint_{\Omega_R} \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} F_{n,m,k}^i e^{-\lambda_{n,m,k} U_{ci} \gamma_i [K_i(t-\tau)]} \frac{8}{abc} \times \left(\sin \frac{n\pi x}{a} \sin \frac{n\pi \eta}{a} \right) \left(\sin \frac{m\pi y}{b} \sin \frac{m\pi \eta}{b} \right) \left(\sin \frac{k\pi z}{c} \sin \frac{k\pi \beta}{c} \right) \right\} - \left[U_{oi} \pm \int_0^{\tau} \sum_{l=1}^L \frac{P_{il}(t)}{|q^z(t)|} dt \right] d\xi d\eta d\beta d\tau - U_{oi}(Z, t) \right\rangle}. \quad (4)$$

For the integration constant $F_{n,m,k}^i$ and the function $[K_i(t-\tau)]$ see chapter A.4.1.5.

4.1.7. THE DISTRIBUTION OF THE TOTAL ELECTRIC CHARGE IN THE CONDUCTIVITY PHENOMENON

On the basis of the considerations of the chapter A.4.1.5, keeping in the mind the realtions $\varrho = \sum_{i=1}^I \varrho_i$ and $U = \sum_{i=1}^I U_i$, we can obtain from eq. (VIII³.2).

$$\frac{\partial \varrho}{\partial t} = \sum_{i=1}^I \gamma_i \nabla^2 U_i. \quad (1)$$

Now, we can use eq. (VIII³.3)

$$\varrho \frac{\partial U}{\partial t} = \sum_{i=1}^I U_{ci} \gamma_i \nabla^2 U_i. \quad (2)$$

Combining eq. (1) with eq. (2), it is possible to obtain

$$\frac{1}{\varrho} \frac{\partial \varrho}{\partial t} = \sum_{i=1}^I \frac{1}{U_{ci}} \frac{\partial U_i}{\partial t}. \quad (3)$$

The solution of the problem presented by eq. (3) has the form

$$\begin{aligned} \varrho_i(Z, t) = \varrho_0(Z, t) e & \sum_{i=1}^I \iiint_{\Omega_R} \left\{ \sum_{n=1}^{\infty} F_{n,m,k}^i e^{-\lambda_{n,m,k} U_{ci} \kappa_i [K_i(t-\tau)]} \frac{8}{abc} \right. \\ & \cdot \left(\sin \frac{n\pi x}{a} \sin \frac{n\pi y}{a} \right) \left(\sin \frac{m\pi y}{b} \sin \frac{m\pi \eta}{b} \right) \left(\sin \frac{k\pi z}{c} \sin \frac{k\pi \beta}{c} \right) \Big\} - [U_{ci} \pm \\ & \pm \int_0^{t-\tau} \sum_{i=1}^I \frac{P_n(t)}{|e^{\tau}(t)|} dt] d\xi d\eta d\beta d\tau \Big\}. \end{aligned} \quad (4)$$

The integration constant $F_{n,m,k}^i$ and $[K_i(t-\tau)]$ are defined in the chapter A.4.1.5.

4.1.8. THE DISRTIBUTED RELATION BETWEEN THE TOTAL ELECTRIC CHARGE AND THE POTENTIAL FORM OF THE FLOW OF THE EXISTING ELECTRIC CHARGE

The is the general potential formula of the field vector velocity of the existing total electric charge which is independent to all other physical phenomena and the total voltage determining this case of multi-component electric fields

$$f \frac{\partial \mathbf{v}}{\partial t} = \text{grad } U. \quad (1)$$

4.2. THE PARTIAL DIFFERENTIAL EQUATIONS OF THE ROTATIONAL FIELD, (variable point $Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q(\xi, \eta, \beta, \tau)$ [1].

4.2.1. THE DISTRIBUTION OF THE FIELD VECTOR VELOCITY OF THE TOTAL ELECTRIC CHARGE

Let us take into consideration the set of partial differential equations, in the forms of the following derivatives (IX) belonging to the variable point, $Q'(\xi', \eta', \beta', \tau')$. We need to verify the reciprocal influences of the state vectors involved to the pertinent points:

point $Z(x, y, z, t)$ — constant space-time coordinates and point $Q'(\xi', \eta', \beta', \tau')$ — variable space-time coordinates

$$\begin{array}{ccc} \left[\begin{array}{c} q_i(x, y, z, t) \\ U_i(x, y, z, t) \\ \varrho(x, y, z, t) \\ U(x, y, z, t) \\ \mathbf{v}(x, y, z, t) \end{array} \right] & \begin{array}{c} \rightarrow (\text{connection}) \leftarrow \\ \leftarrow \\ \text{Reciprocal mutual} \\ \text{relation.} \end{array} & \left[\begin{array}{c} q_i(\xi', \eta', \beta', \tau') \\ U_i(\xi', \eta', \beta', \tau') \\ \varrho(\xi', \eta', \beta', \tau') \\ U(\xi', \eta', \beta', \tau') \\ \mathbf{v}(\xi', \eta', \beta', \tau') \end{array} \right] \end{array} \quad (1)$$

point $Z(x, y, z, t) \in \bar{\Omega}_{(a,b,c)t}$ point $Q'(\xi', \eta', \beta', \tau') \in \Omega_{R'}$. Now, we can investigate the system (IX) from the point of view of its properties

$$\frac{Dq_i}{Dt} = -q_i \operatorname{div} \mathbf{v} \quad (2)$$

$$\frac{D\varrho}{Dt} = -\varrho \operatorname{div} \mathbf{v} \quad (3)$$

$$\varrho \frac{DU}{Dt} = -\varrho U \operatorname{div} \mathbf{v} \quad (4)$$

$$f\varrho \frac{D\mathbf{v}}{Dt} = \hat{a} \nabla^2 \mathbf{v} - f\varrho \mathbf{v} \operatorname{div} \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \quad (5)$$

where: $\frac{D}{Dt} = \frac{\partial}{\partial t} + \mathbf{v} \operatorname{grad}$ [4], [1], [5], [10–12].

In accordance with the previous considerations the multi-component electric charges $\varrho = \sum_{i=1}^I q_i$ and their voltages fulfil $U = \sum_{i=1}^I U_i$. It is necessary to introduce here the multi-component quantities “ f_i ” being the “ i ” component coefficient of the “ i ” mass of the “ i ” component electric charge. In consequence of the use of this point of view

$$f\varrho = \sum_{i=1}^I f_i q_i \quad (6)$$

and eq. (5) can be modified to

$$f\varrho \frac{D\mathbf{v}}{Dt} + f\varrho \mathbf{v} \operatorname{div} \mathbf{v} = \hat{a} \nabla^2 \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \quad (7)$$

Proceeding in this way, from the application of eq. (3) we have

$$f\varrho \frac{D\mathbf{v}}{Dt} - f\mathbf{v} \frac{D\varrho}{Dt} = \hat{a} \nabla^2 \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v}. \quad (8)$$

From the assumption of the reciprocal constant image in the static state, we can present eq. (8) as [3–4], [1], [5], [10–12]

$$\hat{a} \nabla^2 \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} = 0 \quad (9)$$

and

$$\varrho \frac{D\mathbf{v}}{Dt} - \nu \frac{D\varrho}{Dt} = 0. \quad (10)$$

Applying the mathematical rule [4], [6], [8], [5], [10—12]

$$\text{rot rot } \mathbf{v} = \text{grad div } \mathbf{v} - \nabla^2 \mathbf{v} \quad (11)$$

to eq. (9), we can obtain from this equation

$$(\hat{a} + \hat{b}) \hat{a} \text{grad div } \mathbf{v} - \hat{a} \text{rot rot } \mathbf{v} = 0. \quad (12)$$

Eq. (12) can be represented in the form

$$\hat{\nabla}^2 \mathbf{v} = 0 = \hat{a} \text{grad div } \mathbf{v} - \hat{\beta} \text{rot rot } \mathbf{v}. \quad (13)$$

On the basis of the reciprocal constant image, the static Green tensor for the volume element $\bar{\Omega}_{(a,b,c)}$ is [4—5], [10—12]

$$\Gamma_G(Z, Q') = \frac{1}{8\pi} \left[\frac{1}{\hat{\alpha}} \Gamma_{\hat{\alpha}}(Z, Q') + \frac{1}{\hat{\beta}} \Gamma_{\hat{\beta}}(Z, Q') \right] \quad (14)$$

where: $\hat{\alpha} = \hat{a} + \hat{b}$ and $\hat{\beta} = \hat{a}$
— for the fluid phase;

$$\hat{a} = \eta_w \quad \text{and} \quad \hat{b} = \frac{1}{3} \eta_w$$

— for the crystal phase;

$$\hat{a} = \frac{h}{p} \quad \text{and} \quad \hat{b} = \frac{h\bar{k}}{p(\bar{k}-2)},$$

where

$$\Gamma_{\hat{\alpha}} = \frac{1}{|\mathbf{r}|} I - \frac{\mathbf{r} \times \mathbf{r}}{|\mathbf{r}|^3} \quad \text{and} \quad \Gamma_{\hat{\beta}} = \frac{1}{|\mathbf{r}|} I + \frac{\mathbf{r} \times \mathbf{r}}{|\mathbf{r}|^3}$$

by $\mathbf{r} = \mathbf{r}_Z - \mathbf{r}_Q$,

In my paper devoted to the single-component electric field [5], and in my other papers containing the physicochemical modelling of real processes [10—12], the anisotropic Green tensor for the cuboidal selected volume element $\bar{\Omega}_{(a,b,c)}$, was deduced

$$\Gamma_{GA}(Z, Q') = \Gamma_{ln}(Z, Q') = \Gamma_{lp}(Z, Q') * \hat{\Gamma}_{pq}(Z, Q') \Gamma_{qn}(Z, Q') \quad (15)$$

for $k = 1(1)w$,

where:

$\Gamma_{lp}(Z, Q')$ — the static isotropic Green tensor for the material conditions of the cuboidal locally selected volume element $\bar{\Omega}_{(a,b,c)}$ [4—6], [10—12],

$\hat{\Gamma}_{pq}(Z, Q')$ — the modified tensor of static deformations of the cuboidal selected

volume element $\bar{\Omega}_{(a,b,c)}$ for the first order power of the small parameter [5], [10—12],

$\Gamma_{qn}^{(k-1)}(Z, Q')$ — the statical Green tensor of the $(k-1)$ deformations pertinent to the small parameter [5], [10—12],

w — number of the power of the used small parameter.

The consequence of the assumption that $\Gamma_G(Z, Q') \equiv \Gamma_{GA}(Z, Q') \equiv \Gamma_G^W(Z, Q')$ is the possibility of an universal treatment of the considerations of this chapter. The static state of the point $Z(x, y, z, 0)$ leads to the velocity vector $v_0(Z, 0) = v_0$ and the form of the static distribution of the field vector velocity is

$$v(Z, 0) = \iiint_{\Omega_R} \Gamma_G^W(Z, Q') v_0(Q') d\xi' d\eta' d\beta'. \quad (16)$$

Consequently to eq. (10) we have

$$\ln|N| + \ln|v| = \ln|q| \quad (17)$$

with its resultant formula

$$v(t) = \frac{q^z(t)}{N_1} \quad (18)$$

where:

N_1 — constant having the form of eq. (16), with $N_1 = v(Z, 0)$,

$q^z(t)$ — the source function of the total electric charge of the multi-component electric fields at the point $Z(x, y, z, t)$ — see eq. (III.2),

$v(t)$ — the time element of the distribution of the field vector velocity of the total electric charge of the multi-component electric fields.

With the above considerations at hand the distribution of the field vector velocity of the whole electric charge takes its final form

$$v(Z, t) = \int_0^t \iiint_{\Omega_R} \left\{ \frac{1}{N_1} \left[q_0 \pm \int_0^{t-(t-\tau)} \sum_{i=1}^I \sum_{l=1}^L J_{il}(t) dt \right] \Gamma_G^W(Z, Q') \right\} \left[v_0 \pm \int_0^{t-\tau} \sum_{i=1}^I \sum_{l=1}^L \frac{T_{il}(t)}{f_i |q_i^z(t)|} dt \right] d\xi, d\eta, d\beta, d\tau. \quad (19)$$

5. THE SOLUTION OF THE MATHEMATICAL MODEL OF THE DISTRIBUTED PARAMETERETS OF THE MULTI-COMPONENT ELECTRIC FIELDS

Applying all phenomenological solutions of the chapter 4, the structure of the complete solution of the mathematical model of the distributed parameters of the multi-component electric fields is bound with summations of the phenomenological

effects pertinent to potential fields and the rotational field. The structure of this solution has the complete form:

The complete form of
the solution

The multi-component transformations
of multi-component electric fields
at point $Z(x, y, z, t)$

Density of the total
charge
Total voltage
Total field vector
velocity

=

Density of total charge changes
(formula A.4.III.2)
Total voltage changes
(formula A.4.III.3)
Total field vector velocity changes
(formula A.4.III.4)

+

Diffusion
 $Q \rightarrow Z$

Permeability
 $Q \rightarrow Z$

Density of total charge transport
(formula A.4.1.2.10)
Total voltage transport
(formula A.4.1.3.4)
 $\times$

+

$\times$
Total voltage transport
(formula A.4.1.4.11)
 $\times$

+

Conductivity
 $Q \rightarrow Z$

Electrostatic force
 $Q \rightarrow Z$

Density of total charge transport
(formula A.4.1.7.4)
Total voltage transport
(formula A.4.1.5.11)
 $\times$

+

$\times$
 $\times$
Total field vector velocity
distribution
(formula A.4.1.8.1)

+

Field vector velocity

 $Q \rightarrow Z$

$$+ \left[\begin{array}{l} \text{Density of total charge continuity} \\ \text{(formula A.}\mathscr{A}.4.2.1.3) \\ \text{Total voltage continuity} \\ \text{(formula A.}\mathscr{A}.4.2.1.4) \\ \text{Total field vector velocity transport} \\ \text{(formula A.}\mathscr{A}.4.2.1.19) \end{array} \right] \quad (1)$$

$\mathbb{X}$ — coordinates rid of relations between phenomena for all components of the multi-component electric fields.

The structure (1) represents a novelty owing to the solution of set of partial differential equations of the mass/charge, energy and momentum balances of the mathematical physics. While the subject set of partial differential constitutive state equations (I) is already an important novelty, the originality of these considerations can be extended over whole approach to the topic of distributed parameter modelling of the multi-component electric fields. The presented above structure (1) represents the dynamical working space-time characteristics in the sense of the distributed parameter control theory [17—19], [10—12].

B. THE EXISTENCE OF THE INITIAL AND BOUNDARY CONDITIONS FOR THE MULTI-COMPONENT ELECTRIC FIELDS.

β . THE DISTRIBUTED PARAMETER DYNAMICAL WORKING SPACE-TIME CHARACTERISTICS AND PHENOMENOLOGICAL CONTROL CONCEPT BY APPLICATION OF THE INITIAL AND BOUNDARY CONDITIONS.

The procedure of the construction of the dynamical working space-time characteristics by the existence of the initial and boundary condition is the same as in eq. (A.5.1.). The phenomenological solutions for the case of the existence of the initial and boundary conditions are involved with the properties of the phenomenological Green functions, listed in my article [5]. These phenomenological solutions can be described as [4], [6], [8]:

$$W(Z, t) = \int_0^t \iiint_{\Omega_R} G(Z, t; Q^R \tau(m) Q^R, \tau) d\Omega_R(\xi^R, \eta^R, \beta^R, \tau) d\tau + \\ + \int_0^t \oint_{F_R} G_{s0}(Z, t; Q^R, \tau(m_s) Q^R, \tau) dF_R(\xi^R, \eta^R, \beta^R) d\tau \quad (1)$$

where:

$(\xi^R, \eta^R, \beta^R, \tau) \equiv Q(\xi, \eta, \beta, \tau)$ — potential fields, $Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q'(\xi', \eta', \beta', \tau)$ — rotational field,

$W(Z, t)$ — the resultant state variable for the pertinent physical phenomenon,

- $G(Z, t; Q^R, \tau)$ — the phenomenological Green function for the working space volume-time of equipment with multi-component electric fields,
 $m(Q^R, \tau)$ — the source function existing inside the working space volume-time of equipment with the multi-component electric fields,
 $G_{s0}(Z, t; Q^R, \tau)$ — the phenomenological Green function for the working surface-time of equipment with the multi-component electric fields,
 $m_s(Q^R, \tau)$ — the source function acting on the working surface-time of equipment with the multi-component electric fields.

The following relation is valid for the phenomenological Green functions

$$G_{s0}(Z, t; Q^R, \tau) = G(Z, t; Q^R, \tau)|_{F_R} \quad (2)$$

Another important information is that $W(Z, t)$ has all the properties of the phenomenological Green function [6], [4—5], [8], [10—12]. The phenomenological solutions of physical phenomena of the mathematical model of the distributed parameters of the multi-component electric fields can be interpreted in the field of the distributed parameter control theory [17—19], [10—12].

Owing to the possibility of the disposition of the subject system (I) in the partial differential equations of the potential fields and the partial differential equations, in the forms of the following derivatives of the rotational field we can consider the separate influences of every physical phenomenon of multi-component electric fields on the yield and quality aspects of the multi-component transformations of these fields. The phenomenological distributed parameter control concept with the application of the initial and boundary conditions is presented in the figure 1.

Figure 1 displays the possibility of the presentation of exerted by single physical phenomena belonging to the potential fields and the rotational field by applying their phenomenological Green functions — within the brackets $\{ \}$, on the yield and quality aspects of the multi-component transformations of the multi-component electric fields with the application of:

- the initial conditions, and
- the initial and boundary conditions, of these electric fields.

THE APPLICATION OF THE SOLUTIONS OF THE SYSTEM (I) — CASES A AND B TO THE PERTINENT STATICAL WORKING SPACE CHARACTERISTICS.

For the analysis, one can take the solutions of case A, from the chapter A.5. with the elimination of the time coordinate according to the condition $t = \tau = 0$. Applying the procedure (A.5.1) it is possible to obtain the static working space characteristics. The same considerations are valid for the presence of the initial and boundary conditions of the multi-component electric fields. The static phenomenological solutions of the potential fields and the rotational field fulfil the rule [4], [6], [8]:

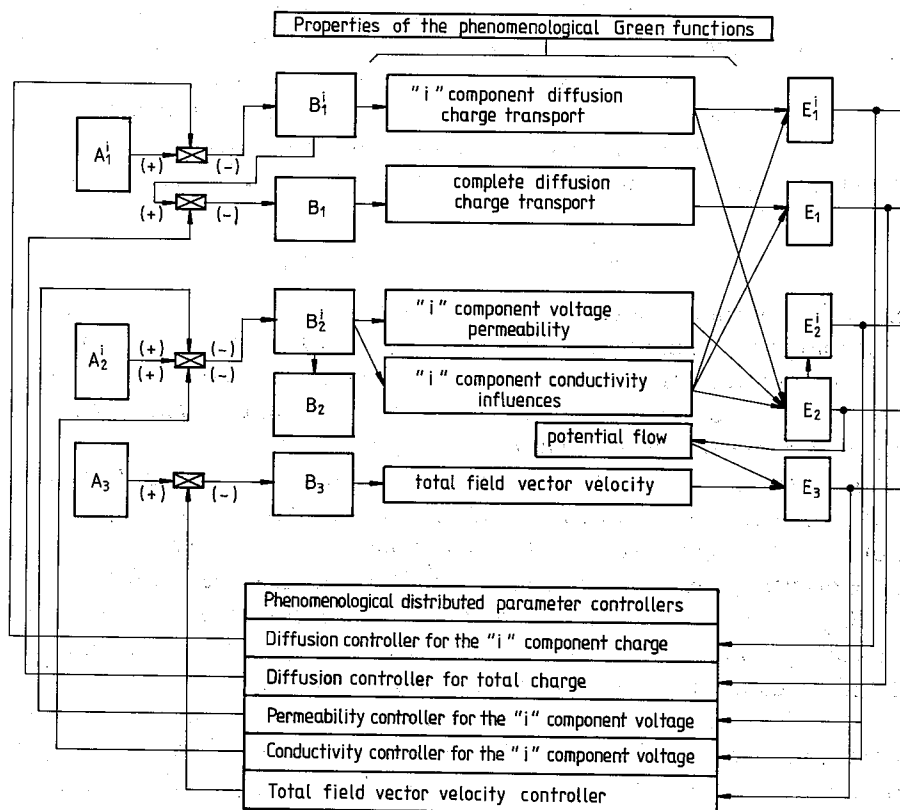


Fig. 1. A concept of the phenomenological distributed parameter control for multi-component electric fields:

A_1^i — INPUT "i" component electric charge,

A_2^i — INPUT "i" component electric voltage,

A_3 — INPUT field vector velocity for all components of multi-component electric fields,

B_1^i — SOURCE of the "i" component electric charge,

B_1 — SOURCE of the total electric charge,

B_2^i — SOURCE of the "i" component electric voltage,

B_2 — SOURCE of the total electric voltage,

B_3 — SOURCE of the field vector velocity for all components of multi-component electric fields,

E_1^i — OUTPUT "i" component electric charge,

E_1 — OUTPUT total electric charge,

E_2^i — OUTPUT "i" component electric voltage,

E_2 — OUTPUT total electric voltage,

E_3 — OUTPUT field vector velocity for all components of the multi-component electric fields.

$$\begin{aligned}
 W(Z) = & \iiint_{\Omega_R} G(Z; Q^R) m(Q^R) d\Omega_R(\xi^R, \eta^R, \beta^R) + \\
 & + \oint_{F_R} G_{s0}(Z; Q^R) m_s(Q^R) dF_R(\xi^R, \eta^R, \beta^R)
 \end{aligned} \quad (1)$$

by: $Z(x, y, z) = Z(x, y, z, 0)$ — the point of the constant working space coordinates, $Q^R(\xi^R, \eta^R, \beta^R) = Q^R(\xi^R, \eta^R, \beta^R, 0)$ — the point of the variable working space coordinates of the equipment with multi-component electric fields. Additionally we know that $Q^R(\xi^R, \eta^R, \beta^R) \equiv Q(\xi, \eta, \beta)$ — for the physical phenomena of the potential fields and $Q^R(\xi^R, \eta^R, \beta^R) \equiv Q'(\xi', \eta', \beta')$ — for the physical phenomena of the rotational field. We can now determine the elements of eq. (1);

- $W(Z)$ — the static from of the resultant state variable for the adequate physical phenomenon,
 $G(Z; Q^R)$ — the phenomenological Green function for the working space volume of the apparatuses of multi-component electric fields,
 $m(Q^R)$ — the source function acting inside the working space volume of equipment with multi-component electric fields,
 $G_{s0}(Z; Q^R)$ — the phenomenological Green function for the working surface of equipment with multi-component electric fields,
 $m_s(Q^R)$ — the source function acting on the working surface of equipment with multi-component electric fields.

The phenomenological Green functions of the working space volume and the working surface problems are mutually related

$$G_{s0}(Z; Q^R) = G(Z; Q^R)|_{F_R} \quad (2)$$

Applying formula (1) for the phenomena of the potential fields and the phenomena of the rotational field we have a basis for the construction of the static working space characteristics at the existence of the initial and boundary conditions, possible to be obtain as a consequence of the use of their structure (A. 5.1) [17—19], [10—12].

6. THE INPUT-SOURCE-ELECTRICAL PROCESS-OUTPUT FOR THE MULTI-COMPONENT ELECTRIC FIELDS

The novelty of the approach proposed in the paper to the solution of the system (I), describing in a complete from the physical aspects of multi-component electric fields leads to an original interpretation of the INPUT-SOURCE-ELECTRICAL PROCESS-OUTPUT relation. We need consider for the multi-component electric fields:

— INPUT; the multicomponent electric charges $q = \sum_{i=1}^I q_i$ and their voltages

$U = \sum_{i=1}^I U_i$, by the general field vector velocity for all components “v”,

— SOURCE; the dynamics of the multi-component transformations at every point $Z(x, y, z)$ is described by the time differential equations of the system (II) which is integrated over the time interval $[0, t]$, to the source having the form of the solutions (III),

— ELECTRICAL PROCESS; the homogeneous form of the system (I) is analysed in the chapter A.2, from the point of view of the theory of the potential and rotational fields [4—6], [8]. As a consequence of the use of this theory we have: for the potential fields — the system (VIII) and for the rotational field — the system (IX). The system (VIII) and its subsequently obtained (VIII¹), (VIII²) and (VIII³) phenomenological forms necessary for the suitable consideration of a single physical phenomena of these electric fields.

— OUTPUT; we extend this part of the considerations to the analytical solutions of the phenomenological partial differential equations of the potential fields (systems (VIII¹), (VIII²) and (VIII³) and the partial differential equations, in the forms of the following derivatives of the rotational field system (IX) not only for case A° but also over the range of case B.

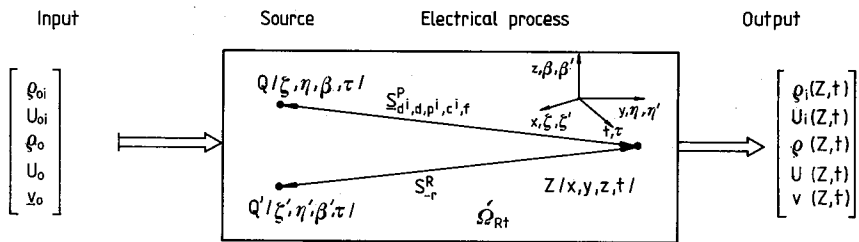


Fig. 2. The phenomenological interpretation of the elements of processing of the multi-component electric fields

$Z(x, y, z, t)$ — point of multi-component transformations of the multi-component electric fields — constant working space-time coordinates of equipment with multi-component electric fields of their volume-time Ω_{Rt} ,

$Q(\xi, \eta, \beta, \tau)$ — variable point of the potential fields of Ω_{Rt} ,

$S_{d^i, d, p^i, c^i, f}^p$ — the phenomenological fluxes of the potential fields (d^i — “i” component diffusion, d — total diffusion, p^i — “i” component permeability, c^i — “i” component conductivity, f — potential flow),

$Q'(\xi', \eta', \beta', \tau)$ — variable point of the rotational field of Ω_{Rt} ,

S_r^R — the phenomenological flux of the rotational field (r — rotational flow),

q_{oi} — the initial “i” component electric charge,

U_{oi} — the initial “i” component electric voltage,

q_o — the initial total electric charge,

U_o — the initial total electric voltage,

v_o — the initial field vector velocity for all components of multi-component electric fields,

$q_i(Z, t)$ — the resultant “i” component electric charge,

$U_i(Z, t)$ — the resultant “i” component electric voltage,

$q(Z, t)$ — the resultant total electric charge,

$U(Z, t)$ — the resultant total electric voltage,

$v(Z, t)$ — the resultant field vector velocity for all components of multi-component electric fields.

According to the above elements it is easy to prove their suitability to what is displayed in figure 2.

Figure 2 shows the reciprocal interpretation of the superposition of the potential fields — the physical phenomena: diffusion, conductivity, electrostatics, to the rotational field — convection and excitation phenomena with the description of the material properties [1]. We have in figure 2 the mutual relation between the constant and variable coordinates points:

$$Q(\xi, \eta, \beta, \tau) \leftrightarrow Z(x, y, z, t) \leftrightarrow Q'(\xi', \eta', \beta', \tau) \quad (1)$$

in Ω_R in $\Omega_{(a,b,c)t}$ in Ω_R

It subsequently follows from the reciprocal interpretation that for the state vectors occurs

$$\begin{bmatrix} q_i(\xi, \eta, \beta, \tau) \\ U_i(\xi, \eta, \beta, \tau) \\ \varrho(\xi, \eta, \beta, \tau) \\ U(\xi, \eta, \beta, \tau) \\ v(\xi, \eta, \beta, \tau) \end{bmatrix} \xleftrightarrow{M_1} \begin{bmatrix} q_i(x, y, z, t) \\ U_i(x, y, z, t) \\ \varrho(x, y, z, t) \\ U(x, y, z, t) \\ v(x, y, z, t) \end{bmatrix} \xleftrightarrow{M_2} \begin{bmatrix} q_i(\xi', \eta', \beta', \tau) \\ U_i(\xi', \eta', \beta', \tau) \\ \varrho(\xi', \eta', \beta', \tau) \\ U(\xi', \eta', \beta', \tau) \\ v(\xi', \eta', \beta', \tau) \end{bmatrix} \quad (2)$$

$M_{1,2}$ — the mutual reciprocal relations of the state vectors M_1 — for the potential fields, M_2 — for the rotational field.

For this part named ELECTRICAL PROCESS an important role is played by the properties of the phenomenological Green functions, especially their symmetry to the constant and variable coordinates [4], [6], [8].

CONCLUSIONS

The originality of this paper involves:

- an original form of the subject system of partial differential constitutive state equations (I), describing wide scope of real processes based on multi-component electric fields,
- an original way of the solution of the system (I), formed by the systems (VIII¹), (VIII²) and (VIII³) — for the potential fields and the system (IX) — for the rotational field,
- an original interpretation of the obtained results with respect to the conditions of the distributed parameter control theory [17—19].

We need to stress that we obtain, in the paper:

- the dynamical working space-time characteristics,
- the static working space characteristics,
- the phenomenological Green functions — within brackets {}, as working space—time transfer functions (with the possibility to obtain the static phenomenal

Green functions by the elimination of the time coordinate), giving a complete information about multi-component electric fields for the phenomenological distributed parameter control [17—19].

Comparing the above presented elements it should be noted that the original concept suggested in the panel of the phenomenological distributed parameter control for the multi-component electric fields based on the effects of a single physical phenomenon fulfils conditions of the control of the yield and quality aspects of real processes [1], [13—16], [17—19]. The fundamental role in the considerations of the multi-component electric fields as continuous media is played by the kind of the mathematical analysis presented by B. Średniawa from Institute of Physics, Jagiellonian University Cracov — Poland. An important position in the Polish literature is the book of S. Węgrzyn [17] where can be found the identification diagram for the physical coefficients of the phenomenological partial differential equations. It is indispensable to mention here the high precision of the considerations of this panel limited only by the precision of determination of the physical coefficients [17].

NOTATION

ρ_{ki} — density of the “ i ” component electric charge in the flux of convection	$\frac{C}{m^3}$
ρ_{wi} — density of the “ i ” component electric charge of the excitation phenomenon	$\frac{C}{m^3}$
ρ_i — complete density of the “ i ” component electric charge $\rho_i = \rho_{ki} + \rho_{wi}$	$\frac{C}{m^3}$
v — the externally generated field vector velocity of the electric charge in the forms of the ions or electrons for all components of the multi-component electric fields	$\frac{m}{s}$
D_{ei} — coefficient of the diffusion of the “ i ” component electric charge	$\frac{m^2}{s}$
γ_i — coefficient of the “ i ” component electric conductivity	$\frac{A}{Vm}$
ε_i — coefficient of the “ i ” component electric permeability	$\frac{C}{Vm}$
U — the total voltage of the multi-component electric fields from the external source	V
U_{di} — the voltage of the multi-component electric fields for the “ i ” component shift current and epidermis phenomenon	V
U_{ci} — the voltage of the multi-component electric fields for the “ i ” component electric conductivity current	V
U_{oi} — the voltage of the electric fields of the “ i ” component electric charge with density	V
“ ρ_i ” $\frac{C}{m^3}$	V
f_i — coefficient of the mass of the “ i ” component electric charge for multi-component electric fields	$\frac{kg}{C}$
$\hat{a}, \hat{b}$ — coefficients of the material properties of the conductor of the electric current	$\frac{Ns}{m^2}$

η_w — coefficient of the dynamic viscosity	$\frac{Ns}{m^2}$
h — the crystal pression module	$\frac{N}{m^2}$
$\bar{k}$ — the Poisson constant of the crystal	$\frac{1}{s}$
p — coefficient of the time transformation of the multi-component electric fields	$\frac{A}{m^3}$
$J_{ii}(t)$ — the external source of the "i" component electric charge undergoing to "i" processing	$\frac{W}{m^3}$
$P_{ii}(t)$ — the external source of the "i" component electric energy undergoing to "i" processing	$\frac{N}{m^3}$
$T_{ii}(t)$ — the external source of the "i" component electric momentum undergoing to "i" processing	$\frac{N}{m^3}$
$\bar{\Omega}_{(x,y,z)t}$ — the locally selected volume-time element (phenomenologically closed — every phenomenon can be treated in a separate way (around point $Z(x, y, z, t)$ of the processing of multi-component electric fields	$m^3 - s$
Ω_{Rt} — the working space volume-time of equipment with multi-component electric fields	$m^3 - s$
$(a, b, c)t$ — the dimensions of the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$, by: $x = a$, $y = b$, $z = c$	m, m, m, s
$(x, y, z)t$ — the constant space-time coordinates of the working space wolume-time of equipment with multi-component electric fields	m, m, m, s
$F_{(a,b,c)t}$ — the externally oriented surface the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}(F_{(x,y,z)t}$ for $x = a$, $y = b$, $z = c$)	$m^2 - s$
$n^{(z)}$ — the normal external suface orientation vector of the surface $F_{(a,b,c)t}$	
F_{Rt} — the externally oriented surface of the working space volume-time of equipment with multi-component electric fields Ω_{Rt}	$m^2 - a$
$n_R^{(z)}$ — the normal external surface orientation vector of the surface F_{Rt}	s
t — time	

REFERENCES

1. W. Niemiec: *On the constitutive theory of the distributed parameter modelling for the phenomenal control of the multi-component electric fields. Part I. The partial differential constitutive state equations describing the multi-component electric fields*, this number
2. C.A. Coulson, A. Jeffrey: *Fale modele matematyczne*. Warszawa: WNT, 1982
3. W. Nowacki: *Dynamiczne zagadnienia termosprężystości*. Warszawa: PWN, 1966
4. B. Średniawa: *Hydrodynamika i teoria sprężystości*. Warszawa: PWN, 1977
5. W. Niemiec: *On the constitutive distributed parameter modelling of the electric field. Part II: The solution of partial differential constitutive state equations of the electric field*. Rozprawy Elektrotechniczne, 1986, z. 4, ss. 1109—1130
6. T. Trajdos: *Matematyka dla inżynierów*. Warszawa: PWN, 1974
7. A.J.A. Morgan: *On the construction of constitutive equations for continuous media*. Archiwum Mechaniki Stosowanej, vol. 17, No. 1., 1965, pp. 145—174
8. A. Tichonow, A. Samarski: *Równania fizyki matematycznej*. Warszawa: PWN, 1963
9. L.D. Landau, E.M. Lifszyc: *Krótki kurs fizyki teoretycznej*. Tom 1. Mechanika. Elektrodynamika, Warszawa: PWN, 1986
10. W. Niemiec: *Mathematical model for distributed parameter control of multicomponent chemical processes in fluid and solid phases*, 2nd Joint ASCE/ASME Mechanics Conference, June 23—26, 1985, University of New Mexico, Albuquerque, NM, USA, Session FWA. 2

11. W. Niemiec: *On the constitutive theory of modelling and information for distributed parameter control for the continuous mass crystallization process: Part II: Complete information for phenomenal distributed parameter control of continuous mass crystallization process*. Third International Conference on Liquid Metal Engineering and Technology in Energy Production, 9—13 April, 1984, Oxford, England, No. 192
12. W. Niemiec: *The constitutive theory of modelling and information for phenomenal distributed parameter control of multicomponent chemical processes in gas, fluid and solid phase. Part II. The complete information for phenomenal distributed parameter control of multicomponent chemical processes in gas, fluid and solid phase*, 7th Miami International Conference on Alternative Energy Sources, 9—11 December, Miami Beach, Florida, USA, Hydrocarbons Energy Transfer, pp. 589—598
13. R. Feynman, R. Leighton, M. Sands: *Feynmana wykłady z fizyki*. Warszawa: PWN, 1974
14. R. Sikora: *Teoria pola elektromagnetycznego*. Warszawa: WNT, 1977
15. T. Cholewicki: *Elektrotechnika teoretyczna*. t. 1 i 2, wyd. 2, Warszawa: WNT, 1971
16. J. Turowski: *Elektrodynamika techniczna*. Warszawa: WNT, 1968
17. S. Węgrzyn: *Podstawy automatyki*. Warszawa: PWN, 1974
18. A. G. Butkovskiy: *Distributed control systems*. American Elsevier Publishing Comp., New York 1969
19. E. D. Gilles: *Zur Theorie der Regelung von Systemen mit örtlich verteilten Parametern*. Regelungstechnik, Haft 10, 1966, 14. Jahrgang, s. 461—469

W. NIEMIEC

O KONSTITUTYWNEJ TEORII ROZŁOŻONEGO PARAMETRYCZNIE MODELOWANIA DLA ZJAWISKOWEGO STEROWANIA WIELOSKŁADNIKOWYCH PÓL ELEKTRYCZNYCH

CZĘŚĆ II

FIZYKALNE WYDAJNOŚCIOWE I JAKOŚCIOWE ZAGADNIENIA STEROWANIA DLA PRZETWARZANIA WIELOSKŁADNIKOWYCH PÓL ELEKTRYCZNYCH

Streszczenie

Wyprowadzony w Części I artykuł pt. „RÓWNANIA RÓŻNICZKOWE CZĄSTKOWE KONSTITUTYWNE STANU OPISUJĄCE WIELOSKŁADNIKOWE POLA ELEKTRYCZNE” konstytutywny rozłożony parametrycznie model matematyczny opisujący rozkład zmiennych stanu w przestrzeni roboczej i czasie dla aparatów elektrycznych wieloskładnikowych pól elektrycznych został rozwiązany analitycznie z punktu widzenia pełnej informacji dla rozłożonego parametrycznie sterowania przetwarzaniem tego pola przy pomocy jego zjawisk fizycznych — zgodnie z klasyczną teorią sterowania. Rozważono dwa przypadki omawianego rozwiązania:

A — istnienie tylko warunków początkowych,

B — istnienie warunków początkowych i brzegowych,

dla zjawisk fizycznych wieloskładnikowych pól elektrycznych. W oparciu o powyższe rozważania wyprowadzono i przedstawiono oryginalną strukturę rozłożonego parametrycznie sterowania zjawiskami fizycznymi wieloskładnikowych pól elektrycznych oraz odpowiadający przetwarzaniu tych pól elektrycznych układ: WEJŚCIE → ŹRÓDŁO → PROCES ELEKTRYCZNY → WYJŚCIE.

On the local stability of uniform sampling phase — locked loops

MARIUSZ ŻÓŁTOWSKI

Instytut Telekomunikacji, Politechnika Gdańska

Received 1989.1.6

Authorized 1989.5.12

The local stability of the uniform sampling digital phase-locked loops [1] is studied in rigorous way. The nonstationary discrete-time functional equation obtained by linearization is used to describe the phase error in the loop. The anticipations of the devised methods are compared and justified by numerical results from the computer program which simulates the loop. New uniform sampling digital phase-locked loop is considered.

1. INTRODUCTION

Digital phase-locked loops (DPLL's)* are the area of interest which is motivated by the continuous progress in increasing performance, speed, reliability and the simultaneous reduction in size and cost of integrated circuits. Generally DPLL's may be referred to uniform and nonuniform sampling schemas [1]. The sampling device acts as a phase detector providing thus information about the deviation of incoming signal versus the local reference signal generated in the digitally controlled oscillator (DCO) of the loop in the case of nonuniform sampling DPLL's. On the other hand the sampling of the incoming input signal is accomplished at constant rate in the uniform sampling phase-locked loops. An uniform sampling phase-locked loop may be advantageous while the high enough sampling rate is required because of signal to noise and acquisition considerations concerning the synchronization system. The rigorous approach to the local stability of the uniform sampling phase-locked loop has not appeared. See the paper of Garodnick, Greco and Schilling for the origin of this problem [2].

2. THE MODEL OF THE UNIFORM SAMPLING PHASE-LOCKED LOOP

The derivation of the theoretical model of the uniform sampling digital phase-locked loop of fig. 1 is given in appendix 1. The basis of the considerations of this paper, the

* Presented at ECCTD 89 in BRIGHTON [8]

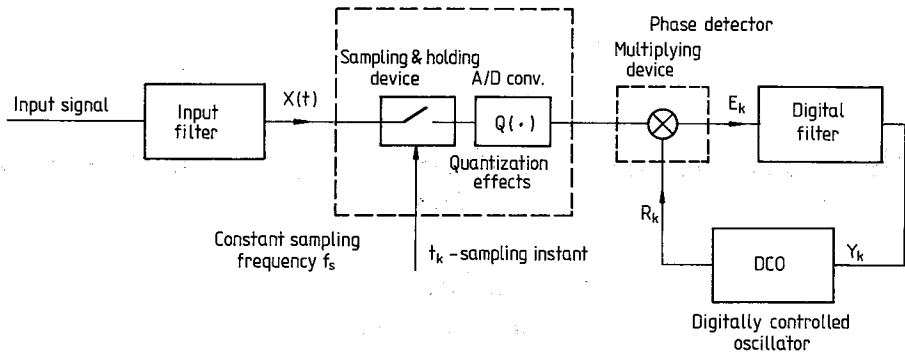


Fig. 1. The block scheme of the uniform sampling digital phase-locked loop

linearized model, is given in appendix 2. The following simplified model is used to study the problem of 1:1 phase entrainment:

- 1) The input signal $x(t)$ is detuned in phase and frequency from the local reference generated by the DCO of the loop.
- 2) No noise is present.
- 3) The effects of quantizations are negligible.

Either effect of noise or the effect of quantization are similar. They results in additional oscillations [1, 2] and are not aimed to be considered in this paper. The following linearized equation has been derived to describe the loop's dynamics;

$$\begin{aligned} \varphi(n) = & \varphi(0)k_1(n) + \Delta \sum_{i=0}^{n-1} k_1(i) + \sum_{i=0}^{n-1} k_2(n-i)\varphi(i) \cos(i\tilde{\omega}_0 + 2\Theta_0) + \\ & - \sum_{i=0}^{n-1} k_2(n-i) \sin(i\tilde{\omega}_0 + 2\Theta_0) \quad n = 0, 1, 2, 3, \dots \sum_{i=0}^{-1} () \equiv 0 \end{aligned} \quad (1)$$

The discrete time waveform $\varphi(n)$, $n = 0, 1, \dots$ of (1) is the phase error process characterizing the deviation from the state of phase entrainment. The discrete time functions $k_1(n)$ and $k_2(n)$, $n = 0, 1, \dots$ are the unit pulse responses of the linearized model of the loop. The parameter Δ depends on the frequency deviation of the tracked waveform. The first part of the right hand side of (1) represents the response of the loop due to the initial condition $\varphi(0)$ while the second one is the response due to the frequency deviation which is characterized by Δ . On the other hand the two last oscillatory components of (1) are not desired from the practical point of view. The loop should be designed to let the following conditions hold true

$$\lim_{n \rightarrow \infty} \varphi(0)k_1(n) = 0 \quad (2a)$$

and

$$\lim_{n \rightarrow \infty} \Delta \sum_{i=0}^{n-1} k_1(i) \triangleq \varphi_s \simeq 0. \quad (2b)$$

Also the oscillatory components of (1) should be suppressed i.e

$$\lim_{N \rightarrow \infty} \sup_{n > N} \left| \sum_{i=0}^{n-1} k_2(n-i) \varphi(i) \cos(i\tilde{\omega}_0 + 2\Theta_0) \right| < \varepsilon \quad (2c)$$

and

$$\lim_{N \rightarrow \infty} \sup_{n > N} \left| \sum_{i=0}^{n-1} k_2(n-i) \sin(i\tilde{\omega}_0 + 2\Theta_0) \right| < \varepsilon \quad (2d)$$

and $\varepsilon \simeq 0$.

We shall call the state of phase entrainment locally stable if these conditions are met. The parameter φ_s of (2b) is the steady state phase error due to the frequency offset. In fact the problem considered in this paper is equivalent to the question how the oscillatory terms of (1) affect the time response of the loop.

3. THE STEADY STATE OF THE PHASE-LOCKED LOOP

Let us consider the equation (1) in a slightly changed equivalent form first in aim to solve the problem of attenuation of the undesired harmonic components in the discrete-time phase-locked loop;

$$x(n) = \sum_{i=1}^{n-1} k_2(n-i) x(i) f(i) + z(n), \quad (3)$$

$$x(n) \triangleq \varphi(n) - \varphi_s, \quad (4a)$$

$$z(n) \triangleq \varphi_s \sum_{i=0}^{n-1} k_2(n-i) \cos(i\tilde{\omega}_0 + 2\Theta_0) - \sum_{i=0}^{n-1} k_2(n-i) \sin(i\tilde{\omega}_0 + 2\Theta_0) + \varphi(0) k_1(n) \quad (4b)$$

$$f(i) \triangleq \cos(i\tilde{\omega}_0 + 2\Theta_0), \quad n = 0, 1, 2, \dots \quad (4c)$$

The equation (3) may be called the discrete time Volterra equation. Next, looking for the steady state of the loop the initial value of time index i of (3) should be $i = -\infty$. Thus

$$x(n) = \sum_{i=-\infty}^{n-1} k_2(n-i) x(i) f(i) + z(n) \quad n = \dots -1, 0, 1, \dots \quad (5)$$

describes the steady state of the loop. We call the function

$$K(e^{j\tilde{\omega}}) \triangleq \sum_{l=0}^{\infty} k(l) e^{-j\tilde{\omega}l}, \quad (6a)$$

the frequency transfer function of the discrete time linear system described by the convolution

$$y(n) = \sum_{m=-\infty}^n k(n-m) x(m), \quad n = \dots -1, 0, 1, \dots \quad (6b)$$

if the series of (6a) converges; e.g. $\sum_{l=0}^{\infty} |k(l)| < \infty$, $\tilde{\omega}$ stands for the radian frequency.

Now, from the first glimpse of eye at the problem of the oscillatory terms of (5) and (4) it follows that the transfer function

$$K_2(e^{j\tilde{\omega}}) = \sum_{n=0}^{\infty} k_2(n) e^{-j\tilde{\omega}n} \quad (7)$$

should suppress every harmonic waveform of the fundamental frequency $\tilde{\omega}_0$, but not only. Also the harmonic components of frequencies being the multiplicity of $\tilde{\omega}_0$ should be suppressed. It is enough in the case of the continuous time systems if the frequency transfer function has the properties of the low pass function and the fundamental frequency is out of the effective frequency bandwidth. Then the harmonics at the frequencies being the multiplicities of the fundamental frequency are suppressed either.

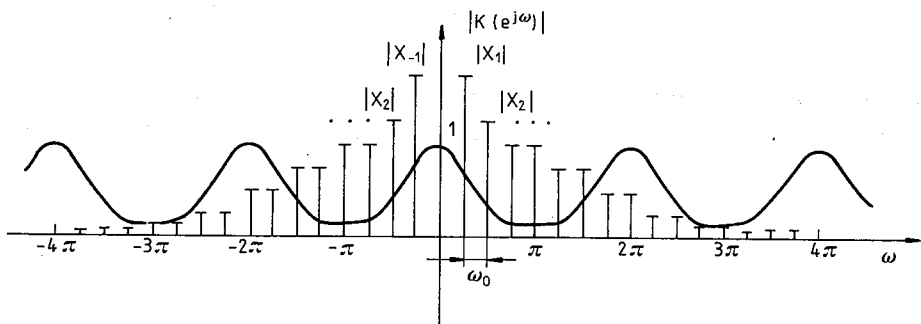


Fig. 2. The modulus of the typical low pass frequency transfer function of the discrete time system and the amplitudes of $x_h(n)$ waveform

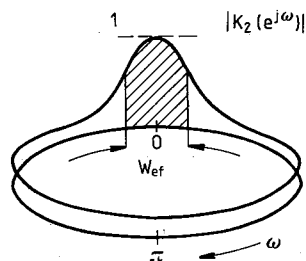
However the last statement is not true for the linear discrete time system because the frequency transfer function is periodic with 2π being a period in this case. See fig. 2 for the typical low pass frequency transfer function of the discrete time linear system and the waveform

$$x_h(n) \triangleq \sum_{m=-\infty}^{\infty} x_m e^{jm\omega_0 n} \quad x_{-m} = x_m^*, \quad n = \dots -1, 0, 1, \dots \quad (8)$$

* denotes the conjugation of complex number.

The waveform of (8) is the composition of the discrete time harmonic waveforms. One can notice that the suppression of the harmonic components of $x_h(n)$ is limited because of the periodicity of the frequency transfer function. Let us investigate this phenomenon in more details. First of all the frequency transfer function $K_2(e^{j\omega})$ of fig. 2 and its effective bandwidth may be visualized on the circle of fig. 3.

Fig. 3. The modulus of the typical low pass frequency transfer function $K_2(e^{j\omega})$ shown on the circle. The shaded region determines the effective bandwidth W_{ef}



Also, the geometrical places of the radian frequencies of the harmonic components of the waveform $x_h(n)$ can be marked in the circle. Let us consider two cases;

- $\tilde{\omega}_0/2\pi$ is a rational number fig. 4
- $\tilde{\omega}_0/2\pi$ is an irrational number fig. 5

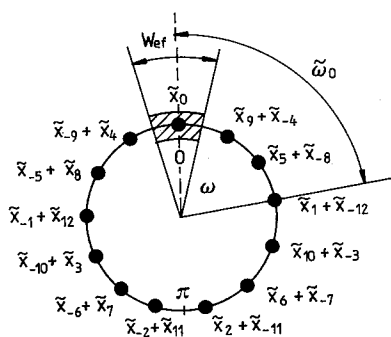


Fig. 4. The geometrical places of the radian frequencies of the harmonic components of the waveform $x_h(n)$ while $\tilde{\omega}_0/2\pi$ is a rational number ($\tilde{\omega}_0/2\pi = 3/13$)

$$N_0 = 13, k_0 = 3$$

$$\tilde{\omega}_0 = \frac{2\pi}{N_0} k_0$$

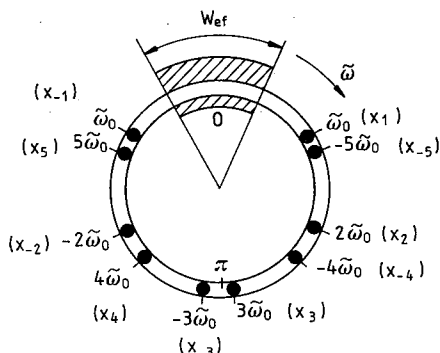


Fig. 5. The geometrical places of the radian frequencies of the harmonic components of the waveform $x_h(n)$, while $\tilde{\omega}_0/2\pi$ is an irrational number

Note that in the case a)

$\tilde{\omega}_0/2\pi = k'/N_0$, the integers k', N_0 have no divisors in common but 1, and

$$x_h(n) = \sum_{l=-N_0+1}^{N_0-1} \tilde{x}_l e^{j l \tilde{\omega}_0 n} \quad (9)$$

$$\tilde{x}_1 \triangleq \sum_{k=0}^{\infty} x_{l+kN_0}, \quad 0 < l \leq N_0 - 1 \quad (9a)$$

$$\tilde{x}_0 \triangleq \sum_{k=-\infty}^{\infty} x_{0+kN_0} \quad l = 0 \quad (9b)$$

$$\tilde{x}_l \triangleq \sum_{k=0}^{\infty} x_{l-kN_0}, \quad -N_0 + 1 \leq l < 0 \quad (9c)$$

Note also, that $\tilde{x}_l$ of (9) is well defined if

$$\sum_{m=-\infty}^{\infty} |x_m| < \infty \quad (10)$$

It is evident that only $\tilde{x}_0$ of (9) is not suppressed in a) case (see fig. 4) and will contribute to the response of the system. On the other hand the geometrical place of the all points which represent the radian frequencies of the harmonic components of $x_h(n)$ has no periodic distribution in the circle any more in b) case (see fig. 5). These points are dense in the circle due to the irrationality of $\tilde{\omega}_0/2\pi$. All harmonic components which are dense within the effective bandwidth W_{ef} will not be suppressed. Now, let us proceed with the solution of (5). Its linear form enables to look for the solution of form

$$x(n) = \sum_{m=-\infty}^{\infty} x_m e^{jm\tilde{\omega}_0 n}, \quad x_{-m} = x_m^* \quad (11)$$

$$\text{Let } z(n) \text{ and } f(n) \text{ of (5) be given by } z(n) = z_1 e^{j\tilde{\omega}_0 n} + z_{-1} e^{-j\tilde{\omega}_0 n} \quad (12a)$$

$$z_1 = \frac{1}{2} K_2(e^{j\tilde{\omega}_0}) (\varphi_s + j) e^{j2\theta_0} \quad (12b)$$

$$\begin{aligned} z_1 &= z_{-1}^* \\ f(n) &= f_1 e^{j\tilde{\omega}_0 n} + f_{-1} e^{-j\tilde{\omega}_0 n} \end{aligned} \quad (12c)$$

$$f_1 \triangleq \frac{1}{2} e^{j2\theta_0}$$

$$f_1 = f_{-1}^* \quad (12d)$$

Then the equation (5) is equivalent to the system of linear equations of infinite dimension;

$$x_r = K_r \sum_{p=-1,1} f_p x_{r-p} + z_r \quad (13)$$

or

$$x_r = K_r (x_{r+1} f_{-1} + x_{r-1} f_1) + z_r \quad (13')$$

$$K_r \triangleq K_2(e^{j\tilde{\omega}_0 r}) \quad r = \dots -1, 0, 1, \dots \quad (13a)$$

The matrix form of (13) is

$$x = Ax + z \quad (14a)$$

$$x \triangleq \text{col}(\dots x_1, x_0, x_{-1} \dots) \quad (14b)$$

$$z \triangleq \text{col}(\dots 0, z_1, z_0, z_{-1}, 0 \dots), \quad (14c)$$

$$A = \begin{bmatrix} \dots & 0 & K_3 f_1 & 0 & \dots \\ 0 & K_2 f_{-1} & 0 & K_2 f_1 & 0 & \dots \\ \dots & 0 & K_1 f_{-1} & 0 & K_1 f_1 & 0 & \dots \\ & \dots & 0 & K_0 f_{-1} & 0 & K_{-0} f_1 & 0 & \dots \\ & & \dots & 0 & K_{-1} f_{-1} & 0 & K_{-2} f_1 & 0 & \dots \\ & & & \dots & 0 & K_{-2} f_{-1} & 0 & \dots \end{bmatrix} \quad (14d)$$

Now, let us make use of Banach theorem to establish the existence and uniqueness of the solution of (14). The l^1 space (the space of any z possessing the property; $\sum_i |z_i| < \infty$) is considered to be consistent with (10).

First of all;

1) The matrix operator A of (14) acting upon any element $z \in l^1$ transforms it into an element of the same space l^1

Secondly;

2) The norm of A is

$$\|A\|_{l^1} = \sup_{i \in Z} |f_1| (|K_i| + |K_{i+2}|) \quad (15)$$

Z is the set of integers and $K_i \triangleq K_2(e^{j\tilde{\omega}_0 i})$

One should assert

$$\|A\|_{l^1} < 1 \quad (16)$$

Notice that it follows from the previous considerations that $K_2(e^{j\omega})$ should be the frequency transfer function of the low pass filter. Thus;

3) Let $K_2^\alpha(e^{j\tilde{\omega}_0})$ possess the properties (see also appendix 6)

$$1^\circ \sup_{\tilde{\omega}_0 \in [0, \pi]} |K_2^\alpha(e^{j\tilde{\omega}_0})| = 1$$

$$2^\circ |K_2^\alpha(e^{j\tilde{\omega}_0})| \text{ decreases as a function of the radian frequency } \tilde{\omega}_0 \text{ while } \tilde{\omega}_0 \in [0, \pi].$$

$$3^\circ \text{ the selectivity of } |K_2^\alpha(e^{j\tilde{\omega}_0})| \text{ increases while the parameter } \alpha \text{ tends to } 0 \text{ (the effective bandwidth diminishes).}$$

The function $K_2^\alpha(e^{j\omega})$ with the stated properties is shown in fig. 6.

Let also (see fig. 7)

$$4^\circ \tilde{\omega}_0 \in [\beta, \pi - \beta] \cup [-\beta, -\pi + \beta] \triangleq I_\beta \quad \beta > 0 \text{ and } \beta \simeq 0.$$

This assumption is not restrictive.

Then the distance between the radian frequencies of any K_i and K_{i+2} ($i \in Z$) is equal $|2\tilde{\omega}_0|$ and satisfies the relations

$$2\beta \leq |2\tilde{\omega}_0| \leq 2(\pi - \beta) \quad (17)$$

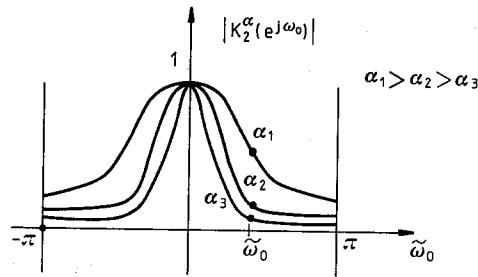


Fig. 6. The dependence of $K_2^\alpha(e^{j\omega})$ on the α parameter

Now it is obvious (see fig. 7) that the event

$$|K_i| = 1 \quad \text{and} \quad |K_{i+2}| = 1 \quad (18)$$

is not possible.

Thus while 1° to 4° hold true, then

$$\|A\|_1 = \sup_{i \in \mathbb{Z}} \frac{1}{2} (|K_i| + |K_{i+2}|) = C_{\alpha}^{\tilde{\omega}_0} \leq M < 1 \quad (19)$$

for all $\tilde{\omega}_0 \in I_\beta$ and α of interest.

Thus the norm of A is bounded uniformly and less than one. According to Banach theorem the solution x of (14) exists, is unique and belongs to l^1 space.

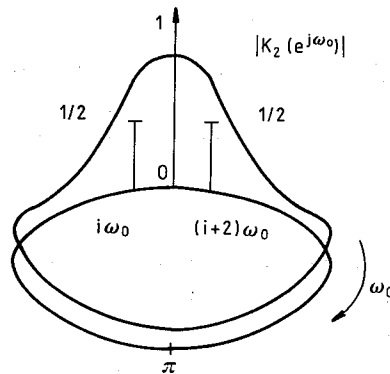


Fig. 7. The illustration of the (19) norm relation. $|K_2^\alpha(e^{j\omega_0})|$ decreases for $\omega_0 \in [0, \pi]$

4. SUPPRESSION OF HARMONIC COMPONENTS IN THE LOOP

Now, we are ready to discuss the problem of the suppression of harmonic components in the discrete time phase locked loop. From the previous chapter it follows that

$$\|x\|_{l^1} \leq \frac{\|z\|_{l^1}}{1-M}. \quad (20)$$

But (see (12))

$$\|z\|_{l^1} = 2|z_1| = |K_2^\alpha(e^{j\bar{\omega}_0})| \sqrt{\varphi_s^2 + 1} \quad (21a)$$

and according to the property 3 of the previous chapter

$$\|z\|_{l^1} \xrightarrow{\alpha \rightarrow 0} 0. \quad (21b)$$

Thus one gets the desired qualitative result; the oscillations can be reduced by any factor after proper reduction of the loop's bandwidth. Note that (21b) is not true when $\bar{\omega}_0 = 0$. The oscillations of the phase error process if of great amplitude can give rise to the cycle slipping phenomenon [6]. The reduction of the effective bandwidth of the loop is always at cost of the prolonged acquisition transients in the loop. That is why the quantitative analysis of the oscillations is also of interest.

5. THE QUANTITATIVE ANALYSIS OF OSCILLATIONS IN THE LOOP

Let us measure the oscillations in the phase-locked loop by the norm of the l^2 space;

$$\|x\|_{l^2} = \sqrt{\sum |x_i|^2} \quad (22)$$

which is related to the notion of energy.

Every solution x of (14) which is the element of l^1 space belongs obviously to l^2 space;

$$\|x\|_{l^2} \leq \sqrt{\sum |x_i|} = \|x\|_{l^1} \quad (23)$$

Two approximate methods of evaluation of the l^2 space norm of the oscillations of the phase error process in the discrete time phase-locked loop shall be presented. The first one, method 1 is based on the relation (20) while the second one, method 2 on the solution of the truncated system of equations obtained from the original system (14).

The method 1. Upper bound on the mean squared value of oscillations.

Let us consider two cases;

$$1^\circ \quad z_1 \neq 0, \quad z_{-1} = 0$$

$$2^\circ \quad z_1 = 0, \quad z_{-1} \neq 0$$

Then (see (14) and (19))

$$\|x^i\|_{l^1} \leq \frac{1}{(20) \quad 1 - C^{\bar{\omega}_0}} \|z_i\|_{l^1}, \quad i = -1, 1. \quad (24)$$

Next, by the superposition principle (the equation (14) is of the linear form) and (4a) and (12b) l^2 norm of the oscillations is approximated as follows

$$\|x\|_{l^2} \leq \sqrt{\|x^1\|_{l_1}^2 + \|x^{-1}\|_{l_1}^2 + \|\varphi_s\|^2} \leq \sqrt{\frac{1}{2(1-C^{\tilde{\omega}_0})}(|\varphi_s|^2 + 1)|K_2(e^{j\tilde{\omega}_0})| + |\varphi_s|^2} \quad (24a)$$

and

$$C^{\tilde{\omega}_0} = \sup_{\omega \in [-\pi, \pi]} \frac{1}{2} \{|K_2(e^{j\omega})| + |K_2(e^{j(\omega+2\tilde{\omega}_0)})|\} \quad (24b)$$

The method 2. Approximate mean squared value of oscillations

The truncated system of equations is solved instead of (14). Namely the following equation is considered

$$x^{2N+1} = A^{2N+1} x^{2N+1} + z^{2N+1} \quad (25a)$$

$$x^{2N+1} \triangleq \text{col}(\hat{x}_N, \dots, \hat{x}_1, \hat{x}_0, \hat{x}_{-1}, \dots, \hat{x}_{-N})_{2N+1} \quad (25b)$$

$$z^{2N+1} \triangleq \text{col}(0, \dots, 0, z_1, z_0, z_{-1}, 0, \dots, 0)_{2N+1} \quad (25c)$$

$$A^{2N+1} = \begin{bmatrix} 0 & K_N f_1 & 0 & & & & \\ K_{N-1} f_{-1} & 0 & K_{N-1} f_1 & 0 & & & \\ 0 & \dots & & & & & \\ & 0 & K_0 f_{-1} & 0 & \dots & & \\ & & & K_{-N+1} f_{-1} & 0 & K_{-N+1} f_1 & \\ & & & 0 & K_{-N} f_{-1} & 0 & \end{bmatrix}_{2N+1 \times 2N+1} \quad (25d)$$

Now, let $x(n) \triangleq \varphi(n)$ (please compare with A2-6 of Appendix 2, 4a and 12). Thus

$$z_1 \triangleq \frac{1}{2j} K_2(e^{j\tilde{\omega}_0}) e^{j2\theta_0} \quad (26a)$$

$$z_{-1} = z_1^* \quad (26b)$$

$$z_0 = \varphi_s \quad (26c)$$

$$f_1 = \frac{1}{2} e^{j2\theta_0} \quad (26e)$$

$$f_{-1} = f_1^* \quad (26c)$$

$$K_m = K(e^{j\tilde{\omega}_0 m}) \quad m = N, \dots, 1, 0, -1, \dots, -N \quad (26f)$$

The mean squared value of oscillations (l^2 norm) may be approximated according to

$$\|x\|_2 \simeq \sqrt{\sum_{i=-N}^N |\hat{x}_i|^2} \quad (27)$$

Also, the upper bound to the error of truncation Δx^{2N+1} can be evaluated (see appendix 3, (19) and (23), (9–10));

$$\|\Delta x^{2N+1}\|_1 \leq \frac{2}{1-C^{\tilde{\omega}_0}} |K_{N+1} f_1| |\hat{x}_N| \quad (28)$$

and $C^{\tilde{\omega}_0}$ is given by (24b), K_{N+1} by (13a), f_1 by (12c) and $\hat{x}_N$ is obtained from the solution of the truncated system of equations (25a). The computer time efficient algorithm for solving the truncated system of equations is presented in appendix 4.

6. THE NUMERICAL RESULTS

The original equation of the loop of the appendix 1 has been simulated by using the personal computer and the required parameters of the solution have been calculated from the experimental numerical data. The discrete-time phase locked loop of the first order of the appendix 2 has been considered with

$$AG_0 = 1; \quad D(z) = K_1 \quad (29a)$$

$$K_1(z) = \frac{z}{z-(1-K_1)}; \quad K_2(z) = \frac{K_1}{z-(1-K_1)} \quad (29b)$$

$$\varphi_s = \lim_{z \rightarrow 1} (z-1) \frac{z^{-1}}{1-z^{-1}} K_1(z) = \frac{\Delta}{K_1} \quad (29c)$$

The obtained results for comparison purpose of the presented approaches are shown in fig. 8—fig. 10. The curves obtained from the computer simulation are referred to the method 3 in these figures, $N = 5$ was assumed for method 2. Note that the parameter

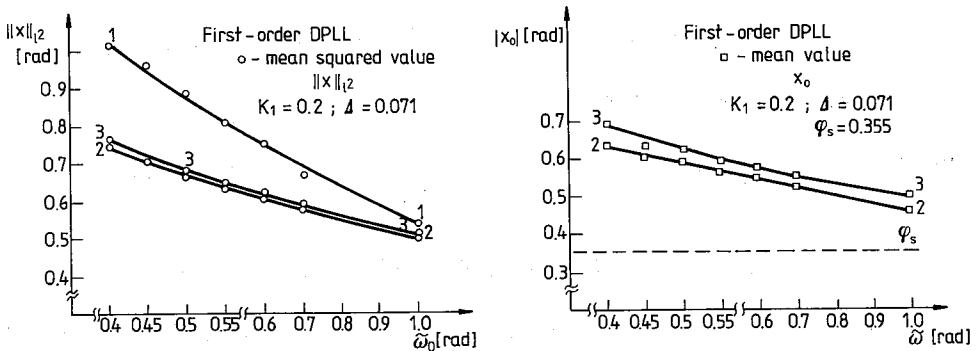


Fig. 8. The results of numerical calculations by the methods 1—3. The first-order DPLL; $K_1 = 0.2$; $\Delta = 0.071$; $\varphi_s = 0.355$

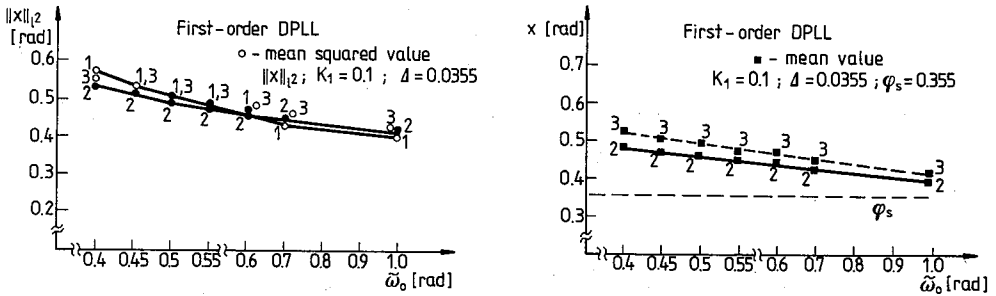


Fig. 9. The results of numerical calculations by the methods 1—3. The first-order DPLL; $K_1 = 0.1$; $\Delta = 0.0355$; $\varphi_s = 0.355$

α of chapter III which parametrize the frequency transfer function $K_2(e^{j\omega_0})$ may be chosen

$$\alpha = K_1 \quad (30)$$

while the loop of first order is considered. When $\alpha \rightarrow 0$ then the effective bandwidth of $K_2(e^{j\omega_0})$ tends to zero. The estimated mean squared value $\|x\|_2$ and mean value x_0 of oscillations are plotted versus the radian frequency $\tilde{\omega}_0$ for different values of parameters of the loop. The parameter K_1 is the most important one while concerning the presented figures. It is directly related to the loop's bandwidth while is determining by the low pass properties of $K_1(z)$. When the effective bandwidth of the loop is not small enough (see fig. 8) the method 1 which is a rough one gives overestimated results. This overestimation decreases while $\tilde{\omega}_0$ increases. On the other hand the presented results are in accord in other cases while $\alpha = K_1$ decreases to 0.

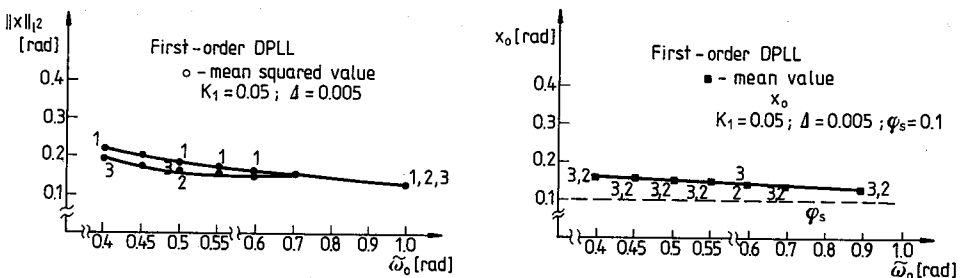


Fig. 10. The results of numerical calculations by the methods 1—3. The first-order DPLL; $K_1 = 0.05$; $\Delta = 0.005$; $\varphi_s = 0.1$

7. SOME MATHEMATICAL JUSTIFICATIONS

The problem of the evaluation of the norm of l^2 space shall be considered carefully. First of all let us state the continuity in $\tilde{\omega}_0$ of the solution of the equation (14) of infinite dimension in l^1 space;

$$\forall_{\varepsilon > 0} \exists_{\delta > 0} \forall_{\tilde{\omega}_1, \tilde{\omega}_2 \in I_\beta} |\tilde{\omega}_1 - \tilde{\omega}_2| < \delta \Rightarrow \|x^{\tilde{\omega}_1} - x^{\tilde{\omega}_2}\|_{l^1} < \varepsilon. \quad (31)$$

The proof is given in appendix 5. Secondly let “ $\sim$ ” stands for the operation of (9a—9c) from the chapter III and

$$x_l^{N_0} \stackrel{\text{df}}{=} \tilde{x}_l + \tilde{x}_{l-N_0} \quad l = 1, \dots, N_0 - 1 \quad (32)$$

$$x_0^{N_0} \stackrel{\Delta}{=} \tilde{x}_0 \quad (32b)$$

$$f_p^{N_0} \stackrel{\Delta}{=} \tilde{f}_p + \tilde{f}_{p-N_0}, \quad p = -1, 1 \quad (32c)$$

$$z_r^{N_0} \stackrel{\Delta}{=} \tilde{z}_r + \tilde{z}_{r-N_0}, \quad r = -1, 0, 1 \quad (32d)$$

while $\tilde{\omega}_0/2\pi \in I_\beta$ is a rational number; $\tilde{\omega}_0/2\pi = k'/N_0$. The integers k' and N_0 have no divisors in common but 1. Generally $a_{l \pm N_0}^{N_0} = a_l^{N_0}$ for any $a_l^{N_0}$ given by the left side of (32). Then any waveform given by (8) being equivalent to (9) is equivalent also to

$$x_h(n) = \sum_{l=0}^{N_0-1} x_l^{N_0} e^{jl\tilde{\omega}_0 n} \quad (33)$$

with $x_l^{N_0}$ given by (32a—32b).

Note, that the expression (33) defines what we call Discrete Fourier Transform DFT of time waveform. The effect of aliasing in the frequency domain is described by (9a—9c) and (32). As a matter of fact the equation (13) is equivalent now to ($\tilde{\omega}_0/2\pi$ is rational)

$$x_r^{N_0} = K, \sum_{m=0}^{N_0-1} f_{(r-m) \bmod N_0}^{N_0} x_m^{N_0} + z_r^{N_0} \quad r = 0, 1, \dots, N_0 - 1, \quad (34)$$

which is the system of equations of finite dimension N_0 . The operation of mod N_0 of (34) can be defined as follows;

Let k be the integer such that

$$l + kN_0 = l_0 \in \{0, 1, \dots, N_0 - 1\}$$

$$\text{then} \quad l \bmod N_0 \stackrel{\Delta}{=} l_0. \quad (35)$$

Also $z_r^{N_0-1}$ and $f_r^{N_0}$ of (34) are given from (32c—32d).

On the other hand the representation (8) still holds while $\tilde{\omega}_0/2\pi$ is an irrational number.

Now the problem of the evaluation of the l^2 space norm of the solution can be considered in more details. The correct expression for the l^2 norm of the oscillations is

$$\|x\|_{l^2} = \sqrt{\frac{1}{N_0} \sum_{n=0}^{N_0-1} |x(n)|^2} = \sqrt{\sum_{l=0}^{N_0-1} |x_l^{N_0}|^2} \quad (36)$$

when $\tilde{\omega}_0/2\pi$ is a rational number, and

$$\|x\|_{l^2} = \lim_{N_0 \rightarrow \infty} \sqrt{\frac{1}{N_0} \sum_{n=0}^{N_0-1} |x(n)|^2} \quad (37)$$

when $\tilde{\omega}_0/2\pi$ is irrational. The expression (36) is what we call Parseval relation for the periodic discrete time series [3]. Now, let us state the equality (37) in the frequency domain. Note that

$$\begin{aligned} \frac{1}{N_0} \sum_{n=0}^{N_0-1} |x(n)|^2 &= \frac{1}{N_0} \sum_{n=0}^{N_0-1} \sum_{l_1=-\infty}^{\infty} x_{l_1} e^{j l_1 \tilde{\omega}_0 n} \sum_{l_2=-\infty}^{\infty} x_{l_2} e^{-j l_2 \tilde{\omega}_0 n} = \\ &= \sum_{l=-\infty}^{\infty} |x_l|^2 + \frac{1}{N_0} \sum_{n=0}^{N_0-1} \sum_{\substack{l_1, l_2=-\infty \\ l_1 \neq l_2}}^{\infty} x_{l_1} x_{l_2}^* e^{j(l_1-l_2)\tilde{\omega}_0 n}. \end{aligned} \quad (38)$$

Next, it can be proved that [4]

$$\lim_{N_0 \rightarrow \infty} \frac{1}{N_0} \sum_{n=0}^{N_0-1} e^{j(m\tilde{\omega}_0)n} = \frac{1}{2\pi} \int_0^{2\pi} e^{j\tau} d\tau = 0, \quad (39)$$

while $\tilde{\omega}_0/2\pi$ is irrational. The integral of (39) is in Riemann sense. Thus

$$\|x\|_{l^2} = \sqrt{\lim_{N_0 \rightarrow \infty} \frac{2}{N_0} \sum_{n=0}^{N_0-1} |x(n)|^2} = \sqrt{\sum_{l=-\infty}^{\infty} |x_l|^2} \quad (40)$$

while $\tilde{\omega}_0/2\pi$ is irrational.

It is well known that the rational numbers are dense in the set of real numbers. One should expect that the numbers which are the values of the norms of the solutions in l^2 space while $\tilde{\omega}_0/2\pi$ is rational form dense set in the set of all numbers being the values of the norms of the solutions in l^2 space (while $\tilde{\omega}_0/2\pi$ is either rational or irrational). In other words the problem is whether solutions for rational $\tilde{\omega}_0/2\pi$ are dense in the set of all solutions in respect to l^2 space. The answer is positive. The short proof is as follows; It is well known that $N_0 \rightarrow \infty$ while a rational number $\tilde{\omega}/2\pi = k'/N_0$ is approximating an irrational number $\tilde{\omega}_0/2\pi$ [7]. But from (36), (9), (32a)

$$\frac{1}{N_0} \sum_{n=0}^{N_0-1} x_{\tilde{\omega}}^2(n) = \sum_{l=0}^{N_0-1} |x_l^{N_0}|^2 = \sum_{l=-(N_0-1)/2}^{(N_0-1)/2} |x_l^{N_0}|^2 \quad (41)$$

and N_0 odd without the loss of generality is assumed. See the remark after (32) for the definition of $x_l^{N_0}$ for negative l .

Next, from (41)

$$\begin{aligned} \lim_{\tilde{\omega} \rightarrow \tilde{\omega}_0} \|x_{\tilde{\omega}}\|_{l^2} &\rightarrow \lim_{N_0 \rightarrow \infty} \sqrt{\sum_{l=-\frac{N_0-1}{2}}^{\frac{N_0-1}{2}} |x_l^{N_0}|^2} \rightarrow \\ &\xrightarrow{(9)(32a)} \sqrt{\sum_{l=-\infty}^{\infty} |x_l|^2} \xrightarrow{(31)} \sqrt{\sum_{l=-\infty}^{\infty} |x_l^{\tilde{\omega}_0}|^2} = \|x_{\tilde{\omega}_0}\|_{l^2} \end{aligned} \quad (42)$$

The last limit is the consequence of the inequality;

$$\|x^{\tilde{\omega}} - x^{\tilde{\omega}_0}\|_{l_2} \leq \|x^{\tilde{\omega}} - x^{\tilde{\omega}_0}\|_{l_1} \quad (43)$$

Thus it is even enough to deal with the system (34) of finite dimension only and compute the mean squared value of oscillations according to (36).

8. CONCLUSIONS

The methods for studying the local stability of the uniform sampling digital phase-locked loops have been provided and the numerical results have been given which are in accord with experimental data. The complexity of analysis does not essentially depend on the order of the loop because of the convolutional form of the equation under consideration. Thus the presented approach can be easily extended for the loops of higher order by considering $K_2(e^{j\tilde{\omega}})$ of higher order only instead of the first order one. See the appendix 6 for more general $K_2(e^{j\tilde{\omega}})$. The parameter K_1 of chapter 6 is affected by the necessity of attenuation of higher frequency terms in the loop. In other case the cycle slipping phenomenon due to these terms will occur. That is why the frequency acquisition region which depends on the value of K_1 (K_1 should be small) of the uniform sampling phase-locked loop is essentially reduced in comparison to that which can be gained with the nonuniform sampling digital phase-locked loop, while the sampling frequencies both of the loops are of the same order of value [5].

While carrying out the numerical calculations of the mean squared value of oscillations it is sufficient to consider the system of equations of finite dimension. The efficient algorithm does not require very high dimensions. This seems to be similar to the problem while the continuous time system described by continuous time Volterra equation is considered. However it is not this case because of the periodicity of the frequency transfer function of the discrete time phase-locked loop. It is possible to avoid multiplication and measure the phase error from the uniformly sampled data without use of multiplier. The oscillatory terms will not be produced this way. This new loop is being considered together with the loop's filter which changes due to the adaptive algorithm [9]. Having the theoretical insight into the problem of the local stability of the uniform sampling DPLL one can deal with any other aspects within computer based modelling.

REFERENCES

1. W. C. Lindsey, C. M. Chie: *A survey of digital phase-locked loops*, Proceedings of the IEEE, vol. 69, no 4, April 1981, pp. 410—431
2. J. Garodnick, J. Greco, D. L. Schilling: *Response of an all digital phase-locked loop*, IEEE Transactions on Communication, vol. COM-22, no 6, June 1974, pp. 751—764
3. A. V. Oppenheim, R. W. Schaffer: *Cyfrowe przetwarzanie sygnałów*, WKŁ, Warszawa 1979, polish translation from english

4. W. Sierpiński: *Arytmetyka Teoretyczna*, PWN, Warszawa 1959
5. W. Szlenk: *Wstęp do teorii gładkich układów dynamicznych*, PWN, Warszawa 1980
6. M. Żółtowski: *On the comparison of the first-order digital phase-locked loops (DPLL-s) with uniform and nonuniform sampling*. Proceedings of the seventh colloquium on microwave communication, Budapest, 6—10 September 1982
7. M. Żółtowski: *On the acquisition behaviour of a first-order digital phase-locked loop with negligible and binary quantization effects*. Rozprawy Elektrotechniczne, 1984, v. 30, z. pp. 401—412
8. M. Żółtowski: *On the local stability of uniform sampling phase-locked loops*. European Conference on Circuit Theory and Design, ECCTD 89, Brighton UK, 5—8 September 1989
9. M. Żółtowski: *Adaptive uniform sampling phase-locked loop*. (in preparation)

APPENDIX 1

The model of the uniform sampling phase-locked loop of fig. 1

The detailed derivation is given in aim to allow further studies of this loop without looking for the model once more. However only the linearized model of appendix 2 is explored in this paper. The 1:1 phase entrainment is studied.

The input waveform $X(t)$ before sampling is

$$X(t) = A \sin(\omega_0 t + \Theta(t)) + N(t); \quad (\text{A1-1})$$

A stands for amplitude ($A = \sqrt{2P}$, P — power of the noiseless waveform $X(t)$), ω_0 stands for the radian frequency, $\Theta(t)$ for the phase of $X(t)$, $N(t)$ for the noise disturbance. The negligible quantization effects are assumed i.e. $Q(x) = x$. The sampling frequencies are;

$$1^\circ f_s = (\tilde{\omega}_0 / 2\pi N_s), \quad N_s > 2, N_s \text{ integer}, \quad (\text{A1-2a})$$

$$2^\circ f_s = (\tilde{\omega}_0 / 2\pi) \cdot 1/(N_s + a), \quad N_s \geq 1, \quad a \simeq 0, \quad a > 0. \quad (\text{A1-2b})$$

The parameter $a > 0$ is necessary for correct loop's behaviour. It also results from the sampling theorem for bandpass signals [2]. See (A1-8), (A1-21).

The discrete time waveform $R_k (k = 0, 1, 2, 3, \dots)$ from the DCO (see fig. 1) is a discrete time equivalent of the continuous time waveform

$$R(t) = 2 \cos(\omega_0 t + G_0 \int_{-\infty}^t y(\tau) d\tau) \quad (\text{A1-3})$$

and has the form

$$R_k = 2 \cos(\omega_0 t_k + G_0 \sum_{j=0}^{k-1} Y_j + \tilde{\varphi}_0) \quad k = 0, 1, 2, 3, \dots \quad (\text{A1-4})$$

G_0 stands for the gain of the voltage controlled oscillator or digitally controlled oscillator respectively. The controlling waveform is $y(t)$ in the continuous time case and $Y_k (k = 0, 1, 2, 3, \dots)$ in the discrete time one. The sampling instants are $t_k (k = 0, 1, 2, 3, \dots)$ and the initial phase of R_k is $\tilde{\varphi}_0$. One can set the initial point of time scale arbitrary. Thus while $t_0 = 0$ the sampling instants are

$$1^\circ \quad t_k = k2\pi/(\omega_0 N_s), \quad (\text{A1-5a})$$

$$2^\circ \quad t_k = (k2\pi/\omega_0) \cdot (N_s + a). \quad (\text{A1-5b})$$

Now, from A1-A5 the sampled input waveform and output waveform from the DCO are

$$X_k = A \sin(2\pi b k + \Theta_k) + N_k \quad (\text{A1-6})$$

$$R_k = 2 \cos(2\pi b k + \tilde{\varphi}_k) \quad (\text{A1-7})$$

$$b = \begin{cases} 1/N_s & \text{for the case } 1^\circ \text{ (A1-2a)} \\ a & \text{for the case } 2^\circ \text{ (A1-2b)} \end{cases}$$

$$\tilde{\varphi}_k = G_0 \sum_{j=0}^{k-1} Y_j + \tilde{\varphi}_0. \quad (\text{A1-8})$$

The following simplified notation has been introduced

$$X(t_k) = X_k, \quad N(t_k) = N_k, \quad \Theta(t_k) = \Theta_k, \quad k = 0, 1, 2, 3, \dots \quad (\text{A1-9})$$

After multiplication (see fig. 1)

$$E_k = X_k R_k = A \sin \varphi_k + A \sin(k4\pi b + 2\Theta_k - \varphi_k) + Z_k \quad (\text{A1-10})$$

$$\varphi_k \stackrel{\Delta}{=} \Theta_k - \hat{\varphi}_k, \quad k = 0, 1, 2, 3, \dots \quad (\text{A1-11})$$

is the phase error process

$$Z_k \stackrel{\Delta}{=} 2N_k \cos(k2\pi b + \Theta_k - \varphi_k), \quad k = 0, 1, 2, 3, \dots \quad (\text{A1-12})$$

The recursive equation may be obtained from A1-8 and A1-11

$$\varphi_{k+1} = \varphi_k + \Theta_{k+1} - \Theta_k - G_0 Y_k. \quad (\text{A1-13})$$

The samples Y_k $k = 0, 1, 2, 3, \dots$ which control the DCO of the loop appear from the DCO of the loop according to the discrete time convolution (the filter is cleared at $t_0 = 0$)

$$Y_k = \sum_{i=0}^k d_{k-i} E_i, \quad k = 0, 1, 2, 3, \dots \quad (\text{A1-14})$$

$\{d_k\}$ is the unit pulse response of the filter [4]. Laurent series of $\{d_k\}$ (or equivalently $\mathcal{Z}$ transform) is the transfer function of the filter;

$$\mathcal{Z}\{d_k\} = D(z) \quad (\text{A1-15})$$

The equation of the loop is obtained from A1-10—A1-14;

$$\begin{aligned} \varphi_{k+1} = & \varphi_k + \Theta_{k+1} - \Theta_k - AG_0 \sum_{i=0}^k d_{k-1} \sin \varphi_i + \\ & - AG_0 \sum_{i=0}^k d_{k-1} \sin(i4\pi b + 2\Theta_i - \varphi_i) - G_0 \sum_{i=0}^k d_{k-1} Z_i. \end{aligned} \quad (\text{A1-16})$$

Now, let us consider the special case of A1-16 while the input waveform $X(t)$ is detuned in phase Θ_0 and radian frequency Ω_0 versus the local reference from DCO i.e;

$$\Theta(t) = \Omega t + \Theta_0 \quad (\text{A1-17})$$

Thus

$$\Theta_k = \Delta k + \Theta_0, \quad k = 0, 1, 2, 3, \dots \quad (\text{A1-18})$$

$$A = \begin{cases} \frac{2\pi}{\omega_0 N_s} \Omega_0 & \text{for the case A1-2a} \\ \frac{2\pi(N_s + a)}{\omega_0} \Omega_0 & \text{for the case A1-2b.} \end{cases} \quad (\text{A1-19})$$

Now the equation A1-16 takes form

$$\begin{aligned} \varphi_{k+1} = & \varphi_k - AG_0 \sum_{i=0}^k d_{k-i} \sin \varphi_i + \Delta - AG_0 \sum_{i=0}^k d_{k-i} \sin(i\tilde{\omega}_0 - \varphi_i + 2\Theta_0) + \\ & - G_0 \sum_{i=0}^k d_{k-i} Z_i; \quad k = 0, 1, 2, 3, \dots \end{aligned} \quad (\text{A1-20})$$

$$\tilde{\omega}_0 = 4\pi b + 2\Delta. \quad (\text{A1-21})$$

APPENDIX 2

The linearized equation of the uniform sampling phase-locked loop

The linearized equation of the loop is derived from the equation A1-20 of Appendix 1 while:

- 1) the effects of noise are negligible
- 2) the state of the phase entrainment is characterized by the oscillations of small amplitude and small mean value that the following approximations hold true;

$$\sin x \simeq x \quad \text{and} \quad \cos x \simeq 1. \quad (\text{A2-1})$$

On the other hand the conditions (A2-1) should be met in order of proper work of the DPLL. Thus the linearized equation of the loop is obtained

$$\begin{aligned} \varphi_{k+1} \simeq & \varphi_k - AG_0 \sum_{i=0}^k d_{k-i} \varphi_i + AG_0 \sum_{i=0}^k d_{k-i} \varphi_i \cos(i\tilde{\omega}_0 + 2\Theta_0) + \\ & - AG_0 \sum_{i=0}^k d_{k-i} \sin(i\tilde{\omega}_0 + 2\Theta_0) + \Delta. \end{aligned} \quad (\text{A2-2})$$

Laurent series ($\mathcal{Z}$ transforms) of the both sides of A2-2 are

$$\Phi(z) = \frac{z\varphi_0}{z-1+AG_0D(z)} + \frac{\mathcal{Z}\{\Delta\}}{z-1+AG_0D(z)} - \frac{AG_0D(z)}{z-1+AG_0D(z)} \mathcal{Z}\{\sin(i\tilde{\omega}_0 + 2\Theta_0)\} +$$

$$+ \frac{AG_0 D(z)}{z-1+G_0 D(z)} \mathcal{Z} \{ \varphi_i \cos(i\tilde{\omega}_0 + 2\Theta_0) \} \quad (\text{A2-3})$$

$$X(z) \triangleq \mathcal{Z} \{ x_k \} \triangleq \sum_{k=0}^{\infty} x_k z^{-k}, \quad z \in C \quad C — \text{the set of complex numbers}$$

$$\Phi(z) \triangleq \mathcal{Z} \{ \varphi_k \}, \quad D(z) \triangleq \mathcal{Z} \{ d_k \}.$$

Let

$$k_1(n) \triangleq \mathcal{Z}^{-1} \{ K_1(z) \} = \mathcal{Z}^{-1} \left\{ \frac{z}{z-1+AG_0 D(z)} \right\}, \quad (\text{A2-4})$$

$$k_2(n) \triangleq \mathcal{Z}^{-1} \{ K_2(z) \} = \mathcal{Z}^{-1} \left\{ \frac{AG_0 D(z)}{z-1+AG_0 D(z)} \right\} \quad (\text{A2-5})$$

$$n = 0, 1, 2, 3, \dots$$

Then the equation A2-3 after the reciprocal $\mathcal{Z}^{-1}$ transformation takes form

$$\begin{aligned} \varphi(n) = \varphi_0 k_1(n) + \bigwedge \sum_{i=0}^{n-1} k_1(i) + \sum_{i=0}^{n-1} k_2(n-i) \varphi(i) \cos(i\tilde{\omega}_0 + 2\Theta_0) + \\ - \sum_{i=0}^{n-1} k_2(n-1) \sin(i\tilde{\omega}_0 + 2\Theta_0), \quad (\text{A2-6}) \\ \varphi(n) = \varphi_n, \quad n = 0, 1, 2, 3, \dots \end{aligned}$$

The equation A2-6 is the linearized equation of the uniform sampling phase-locked loop.

APPENDIX 3

The upper bound on the error of truncation of the infinite system of equations (14)

Let us consider the original system of equations

$$x = Ax + z \quad (\text{A3-1})$$

and the truncated one

$$x^{2N+1} = Ax^{2N+1} + z^{2N+1}. \quad (\text{A3-2})$$

Next, let us find the equation satisfied by the error of approximation

$$\Delta x^{2N+1} \triangleq x - x^{2N+1}. \quad (\text{A3-3})$$

One gets

$$\Delta x^{2N+1} = A\Delta x^{2N+1} + (A - A^{2N+1})x^{2N+1} + \Delta z^{2N+1} \quad (\text{A3-4})$$

and

$$\Delta z^{2N+1} \triangleq z - z^{2N+1} = 0. \quad (\text{A3-5})$$

Thus

$$\begin{aligned} \|\Delta x^{2N+1}\|_\mu &\leq \frac{1}{1 - \|A\|_\mu} \|(A - A^{2N+1})x^{2N+1}\|_\mu = \\ &= \frac{|K_{N+1}f_1||\hat{x}_N| + |K_{-(N+1)}f_{-1}||\hat{x}_{-N}|}{1 - \|A\|_\mu} = \frac{2}{1 - \|A\|_\mu} |K_{N+1}f_1||\hat{x}_N|. \end{aligned} \quad (\text{A3-6})$$

The $\hat{x}_N$ of A3-6 is obtained from the solution of the truncated system A3-2.

APPENDIX 4

The time efficient algorithm for solving the truncated system of equations (25)

Let

$$a_1^m \triangleq K_m f_1 \quad \text{and} \quad a_{-1}^m \triangleq K_m f_{-1} \quad (\text{A4-1})$$

$$m = -N, \dots, -1, 0, 1, \dots, N$$

Then the solution $x^{2N+1} = \text{col}(\hat{x}_N, \dots, \hat{x}_1, \hat{x}_0, \hat{x}_{-1}, \dots, \hat{x}_{-N})$ of (25) can be computed as follows

$$1) \ b_N = a_1^N$$

$$2) \ b_{n-1} = \frac{a_1^{n-1}}{1 - a_{-1}^{n-1} b_n}; \quad \text{for } n = N \text{ down to } 1$$

$$3) \ c_1 = z_1 b_1 / a_1^1, \quad c_2 = (a_{-1}^0 c_1 + z_0) \cdot b_0 / a_1^0$$

$$4) \ x_1 = (b_1 c_2^* + c_1) / (1 - b_1 b_0^*)$$

$$5) \ \hat{x}_0 = (\hat{x}_1 - c_1) / b_1$$

$$6) \ \hat{x}_n = b_n \hat{x}_{n-1}; \quad \text{for } n = 2 \text{ to } N.$$

The solution x^{2N+1} is calculated in steps 4 to 6.

APPENDIX 5

The proof of the continuity relation (statement (31))

The existence and uniqueness of the solution of (14) for all $\tilde{\omega}_0 \in I_\rho$ follows from Banach theorem because the norm of $A(\|A\|_\mu)$ is less than one. Now, let the following equations be satisfied;

$$x^{\tilde{\omega}_1} = A \tilde{\omega}_1 x^{\tilde{\omega}_1} + z^{\tilde{\omega}_1} \quad (\text{A5-1a})$$

$$x^{\tilde{\omega}_2} = A\tilde{\omega}_2 x^{\tilde{\omega}_2} + z^{\tilde{\omega}_2} \quad (\text{A5-1b})$$

then

$$\begin{aligned} \|\Delta x\|_{l^1} &\leq \frac{1}{1 - \|A\|_{l^1}} (\|A\tilde{\omega}_1 - A\tilde{\omega}_2\|_{l^1} x^{\tilde{\omega}_2} + \|\Delta z\|_{l^1}) \\ &\leq \frac{1}{1 - M} (\|(A\tilde{\omega}_1 - A\tilde{\omega}_2)x^{\tilde{\omega}_2}\|_{l^1} + \|\Delta z\|_{l^1}) \end{aligned} \quad (\text{A5-2})$$

$$\Delta x \triangleq x^{\tilde{\omega}_1} - x^{\tilde{\omega}_2}; \quad \Delta z \triangleq z^{\tilde{\omega}_1} - z^{\tilde{\omega}_2}. \quad (\text{A5-3})$$

The following statements hold true

$$\|x^{\tilde{\omega}_2}\|_{l^1} \leq \infty \quad \text{and} \quad K(e^{j\tilde{\omega}}) \text{ is continuous in } \tilde{\omega}. \quad (\text{A5-4})$$

Thus (see 12 and 14)

$$\|(A\tilde{\omega}_1 - A\tilde{\omega}_2)x^{\tilde{\omega}_2}\|_{l^1} \rightarrow 0 \quad \text{when} \quad \tilde{\omega}_1 \rightarrow \tilde{\omega}_2 \quad (\text{A5-5})$$

and

$$\|\Delta z\|_{l^1} \rightarrow 0 \quad \text{when} \quad \tilde{\omega}_1 \rightarrow \tilde{\omega}_2. \quad (\text{A5-6})$$

Eventually

$$\|\Delta x\|_{l^1} \rightarrow 0 \quad \text{when} \quad \tilde{\omega}_1 \rightarrow \tilde{\omega}_2. \quad (\text{A5-7})$$

APPENDIX 6

Let $K_2(e^{j\omega})$ be the frequency transfer function of the low pass filter. If the $|K_2(e^{j\omega})|$ is not strictly decreasing with ω another proof is required to establish that the solution of the equation (14) belongs to l^1 space. The belongingness to l^1 space makes the all results of this paper valid even in more complicated case. Another conditions which guarantee the belongingness to l^1 space will be presented below.

Let the matrix operator A of the equation (14) be partitioned as follows;

$$A = A^u + A^d \quad (\text{A6-1})$$

$$A^u \triangleq \begin{array}{ccccccc} \dots & 0 & K_3 f_1 & 0 & \dots & & \\ & \dots & 0 & K_2 f_1 & 0 & \dots & \\ & & \dots & 0 & K_1 f_1 & 0 & \dots \\ & & & \dots & 0 & K_0 f_1 & 0 \dots \\ & & & & & 0 & K_{-1} f_1 & 0 \dots \end{array} \quad (\text{A6-2})$$

$$A^d \triangleq \begin{array}{ccccccc} \dots & 0 & K_2 f_{-1} & 0 & \dots & & \\ & \dots & 0 & K_1 f_{-1} & 0 & \dots & \\ & & \dots & 0 & K_0 f_{-1} & 0 & \dots \\ & & & \dots & 0 & K_{-1} f_{-1} & 0 \dots \\ & & & & & 0 & K_{-2} f_{-1} & 0 \dots \end{array} \quad (\text{A6-3})$$

The required solution of (14) may be given by

$$x = z + Az + A^2z + \dots + A^n z + \dots \quad (\text{A6-4})$$

if the series of (A6-4) converges to the element of l^1 space. One can establish the conditions assuring the convergence of (A6-4), based on the well known properties of geometrical progression. Bearing in mind the argument of superposition it is enough to deal with

$$z = (z)_i \triangleq \text{col}(\dots, 0, z_i, 0, 0, \dots), \quad i = -1, 0, 1. \quad (\text{A6-5})$$

Let also denote

$$(a)_i \triangleq \text{col}(\dots, 0, a_i, 0, \dots) \quad (\text{A6-6})$$

and i stands for the coordinate index of the only nonzero element of $(a)_i$. Then the action of the operators A^u and A^d on $(z)_i$ can be established:

$$A^u(z)_i = (K_{i+1} f_1 z)_i \quad (\text{A6-7a})$$

$$A^d(z)_i = (K_{i-1} f_{-1} z)_{i-1}. \quad (\text{A6-7a})$$

Also

$$A^n(z)_i = (A^u + A^d)^n(z)_i = \sum_{n_1 + n_2 = n}^{2^n} A^{u_{n_1}} A^{d_{n_2}}(z)_i \quad (\text{A6-8})$$

and the sum of (A6-8) consists of 2^n different combinations of $A^{u_{n_1}} A^{d_{n_2}}(z)_i$ terms only. But

$$\begin{aligned} A^{u_{n_1}} A^{d_{n_2}}(z)_i &= A^{u_{n_1}}(K_{i-1} K_{i-2} \dots K_{i-n_2} f_{-1}^{n_2} z)_{i-n_2} = \\ &= (K_{i-1} K_{i-2} \dots K_{i-n_2} f_{-1}^{n_2} z)_{i-n_2+n_1} \dots K_{i-n_2+n_1} f_1^{n_1} z)_{i-n_2+n_1}. \end{aligned} \quad (\text{A6-9})$$

Next from the known properties of the arithmetical and the geometric mean one obtains $(|f_1| = |f_{-1}|)$:

$$\begin{aligned} \sqrt[n]{2^n |A^{u_{n_1}} A^{d_{n_2}}(z)_i|} &\leq 2 \cdot |f_1| \sqrt[n]{|z_i|} 2 \cdot \sum_{l=0}^n \frac{|K_{i-l}| + |K_{i+l}|}{n} = \\ &= \sqrt[n]{|z_i|} 2 \cdot \sum_{l=0}^n \frac{|K_{i-l}| + |K_{i+l}|}{n}. \end{aligned} \quad (\text{A6-10})$$

Notice that the upper bound of (A6-10) does not depend on the specific values of the integers n_1 and n_2 . Thus

$$\begin{aligned} \lim_{n \rightarrow \infty} \sqrt[n]{\|A^n(z)_i\|_{l^1}} &\leq \lim_{n \rightarrow \infty} \sqrt[n]{\sum_{n_1 + n_2 = n}^{2^n} |A^{u_{n_1}} A^{d_{n_2}}(z)_i|} \leq \\ &\leq \lim_{n \rightarrow \infty} \sqrt[n]{2^n \sup |A^{u_{n_1}} A^{d_{n_2}}(z)_i|} \leq \\ &\leq 2 \lim_{n \rightarrow \infty} \sqrt[n]{|z_i|} \sum_{l=0}^n \frac{|K_{i-l}| + |K_{i+l}|}{n}. \end{aligned} \quad (\text{A6-11})$$

The limit of (A6-11) can be handled as follows

$$\lim_{n \rightarrow \infty} \sqrt[n]{|z_i|} = 1 \quad (\text{A6-12})$$

and while $\tilde{\omega}_0/2\pi$ is irrational (see (39) of chapter VII) [4]

$$\lim_{n \rightarrow \infty} \sum_{i=0}^n \frac{|K_{i-1}| + |K_{i+1}|}{n} \rightarrow \frac{1}{\pi} \int_{-\pi}^{\pi} |K_2(e^{jt})| dt. \quad (\text{A6-13})$$

Thus for the irrational $\tilde{\omega}_0/2\pi$:

$$q \triangleq \lim_{n \rightarrow \infty} \sqrt[n]{\|A^n(z)\|_1} = \frac{2}{\pi} \int_{-\pi}^{\pi} |K_2(e^{jt})| dt. \quad (\text{A6-14})$$

The value $q < 1$ assures that the solution (A6-4) of the equation (14) exists and belongs to the l^1 space. The value of the parameter q is directly proportional to the measure of the region under $|K_2(e^{jt})|$ and over t axis. While $\tilde{\omega}_0/2\pi$ is rational one can always deal with the system of equations of finite dimension. However when N_0 is large ($\tilde{\omega}_0/2\pi = k'/N_0$) then the limit of (A6-13) is

$$\sum_{i=0}^{N_0-1} \frac{|K_{i-1}| + |K_{i+1}|}{N_0}. \quad (\text{A6-15})$$

While N_0 is large it can be reasonably approximated by the integral of (A6-13) which is in Riemman sense. Thus the results of this paper are also valid in the more general case of the low pass filter of the loop.

M. ŻÓŁTOWSKI

O LOKALNEJ STABILNOŚCI PĘTLI FAZOWYCH Z JEDNOSTAJNYM PRÓBKOWANIEM

Streszczenie

W pracy przeprowadzono teoretyczną analizę problemu lokalnej stabilności cyfrowych pętli fazowych z jednostajnym próbkowaniem. Przebieg błędu fazy w pętli opisano przy pomocy zlinearyzowanego funkcjonalnego równania pętli z czasem dyskretnym. Przewidywania podanych nowych metod zostały porównane w świetle wyników z numerycznej symulacji modelu pętli. Pod rozwagę poddano nową pętlę z jednostajnym próbkowaniem.

Układy dwójników zawierające sterowaną przewodność jako funkcjonalne przetworniki sygnału

ZYGMUNT WRÓBEL

*Zakład Elektrotechniki i Automatyki
Instytut Problemów Techniki, Uniwersytet Śląski, Sosnowiec*

IGOR W. GIERASIMOW

*Katedra Techniki Obliczeniowej
Leningradzki Instytut Elektrotechniki, Leningrad*

Otrzymano 1988.03.12

Autoryzowano do druku 1989.10.16

W pracy przedstawiono próbę wykorzystania otrzymanych struktur kanonicznych bi-grafów układów dwójników [37, 38], jako funkcjonalnych przetworników informacji w których sygnał wyjściowy $\{U_{wy}\}$ jest zadaną funkcją sygnału wejściowego $\{U_{we}\}$ [$U_{wy} = f(U_{we})$].

Przedstawiono ograniczenia nakładane na funkcję wymierną $R_{m/n}(x)$, aproksymującą daną zależność $y = f(x)$, oraz warunki realizowalności funkcji $R_{m/n}(x)$ w postaci admitancji wejściowej $Y_{m/n}(q)$ układu dwójnika, składającego się z liniowo sterowanych przewodności $\{G_\alpha \cdot q\}$ i niesterowanych przewodności $\{g_\beta\}$.

Przedstawiono wpływ struktury topologicznej układu dwójnika realizującego funkcję $Y_{m/n}(q)$ na wartości tworzących go elementów $y_i \in \{G_\alpha \cdot q, g_\beta\}$.

1. WSTĘP

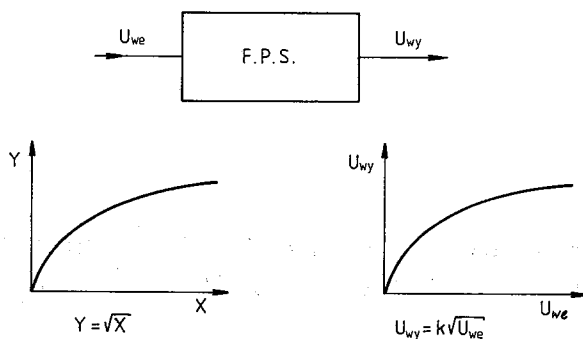
Funkcjonalne przetworniki sygnału [w skrócie F.P.S.] zwane również układami nieliniowymi [17, 23] lub generatorami funkcji [20, 25], to układy w których sygnał wyjściowy $\{y - U_{wy}\}$ jest zadaną funkcją sygnału wejściowego $\{x - U_{we}\}$ (rys. 1), co można zapisać w postaci:

$$y = f(x), \quad (1)$$

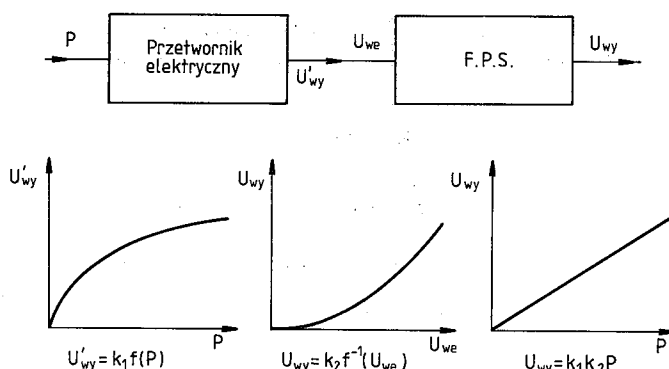
lub w postaci:

$$U_{wy} = f(U_{we}). \quad (2)$$

a)



b)



Rys. 1a. Symbol funkcjonalnego przetwornika sygnału i jego przykładowa charakterystyka

Rys. 1b. Funkcjonalny przetwornik sygnału w układzie linearyzatora statycznych charakterystyk przetworników pomiarowych

W ostatnich latach, oprócz tradycyjnych analogowych i cyfrowych układów F.P.S., w których sygnał informacyjny ma tylko jedną postać — wartości napięcia lub kodu, szybko rozwijają się nowe rodzaje układów F.P.S. w których sygnał informacyjny ma postać impulsowo-czasową [9, 10, 25, 29].

Przedstawienie sygnału informacyjnego w postaci impulsowo-czasowej zwiększa zarówno możliwości wyboru parametrów sygnału, którymi mogą być: amplituda, szerokość impulsu, częstotliwość itp., jak również wyraźnie zwiększa się odporność tego typu przetworników na zakłócenia zewnętrzne [9, 25, 29, 30, 31].

Funkcjonalne przetworniki sygnału znajdują zastosowania między innymi:

- w analogowych systemach przetwarzania informacji, jako człony operacyjne o żądanej charakterystyce transmitancyjnej [7, 8, 9, 19, 20, 22, 24],
- w systemach sterowania, jako człony operacyjne o zadanej a priori charakterystyce transmitancyjnej [1, 7, 21],
- w systemach pomiarowych (między innymi w układach linearyzacji nielinio-

wych charakterystyk przetworników pomiarowych), jako człony operacyjne o tak dobranej charakterystyce transmitancyjnej aby system pomiarowy był liniowy [9, 14, 18, 24, 25, 30, 31,].

— w cyfrowych systemach przetwarzania danych, jako nieliniowe przetworniki A/C, C/A, C/C [1, 15].

Jak wynika z wyżej przytoczonych przykładów zastosowań układów F.P.S. we wszystkich tych systemach istnieje konieczność budowy członu nieliniowego, który mógłby realizować a priori zadaną zależność funkcyjną $y = f(x)$. (rys. 1.).

Ponieważ nie każda zależność funkcyjna $y = f(x)$ może być zrealizowana przy pomocy elementów o odpowiedniej charakterystyce, jak to ma miejsce np. dla funkcji logarytmicznej — gdzie w elektronicznych układach logarytmujących wykorzystuje się eksponentyjalną zależność prądu bazy od napięcia baza emiter [19]. Istnieje więc konieczność przybliżenia (aproksymacji) zadanej analitycznie lub empirycznie funkcji $f(x)$ inną „prostsza” funkcją $F(x)$ możliwą do zrealizowania przy pomocy znanych elementów elektronicznych.

Ze względu na sposób aproksymacji analitycznie lub empirycznie zadanej funkcji $f(x)$, można wyróżnić w zasadzie dwa podejścia do syntezy układów F.P.S. składających się ze skupionych liniowych skończonych stacjonarnych i odwracalnych [słssso] elementów. Historycznie pierwszy kierunek syntezy układów nieliniowych, oparty jest na aproksymacji zadanej funkcji $f(x)$ — funkcją odcinkową. Można tutaj wyróżnić [4, 25]:

- układy nieliniowe realizujące funkcję odcinkowo-liniową,
- układy nieliniowe realizujące funkcję odcinkowo-nieliniową,
- układy nieliniowe realizujące funkcję odcinkowo-skokową.

Zaletą metody syntezy wykorzystującej funkcje odcinkowe, jest prostota aproksymacji (szczególnie dla funkcji odcinkowo-nieliniowej i odcinkowo-skokowej), oraz stosunkowo niski koszt realizacji układów [4, 25]. Wadą takich układów jest zazwyczaj mała dokładność realizowanych charakterystyk oraz duża wrażliwość ich parametrów na zmiany wartości ich elementów.

Drugi intensywnie obecnie rozwijany kierunek syntezy układów nieliniowych, oparty jest na aproksymacji zadanej funkcji $f(x)$ — funkcją ciągłą, jednej lub wielu zmiennych [9, 25, 29]. Można tutaj wyróżnić:

— układy nieliniowe realizujące funkcję ciągłą w postaci wielomianu (najczęściej wykładniczego lub potęgowego),

— układy nieliniowe realizujące funkcję ciągłą w postaci — funkcji wymiernej. Podejście to ma szereg zalet w stosunku do poprzedniego. Zapewnia zazwyczaj bardzo wysoką dokładność realizacji charakterystyk nieliniowych, ułatwia analizę ich wrażliwości. Wadą metody jest wyższy zazwyczaj koszt realizacji układu.

Podstawowymi elementami F.P.S. są zazwyczaj elementy których wartość jest w większości przypadków liniową funkcją sygnału sterującego. W tradycyjnym ujęciu są to elementy typu $G(x)$, $R(x)$, $C(x)$, $L(x)$ sterowane napięciem lub prądem [4]. Obecnie szybko rozwijają się nowe typy elementów sterowanych w których sygnał informacyjny $\{x\}$ ma postać impulsowo-czasową [9, 10, 11, 25, 29, 32].

Celem niniejszej pracy, jest przedstawienie możliwości wykorzystania układów dwójników składających się z dwóch rodzajów [slsso] elementów w postaci tzw. sterowanej przewodności $\{G_{\alpha}q\}$ i stacjonarnej przewodności $\{g_{\beta}\}$, jako funkcjonalnych przetworników sygnału.

2. STEROWANA PRZEWODNOŚĆ

Jednym z podstawowych elementów wykonawczych w czasowo-impulsowych układach przetworników sygnału jest sterowana przewodność. Sterowaną przewodność w zależności od formy sygnału informacyjnego, można podzielić na trzy podstawowe grupy:

- impulsowo sterowaną przewodność,
- kodowo sterowaną przewodność,
- impulsowo-kodowo sterowaną przewodność.

Niżej omówiono podstawowe parametry poszczególnych rodzajów sterowanej przewodności.

2.1. IMPULSOWO STEROWANA PRZEWODNOŚĆ

Sygnał informacyjny dla układu impulsowo sterowanej przewodności ma postać ciągu impulsów prostokątnych w okresie T i czasie trwania t_x zależnym od wartości sygnału wejściowego. Postać sygnału informacyjnego pokazano na rysunku 2.a.

Schemat ideowy impulsowo sterowanej przewodności pokazano na rysunku 2.b. Podając na klucz elektroniczny K impulsy prostokątne o okresie T i czasie trwania t_x , średnią wartość przewodności $g(t)$ można zapisać równaniem:

$$g(t) = \begin{cases} \frac{1}{R}; & \text{dla } k \cdot T < t < k \cdot T + t_x \\ 0; & \text{dla } k \cdot T + t_x < t < (k+1) \cdot T \end{cases} \quad (3)$$

gdzie:

$$k = 1, 2, \dots, \infty$$

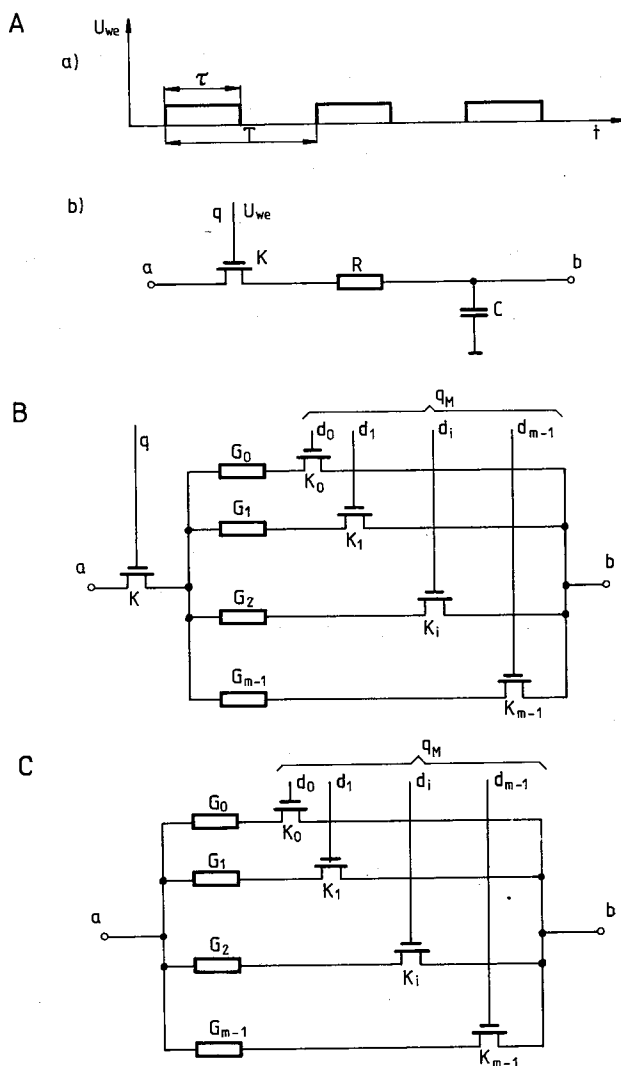
Periodyczne zmiany przewodności elementu można analitycznie przedstawić rozwijając w szereg Fouriera, w postaci:

$$g(t) = G \frac{t_x}{T} + \sum_{k=0}^{\infty} \frac{2G}{k} \cdot \sin\left(\frac{2k\pi}{T} \cdot t\right) \cdot \cos\left(\frac{2k\pi}{T} \cdot t\right) \quad (4)$$

gdzie:

$$G = \frac{1}{R}$$

Jak wynika z powyższego wzoru, średnią za okres T wartość przewodności $\bar{G}$ elementu można zapisać wzorem:



Rys. 2. Schematy ideowe: A) impulsowo sterowanej przewodności, B) Kodowo sterowanej przewodności, C) impulsowo i kodowo sterowanej przewodności

$$\bar{G} = G \cdot q \quad (5)$$

gdzie:

$$q = \frac{t_x}{T}. \quad (6)$$

Włączenie do układu szeregowo połączonych rezystora R i klucza elektronicznego K , kondensatora C powoduje odfiltrowanie wyższych harmonicznych napięcia. Element taki nazywamy liniową impulsowo sterowaną przewodnością, ponieważ

średnia wartość przewodności elementu jest liniową funkcją parametru informacyjnego q , sterującego kluczem elektronicznym [10, 11, 29, 33].

2.2. KODOWO STEROWANA PRZEWODNOŚĆ

Sygnal informacyjny dla układu kodowo sterowanej przewodności ma postać kodu cyfrowego $M(d_0, d_1, \dots, d_{m-1})$ gdzie $d_i \in (0, 1)$. Schemat ideowy kodowo sterowanej przewodności pokazano na rysunku 2.c.

Podając na klucze elektroniczne K_i sygnał informacyjny w postaci kodu dwójkowego, przewodności poszczególnych elementów układu mają postać:

$$G(d_i) = \begin{cases} G_i; & \text{jeśli } d_i = 1 \\ 0; & \text{jeśli } d_i = 0 \end{cases} \quad (7)$$

gdzie:

$$i = 0, 1, \dots, m-1.$$

Dla układu kodowo sterowanej przewodności (rys. 2c) przewodność całkowita wyraża się wzorem:

$$G(M) = \sum_{i=0}^{m-1} G(d_i) \cdot d_i. \quad (8)$$

Jeśli:

$$G(d_i) = G_0 \cdot 2^i \quad (9)$$

to

$$G(M) = G_0 \sum_{i=0}^{m-1} 2^i \cdot d_i = G_{\max} \cdot \frac{M(d_i)}{M_{\max}} \quad (10)$$

gdzie:

$G(M)$ — wartość przewodności układu
 G_0 — waga przewodności minimalnej
 $M(d_i)$ — słowo kodu cyfrowego

$$G_{\max} = G_0 \sum_{i=0}^{m-1} 2^i = G_0 \cdot (2^m - 1). \quad (11)$$

Powyższą zależność można również zapisać w postaci:

$$G(M) = G_{\max} \cdot q_M; \quad (12)$$

$$q_M = \frac{M(d_i)}{M_{\max}}. \quad (13)$$

2.3. IMPULSOWO-KODOWO STEROWANA PRZEWODNOŚĆ

Sygnal informacyjny dla impulsowo-kodowo sterowanej przewodności jest kompilacją sygnału w postaci impulsu prostokątnego o okresie T i czasie trwania t_x ,

sygnału w postaci kodu cyfrowego $M(d_0, d_1, \dots, d_{m-1})$ gdzie $d_i \in (0,1)$. Schemat ideowy impulsowo-kodowo sterowanej przewodności pokazano na rysunku 2.d.

Jeśli częstotliwość zmian kodu cyfrowego M sterującymi kluczami elektronicznymi K_i jest wielokrotnie większa od częstotliwości komutacji klucza K (co najmniej o dwa rzędy wielkości) to średnia wartość przewodności układu prezentowanego na rys. 2.d, może być zapisana w postaci:

$$G(M, q) = \frac{t_x}{T} \cdot G_{\max} \cdot \frac{M(d_i)}{M_{\max}}. \quad (14)$$

Powyższą zależność można w skrócie zapisać w postaci:

$$G(q, q_M) = G_{\max} \cdot q \cdot q_M. \quad (15)$$

3. ZBIORY STRUKTUR KANONICZNYCH BI-GRAFÓW UKŁADÓW DWÓJNIKÓW SKŁADAJĄCYCH SIĘ Z DWÓCH RODZAJÓW [SLSO] ELEMENTÓW

Jak już wspomniano we wstępie w wielu dziedzinach techniki pojawia się następujący problem: jak mając *a priori* zadaną nieliniową zależność między sygnałem wyjściowym a wejściowym, można zbudować układ elektroniczny realizując z pewnym przybliżeniem zadaną zależność funkcyjną $y = f(x)$.

Jak pokazano w pracy [37, 38] admitancja wejściowa $Y_{m/n}(t_1, t_2)$ dwójnika składającego się z dwóch rodzajów [soso] elementów jest funkcją wymierną parametru $\{t_1\}$ lub $\{t_2\}$.

$$Y(t_1, t_2) = \frac{1}{Z(t_1, t_2)} = \frac{\sum_{n_{11}=0}^{n_{11}=v-1} A_{n_{11}, n_{21}} \cdot t_1^{n_{11}} \cdot t_2^{n_{21}}}{\sum_{n_{12}=0}^{n_{12}=v-2} B_{n_{12}, n_{22}} \cdot t_1^{n_{12}} \cdot t_2^{n_{22}}}. \quad (16)$$

Jeśli:

$$Y_{1i} t_1 = G_\alpha \cdot q, \quad (17)$$

$$Y_{2j} t_2 = g_\beta, \quad (18)$$

to wyrażenie (16) można zapisać w postaci:

$$Y_{m/n}(q) = \frac{1}{Z_{m/n}(q)} = \frac{\sum_{i=0}^m A_i \cdot q^i}{\sum_{j=0}^n B_j \cdot q^j}. \quad (19)$$

Z pracy [37] wynika, że funkcja wymierna $Y_{m/n}(q)$ musi spełniać jeden z czterech niżej wymienionych warunków:

$$1: m = v-1, \quad a_i > 0, \quad \text{dla } i = 0, 1, \dots, m \quad (20a)$$

$$n = v-2, \quad b_j > 0, \quad \text{dla } j = 0, 1, \dots, n \quad (20b)$$

$$2: m = v-1, \quad a_i > 0, \quad \text{dla } i = 1, 2, \dots, m \quad (a_0 = 0) \quad (21a)$$

$$n = v-2, \quad b_j > 0, \quad \text{dla } j = 0, 1, \dots, n \quad (21b)$$

$$3: m = v-2, \quad a_i > 0, \quad \text{dla } i = 0, 1, \dots, m-1 \quad (a_m = 0) \quad (22a)$$

$$n = v-2, \quad b_j > 0, \quad \text{dla } j = 0, 1, \dots, n \quad (22b)$$

$$4: m = v-2, \quad a_i > 0, \quad \text{dla } i = 1, 2, \dots, m-1 \quad (a_0 = 0, a_m = 0) \quad (23a)$$

$$n = v-2, \quad b_j > 0, \quad \text{dla } j = 0, 1, \dots, n \quad (23b)$$

gdzie: v — liczba wierzchołków w układach dwójnika.

Wynika stąd, że aby układ dwójnika składającego się z dwóch rodzajów [slsso] elementów w postaci elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$, którego admitancja wejściowa $Y_{m/n}(q)$ jest funkcją wymierną parametru sterującego $\{q\}$, mógł realizować a priori zadaną analitycznie lub empirycznie funkcję $f(x)$, należy aproksymować tę zależność funkcją wymierną typu:

$$R_{m/n}(x) = \frac{\sum_{i=0}^m a_i x^i}{\sum_{j=0}^n b_j x^j} \quad (24)$$

4. APROKSYMACJA FUNKCJI $f(x)$ FUNKCJĄ WYMIERNĄ $R_{m/n}(x)$

Zagadnienie najlepszego przybliżenia z żadaną dokładnością, zadanej analitycznie lub empirycznie funkcji $y = f(x)$ — funkcją wymierną $R_{m/n}(x)$ jest „czysto” matematycznym zagadnieniem przedstawionym na przykład w pracy [5, 6], którego bliższe rozpatrzenie wychodzi poza ramy niniejszej pracy. Niżej omówiono tylko pewne ograniczenia nakładane na funkcję wymierną $R_{m/n}(x)$ wynikające z fizycznej realizacji układów dwójników.

Pierwsze ograniczenie nakładane na funkcję wymierną $R_{m/n}(x)$ wynika z faktu, że aby układ dwójnika składający się z elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$, mógł realizować z pewnym przybliżeniem daną zależność funkcyjną $y = f(x)$, aproksymującą ją funkcja wymierna $R_{m/n}(x)$ musi mieć analogiczną postać do jednego z czterech typów funkcji wymiernej $Y_{m/n}(q)$ jakie można realizować w postaci admitancji układ dwójnika (zależności 20÷23). W tabeli 1 zestawiono przykłady funkcji wymiernych $R_{m/n}(x)$ aproksymujących wybrane analityczne funkcje $f(x)$, i spełniające powyższe ograniczenia (20÷23).

Drugie ograniczenie nakładane na funkcję wymierną $R_{m/n}(x)$ wynika z faktu, że z jednej strony przybliżając daną funkcję (np. $y = \sqrt{x}$ w przedziale jej zmienności $y \in (0 - \infty)$) trudno spodziewać się małego błędu w całym przedziale. Z drugiej zaś

Tabela 1

Aproksymacja funkcji $f(x)$ funkcją wymierną $R_{m/n}(x)$

$f(x)$	x	m/n	$R_{m/n}(x)$	$[\varepsilon]\%$
$\sqrt{x}$	$0,1 \div 1$	$2/1$	$\frac{0.41x^2 + 0.867x + 0.041}{x + 0.316}$	$1.2 \cdot 10^{-1}$
	$0,1 \div 1$	$2/2$	$\frac{3.046x^2 + 2.007x + 0.056}{x^2 + 3.569x + 0.542}$	$2 \cdot 10^{-2}$
	$0,1 \div 1$	$2/2$	$\frac{3.1314x^2 + 2.1883x + 6.4905 \cdot 10^{-2}}{x^2 + 3.7771x + 6.0824 \cdot 10^{-1}}$	$1.3 \cdot 10^{-2}$
	$0,1 \div 1$	$3/2$	$\frac{2.6662 \cdot 10^{-1}x^3 + 1.5165x^2 + 5.0703 \cdot 10^{-1}x + 9.9136 \cdot 10^{-3}}{x^2 + 1.1884x + 1.1163 \cdot 10^{-1}}$	$1.8 \cdot 10^{-3}$
	$0,1 \div 1$	$4/3$	$\frac{2.0096 \cdot 10^{-1}x^4 + 2.1306x^3 + 1.8449x^2 + 2.3289 \cdot 10^{-1}x + 2.3920 \cdot 10^{-3}}{x^3 + 2.5367x^2 + 8.3892 \cdot 10^{-1}x + 3.6095 \cdot 10^{-2}}$	$3.3 \cdot 10^{-5}$
	$0,5 \div 1$	$4/3$	$\frac{0.14735x^4 + 2.9629x^3 + 5.2748x^2 + 1.4898x + 3.7251 \cdot 10^{-2}}{x^3 + 5.0059x^2 + 3.5497x + 0.35654}$	$5 \cdot 10^{-9}$
$\operatorname{arcsh} \sqrt{x}$	$0,1 \div 1$	$2/2$	$\frac{0.1444 \cdot 10^1 x^2 + 0.1916x + 0.6412 \cdot 10^{-4}}{x^2 + 0.852x + 0.1371 \cdot 10^{-1}}$	$7 \cdot 10^{-3}$
$\operatorname{sh} \sqrt{x}$	$0,1 \div 1$	$2/2$	$\frac{0.3643 \cdot 10^1 x^2 + 0.7020x + 0.392 \cdot 10^{-3}}{x^2 + 0.566 \cdot 10^{-1}x + 0.266 \cdot 10^1}$	$3 \cdot 10^{-3}$
$\operatorname{tg} \frac{\pi}{4} \sqrt{x}$	$0,1 \div 1$	$2/2$	$\frac{0.578 \cdot 10^1 x^2 + 0.135 \cdot 10^1 x + 0.952 \cdot 10^{-3}}{x^2 + 0.603 \cdot 10^1 x + 0.1483}$	$2 \cdot 10^{-2}$

strony parametr sterujący $\{q\}$ może zmieniać się tylko w ściśle określonych granicach wynikających z fizycznej realizowalności, co determinowało wybór przedziałów zmiennej x (tabela 1).

5. WARUNKI REALIZOWALNOŚCI FUNKCJI $R_{m/n}(x)$ W POSTACI ADMITANCJI WEJŚCIOWEJ $Y_{m/n}(q)$ UKŁADU DWÓJNIKA

Warunkiem koniecznym realizowalności danej funkcji $y = f(x)$ aproksymowanej funkcją wymierną $R_{m/n}(x)$, poprzez układ dwójnika składającego się z elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$ jest, aby jego admitancja wejściowa $Y_{m/n}(q)$ (wzór 19) miała analogiczną postać jak funkcja $R_{m/n}(x)$ (wzór 24), co można zapisać w postaci następującego układu równań:

$$a_i = h \cdot A_i \quad (25)$$

$$b_j = h \cdot B_j \quad (26)$$

$$x = h \cdot q \quad (27)$$

gdzie:

h — współczynnik proporcjonalności.

Warunkiem dostatecznym realizowalności danej funkcji $R_{m/n}(x)$ w postaci admitancja wejściowa $Y_{m/n}(q)$ układu dwójnika, jest istnienie rozwiązania powyższego układu równań nieliniowych. Niżej na przykładzie analitycznie zadanej funkcji $y = \sqrt{x} : x \in (0.1 \div 1)$, zilustrowano wypełnienie warunków koniecznych i dostatecznych realizowalności danej funkcji $y = \sqrt{x}$ aproksymowanej funkcją wymierną $R_{3/2}(x)$ w postaci admitancji wejściowej $Y_{3/2}(q)$ układu dwójnika składającego się z elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$.

Jak wynika z tabeli 1 funkcję $y = \sqrt{x}$ można aproksymować funkcją wymierną, na przykład w postaci:

$$\sqrt{x} \approx \frac{2.6662 \cdot 10^{-1} \cdot x^3 + 1.5165 \cdot x^2 + 5.0703 \cdot x + 9.9136 \cdot 10^{-3}}{x^2 + 1.1884 \cdot x + 1.1163 \cdot 10^{-1}} \quad (28)$$

Z warunków (25 ÷ 27) wynika, że admitancja wejściowa układu dwójnika składającego się z elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$ musi mieć postać:

$$Y_{3/2}(q) = \frac{A_3 \cdot q^3 + A_2 \cdot q^2 + A_1 \cdot q^1 + A_0 \cdot q^0}{B_2 \cdot q^2 + B_1 \cdot q^1 + B_0 \cdot q^0}, \quad (29)$$

czyli zgodnie z twierdzeniem 1.1 z pracy [37], układ dwójnika musi być sumą logiczną dendrytu 1-drzewowego składającego się tylko z elementów sterowanych $\{G_\alpha \cdot q\}$ i dendrytu 1-drzewowego składającego się tylko z elementów niesterowanych $\{g_\beta\}$.

Z pracy [38] wynika, że istnieją 32 klasy nierównoważnych między sobą struktur kanonicznych bi-grafów układów dwójników składających się z elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$ których admitancja wejściowa $Y_{3/2}(q)$ ma postać funkcji wymiernej opisanej wzorem 29.

Przykład układu dwójnika składającego się z trzech elementów sterowanych $\{G_\alpha \cdot q\}$ i trzech elementów niesterowanych $\{g_\beta\}$ którego admitancja wejściowa ma postać funkcji $Y_{3/2}(q)$, pokazano na rysunku 3. Dla prezentowanego układu dwójnika, korzystając z tabeli na rysunku gdzie zestawiono wszystkie dendryty 1- i 2-drzewowe istniejące w danym układzie, sformułowane wyżej warunki realizowalności można zapisać w postaci:

$$2.6662 \cdot 10^{-1} = h \cdot (G_1 \cdot G_2 \cdot G_5) \quad (30)$$

$$1.6165 = h \cdot (g_1 \cdot G_2 \cdot G_5 + g_3 \cdot G_1 \cdot G_5 + g_3 \cdot G_2 \cdot G_5 + g_6 \cdot G_1 \cdot G_2 + g_6 \cdot G_2 \cdot G_5) \quad (31)$$

$$5.0703 = h \cdot (g_1 \cdot g_3 \cdot G_5 + g_1 \cdot g_6 \cdot G_2 + g_3 \cdot g_6 \cdot G_1 + g_3 \cdot g_6 \cdot G_2 + g_3 \cdot g_6 \cdot G_5) \quad (32)$$

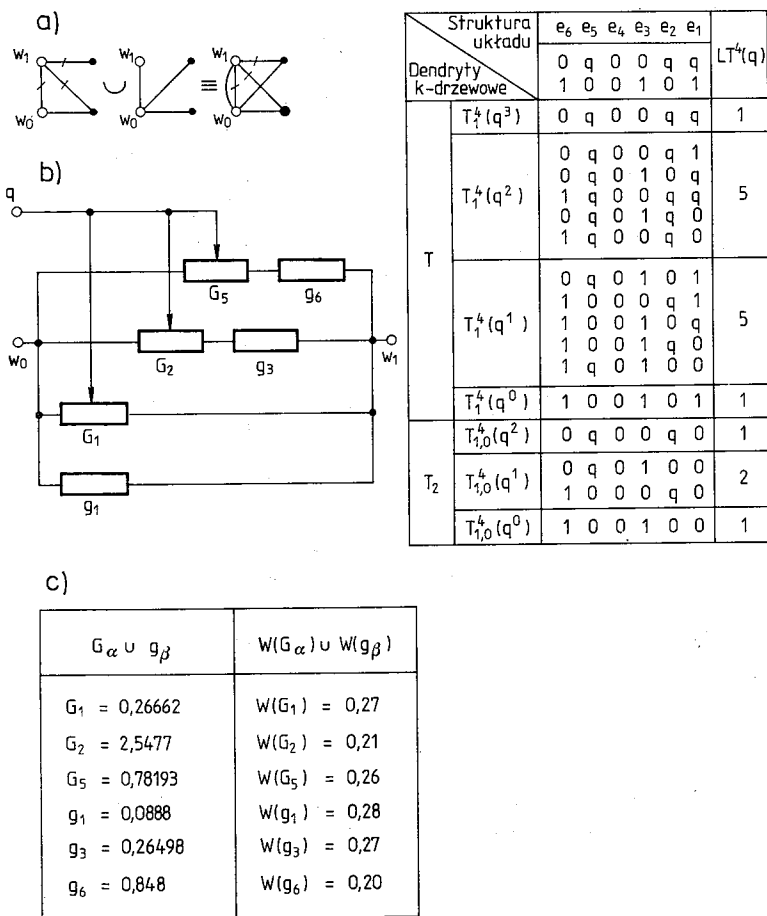
$$9.9136 \cdot 10^{-3} = h \cdot (g_1 \cdot g_3 \cdot g_6) \quad (33)$$

$$1.0000 = h \cdot (G_2 \cdot G_5) \quad (34)$$

$$1.1884 = h \cdot (g_3 \cdot G_5 + g_6 \cdot G_2) \quad (35)$$

$$1.1163 \cdot 10^{-1} = h \cdot (g_3 \cdot g_6). \quad (36)$$

Rozwiązując powyższy układ równań nieliniowych [9] otrzymano wartości elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$, które zestawiono w tabeli na rysunku 3.



Rys. 3. Przykład układu dwójnika zawierającego sterowane i niesterowane przewodności, a) jego bi-graf, b) schemat ideowy, c) unormowane wartości G_α i g_β oraz wrażliwość względne $W(G_\alpha)$ i $W(g_\beta)$ dwójnika realizującego funkcję $Y_{3/2}(q) \sim \sqrt{q}$, tabela — dendryty k-drzewowe konieczne do określenia parametrów dwójnika

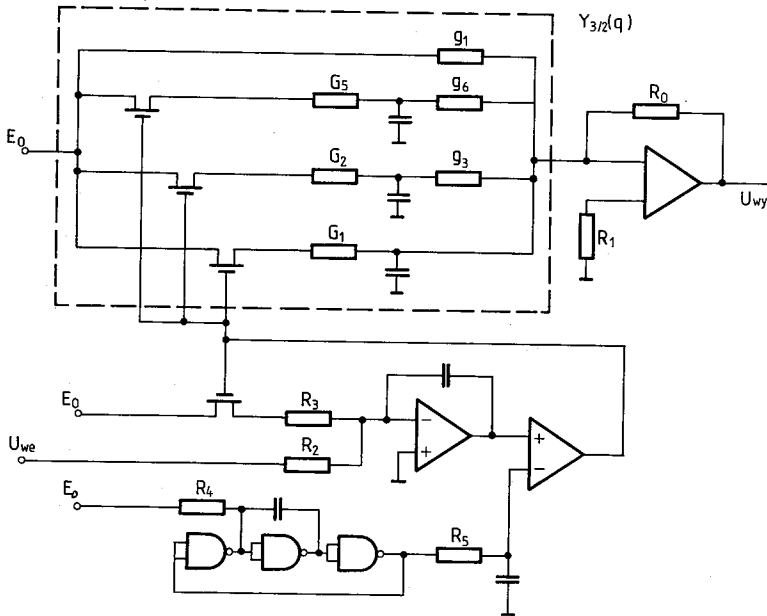
6. SCHEMAT IDEOWY FUNKCJONALNEGO PRZETWORNIKA SYGNAŁU

Schemat ideowy F.P.S. zawierającego człon nieliniowy w postaci układu dwójnika składającego się z elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$ pokazano na rysunku 4.

Sygnał na wyjściu wzmacniacza operacyjnego można zapisać w postaci:

$$U_{wy} = -R_o \cdot E_o \cdot Y_{3/2}(q). \quad (37)$$

Człon nieliniowy w postaci dwójnika o transmitancji $Y_{3/2}(q)$ jest sterowany z modulatora szerokości impulsu, który przetwarza sygnał wejściowy U_{we} na szerokość impulsu q sterującego tzw. przewodność sterowaną $\{G_\alpha \cdot q\}$ [29]. Zależność



Rys. 4. Schemat ideowy funkcjonalnego przetwornika sygnału realizującego funkcję $U_{wy} = p\sqrt{U_{we}}$

parametru sterującego q od napięcia wejściowego U_{we} modulatora szerokości impulsu, można wyznaczyć korzystając z zasady zachowania ładunku. Otrzymujemy wtedy:

$$\frac{U_{we}}{R_2} \cdot T - \frac{E_0}{R_3} t_x = 0. \quad (38)$$

Stąd korzystając z zależności (6) otrzymujemy:

$$q = \frac{R_3}{R_2 \cdot E_0} \cdot U_{we}. \quad (39)$$

Podstawiając otrzymaną zależność do wzoru (37) otrzymujemy:

$$U_{wy} = -R_0 \cdot E_0 \cdot Y_{3/2} \left(\frac{R_3}{R_2 \cdot E_0} \cdot U_{we} \right). \quad (40)$$

Ponieważ transmitancja układu dwójnika jako członu nieliniowego jest funkcją wymierną $Y_{3/2}(q)$ parametru q , stąd sygnał wyjściowy układu F.P.S. (rys. 4) korzystając z zależności (19) można zapisać w postaci:

$$U_{wy} = -R_0 \cdot E_0 \cdot \frac{\sum_{i=0}^m A_i \cdot \left(\frac{R_3}{R_2 \cdot E_0} \cdot U_{we} \right)^i}{\sum_{j=0}^n B_j \cdot \left(\frac{R_3}{R_2 \cdot E_0} \cdot U_{we} \right)^j}. \quad (41)$$

Ponieważ wyrażenia (28) i (41) mają analogiczną postać, wynika stąd że sygnał wyjściowy U_{wy} układu F.P.S. pokazanego na rys. 4 jest funkcją pierwiastkową sygnału wejściowego U_{we}

$$U_{wy} \approx p \cdot \sqrt{U_{we}} \quad (42)$$

gdzie: p — współczynnik proporcjonalności.

Ważnym parametrem określającym jakość członu nieliniowego układu F.P.S., jest wrażliwość jego całkowitej admitancji $Y_{m/n}(q)$ na zmiany wartości tworzących go elementów $y_i \in \{G_\alpha \cdot q, g_\beta\}$.

Wrażliwość względną $W(y_i)$ funkcji $Y(q) = f(y_i)$ na zmiany wartości tworzących go elementów $y_i \in \{G_\alpha \cdot q, g_\beta\}$ to zgodnie z pracami [3, 16, 28].

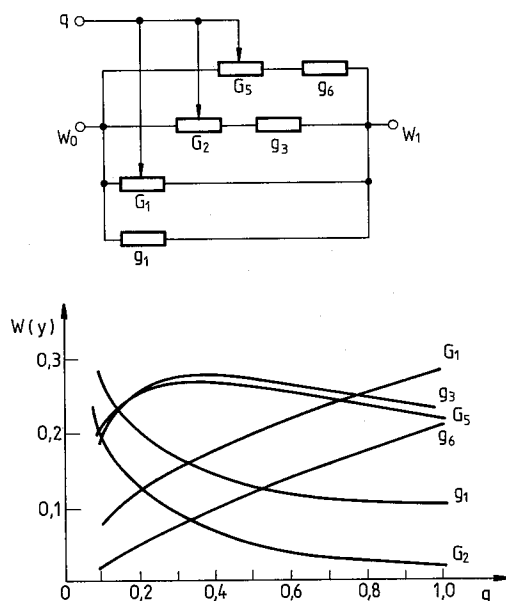
$$W(y_i) = \frac{y_i}{Y(q)} \cdot \frac{\delta Y(q)}{\delta y_i} \quad (42)$$

Praktycznie wrażliwość względną $W(y_i)$ układu dwójnika na zmiany wartości tworzących go elementów y_i obliczono korzystając ze wzoru przybliżonego

$$W(y_i) = \frac{y_i}{Y(q)} \cdot \frac{\Delta Y(q)}{\Delta y_i}, \quad (43)$$

przy założeniu

$$\frac{\Delta y_i}{y_i} = 0.01. \quad (44)$$



Rys. 5. Układ dwójnika oraz zależność wrażliwości jego admitancji $W(y)$ od parametru sterującego q

Tabela 2

Parametry funkcjonalnych przetworników sygnału, których admitancja ma postać $Y_{3/2}(q) \approx \sqrt{q}$

N°	Układ dwójnika			Wartość elementów sterowanych $\{G_\alpha\}$ i niesterowanych $\{g_\beta\}$	Wrażliwość względna w % na 1% zmiany wartości elementów $\{G_\alpha\}$ i $\{g_\beta\}$
	unigraf	bi-graf	schemat ideowy		
A	B	C	D	E	F
1				$G_2 = 1.391$ $G_3 = 2.205$ $G_5 = 0.3033$ $g_1 = 0.0888$ $g_4 = 0.3524$ $g_6 = 1.1057$	$W(G_2) = 0.32$ $W(G_3) = 0.19$ $W(G_5) = 0.22$ $W(g_1) = 0.28$ $W(g_4) = 0.095$ $W(g_6) = 0.39$
2				$G_2 = 2.5808$ $G_3 = 0.2973$ $G_5 = 1.614$ $g_4 = 0.4016$ $g_5 = 0.0959$ $g_6 = 1.195$	$W(G_2) = 0.19$ $W(G_3) = 0.22$ $W(G_5) = 0.32$ $W(g_4) = 0.096$ $W(g_5) = 0.17$ $W(g_6) = 0.46$
3				$G_1 = 0.2666$ $G_2 = 1.301$ $G_4 = 2.029$ $g_1 = 0.0888$ $g_3 = 1.111$ $g_5 = 0.2653$	$W(G_1) = 0.27$ $W(G_2) = 0.31$ $W(G_4) = 0.14$ $W(g_1) = 0.28$ $W(g_3) = 0.37$ $W(g_5) = 0.11$
4				$G_1 = 0.26662$ $G_2 = 2.5477$ $G_5 = 0.78193$ $g_1 = 0.088815$ $g_3 = 0.26198$ $g_6 = 0.84887$	$W(G_1) = 0.27$ $W(G_2) = 0.21$ $W(G_5) = 0.26$ $W(g_1) = 0.28$ $W(g_3) = 0.27$ $W(g_6) = 0.20$
5				$G_2 = 0.9541$ $G_3 = 0.9377$ $G_5 = 0.2945$ $g_2 = 0.099$ $g_5 = 0.3197$ $g_6 = 2.815$	$W(G_2) = 0.26$ $W(G_3) = 0.26$ $W(G_5) = 0.20$ $W(g_2) = 0.21$ $W(g_5) = 0.27$ $W(g_6) = 0.27$

c.d. tabl 2

A	B	C	D	E	F
6				$G_1 = 0.26662$ $G_2 = 2.1356$ $G_4 = 3.3299$ $g_1 = 0.0888$ $g_2 = 0.7146$ $g_6 = 1.1109$	$W(G_1) = 0.27$ $W(G_2) = 0.12$ $W(G_4) = 0.38$ $W(g_1) = 0.28$ $W(g_2) = 0.11$ $W(g_6) = 0.37$
7				$G_2 = 1.6318$ $G_3 = 0.3187$ $G_4 = 4.1658$ $g_3 = 1.899$ $g_4 = 0.1242$ $g_5 = 0.3816$	$W(G_2) = 0.44$ $W(G_3) = 0.078$ $W(G_4) = 0.15$ $W(g_3) = 0.42$ $W(g_4) = 0.044$ $W(g_5) = 0.28$
8				$G_3 = 2.402$ $G_4 = 0.3187$ $G_5 = 5.087$ $g_3 = 0.8035$ $g_4 = 1.699$ $g_5 = 0.106$	$W(G_3) = 0.23$ $W(G_4) = 0.078$ $W(G_5) = 0.43$ $W(g_3) = 0.23$ $W(g_4) = 0.42$ $W(g_5) = 0.088$

Wybór wartości 0.01 wynika z faktu, że w praktyce do realizacji przewodności sterowanej $\{G_\alpha \cdot q\}$ i niesterowanej $\{g_\beta\}$ używano rezystorów klasy 1.

Jak pokazano na rysunku 5 wrażliwość względna poszczególnych elementów $\{G_\alpha \cdot g_\beta\}$ układu dwójnika, jest również funkcją parametru sterującego q . Dla rozpatrywanych dalej układów dwójników będą podawane tylko maksymalne wrażliwości względne $W(G_\alpha)$ i $W(g_\beta)$.

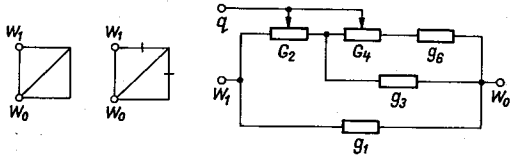
7. PRZYKŁADY F.P.S. W POSTACI DWÓJNIKÓW SKŁADAJĄCYCH SIĘ Z DWÓCH RODZAJÓW ELEMENTÓW

Jak pokazano w pracy [37, 38] dwójnik składający się z dwóch rodzajów [s]lso elementów, może realizować transmitancję w postaci czterech typów funkcji wymiernej. Aby dany dwójnik mógł realizować w postaci transmitancji $Y_{m/n}(q)$, analitycznie lub empirycznie zadaną funkcję $y = f(x)$, musi być ona aproksymowana jednym z czterech typów funkcji wymiernej $R_{m/n}(x)$ (wzory 20÷23) oraz musi istnieć rozwiązanie układu równań 25÷27.

Ponieważ równania 25 i 26 są równaniami nieliniowymi, rozwiązania ich szukano metodami iteracyjnymi. W pracy [9] opracowano algorytm, który pozwalał uzyskiwać rozwiązania układów równań 25÷27, dla wielu cennych z praktycznego punktu

Tabela 3

Parametry funkcjonalnych przetworników sygnału, których admitancja ma postać $Y_{2/2}(q) \sim f(q)$

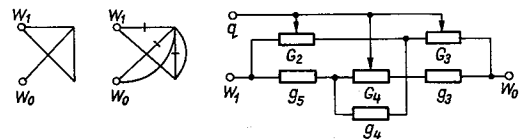
			
No	Realizowana funkcja $Y_{2/2}(q) \sim f(q)$	Wartość elementów sterowanych $\{G_\alpha\}$ i niesterowanych $\{g_\beta\}$	Wrażliwość względna w % na 1% zmiany wartości elementów $\{G_\alpha\}$ i $\{g_\beta\}$
A	B	C	D
1	$Y_{2/2}(q) \approx \sqrt{q}$	$G_3 = 2.199$ $G_5 = 0.7355$ $g_1 = 0.1067$ $g_2 = 0.3707$ $g_6 = 2.654$	$W(G_3) = 0.27$ $W(G_5) = 0.45$ $W(g_1) = 0.34$ $W(g_2) = 0.29$ $W(g_6) = 0.12$
2	$Y_{2/2}(q) \approx \operatorname{tg} \frac{\pi}{2} \sqrt{q}$	$G_3 = 7.954$ $G_5 = 0.928$ $g_1 = 0.00642$ $g_2 = 0.1963$ $g_6 = 5.578$	$W(G_3) = 0.12$ $W(G_5) = 0.69$ $W(g_1) = 0.025$ $W(g_2) = 0.49$ $W(g_6) = 0.11$
3	$Y_{2/2}(q) \approx \operatorname{sh} \sqrt{q}$	$G_3 = 1.289$ $G_5 = 10.779$ $g_1 = 0.00692$ $g_2 = 3.405$ $g_6 = 0.2312$	$W(G_3) = 0.60$ $W(G_5) = 0.10$ $W(g_1) = 0.022$ $W(g_2) = 0.22$ $W(g_6) = 0.49$
4	$Y_{2/2}(q) \approx \operatorname{arc sh} \sqrt{q}$	$G_3 = 12.421$ $G_5 = 1.496$ $g_1 = 0.0042$ $g_2 = 0.1986$ $g_6 = 1.236$	$W(G_3) = 0.076$ $W(G_5) = 0.4$ $W(g_1) = 0.014$ $W(g_2) = 0.48$ $W(g_6) = 0.42$

widzenia funkcji analitycznych (np. $y = \sqrt{x}$, $y = \log x$, $y = \operatorname{tg} x$), oraz empirycznie zadanych funkcji $y = f(x)$ [31] aproksymowanych funkcją wymierną $R_{m/n}(x)$ spełniającą zależności 25 ÷ 27.

W procesie wyboru typu struktury kanonicznej bi-grafu układu do realizacji F.P.S., można wyróżnić trzy grupy zagadnień. Pierwsza grupa zagadnień związanych z kryteriami wyboru typu struktury kanonicznej bi-grafu, wynika z faktu, że daną

Tabela 4

Parametry funkcjonalnych przetworników sygnału, których admitancja ma postać $Y_{2/2}(q) \sim f(q)$

			
N°	Realizowana funkcja $Y_{2/2}(q) \sim f(q)$	Wartość elementów sterowanych $\{G_\alpha\}$ i niesterowanych $\{g_\beta\}$	Wrażliwość względna w % na 1% zmiany wartości elementów $\{G_\alpha\}$ i $\{g_\beta\}$
A	B	C	D
1	$Y_{2/2}(q) \approx \sqrt{q}$	$G_2 = 3.706$ $G_6 = 0.7328$ $g_3 = 0.4885$ $g_4 = 0.1444$ $g_5 = 2.498$	$W(G_2) = 0.29$ $W(G_6) = 0.44$ $W(g_3) = 0.45$ $W(g_4) = 0.089$ $W(g_5) = 0.12$
2	$Y_{2/2}(q) \approx \operatorname{tg} \frac{\pi}{2} \sqrt{q}$	$G_2 = 9.356$ $G_6 = 1.021$ $g_3 = 0.2496$ $g_4 = 5.525$ $g_5 = 0.0066$	$W(G_2) = 0.27$ $W(G_6) = 0.59$ $W(g_3) = 0.52$ $W(g_4) = 0.11$ $W(g_5) = 0.002$
3	$Y_{2/2}(q) \approx \operatorname{sh} \sqrt{q}$	$G_2 = 12.653$ $G_6 = 1.409$ $g_3 = 0.2947$ $g_4 = 3.341$ $g_5 = 0.0071$	$W(G_2) = 0.25$ $W(G_6) = 0.50$ $W(g_3) = 0.51$ $W(g_4) = 0.21$ $W(g_5) = 0.002$
4	$Y_{2/2}(q) \approx \operatorname{arc sh} \sqrt{q}$	$G_2 = 1.598$ $G_6 = 14.188$ $g_3 = 0.0043$ $g_4 = 1.1831$ $g_5 = 0.2517$	$W(G_2) = 0.36$ $W(G_6) = 0.21$ $W(g_3) = 0.002$ $W(g_4) = 0.39$ $W(g_5) = 0.50$

funkcję $y = f(x)$ aproksymowaną funkcją wymierną $R_{m/n}(x)$ spełniającą zależności 25 ÷ 27 mogą realizować dwójniki o różnej strukturze kanonicznej.

Wpływ postaci struktury kanonicznej bi-grafu układu dwójnika na wartości poszczególnych jego elementów $\{G_\alpha \cdot q, g_\beta\}$ oraz ich wrażliwości $W(G_\alpha)$ i $W(g_\beta)$, zilustrowano na przykładzie funkcji $y = \sqrt{x}$ aproksymowanej funkcją wymierną $R_{3/2}(x)$ (tabela 1). Jak pokazano w pracy [38], istnieją 32 klasy struktur kanonicznych bi-grafów dwójników, których admitancja wejściowa $Y_{3/2}(q)$ ma analogiczną postać funkcji wymiernej $R_{3/2}(x)$.

Tabela 5

Parametry funkcjonalnych przetworników sygnału, których admitancja ma postać $Y_{2/2}(q) \approx \sqrt{q}$

N°	Układ dwójnika			Wartość elementów sterowanych $\{G_\alpha\}$ i niesterowanych $\{g_\beta\}$	Wrażliwość względna w % na 1% zmiany wartości elementów $\{G_\alpha\}$ i $\{g_\beta\}$
	unigraf	bi-graf	schemat ideowy		
A	B	C	D	E	F
1				$G_2 = 1.0179$ $G_4 = 3.089$ $g_3 = 2.996$ $g_5 = 0.5029$ $g_6 = 0.1354$	$W(G_2) = 0.48$ $W(G_4) = 0.22$ $W(g_3) = 0.30$ $W(g_5) = 0.32$ $W(g_6) = 0.13$
2				$G_2 = 2.935$ $G_4 = 0.865$ $g_1 = 0.107$ $g_3 = 0.65$ $g_6 = 2.37$	$W(G_2) = 0.24$ $W(G_4) = 0.04$ $W(g_1) = 0.34$ $W(g_3) = 0.15$ $W(g_6) = 0.08$
3				$G_3 = 2.199$ $G_5 = 0.7355$ $g_1 = 0.1067$ $g_2 = 0.3707$ $g_6 = 2.654$	$W(G_3) = 0.27$ $W(G_5) = 0.45$ $W(g_1) = 0.34$ $W(g_2) = 0.29$ $W(g_6) = 0.12$
4				$G_2 = 3.648$ $G_5 = 0.7354$ $g_2 = 0.1374$ $g_3 = 0.4775$ $g_6 = 2.654$	$W(G_2) = 0.27$ $W(G_5) = 0.45$ $W(g_2) = 0.10$ $W(g_3) = 0.45$ $W(g_6) = 0.12$
5				$G_2 = 1.403$ $G_4 = 2.935$ $g_1 = 0.1067$ $g_2 = 0.828$ $g_6 = 3.025$	$W(G_2) = 0.22$ $W(G_4) = 0.47$ $W(g_1) = 0.34$ $W(g_2) = 0.19$ $W(g_6) = 0.26$
6				$G_4 = 3.146$ $G_5 = 1.503$ $g_2 = 0.8873$ $g_3 = 3.131$ $g_5 = 0.1105$	$W(G_4) = 0.47$ $W(G_5) = 0.22$ $W(g_2) = 0.19$ $W(g_3) = 0.32$ $W(g_5) = 0.28$
7				$G_2 = 4.053$ $G_4 = 1.358$ $g_2 = 0.1258$ $g_4 = 0.9069$ $g_6 = 3.1314$	$W(G_2) = 0.45$ $W(G_4) = 0.26$ $W(g_2) = 0.14$ $W(g_4) = 0.29$ $W(g_6) = 0.32$

Dla wszystkich 32 klas struktur kanonicznych bi-grafów dwójników, uzyskano rozwiązanie układu równań 25—27 z dokładnością 10^{-3} .

W tabeli 2 pokazano osiem przykładów (po jednym z każdej klasy unigrafów z pracy [38]) dwójników, których admitancja wejściowa jest funkcją wymierną typu $Y_{3/2}(q)$ i które z pewnym przybliżeniem realizują funkcję $Y_{3/2}(q)$. Otrzymywane w procesie rozwiązywania układu równań 25—27, wartości admitancji poszczególnych elementów $\{G_\alpha \cdot q, g_\beta\}$ i wpływu ich zmian na całkowitą wrażliwość układu zestawiono w tej samej tabeli.

Druga grupa zagadnień związanych z kryteriami wyboru typu struktury kanonicznej realizującej daną funkcję $y = f(x)$ wynika z faktu, że różne analitycznie lub empirycznie zadane funkcje (tabela 1), mogą być aproksymowane jednym typem funkcji wymiernej $R_{m/n}(x)$ spełniającej jeden z czterech warunków 20—23.

Na przykład funkcje $y = \sqrt{x}$, $y = \operatorname{tg} \frac{\pi}{2} \sqrt{x}$, $y = \operatorname{sh} \sqrt{x}$, $y = \operatorname{arcsh} \sqrt{x}$, mogą być aproksymowane jednym typem funkcji wymiernej $R_{3/2}(x)$ różniącej się tylko wartościami współczynników a_i i b_j . Jeśli dla danego typu struktury kanonicznej bi-grafu dwójnika, oraz danego typu funkcji wymiernej $R_{m/n}(x)$ (różniącej się tylko wartościami współczynników a_i i b_j) istnieje rozwiązanie układu równań nieliniowych 25÷27, to dany typ struktury kanonicznej dwójnika może realizować z pewnym przybliżeniem różne funkcje $y = f(x)$.

W tabelach 3 i 4 pokazano możliwość realizacji funkcji $y = \sqrt{x}$, $y = \operatorname{tg} \frac{\pi}{2} \sqrt{x}$, $y = \operatorname{sh} \sqrt{x}$, $y = \operatorname{arcsh} \sqrt{x}$, przez dwie losowo wybrane struktury kanoniczne bi-grafów dwójników. Wynika stąd, że jedna i ta sama struktura kanoniczna bi-grafu układu dwójnika składający się z elementów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$, może realizować różne funkcje $f(x)$ jeśli tylko istnieje rozwiązanie układu równań nieliniowych 25÷27.

Trzecia grupa zagadnień związanych z kryteriami wyboru typu struktury kanonicznej realizującej daną funkcję $y = f(x)$, wynika z faktu, że różne analitycznie lub empirycznie zadane funkcje (tabela 1), mogą być przybliżone z określonym błędem aproksymacji — funkcją wymierną $R_{m/n}(x)$ o różnych stopniach licznika i mianownika.

Zilustrujemy to na przykładzie funkcji $y = \sqrt{x}$. Z tabeli 1 wynika, że funkcja $y = \sqrt{x}$ może być aproksymowana między innymi następującymi funkcjami wymiernymi $R_{2/2}(x)$, $R_{3/2}(x)$, $R_{4/3}(x)$. W tabeli 5 pokazano siedem przykładów (po jednym przykładzie z każdej klasy struktur kanonicznych unigrafów z pracy [38]), których admitancja wejściowa ma postać funkcji wymiernej $Y_{2/2}(q)$, i które realizują funkcję $Y(q) \approx \sqrt{q}$.

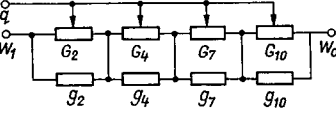
Przykłady układów dwójników, których admitancja wejściowa ma postać funkcji wymiernej typu $Y_{3/2}(q)$, i które realizują funkcję $Y(q) \approx \sqrt{q}$ pokazano w tabeli 2. Natomiast w tabeli 6 pokazano pięć przykładów układów dwójników, których admitancja wejściowa ma postać funkcji wymiernej typu $Y_{4/3}(q)$, i które realizują

Tabela 6

Parametry funkcjonalnych przetworników sygnału, których admitancja ma postać $Y_{4/3}(q) \approx \sqrt{q}$

N°	Układ dwójnika			Wartość elementów sterowanych $\{G_\alpha\}$ i niesterowanych $\{g_\beta\}$	Wrażliwość względna w % na 1% zmiany wartości elementów $\{G_\alpha\}$ i $\{g_\beta\}$
	unigraf	bi-graf	schemat ideowy		
A	B	C	D	E	F
1				$G_1 = 0.2009$ $G_5 = 1.2182$ $G_7 = 1.5052$ $G_8 = 5.1503$ $g_2 = 0.41816$ $g_4 = 0.08569$ $g_6 = 0.9729$ $g_{10} = 0.64787$	$W(G_1) = 0.20$ $W(G_5) = 0.34$ $W(G_7) = 0.005$ $W(G_8) = 0.16$ $W(g_2) = 0.28$ $W(g_4) = 0.04$ $W(g_6) = 0.26$ $W(g_{10}) = 0.14$
2				$G_1 = 0.2009$ $G_2 = 0.5003$ $G_5 = 0.9386$ $G_9 = 3.272$ $g_1 = 0.0663$ $g_3 = 1.078$ $g_6 = 0.3105$ $g_{10} = 0.1656$	$W(G_1) = 0.2$ $W(G_2) = 0.24$ $W(G_5) = 0.18$ $W(G_9) = 0.1$ $W(g_1) = 0.21$ $W(g_3) = 0.11$ $W(g_6) = 0.18$ $W(g_{10}) = 0.24$
3				$G_2 = 0.2133$ $G_3 = 3.4729$ $G_5 = 0.9386$ $G_{10} = 0.5637$ $g_2 = 0.1866$ $g_6 = 0.3105$ $g_9 = 1.1446$ $g_{10} = 0.0703$	$W(G_2) = 0.17$ $W(G_3) = 0.15$ $W(G_5) = 0.18$ $W(G_{10}) = 0.24$ $W(g_2) = 0.24$ $W(g_6) = 0.18$ $W(g_9) = 0.15$ $W(g_{10}) = 0.18$
4				$G_2 = 6.5064$ $G_4 = 3.8366$ $G_9 = 1.1515$ $G_{10} = 0.2434$ $g_2 = 0.0868$ $g_4 = 0.7069$ $g_7 = 0.5767$ $g_{10} = 2.3528$	$W(G_2) = 0.18$ $W(G_4) = 0.053$ $W(G_9) = 0.44$ $W(G_{10}) = 0.036$ $W(g_2) = 0.023$ $W(g_4) = 0.088$ $W(g_7) = 0.26$ $W(g_{10}) = 0.35$

c.d. tabl. 6

A	B	C	D	E	F
5				$G_2 = 7.1256$ $G_4 = 4.6369$ $G_7 = 1.9515$ $G_{10} = 0.2434$ $g_2 = 0.0824$ $g_4 = 0.6440$ $g_7 = 1.5365$ $g_{10} = 2.3528$	$W(G_2) = 0.36$ $W(G_4) = 0.19$ $W(G_7) = 0.16$ $W(G_{10}) = 0.036$ $W(g_2) = 0.040$ $W(g_4) = 0.17$ $W(g_7) = 0.19$ $W(g_{10}) = 0.35$

funkcję $Y(q) \approx \sqrt{q}$. Oczywiście im wyższy stopień licznika i mianownika funkcji wymiernej $R_{m/n}(x)$, to z jednej strony mniejszy błąd aproksymacji danej funkcji $y = f(x)$ — funkcją wymierną typu $R_{m/n}(x)$, z drugiej strony wyższy stopień licznika i mianownika funkcji $Y_{m/n}(q)$, to zmniejszenie wrażliwości admitancji układu dwójnika na zmiany parametrów sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$ co ilustrują tabele 2 i 6.

W praktyce konieczność stosowania F.P.S. w postaci dwójników, których admitancja wejściowa ma postać funkcji wymiernej typu $Y_{4/3}(q)$ a nawet $Y_{5/4}(q)$ wynika z faktu, że wrażliwość admitancji $W(y_i)$ układu dwójnika na 1% zmiany wartości przewodności sterowanych $\{G_\alpha \cdot q\}$ i niesterowanych $\{g_\beta\}$, może mieć w tych układach wartość $W(y_i) < 0,1\%$.

8. PODSUMOWANIE

Z rozważań niniejszej pracy wynika, że jeśli zadaną analitycznie lub empirycznie zależność $y = f(x)$ można aproksymować funkcją wymierną $R_{m/n}(x)$ spełniającą warunki 20—23 oraz jeśli istnieje rozwiązanie układu równań nieliniowych 25—27, to można zbudować funkcjonalny przetwornik sygnału zawierający człon nieliniowy w postaci układu dwójnika składającego się z przewodności sterowanej $\{G_\alpha \cdot q\}$ i niesterowanej $\{g_\beta\}$, którego admitancja $Y_{m/n}(q)$ realizuje z określoną dokładnością zadaną zależność funkcyjną $y = f(x)$.

BIBLIOGRAFIA

1. K. Badźmirowski, H. Karkowska, Z. Karkowski: *Cyfrowe systemy pomiarowe*. Warszawa: WNT, 1979
2. M. Białko, A. Guziński, W. Sienko, J. Zaruda: *Filtry aktywne RC*. Warszawa: WNT, 1979
3. J. Chocjan, L. Lasek: *Metody analizy wrażliwości układów elektronicznych*. Gliwice: Wyd. Pol. Śląskiej, 1980
4. A. Cichocki: *Synteza układów nieliniowych przy użyciu wzmacniaczy operacyjnych i elementów sterowanych*. Prace naukowe Elektryka, z. 67, Pol. Warszawska, Warszawa: 1982
5. G. Dahlquist, A. Bjorek: *Metody numeryczne*. Warszawa: PWN, 1983
6. Z. Fortuna, B. Mackow, J. Wąsowski: *Metody numeryczne*. Warszawa: WNT, 1981
7. P. H. Garrett: *Układy analogowe w systemach cyfrowych*. Warszawa: WNT, 1981

8. I. W. Gierasimow, Z. Wróbel: *Elektroniczny aproksymator funkcji $y = \sqrt{x}$* . Archiwum Elektrotechniki, 1985, t. XXIV, z. 3, ss. 855—856
9. I. W. Gierasimow: *Teoria projektowania i primienienia wycislitelno-prieobrazowatielnych ciepiej (sintez i primienienia)*. Rozprawa habilitacyjna, Leningradzki Instytut Elektrotechniki, Leningrad: 1986
10. A. Guziński, J. C. Matheau: *Filtry R-przelączane realizowane technologią MOS*. Mat. VI KK, TOiUE, 1983, ss. 325—331, Gliwice, październik, 1983
11. A. Guziński, J. C. Matheau: *Projektowanie filtrów R-przelączanych*. Mat. VII KK, Kazimierz, TOiUE, 1984, ss. 339—345
12. A. Guziński, S. Słotwiński: *Komputerowa analiza układów liniowych z kluczowaną rezystancją*. Mat. X KK, TOiUE, 1987, Gdańsk, ss. 261—266
13. I. W. Guriewicz: *Synteza parametryczeskich funkcjonalnych struktur*. Moskwa: Swiaz, 1977
14. J. Huertas, J. Acha, A. Gago: *Design of general voltage or current controlled resistive elements and their applications to the synthesis of nonlinear networks*. IEEE. Trans. Circuits and Systems, 1980, vol. CAS-27, pp. 92—103
15. J. Janiczek: *Linearyzacja układów pomiarowych zawierających przetwornik A/C*. Praca doktorska, Pol. Wrocławska, 1975
16. J. Jaworska: *Synteza układów optymalnych mało-wrażliwych na zmiany wybranych parametrów*. Rozprawy Elektrotechniczne, 1985, t. XXXIV, ss. 95—99
17. E. Kącki, J. Woźniakowski: *Modelowanie analogowe i hybrydowe oraz cyfrowa symulacja maszyn analogowych*. Warszawa; PWN, 1979
18. A. A. Khan: *Linearization of thermistor thermometer*. Int. J. Electron, 1985, pp. 129—139
19. Z. Kulka, M. Nadachowski: *Liniowe układy scalone i ich zastosowania*. Warszawa; WKiŁ, 1974
20. J. Mędrzycki: *Technika analogowa i hybrydowa*. Warszawa; WNT, 1974
21. A. Niderliński: *Systemy komputerowe automatyki przemysłowej*. Warszawa: WNT, 1984
22. W. Nowakowski: *Podstawowe układy elektroniczne — układy impulsowe*. Warszawa: WKiŁ, 1982
23. J. Pawłowski: *Podstawowe układy elektroniczne — nieliniowe układy analogowe*. Warszawa: WKiŁ, 1976
24. S. Plesiewicz: *Cyfrowa metoda linearyzacji charakterystyk statycznych torów pomiarowych w układzie z blokiem funkcji odwrotnej*. Archiwum Elektrotechniki, 1985, t. XXXIV, z. 3/4, ss. 815—830
25. W. B. Smolow: *Funkcjonalnyye prieobrazowatieli informacii*. Leningrad: Energoizdat, 1981
26. W. B. Smolow: *Kompleks impulsno-analogowych wycislitelnych ustrojstw dla sistem uprawlenia i modilirowania*. Pribory i sistemy uprawlenia, 1976, t. 5, ss. 30—36
27. W. B. Smolow: *Analogowo wycislitelnyje maszyny*. Leningrad: Energoizdat, 1981
28. M. Stybliński: *Metody analizy i optymalizacja tolerancji parametrów układów elektronicznych*. Warszawa: WNT, 1981.
29. *Wremia — impulsnyje wycislitelnyje ustrojstwa* pod redakcją W. B. Smołowa i E. P. Ugriumowa: Moskwa, Radio i swiaź, 1983
30. Z. Wróbel, I. W. Gierasimow: *Elektroniczny aproksymator funkcji logarytmicznej*. PAK. Nr 8, 1986, ss. 180—181.
31. Z. Wróbel, I. W. Gierasimow: *Sterowana przewodność w układzie linearyzatora charakterystyk termopar*. PAK. Nr 11, 1988, ss. 258÷260.
32. Z. Wróbel: *Elektroniczny potencjometr*, Patent PRL nr 137437
33. Z. Wróbel: *Przetwornik informacji z potencjometrem elektronicznym*, Patent PRL nr 135893
34. Z. Wróbel: *Układ do linearyzacji charakterystyk przetworników pomiarowych*, Patent PRL nr 133313
35. Z. Wróbel: *Funkcje skrukturalne grafów układów dwójników i czwórników składających się z dwóch rodzajów elementów*. Część I. Homeomorfizm unigrafów układów. Kwart. Elektr. i Telekom. 1991, t. 37, z. 1—2

36. Z. Wróbel: *Funkcje strukturalne grafów układów dwójników i czwórników składających się z dwóch rodzajów elementów. Część II. Homeomorfizm bi-grafów układów.* Kwart. Elektr. i Telekom. 1991, t. 37, z. 1—2
37. Z. Wróbel: *Synteza struktur kanonicznych bi-grafów układów dwójników. Część I. Metoda przeglądu bezpośredniego.* Kwart. Elektr. i Telekom. 1990, t. 199, z. 4
38. Z. Wróbel: *Synteza struktur kanonicznych bi-grafów układów dwójników. Część II. Metodą przeglądu hierarchicznego.* Kwart. Elektr. i Telekom. 1991, t. 37, z. 1—2
39. J. Zakrzewski: *Metodyka syntezy układów linearyzujących nieliniowe charakterystyki statyczne przetworników pomiarowych.* Prace naukowe Pol. Śląskiej, Elektryka z. 65, Gliwice, 1979

Z. WRÓBEL, I. W. GIERASIMOW

ONE-PORT CONSISTING SWITCHED CONDUCTIVITY AS A FUNCTIONAL SIGNAL CONVERTERS

Summary

In the paper has been shown the test of the use received canonical structures of bi-graph one — port [37, 38], as the functional converters of the information where the output signal $\{U_{wy}\}$ is the function input signal $\{U_{we}\}$ [$U_{wy} = f(U_{we})$].

It has been shown restrictions imposed on the rational function $R_{m/n}(x)$ known approximating function dependence $y = f(x)$ and conditions of realizeability of the function $R_{m/n}(x)$ in the form of the input admittance $Y_{m/n}(q)$ one-port consisting switched conductivity $\{G_\alpha \cdot q\}$ and linear conductivity $\{g_\beta\}$.

It has been shown the structure topological effect of one-port with realize function $Y_{m/n}(q)$ on the values consisting its elements $y_i \in \{G_\alpha \cdot q, g_\beta\}$.

Elektrotermiczna, dwuwymiarowa analiza procesu załączania w wielowarstwowych przyrządach półprzewodnikowych

BOGUSŁAW WIĘCEK

Instytut Elektroniki, Politechnika Łódzka

Otrzymano 1989.12.29

Autoryzowano do druku 1990.05.16

W pracy przedstawiono dwuwymiarowy, elektrotermiczny model wielowarstwowych przyrządów półprzewodnikowych ze szczególnym uwzględnieniem zjawisk zachodzących przy załączaniu. Uwzględniono wpływ gęstości nośników na ruchliwość. Zamodelowano efekt nierównomiernego rozkładu gęstości prądu. Do obliczeń wykorzystano pomiary powierzchni obszaru przewodzącego przy użyciu promieniowania rekombinacyjnego. Model zweryfikowano zaciskowymi pomiarami napięcia przewodzenia w czasie załączania.

1. WSTĘP

Proces załączania wielowarstwowej, półprzewodnikowej struktury mocy (tyrystora, tyrystora GATT, GTO), szczególnie przy impulsach prądowych o dużej stromości charakteryzuje wydzielanie się znacznej mocy. Wynika to ze zlokalizowanego przewodzenia i skończonej szybkości rozprzestrzeniania się przewodzącego kanału. Wzrost temperatury wpływa na parametry fizyczne półprzewodnika: ruchliwość, koncentracje ładunków, generację i rekombinację nośników itd. Wpływ temperatury oraz przestrzenny charakter opisywanych zjawisk powoduje, iż dokładna analiza tego procesu musi wykorzystywać wielowymiarowe i elektrotermiczne modelowanie przyrządów półprzewodnikowych.

Analizie zjawisk zachodzących w przyrządach półprzewodnikowych w stanach statycznych i dynamicznych poświęcono prace [1—26], których wyniki umożliwiają oszacowanie przyrostów temperatury w strukturach dla określonych pobudzeń. Dużą grupę stanowią prace dotyczące modelowania i badania zjawisk w tyrystorach [1—2, 4, 8—10, 12—15, 18—26]. Prezentowane prace dotyczą przede wszystkim jednowymiarowego modelowania elektrycznego, które pozwala ocenić wpływ parametrów technologicznych na prędkość rozprzestrzeniania się przewodzącego kanału w tyrystorze [2, 15, 18—22, 26]. Inne prace związane są z zastosowaniem elektrycznych

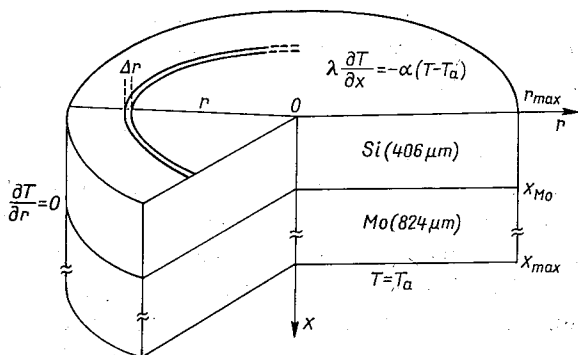
i optycznych metod pomiarowych zarówno powierzchni obszaru początkowo załączonego jak i powierzchni obszaru przewodzącego w czasie załączania [8—9, 18, 20—23]. Przedstawione wyniki jedynie w przybliżeniu pozwalają ocenić przyrost temperatury w stanach dynamicznych w strukturze. Próby połączenia modeli elektrycznych i termicznych kończyły się najczęściej zawężeniem zagadnienia do przypadku jednowymiarowego [19, 13—14] lub dotyczyły wielowymiarowego opisu zjawisk w stanach ustalonych [16]. Znana jest uproszczona elektrotermiczna analiza zjawisk dynamicznych w tranzystorze mocy [5—6], w której wykorzystano quasidwuwymiarowy model elektryczny przyrządu. Modelowanie polega tu na podziale elementu na skończoną liczbę struktur elementarnych, dla których obowiązują modele jednowymiarowe. Podobne podejście prezentowane jest w pracy [16]. Modele quasidwuwymiarowe, przy dobrym odwzorowaniu zjawisk zachodzących w przyrządach, zapewniają znaczne skrócenie czasu obliczeń. Dwuwymiarowy model cieplny wraz z optycznym pomiarem szerokości kanału przewodzącego zaproponowano w pracy [24], uzyskując przybliżone rozkłady temperatur wewnątrz i na powierzchni pastylki tyrystora.

2. ZAŁOŻENIA

1. Model dotyczy struktury z centryczną bramką. Jest także słuszny dla produkowanych obecnie na szeroką skalę przyrządów z bramką wzmacniającą. Może być również łatwo przystosowany do struktur z centryczną katodą. Do obliczeń przyjęto wartości parametrów konstrukcyjnych tyrystora produkcji krajowej TOO-40.

2. Ze względu na symetrię geometryczną struktury z bramką centryczną, model przepływu ciepła opisany w trójwymiarowym układzie współrzędnych walcowych zredukowano do postaci dwuwymiarowej o osiach umieszczonych w środku symetrii górnej powierzchni struktury.

3. Struktura półprzewodnika umieszczona jest na podstawie z molibdenu. Pomiędzy warstwę lutownia na styku krzem-molibden ze względu na niewielką jego grubość



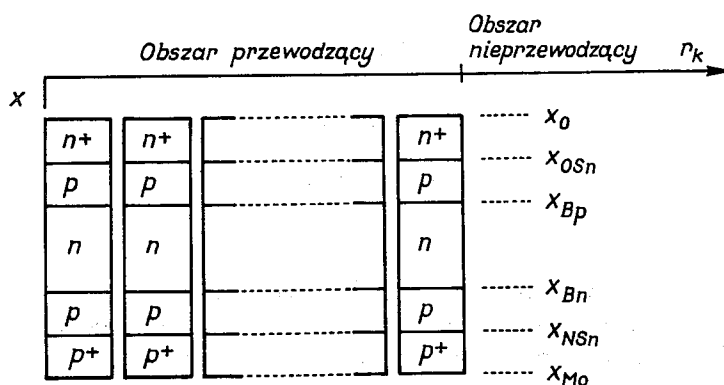
Rys. 1. Przekrój struktury tyrystora

(kilka μm). Z uwagi na znaczną grubość pastylki wraz z molibdenową podstawą (ok. 1200 μm) oraz analizę stanów przejściowych dla czasów krótszych niż 300 μs założono izotermiczne warunki brzegowe na krawędzi molibdenu. Na górnej powierzchni krzemu uwzględniono wpływ konwekcji naturalnej na straty ciepła i rozkład temperatury. Pominęto cienką (kilka ÷ kilkanaście μm) warstwę metalizacji (Al).

4. Proces załączania podzielono na dwie fazy: fazę zaniku ładunku przestrzennego oraz fazę rozprzestrzeniania się przewodzącego kanału [8—9, 13], jako dominujące ze względu na wydzielanie mocy w przyrządzie.

5. Złącza struktury tyrystora są skokowe, co w konsekwencji prowadzi do liniowego rozkładu pola elektrycznego w warstwie ładunku przestrzennego.

6. Przyjęto podział struktury półprzewodnika na dwa obszary: przewodzący i nieprzewodzący. Pominęto strefę przejściową między obszarem załączonym i niezałączonym ze względu na jego niewielką szerokość (rzędu kilku dróg dyfuzji w krzemie) i niewielki udział w generacji ciepła [9, 21].



Rys. 2. Podział tyrystora na struktury elementarne

7. Dla fazy rozprzestrzeniania kanału przewodzącego zaproponowano model quasidwuwymiarowy. Oznacza to, że daną strukturę podzielono wzdłuż osi O_r na skończoną liczbę struktur elementarnych, dla których obowiązują modele jednowymiarowe [6, 16]. Dla uzyskania określonej dokładności liczba struktur elementarnych musi być zwiększana stosownie do prędkości rozprzestrzeniania się przewodzącego kanału.

8. W stanie rozprzestrzeniania się kanału przewodzącego model tyrystora dla struktur elementarnych w części przewodzącej (załączonej) zastąpiono modelem diody p-s-n [4, 12—13].

9. Uwzględniono wpływ gęstości nośników nadmiarowych oraz temperatury na ruchliwość, co jest niezbędne przy dużym poziomie wstrzykiwania.

10. Uwzględniono efekt nierównomiernego rozkładu gęstości prądu, który wynika z lokalnego przegrzania struktury w okolicach bramki. Słuszność tego założenia potwierdziła eksperymentalna weryfikacja modelu.

11. Prądy na granicach obszaru słabo domieszkowanego wynikają z przepływu nośników jednoimiennych [3, 7].

12. W modelu elektrycznym wykorzystano wyniki pomiarów pola powierzchni przewodzącego kanału w tyrystorze przy użyciu mikroskopu podczerwieni do detekcji promieniowania rekombinacyjnego.

3. WPŁYW TEMPERATURY NA FIZYCZNE PARAMETRY PÓŁPRZEWODNIKA

Podstawową zależnością, która łączy właściwości elektryczne i termiczne przyrządów półprzewodnikowych są równania wyrażające związek ruchliwości nośników i temperatury. Dodatkowo przy analizie stanów przejściowych w przyrządach mocy należy uwzględnić wysoki poziom wstrzykiwania.

Wpływ temperatury na ruchliwość nośników przedstawiają równania [16]:

$$\left. \begin{aligned} \mu_{Lp} &= \mu_{p0} (T/T_0)^{-\alpha} \\ \mu_{Ln} &= \mu_{n0} (T/T_0)^{-\alpha} \end{aligned} \right\} \quad (1)$$

przy czym:

$$\begin{aligned} \mu_{p0} &= 0.1430 \text{ m}^2 \text{v}^{-1} \text{s}^{-1}, \\ \mu_{n0} &= 0.495 \text{ m}^2 \text{v}^{-1} \text{s}^{-1}, \\ \alpha &= 2.2, \\ T_0 &= 300 \text{ K}. \end{aligned}$$

Przy dużym poziomie wstrzykiwania w obszarze s_n diody p-s-n ważną rolę odgrywa składowa ruchliwości μ_{CCS} związana z wzajemnym oddziaływaniem nośników (Carrier-Carrier Scattering). W pracy wykorzystano zależność postaci [12, 16]:

$$\mu_{CCS}(x, T) = \frac{1.1 * 10^{17} T^{1.5}}{n(x, T) * \ln[2.9 * 10^{14} T^2 / n(x, T)]} [\text{cm}^2 \text{v}^{-1} \text{s}^{-1}] \quad (2)$$

Ze względu na niewielki poziom domieszkowania obszaru s_n diody p-s-n ($N_D \simeq 10^{14}$) pominięto wpływ składowej ruchliwości związanej z oddziaływaniem nośników ze zjonizowanymi atomami domieszek μ_1 [16].

Przedstawione wyżej ruchliwości tworzą wypadkową ruchliwość dziur i elektronów w słabo domieszkowanym obszarze diody p-s-n:

$$\left. \begin{aligned} \mu_p^{-1} &= \mu_{Lp}^{-1} + \mu_{CCS}^{-1} \\ \mu_n^{-1} &= \mu_{Ln}^{-1} + \mu_{CCS}^{-1} \end{aligned} \right\} \quad (3)$$

Uwzględniono wpływ temperatury na koncentrację nośników w półprzewodniku samoistnym. Wzrost temperatury, poprzez wpływ na koncentracje n_i wywołuje zmiany napięć na złączach struktury p-s-n [7].

$$n_i^2 = AT^3 \exp \left[-\frac{E_g}{kT} \right] \quad (4)$$

przy czym:

$$A = \frac{32\pi^3 k^3}{h^6} (m_n^* m_p^*)^{1.5},$$

m_n^*, m_p^* — efektywne masy: elektronu i dziury,

k, h — stałe: Boltzmanna i Plancka.

Wpływ temperatury na szerokość pasma zabronionego w krzemie przedstawia zależność:

$$E_g = E_{g0} - a(T - T_0) \quad (5)$$

przy czym: $a = (dE_g/dT) = 2.3 \cdot 10^{-4} \text{ eV/K},$

$$E_{g0} = 1.12 \text{ eV},$$

$$T_0 = 300 \text{ K}.$$

4. MODEL ELEKTRYCZNY

Zgodnie z przedstawionymi wyżej założeniami uwzględniono dwie fazy załączania tyrystora: fazę zaniku ładunku przestrzennego i fazę rozprzestrzeniania się kanału przewodzącego.

Podczas pierwszej fazy występuje gwałtowny spadek napięcia na strukturze [8, 13]. Napięcie to odkłada się na warstwie ładunku przestrzennego, którą prawie w całości jest n-baza w przyrządzie. Prąd płynie przez obszar początkowo załączony, którego pole przekroju poprzecznego stanowi kilka procent całkowitego pola pastylki [9-10]. Przewidywania teoretyczne oraz szacunkowe badania i pomiary wykazują, że obszar ten występujący wzdłuż wewnętrznej krawędzi pierścienia katody (konstrukcja bramki centrycznej lub wzmacniającej) ma szerokość $250 \div 550 \text{ }\mu\text{m}$. Pomiary elektryczne prowadzą zwykle do zawyżenia tej wartości.

Model elektryczny dla tej fazy załączania zakłada [8] liniową zmianę napięcia struktury w funkcji gęstości prądu:

$$U(t) = U_0 \left[1 - \frac{J(t)}{qN_D v_s} \right] \quad (6)$$

gdzie: U_0 — napięcie blokowania,

q — ładunek elektronu,

N_D — poziom domieszkowania w n-bazie,

v_s — maksymalna prędkość unoszenia elektronów.

Powyższe wyrażenie może być wykorzystane do określenia pola powierzchni obszaru początkowo załączonego S_0 .

$$S_0 = \frac{I_0}{qN_D v_s} \quad (7)$$

gdzie: I_0 oznacza hipotetyczny prąd, dla którego napięcie na strukturze byłoby równe zeru zgodnie z wyrażeniem (6).

Rozkład gęstość mocy w przyrządzie jako iloczyn gęstości prądu $J(t)$ i pola elektrycznego $E(x, t)$ ma postać [8, 13]:

$$g(x, t) = \frac{2U_0}{W_{Bn}} \left[1 - \frac{J(t)}{qN_D v_s} \right] \left[1 - \frac{x - x_{Bp}}{W_{Bn}} \right] J(t), \quad (8)$$

gdzie: W_{Bn} oznacza szerokość n -bazy w tyrystorze.

W stanie rozprzestrzeniania się kanału przewodzącego, ze względu na podobny profil domieszkowania oraz rozkład nośników nadmiarowych, model tyrystora dla załączonych struktur elementarnych dobrze odpowiada modelowi diody p-s-n. W stanie załączenia, przy dużym poziomie wstrzykiwania ($n(x, t) \gg N_D$) przyjmuje się na podstawie neutralności ładunku równe koncentracje nośników nadmiarowych w słabo domieszkowanym obszarze s_n diody p-s-n [3—4, 13].

$$n(x, t) \simeq p(x, t). \quad (9)$$

Powyższe równanie opisuje ambipolarny stan rozkładu nośników w półprzewodniku przy dużym poziomie wstrzykiwania. Wynika to z faktu, że $\Delta n = n - n_0 \simeq \Delta p = p - p_0$, gdzie: n_0, p_0 są równowagowymi koncentracjami elektronów i dziur w obszarze słabo domieszkowanym. Przyjęcie warunku objętościowej neutralności ładunku do opisu stanów dynamicznych w diodzie p-s-n wynika z krótkiego czasu zaniku zaburzenia obojętności elektrycznej struktury. Przyjmuje się, że stała czasowa zaniku zaburzenia zwana czasem relaksacji dielektrycznej ma postać: $\tau_{rd} = \frac{\varepsilon}{\sigma}$, gdzie: ε ,

$\sigma = q(\mu_n n + \mu_p p)$ oznaczają stałą dielektryczną i konduktywność krzemu.

Przyjęcie typowych danych materiałowych dla struktury tyrystora pozwala oszacować czas relaksacji dielektrycznej na poziomie $10^{-12} \div 10^{-10}$ s, co stanowi wartość o kilka rzędów wielkości mniejszą od najmniejszego kroku czasowego przy iteracyjnej analizie stanów dynamicznych w półprzewodnikowych przyrządach mocy. Jest to podstawą uproszczenia ogólnych równań transportu i ciągłości do postaci ambipolarnej przy dużym poziomie wstrzykiwania.

Skojarzenie równań transportu i ciągłości prowadzi do następujących zależności:

$$\left. \begin{aligned} -R - \frac{\partial}{q \partial x} \left[q \mu_p \left\{ n(x, t) E(x, t) - U_T \frac{\partial n(x, t)}{\partial x} \right\} \right] &= \frac{\partial n(x, t)}{\partial t} \\ -R + \frac{\partial}{q \partial x} \left[q \mu_n \left\{ n(x, t) E(x, t) - U_T \frac{\partial n(x, t)}{\partial x} \right\} \right] &= \frac{\partial n(x, t)}{\partial t} \end{aligned} \right\} \quad (10)$$

gdzie:

$$U_T = \frac{KT}{q} \text{ — potencjał elektrotermiczny,}$$

μ_p — ruchliwość dziur,

μ_n — ruchliwość elektronów,

R — szybkość rekombinacji.

Powyższe równania są łatwe do rozwiązywania przy założeniu stałych (niezależnych od x) ruchliwości μ_p , μ_n . W rzeczywistości w stanie przejściowym, przy załączeniu tyrystora występują znaczne gradienty gęstości nośników nadmiarowych wywołane zarówno dużym poziomem wstrzykiwania jak i nierównomiernym rozkładem temperatury wzdłuż osi OX [3—4, 12]. Przyjęcie w pełni niezależnych funkcji $\mu_p(x)$, $\mu_n(x)$ znacznie komplikuje rozwiązanie równań (10). By uprościć zagadnienie, skrócić czas obliczeń oraz uwzględnić wpływ gęstości nośników nadmiarowych na ruchliwość, założono w niniejszej pracy stały iloraz ruchliwości elektronów i dziur, oznaczając go symbolem mr :

$$mr = \frac{\bar{\mu}_n}{\bar{\mu}_p} \simeq \frac{\mu_n(x)}{\mu_p(x)} \quad (11)$$

przy czym: $\bar{\mu}_n$, $\bar{\mu}_p$ oznaczają średnie wartości ruchliwości dla danej struktury elementarnej.

Współczynnik mr zależy od koncentracji nośników i zmienia się dla krzemu w granicach od 2.6 przy gęstości nośników nadmiarowych 10^{14} cm^{-3} do ok. 1.6 przy gęstości 10^{19} cm^{-3} . Dla mniejszych zmian koncentracji (np. o rząd wielkości) z dobrą dokładnością można przyjąć stałość tego współczynnika. Warunek niewielkiej zmiany koncentracji jest spełniony o ile równania (10) rozwiązywane są metodami iteracyjnymi, a współczynnik „ mr ” uaktualniany jest w każdym kroku.

Relacje (10—11) prowadzą do sformułowania ambipolarnego równania rozkładu nośników w półprzewodniku przy zmiennych ruchliwościach:

$$-R + D(x) \left[\frac{\partial^2 n}{\partial x^2} + \frac{1}{\mu_p} \frac{\partial \mu_p}{\partial x} \frac{\partial n}{\partial x} \right] = \frac{\partial n}{\partial t}, \quad (12)$$

przy czym:

$$D(x) = \frac{2U_T \mu_p(x) \mu_n(x)}{\mu_p(x) + \mu_n(x)} \text{ — ambipolarna stała dyfuzji dla dużego poziomu wstrzykiwania.}$$

Szybkość rekombinacji R dla słabo domieszkowanego obszaru diody p-s-n uwzględnia model rekombinacji zgodny z teorią Shockleya-Reada-Halla oraz rekombinację Augera, która ma istotny wpływ na straty mocy w stanie przewodzenia przy dużym poziomie wstrzykiwania nośników [10]. Duży poziom wstrzykiwania umożliwia uproszczenie współczynnika rekombinacji do postaci [4]:

$$R = \frac{n(x, t)}{\tau} + C_A n^3(x, t) \quad (13)$$

gdzie:

$$\tau = \tau_{p0} + \tau_{n0},$$

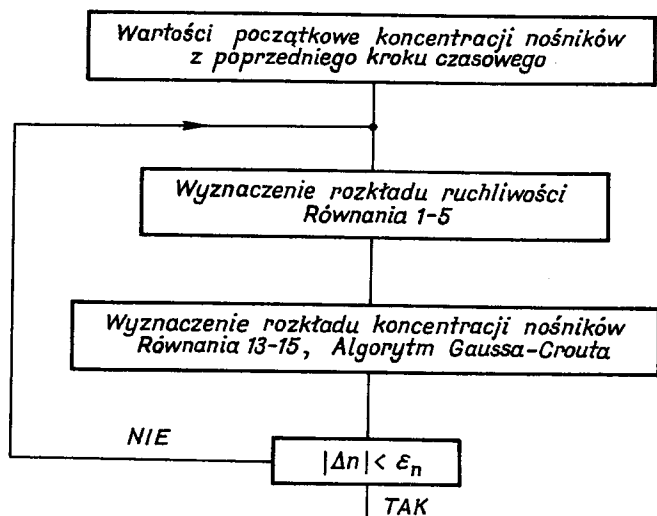
$\tau_{p0}(\tau_{n0})$ — graniczny czas życia dziur (elektronów) w występujący w silnie domieszkowanym półprzewodniku typu n^+ (p^+), parametry modelu rekombinacji Shockleya-Reada-Halla.

Warunki brzegowe wynikają z faktu, że prądy na krańcach obszaru s_n mają charakter prądów dziurowych (od strony obszaru p) i elektronowych (od strony obszaru n), a ich gęstości są jednakowe i wynikają z chwilowej wartości prądu i szerokości obszaru przewodzącego [3, 7].

$$\left. \begin{aligned} \frac{\partial n(x, t)}{\partial x} \Big|_{x=x_{0s_n}} &= -\frac{J(t)}{2q\mu_n U_T} \\ \frac{\partial n(x, t)}{\partial x} \Big|_{x=x_{Ns_n}} &= \frac{J(t)}{2q\mu_p U_T} \end{aligned} \right\} \quad (14)$$

Równania (1—5, 12—14) tworzą układ równań, które rozwiązuje się metodą różnic skończonych. Wymaga to sformułowania układu równań nieliniowych o liczbie niewiadomych określonych liczbą punktów, dla których wyznacza się nieznane koncentracje nośników. W praktyce rozmiary macierzy dochodzą do kilkudziesięciu, a nawet kilkuset wierszy i kolumn, co komplikuje rozwiązanie i wydłuża czas obliczeń. Stosuje się znane metody rozwiązywania układów równań nieliniowych np. metodą Newtona lub siecznych.

W pracy zastosowano inne podejście znacznie skracające czas obliczeń. Zaproponowano przedziałową linearyzację nieliniowej części równania (13) oraz metodę kolejnych przybliżeń dla układu równań (1—5, 12—14). Dużą zbieżność metody zapewnia przyjęcie punktu startu dla każdej iteracji na podstawie danych z kroku poprzedniego. Szybkość rekombinacji Augera przybliżono wykorzystując liniową aproksymację funkcji $n^3(x, t)$ postaci:



Rys. 3. Algorytm wyznaczania koncentracji nośników nadmiarowych metodą kolejnych przybliżeń.
 Δn zmiana koncentracji, ε_n założona dokładność wyznaczania koncentracji

$$n_{l+1}^3(x, t) \simeq 3n_l^2(x, t) [n_{l+1}(x, t) - n_l(x, t)] + n_l^3(x, t), \quad (15)$$

przy czym „ l ” jest symbolem określającym numer iteracji dla metody kolejnych przybliżeń przyjętej do rozwiązania układu (1—5, 12—15).

Procedurę wyznaczania koncentracji nośników nadmiarowych, która uwzględnia wpływ temperatury i gęstości nośników na ruchliwość przedstawia rys. 3. Zastosowano algorytm Gaussa-Crouta z częściowym wyborem elementu podstawowego do rozwiązywania układu równań liniowych z macierzą rzadką.

Ze względu na duże gradienty temperatury istnieje potrzeba uwzględnienia efektu wypierania prądu z obszarów o wyższej temperaturze w quasidwuwymiarowej, elektrotermicznej analizie procesu załączania elementu. Efekt ten wynika z termicznej zmiany ruchliwości, a przez to i zmiany konduktywności obszarów s_n struktur elementarnych (Rys. 2). Obszary o wyższej temperaturze cechuje mniejsza konduktywność, co przy stałym napięciu na wszystkich strukturach elementarnych pociąga za sobą zmniejszenie gęstości prądu.

$$T \rightarrow, (\mu_p, \mu_n) \rightarrow, \sigma = p\mu_p + n\mu_n \rightarrow, I \rightarrow J \rightarrow. \quad (16)$$

W celu wyznaczenia rozkładu gęstości prądu w poszczególnych strukturach elementarnych zastosowano optymalizację gradientową definiując funkcję celu przez rozrzut napięć $\Delta U = \Delta U_{\max} - U_{\min}$ na strukturach elementarnych.

$$\Delta U = f(I_1, I_2, \dots) \Big|_{\sum_k I_k = I} < \varepsilon, \quad (17)$$

przy czym: I_k — prąd w k -tej strukturze, I — prąd całkowity, $U_{\min}$, $U_{\max}$ — minimalny i maksymalny spadek napięcia w zbiorze struktur elementarnych, ε — założony prąd wyznaczenia napięcia.

Do rozwiązania zagadnienia optymalizującego zastosowano algorytm wykorzystujący przybliżenie gradientu wielowymiarowej funkcji celu $\Delta U = f(I_1, I_2, \dots)$ postaci:

$$I_k^{p+1} = I_k^p - \alpha \sigma_k^p (U_k^p - \bar{U}^p) \quad (18)$$

przy czym:

I_k^p, U_k^p — prąd i napięcie k -tej struktury elementarnej w p -tym kroku optymalizacji,

$\bar{U}^p = (\sum_k U_k^p) / M$ — napięcie średnie, M — liczba załączonych struktur elementarnych,

$\sigma_k^p = I_k^p / U_k^p$ — konduktywność zastępcza k -tej struktury,

α — współczynnik określający zbieżność.

By spełnić warunek równości prądu zaciskowego z sumą prądów struktur elementarnych, po osiągnięciu przez funkcję celu wartości minimalnej (z odpowiednią dokładnością) następuje proporcjonalna korekcja wartości prądów. Napięcie na każdej strukturze elementarnej wyznaczono poprzez całkowanie pola elektrycznego [3,7]. Dodatkowo uwzględniono spadki napięć na skrajnych złączach diody p-s-n [7].

$$U = \int_{x_{0Sn}}^{x_{NSn}} \left(\frac{J}{qn(\mu_p + \mu_n)} - \frac{\mu_p - \mu_n}{\mu_p + \mu_n} U_T \frac{1}{n} \frac{\partial n}{\partial x} \right) dx + U_T \ln \left(\frac{n_{0Sn} P_{NSn}}{n_i^2} \right), \quad (19)$$

gdzie:

n_{OSn}, p_{NSn} — brzegowe koncentracje elektronów i dziur w obszarze s_n ,
 n_i — koncentracja elektronów w półprzewodniku samoistnym.

Przedstawiona procedura iteracyjna wykazuje dobrą zbieżność. W praktyce po kilku krokach osiągnano zadowalającą dokładność ($\Delta U/\bar{U} < 1\%$).

Wyznaczenie koncentracji nośników nadmiarowych w słabo domieszkowanym obszarze diody p-s-n umożliwia określanie gęstości mocy traconej w strukturze tyrystora podczas rozprzestrzeniania się kanału przewodzącego. Wykorzystując powyższe rozważania oraz uwzględniając moc strat rekombinacyjnych, całkowitą gęstość wydzielanej mocy w przyrządzie można przedstawić w postaci:

$$g(x, t) = J(t) E(x, t) + R(x, t) E_g, \quad (20)$$

przy czym E_g oznacza szerokość pasa zabronionego w krzemie.

5. ROZPRZESTRZENIANIE SIĘ PRZEWODZĄCEGO KANAŁU W TYRYSTORZE

Jednym z parametrów dynamicznych tyrystorów jest prędkość rozprzestrzeniania się przewodzącego kanału (plazmy) podczas załączania. Efekt rozprzestrzeniania się przewodzącego kanału wynika głównie z gradientu koncentracji nośników nadmiarowych we wszystkich obszarach tyrystora, a w szczególności w p-bazie i n-bazie oraz z faktu istnienia poprzecznego pola elektrycznego, które jest szczególnie silne w obszarze przejściowym (na granicy obszaru przewodzącego i nieprzewodzącego).

Do podstawowych czynników, które istotnie wpływają na prędkość rozprzestrzeniania się kanału przewodzącego można zaliczyć [2, 15, 18—23, 26]:

- gęstość prądu,
- poprzeczne pole elektryczne,
- poziom domieszkowania i szerokość p-bazy i n-bazy,
- czas życia nośników mniejszościowych w n-bazie,
- efekt zwarcia emiterowego.

Dyfuzja nośników nadmiarowych i unoszenie przez silne pole elektryczne są głównymi elementami mechanizmu rozprzestrzeniania się przewodzącego kanału. Dyfuzja nośników zależy od gradientu ich koncentracji, a ta od gęstości prądu w obszarze przewodzącym. Wzrost gęstości prądu powoduje wzrost prędkości rozchodzenia się kanału przewodzącego. Badania eksperymentalne doprowadziły do sformułowania wielu zależności, wśród których często cytowana ma postać [15, 22]:

$$V = C J^n \quad (21)$$

gdzie: C i $n = 2-6$ są stałymi.

Poprzeczne pole elektryczne występuje zarówno w p-bazie [18] jak i w n-bazie [19]. Jest ono zlokalizowane głównie w obszarze przejściowym na granicy obszaru

przewodzącego i nieprzewodzącego. Obszar przejściowy jest wąski. Przyjmuje się, że jego szerokość jest w przybliżeniu równa drodze dyfuzji w krzemie, co oznacza, że przy częściowym przewodzeniu struktury, przy napięciach zaciskowych rzędu kilku woltów poprzeczne pole elektryczne w okolicy złącza kolektorowego może dochodzić do poziomu 10^3 V/cm. Uproszczona analiza poprzecznego rozkładu potencjału w p-bazie prowadzi do sformułowania znanej w literaturze relacji prędkości rozprzestrzeniania się kanału przewodzącego i gęstości prądu:

$$V = A \ln(J) + B, \quad (22)$$

gdzie: A i B są stałymi.

Poziom domieszkowania i szerokość p-bazy wpływa na współczynnik wzmocnienia prądowego tranzystora n-p-n struktury czterowarstwowej [26]. Ze wzrostem domieszkowania maleje wzmocnienie, maleje głębokość dodatniego sprzężenia zwrotnego ($\alpha_1 + \alpha_2 \rightarrow 1$) oraz zmniejsza się prędkość rozprzestrzeniania się przewodzącego kanału. Przeprowadzone badania empiryczne wskazują, że wzrost domieszkowania w granicach $N_A = 10^{17} - 10^{18} \text{ cm}^{-3}$ powoduje zmniejszenie prędkości rozprzestrzeniania się plazmy z 500 m/s do 50 m/s przy gęstości prądu około 400 A/cm^2 .

$$V \sim N_A^{-0.9} \quad (23)$$

Wpływ szerokości bazy obowiązuje zarówno dla p-bazy jak i n-bazy. W obu przypadkach wraz ze zmniejszeniem szerokości obszaru bazy rośnie współczynnik wzmocnienia prądowego (α_1, α_2). Zależność uzyskana z doświadczeń dla struktury domieszkowanej złotem w temperaturze 860°C , ze zwartym emiterem o średnicy zwarć $d = 0.3 \text{ mm}$ i odległości między nimi $D = 1.1 \text{ mm}$, przy gęstości prądu $J = 100 \text{ A/cm}^2$ ma postać [15]:

$$V \sim W_{Bn}^{-1}, \quad \text{gdzie } l \simeq 1. \quad (24)$$

Uwzględniając zależność prędkości od czasu życia nośników mniejszościowych w n-bazie postaci $V \sim \tau_p^{1/2}$ [15, 19] można otrzymać kolejną relację wpływu parametrów struktury tyrystora na prędkość rozprzestrzeniania się przewodzącego kanału:

$$V \sim \frac{\tau_p^{1/2}}{W_{Bn}} \sim \frac{L}{W_{Bn}} \quad (25)$$

gdzie: L oznacza drogę dyfuzji w n-bazie.

Analogiczne rozważania można przeprowadzić dla czasu życia nośników mniejszościowych i szerokości p-bazy.

Wpływ zwarcia emiterowego na szybkość rozprzestrzeniania się przewodzącego kanału wynika głównie ze zmiany wzmocnienia prądowego tranzystora n-p-n, struktury. Definiuje się [15] współczynnik „gęstości zwarcia emiterowego” określony przez średnicę zwarć (d) i ich wzajemne odległości (D). Wzrost „gęstości zwarcia” (zmniejszenie średnic i ich wzajemnych odległości) wywołuje zmniejszenie szybkości rozprzestrzeniania się plazmy.

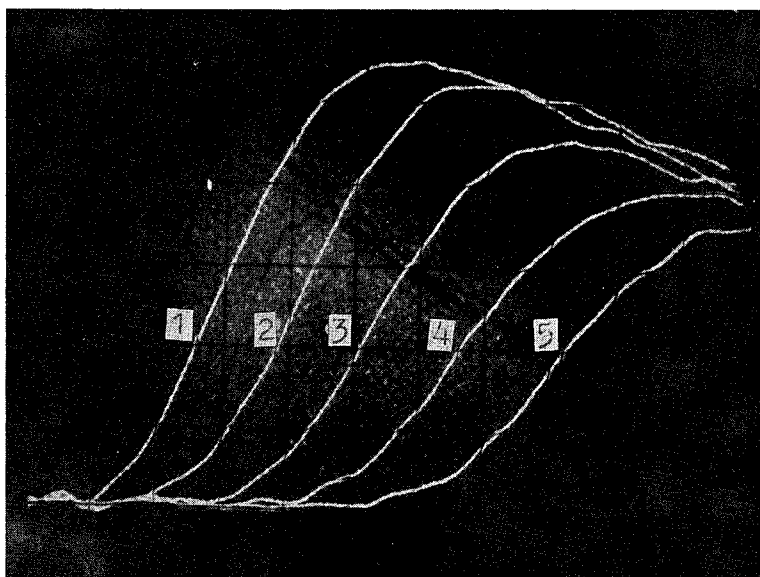
$$\ln(V) \sim \frac{1}{16}(d^2 + D^2[2\ln(D/d) - 1]). \quad (26)$$

Znane są także bezpośrednie badania wpływu wzmocnień prądowych obu tranzystorów jako elementów tyrystora [26] na szybkość rozprzestrzeniania się przewodzącego kanału. Prowadzą one do wniosku o liniowej zależności tej prędkości od logarytmu iloczynu α_1, α_2 , przy czym gęstość prądu jest tu parametrem.

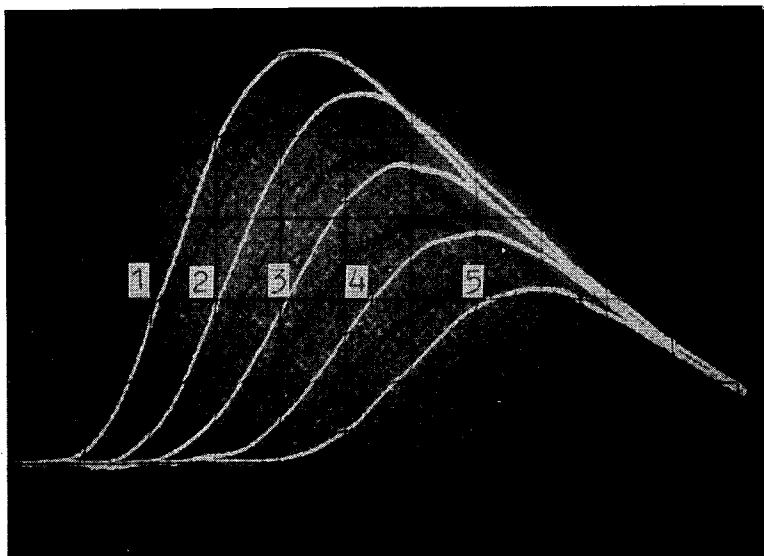
$$V \sim \ln(\alpha_1 \alpha_2). \quad (27)$$

Do dziś nie opracowano spójnej teorii rozprzestrzeniania się plazmy w strukturach czterowarstwowych. Każdy z przedstawionych wyżej modeli wymaga empirycznej identyfikacji parametrów. Wiąże się to z pomiarami elektrycznymi lub optycznymi specjalnie przygotowanych struktur. W pracy dla celów modelowania elektrotermicznego wykonano badania mierząc sygnał wyjściowy detektora promieniowania rekombinacyjnego w pięciu punktach struktury leżących na promieniu pastylki (tyrystor T00-40) w następujących odległościach od jej środka:

- 1—0.30 cm,
- 2—0.35 cm,
- 3—0.40 cm,
- 4—0.45 cm,
- 5—0.50 cm.



Rys. 4. Intensywność promieniowania rekombinacyjnego w różnych punktach promienia pastylki, w jednostkach względnych $I_{TM} = 200$ A, $t_{iT} = 275 \mu s$, $10 \mu s/działkę$



Rys. 5. Intensywność promieniowania rekombinacyjnego w różnych punktach promienia pastylki, w jednostkach względnych $I_{TM} = 500$ A, $t_{iT} = 118$ μ s, 10 μ s/działkę

Promień pastylki tyrystora wynosił $r_{max} = 0.78$ cm. Pomiary przeprowadzono dla impulsów prądowych (półfale sinusoid) o amplitudach w zakresie $I_{TM} = 100$ A $\div$ 500 A przy czasach trwania $t_{iT} = 50 \div 275$ μ s, a przykładowe wyniki przedstawiono na rys. 4 $\div$ 5. Początek wykresów oznacza chwilę załączenia tyrystora. Numeracja na wykresach odpowiada położeniu badanego punktu na promieniu i jest zgodna z powyższym wykazem. Uzyskane charakterystyki posłużyły do wyznaczenia promienia przewodzącej części struktury jako funkcji czasu. Uwzględniono opóźnienie wnoszone przez detektor. Dane empiryczne o szybkości rozprzestrzeniania się kanału wraz z elektrycznym pomiarem pola powierzchni obszaru, początkowo załączonego (Równ. 7) zostały wykorzystane do elektrotermicznego modelowania procesu załączania tyrystora.

6. MODEL CIEPLNY

W modelu termicznym założono wielowarstwowe rozchodzenie się ciepła. Źródłem ciepła są fragmenty pastylki tyrystora, n-baza dla fazy zaniku ładunku przestrzennego oraz obszar $x_{OSi} - x_{NSi}$ dla fazy rozprzestrzeniania się kanału przewodzącego — (Rys. 2). Pominięto ciepło wydzielane w obszarach silnie domieszkowanych. Spadki napięć na tych obszarach przy gęstościach prądu rzędu 1000 A/cm² wynoszą zaledwie około 50 mV [4], co przy napięciach przewodzenia w stanach dynamicznych ok. 2 $\div$ 8 V oznacza niewielki, co najwyżej kilku procentowy, udział strat mocy w tych obszarach w ogólnym bilansie energetycznym.

Dwuwymiarowy model rozptywu ciepła dla warstwy, w której generowane jest ciepło ma postać:

$$\frac{\partial^2 T}{\partial x^2} + \frac{\partial^2 T}{\partial r^2} + \frac{\partial T}{r \partial r} = \frac{\partial T}{\kappa \partial t} + \frac{g(x, r, t)}{\lambda(T)} \quad (28)$$

gdzie: $T = T(x, r, t)$ — temperatura,

t — czas,

λ — przewodność cieplna,

κ — $\lambda/c\rho$, c — ciepło właściwe, ρ — gęstość materiału,

$g(x, r, t)$ — gęstość mocy wydzielanej w strukturze.

Podobne wyrażenie obowiązuje dla warstw, które ciepło przejmują. W tych przypadkach w powyższym równaniu $g(x, r, t) = 0$.

Na górnej powierzchni pastylki tyrystora założono wymianę ciepła poprzez konwekcję naturalną oraz adiabatyczne warunki termiczne dla krawędzi bocznych. Pominęto wpływ promieniowania na straty ciepła [16]. Przyjęto izotermiczne warunki brzegowe na dolnej krawędzi molibdenu, co wynika z grubości podstawy molibdenowej (około 800 μm) i czasu trwania analizowanych elektrotermicznych stanów przejściowych (max. 300 μs).

$$\left. \begin{aligned} T|_{x=x_{\max}, r, t} &= T_a \\ \frac{\partial T}{\partial r}|_{x, r=r_{\max}, t} &= 0 \end{aligned} \right\} \quad (29)$$

przy czym T_a jest temperaturą otoczenia.

Wymianę ciepła poprzez konwekcję opisuje prawo Newtona, które zakłada, że strumień wymienianego ciepła jest wprost proporcjonalny do różnicy temperatur: temperatury na krawędzi pastylki i temperatury otoczenia.

$$\lambda \frac{\partial T}{\partial r} \Big|_{x_0, r, t} = -\alpha [T(x_0, r, t) - T_a] \quad (30)$$

gdzie: α — współczynnik przejmowania ciepła,

T_a — temperatura otoczenia,

$T(x, r, t)$ — temperatura punktu na powierzchni pastylki w chwili t .

Uwzględniono wpływ temperatury na przewodność cieplną λ [5, 16].

$$\lambda(T) = \frac{311000}{T^{1.33}} \text{ [W/mK]} \quad (31)$$

Zastąpienie pochodnych w równaniu (28) różnicami skończonymi określonymi dla dwuwymiarowej siatki (równomiernej) i uzależnienie temperatury danego węzła od temperatur węzłów sąsiadujących dla chwili poprzedzającej prowadzi do sformułowania tzw. jawnego algorytmu rozwiązywania zagadnienia brzegowego.

$$T_{i,k}^{j+1} = T_{i,k}^j + \frac{\Delta t}{cQ} \left(g_{i,k}^j + \lambda \left[\frac{T_{i-1,k}^j - 2T_{i,k}^j + T_{i+1,k}^j}{(\Delta x)^2} + \frac{T_{i,k-1}^j - 2T_{i,k}^j + T_{i,k+1}^j}{(\Delta r)^2} + \frac{T_{i,k+1}^j - T_{i,k-1}^j}{2r_k \Delta r} \right] \right) \quad (32)$$

przy czym:

$$T_{i,k}^j = T(x_i, r_k, t_j),$$

$\Delta x, \Delta r, \Delta t$ — kroki zmiennych x, r, t .

7. PRZEBIEG I WYNIKI SYMULACJI

Procedurę obliczeniową modelowania elektrotermicznego można podzielić na dwa etapy. Oba fragmenty wykonywane są w pętli kolejno, jeden po drugim, aż do zakończenia czasu trwania impulsu grzejnego.

1. Wyznaczenie gęstości mocy w strukturze z modelu elektrycznego (Rys. 3, Równ. 8.20) przy uwzględnieniu:

- pola powierzchni przewodzącego kanału (detekcja promieniowania rekombinacyjnego),
- zjawiska wypierania prądu obszarów o wyższej temperaturze (procedura optymalizacyjna, Równ. 17—18),

2. Wyznaczenie rozkładu temperatury przy użyciu modelu cieplnego.

Dla celów symulacji tyristor podzielono na $M = 100$ struktur elementarnych. Dla każdej z nich przeprowadzono analizę elektryczną w $N = 30$ punktach. Wynika stąd, że dwuwymiarowa symulacja termiczna struktury krzemowej (prócz warstwy Mo) dotyczy 3000 węzłów. Przedstawiono wyniki rozkładu temperatury w płaszczyźnie Oxr tyristora w czasie trwania pierwszego impulsu grzejnego. Dodatkowo wykreślono wybrane rozkłady gęstości prądu wzdłuż osi Or ilustrując efekt wypierania prądu z obszarów o wyższej temperaturze. Wykonano obliczenia dla impulsów o kształtach półfal sinusoid, o amplitudach $I_{TM} = 500$ A i czasach trwania $t_{iT} = 50 \mu s$ i $118 \mu s$.

Do obliczeń przyjęto dane konstrukcyjne tyristora T00-40:

$$x_{Mo} = 406 \mu m,$$

$$x_{max} = 1230 \mu m,$$

$$x_{OSn} = 20 \mu m,$$

$$x_{NSn} = 386 \mu m,$$

$$r_{max} = 780 \mu m,$$

$$W_{Bn} = 246 \mu m,$$

$$N_D = 10^{14} \text{ cm}^{-3},$$

$$\tau = 10 \mu s,$$

$$C_A = 2 \cdot 10^{31} \text{ cm}^{-6} \text{ s}^{-1},$$

$$v_s = 10^7 \text{ cms}^{-1},$$

$$E_{g0} = 1.12 \text{ eV},$$

$$T_a = 296 \text{ K},$$

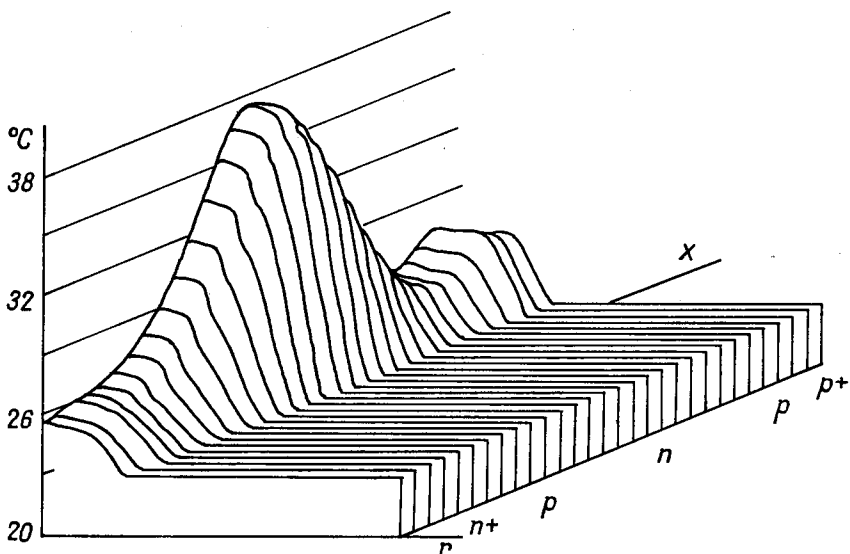
$$\rho = 2.328 \text{ gcm}^{-3}, \text{ gęstość krzemu,}$$

$$c = 0.71 \text{ Jg}^{-1}\text{K}^{-1}, \text{ ciepło właściwe krzemu,}$$

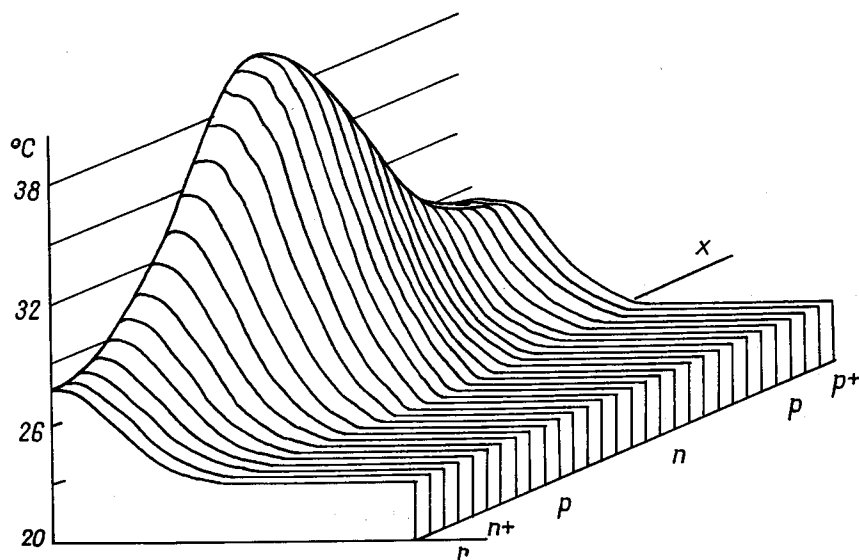
$$\alpha = 0.004 \text{ Wcm}^{-2}\text{K}^{-1}, \text{ współczynnik przejmowania ciepła.}$$

Wykresy temperatur przedstawiono w trójwymiarowym układzie współrzędnych liniowych, przy czym stosunek skal na osiach Ox i Oy wynosi $7800 \mu\text{m}/406 \mu\text{m}$. Uzyskane wyniki rozkładu temperatury wewnątrz i na powierzchni struktury potwierdzają wpływ stromości prądowej w tyrystorze na poziom lokalnego przegrzania w okolicach bramki. Przyrost temperatury na powierzchni przy dużej stromości prądu może być wielokrotnie mniejszy niż w warstwach wewnętrznych. Nierównomierny rozkład gęstości prądu w strukturze wywołany termiczną modulacją przewodności półprzewodnika zmniejsza lokalny wzrost temperatury. Straty rekombinacyjne uwytłaczają się przy dużej gęstości prądu w tyrystorze w okolicach skrajnych złączy struktury. Duży wpływ na poziom wydzielanej mocy ma dokładność wyznaczenia obszaru początkowo załączonego S_0 .

W celu weryfikacji przyjętych założeń porównano napięcie pomierzone na strukturze w stanie dynamicznym podczas rozprzestrzeniania się kanału przewodzącego z wartościami obliczonymi (Rys. 11). Pomiar przeprowadzono metodą oscylograficzną przy następujących parametrach: $I_{TM} = 500 \text{ A}$, $t_{iT} = 118 \mu\text{s}$, $I_{GM} = 0.6 \text{ A}$. Małe różnice między wartościami pomierzonymi i obliczonymi uzyskano po uwzględnieniu, prócz wpływu temperatury, także wpływu gęstości nośników nadmiarowych na ruchliwość oraz zjawisko wypierania prądu z obszarów o wyższej temperaturze.

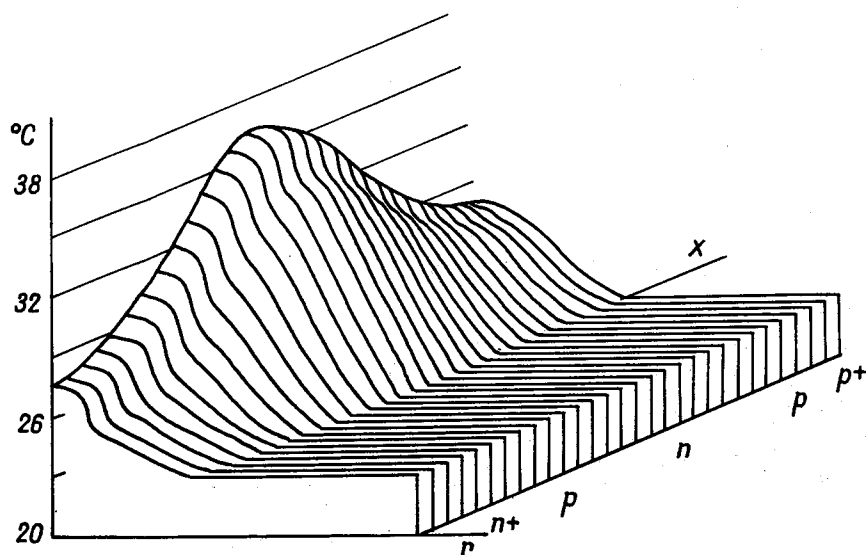


Rys. 6. Rozkład temperatury w przyrządzie $I_{TM} = 500 \text{ A}$, $t_{iT} = 50 \mu\text{s}$, $t = 5 \mu\text{s}$

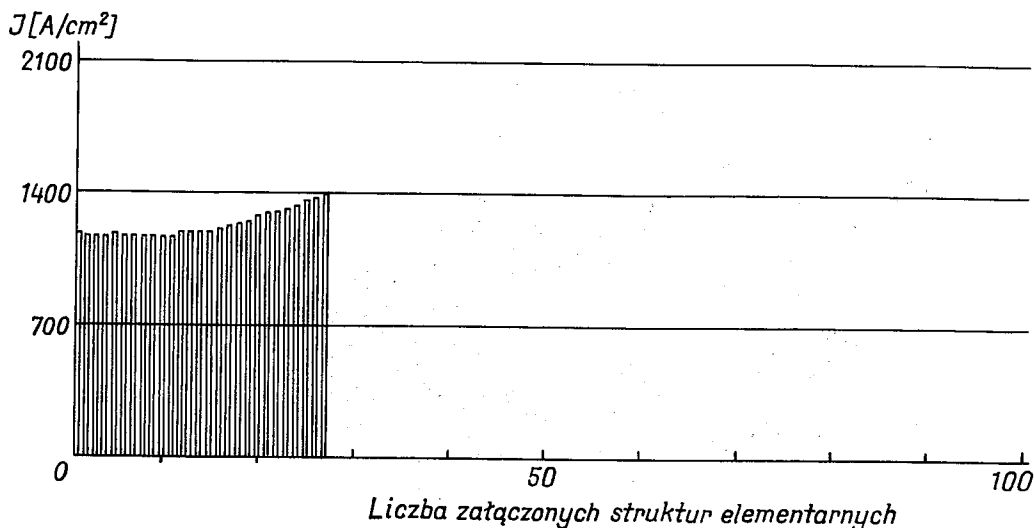


Rys. 7. Rozkład temperatury w przyrządzie $I_{TM} = 500$ A, $t_{iT} = 50$ μ s, $t = 15$ μ s

Przy braku oddziaływań elektrotermicznych (warunki izotermiczne w temperaturze pokojowej) ruchliwość zależy wyłącznie od koncentracji nośników i nie występuje efekt nierównomiernego rozkładu gęstości prądu. Przy symulacji warunków izotermicznych założono stałą temperaturę $T = 296$ K we wzorach 2.1 ÷ 2.5 oraz pominięto procedurę

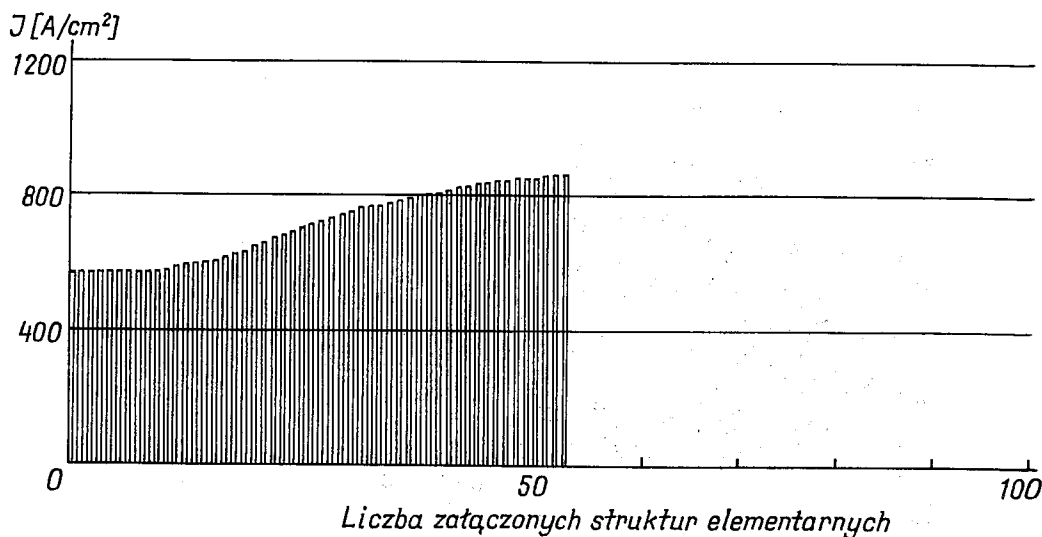


Rys. 8. Rozkład temperatury w przyrządzie $I_{TM} = 500$ A, $t_{iT} = 118$ μ s, $t = 10.16$ μ s

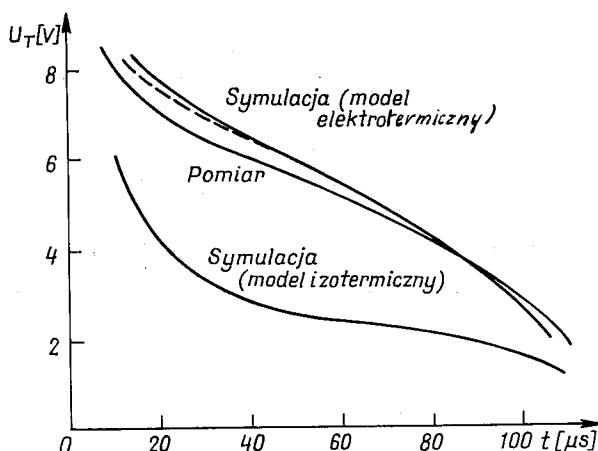


Rys. 9. Rozkład gęstości prądu w strukturze $I_{TM} = 500$ A, $t_{iT} = 50$ μ s, $t = 5$ μ s

optymalizacyjną (Równ. 19). Przeprowadzono obliczenia przy uwzględnieniu opóźnienia układu detekcji promieniowania rekombinacyjnego — linia przerywana. Rozbieżności wynikające ze skończonego czasu reakcji detektora są istotne jedynie w początkowym zakresie otrzymanych wykresów i nie przekraczają kilku procent wartości dokładnej.



Rys. 10. Rozkład gęstości prądu w strukturze $I_{TM} = 500$ A, $t_{iT} = 50$ μ s, $t = 15$ μ s



Rys. 11. Napięcie przewodzenia tyrystora w czasie rozprzestrzeniania się przewodzącego kanału
 $I_{TM} = 500 \text{ A}$, $t_{iT} = 118 \text{ μs}$

W tabelicy 1 przedstawiono wyniki przykładowej iteracji (Równ. 19), która jest elementem procedury optymalizacyjnej modelującej zjawisko wypierania prądu. Średni rozrzut napięć wyrażony jako $\Sigma |U_k - \bar{U}|/M$, gdzie: M — liczba załączonych struktur elementarnych, jest kryterium określającym dokładność analizy. Zgodność prądów: obliczonego ΣI_k oraz płynącego przez przyrząd jest warunkiem pomocniczym, przy czym spełnienie obu warunków kończy obliczenia w danym kroku.

Tablica 1

Przykład przebiegu procedury optymalizacyjnej
 wg równania 2.19.

Krok	Średni rozrzut napięć [V]	ΣI_k [A]	I [A]
1	0.83	315.313	293.892
2	0.48	296.558	"
3	0.29	295.418	"
4	0.18	295.002	"
5	0.12	294.335	"
6	0.08	294.396	"
7	0.06	293.960	"
8	0.04	294.241	"
9	0.02	294.003	"

Dla postawionego problemu (dwuwymiarowego, elektrotermicznego modelowania stanów dynamicznych w tyrystorze) czas obliczeń jest stosunkowo krótki. Program analizy elektrotermicznej tyrystora dla stanu załączenia „THERM” napisano w języku

Tablica 2

Porównanie czasów obliczeń programu „THERM”

IBM-386 Turbo	IBM-XT, 8087	IBM-XT
3.5 h pełny proces załączania	9 h	11' jedna iteracja

„C” przy użyciu kompilatora firmy Microsoft Co. wersja 4.0, stosując optymalizację czasową. Porównanie przedstawia poniższa tablica (cały proces załączania odpowiada 100 iteracjom).

PODSUMOWANIE

Prezentowano elektrotermiczny model struktury czterowarszawowej dla stanu załączania. Część modelu dotycząca zjawisk elektrycznych w przyrządzie ma charakter quasidwuwymiarowy, co z jednej strony zapewnia dużą wierność odwzorowania rzeczywistości, a z drugiej stosunkowo krótki czas obliczeń [16]. Model cieplny jest modelem dwuwymiarowym. Uwzględniono zmienność ruchliwości μ_p , μ_n jako efekt wpływu temperatury i gęstości nośników namiarowych, co ma duże znaczenie przy dużym poziomie wstrzykiwania. Stosując metodę optymalizacji gradientowej uwzględniono zjawisko wypierania prądu z obszarów o wyższej temperaturze.

WNIOSKI

1. Duża szybkość rozprzestrzeniania się przewodzącego kanału ($V > 100$ m/s) w początkowej fazie załączania ($t < 15$ μ s) wyraźnie łagodzi efekt lokalnego przegrzania struktury,
2. Termiczna modulacja konduktywności półprzewodnika powodująca nierównomierny rozkład gęstości prądu zwiększa odporność tyrystora na wymuszenia prądowe o średnich i dużych stromościach [$(di/dt) > 50$ A/ μ s],
3. Duży wpływ na dokładność wyznaczenia maksymalnej temperatury struktury podczas załączania ma precyzja określenia powierzchni obszaru początkowo załączonego,
4. Straty rekombinacyjne mają istotne znaczenie przy dużych gęstościach nośników nadmiarowych, co ma miejsce w początkowej fazie załączania ($t < 15$ μ s), w okolicy skrajnych złączy tyrystora.

BIBLIOGRAFIA

1. M. Adler, V. Temple: *The Dynamics of the Thyristor Turn-On Process*. IEEE Trans. ED, Vol. 27, No. 2, February 1980
2. M. Adler: *Details of the Plasma-Spreading Process in Thyristors*. IEEE Trans. ED, Vol. 27, No. 2, February 1980

3. H. Benda, E. Spenke: *Reverse Recovery Process in Silicon Power Rectifiers*. Proc. IEEE. Vol. 55, No. 8, August 1967
4. A. Blicher: *Thyristor Physics*. Springer-Verlag, New York, 1976
5. S. Gaur: *Safe Operating Area for Bipolar Transistors*. IBM Journal of Research and Development, Vol. 21, No. 5, September 1977
6. S. Gauer, D. Navon: *Two Dimensional Carrier Flow in a Transistor Structure Under Nonisothermal Conditions*. IEEE Trans. ED-23, No. 1, 1976
7. A. Herlet: *The Forward Characteristic of Silicon Power Rectifiers at High Current Densities*. Solid-State Electronics, Vol. 11, Pergamon Press 1968
8. A. Jaecklin: *The First Dynamic Phase at Turn-On of a Thyristor*. Trans. IEEE ED-23, No. 8, August 1976
9. A. Jaecklin: *On the Dynamic Phase of Thyristor Turn-On*. SEV/VSE Bulletin (Brown Boveri). Vol. 68, 1977
10. Jaecklin, M. Lietz: *Analysis of a Power Semiconductor Near Its Threshold of Destruction*. IEEE Power Electronics Specialist Conf., Murray Hill/N.Y., June, 1974
11. W. Janke: *Modele elektrotermiczne elementów półprzewodnikowych*. Zeszyty naukowe Politechniki Gdańskiej, Gdańsk, 1984
12. R. Kokosa: *The Potential and Carrier Distribution of p-n-p-n Device in the ON State*. Proc. IEEE. Vol. 55, No. 8, August 1967
13. Z. Lisik: *Numeryczny model cieplny tyrystora dla stanów dynamicznych*. Rozprawa doktorska, Instytut Elektroniki, Politechnika Łódzka, 1980
14. Z. Lisik: *Model generacji ciepła w tyrystorze podczas załączania*. IV KK „Teoria Obwodów i Układy Elektroniczne”, Drzonków, 1981
15. T. Matsuzawa: *Spreading Velocity of the ON State in High-Speed Thyristors*. Electrical Eng. Japan, Vol. 93. No. 1, 1973
16. A. Napieralski: *Komputerowe projektowanie układów półprzewodnikowych dużej mocy ze szczególnym uwzględnieniem ich właściwości termicznych*. Zeszyty Naukowe Politechniki Łódzkiej, Nr 562, Łódź 1988
17. A. Nowakowski: *Badanie procesów termicznych w przyrządach półprzewodnikowych*. Zeszyty Naukowe Politechniki Gdańskiej, Nr 60, Gdańsk, 1985
18. H. Ruhl: *Spreading Velocity of the Active Area Boundary in a Thyristor*. IEEE Trans. ED, Vol. 17, No. 9, September 1970
19. N. Sawaki, N. Sato, T. Arizumi: *Theory of Plasma Spreading Velocity in a Thyristor*. Japanese Journal of Applied Physics, Vol. 14, No. 12, 1975
20. N. Sawaki: *High and Lateral Current at Turn-on of Large Area Thyristor*. Japanese Journal of Applied Physics, Vol. 16, Supplement 16—1, 1977
21. N. Sawaki, T. Nishinaga: *Plasma Spreading and Turn-On Delay in a Power Thyristor*. Japanese Journal of Applied Physics, Vol. 17, No. 7, July 1978
22. I. Somos, D. Piccone: *Plasma Spread in High-Power Thyristor Under Dynamic and Static Condition*. IEEE Trans. ED, Vol. 17. No. 9, September 1970
23. M. Suzuki, N. Sawaki, K. Iwata, T. Nishinaga: *Current Distribution at the Lateral Spreading of Electron-Hole Plasma in a Thyristor*. IEEE Trans. ED, Vol. 29, No. 8, August 1982
24. Y. Terasawa, H. Fukui: *Calculation of Two-Dimensional Temperature Rise Distribution in a Thyristor*. IEEE Trans. ED, August 1976
25. B. Więcek: *Dwuwymiarowy elektrotermiczny model tyrystora*. IX KK „Teoria obwodów i układów elektronicznych”, Myszkowce-Rzeszów, 1989
26. A. Zekry, W. Gerlach: *Investigation of the Lateral Turn-On Process of Thyristor with an Epitaxial p-Base*. IEEE Trans. ED, Vol. 30, No. 2, February 1983

B. WIĘCEK

TWO-DIMENSIONAL ELECTROTHERMAL ANALYSIS OF TURN-ON PROCESS
IN P-N-P-N POWER DEVICES

Summary

Two-dimensional, electrothermal model of p-n-p-n power devices with taking into the special considerations turn-on process phenomenons is presented. Carriers concentration influence on their mobility is mentioned. Effect of nonuniform current density distribution is modelled. Measurements of conductive area with the help of recombination radiance are used to the computer calculations. The model is verified with the use of thyristor forward voltage.

Metoda formułowania równań hybrydowych nieliniowych obwodów rezystancyjnych

MICHAŁ TADEUSIEWICZ, JADWIGA BEK

Instytut Podstaw Elektrotechniki, Politechnika Łódzka

Otrzymano 1989.3.15

Autoryzowano do druku 1989.11.30

W pracy przedstawiono metodę formułowania równań hybrydowych DC obwodów nieliniowych. Wykazano, że w procesie tworzenia równań występuje konieczność wyznaczania inwersji pewnych macierzy. Struktura tych macierzy umożliwia opracowanie szybkiej i łatwej do zalgorytmizowania procedury. Zaproponowana metoda polega na dekompozycji odpowiednich obwodów i wykonywaniu operacji macierzowych przy wykorzystaniu formuły Woodbury'ego na podobwodach systematycznie powiększanych i łączonych. Opracowano kilka wariantów algorytmu. Wykazano, że metoda może być również bezpośrednio zastosowana do formułowania równań Sandberga-Willsona opisujących DC obwody diodowo-tranzystorowe. Przytoczono przykłady liczbowe potwierdzające efektywność algorytmu.

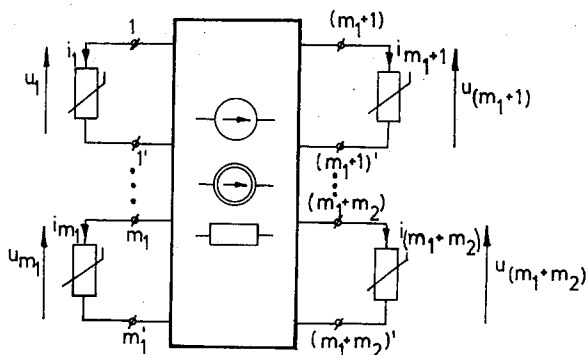
1. WPROWADZENIE

Metoda hybrydowa należy do ogólnych metod formułowania równań obwodów nieliniowych. Jej cechą charakterystyczną jest to, że liczba równań jest taka sama jak liczba elementów nieliniowych. Ponadto w każdym równaniu występuje tylko jedna funkcja nieliniowa jednej zmiennej. Metoda prowadzi więc do małej liczby równań w obwodach zawierających niewiele elementów nieliniowych. Struktura tych równań jest bardzo dogodna zarówno do analizy ilościowej jak i jakościowej. Wystarczy tu nadmienić, że jedyny aktualnie skuteczny algorytm obliczania wszystkich rozwiązań DC obwodów nieliniowych (Chua-Ying [2]) odnosi się do opisu hybrydowego. Również wiele rezultatów dotyczących problemu istnienia i jednoznaczności rozwiązań formułuje się w kategoriach macierzy hybrydowej lub jej szczególnych przypadkach. Dotyczy to w szczególności układów diodowo-tranzystorowych opisanych równaniem Sandberga-Willsona [4—5]. Powyższe fakty czynią z metody hybrydowej bardzo atrakcyjną reprezentację układów nieliniowych ale pewną prze-

szkodą w jej powszechnym stosowaniu jest stosunkowo skomplikowany sposób formułowania równań. Należy ponadto podkreślić, że opis hybrydowy może w pewnych przypadkach prowadzić do większej liczby równań niż np. metoda węzłowa, ale w omawianym obszarze zastosowań jest on niezbędny. Dla złożonych obwodów bardzo czasochłonnym ogniwem procesu formułowania równań hybrydowych jest odwracanie pewnych macierzy oznaczonych w dalszej części przez R i G . Macierze te mają struktury umożliwiające opracowanie szybkiej i łatwej do zalgorytmizowania procedury. W pracy zaproponowano efektywną metodę generacji macierzy R^{-1} i G^{-1} . Wykorzystano pewne koncepcje diakptyki dzięki czemu odpowiednie operacje wykonuje się na macierzach o mniejszych rozmiarach systematycznie powiększanych. Procesowi obliczeniowemu nadano interpretację fizyczną. Omówiono różne warianty oraz przytoczono przykłady liczbowe potwierdzające skuteczność algorytmu.

2. ZALEŻNOŚCI PODSTAWOWE

Rozpatrzmy obwód zawierający rezystory nieliniowe zarówno uzależnione napięciowo jak i prądowo, rezystory liniowe o dodatnich rezystancjach oraz niezależne źródła napięcia i prądu. Wyodrębnimy z obwodu wszystkie rezystory nieliniowe w sposób pokazany na rys. 1, gdzie rezystory o numerach od $1 - m_1$ są uzależnione napięciowo zaś rezystory o numerach od $(m_1 + 1)$ do $(m_1 - m_2)$ są uzależnione prądowo.

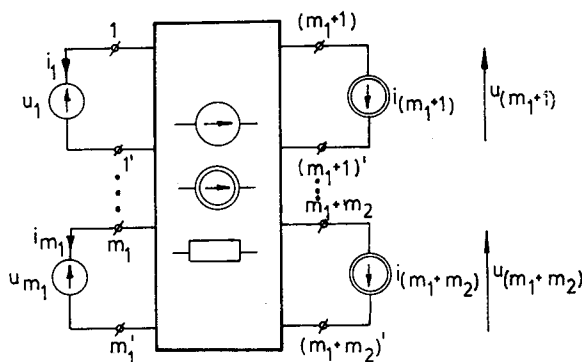


Rys. 1. Obwód z wyodrębnionymi rezystorami nieliniowymi

Następnie zastąpimy rezystory uzależnione napięciowo (prądowo) źródłami napięcia (prądu) (rys. 2). Otrzymujemy w ten sposób $(m_1 + m_2)$ — wrotnik zbudowany z rezystorów liniowych o dodatnich rezystancjach i źródeł niezależnych obciążony niezależnymi źródłami napięcia i prądu. Przyjmijmy, że w rozpatrywanym obwodzie, który oznaczmy przez $\mathcal{N}$, spełnione są następujące założenia:

- 1) źródło napięcia nie tworzą żadnej pętli,
- 2) źródło prądu nie tworzą żadnego rozcięcia.

W takim przypadku istnieje [1] następujący opis hybrydowy:



Rys. 2. Obwód z rys. 1 po wprowadzeniu źródeł napięcia i prądu

$$\begin{bmatrix} i_A \\ u_B \end{bmatrix} = -H \begin{bmatrix} u_A \\ i_B \end{bmatrix} + s, \quad (1)$$

gdzie

$$i_A = \begin{bmatrix} i_1 \\ \vdots \\ i_{m_1} \end{bmatrix}, \quad u_A = \begin{bmatrix} u_1 \\ \vdots \\ u_{m_1} \end{bmatrix}, \quad i_B = \begin{bmatrix} i_{m_1+1} \\ \vdots \\ i_{m_1+m_2} \end{bmatrix}, \quad u_B = \begin{bmatrix} u_{m_1+1} \\ \vdots \\ u_{m_1+m_2} \end{bmatrix},$$

zaś indeksy A i B oznaczają odpowiednio wrota napięciowe i prądowe, $s = [s_1 \dots s_{m_1+m_2}]^T$ jest wektorem zależnym od wymuszeń występujących wewnątrz (m_1+m_2) — wrotnika, $H = [h_{ij}]_{(m_1+m_2) \times (m_1+m_2)}$ jest macierzą hybrydową (m_1+m_2) — wrotnika. Uwzględniając, że dla rezystorów nieliniowych zachodzi:

$$i_A = f_A(u_A) = \begin{bmatrix} f_1(u_1) \\ \vdots \\ f_{m_1}(u_{m_1}) \end{bmatrix}, \quad u_B = f_B(i_B) = \begin{bmatrix} f_{m_1+1}(i_{m_1+1}) \\ \vdots \\ f_{m_1+m_2}(i_{m_1+m_2}) \end{bmatrix},$$

otrzymujemy na podstawie (1) następujący opis hybrydowy obwodu:

$$\begin{bmatrix} f_A(u_A) \\ f_B(i_B) \end{bmatrix} + H \begin{bmatrix} u_A \\ i_B \end{bmatrix} = s. \quad (2)$$

W celu sformułowania reprezentacji hybrydowej dzielimy obwód z rys. 2 na drzewo i dopełnienie zaliczając do drzewa wszystkie źródła napięcia a do dopełnienia wszystkie źródła prądu. Tworzymy następujące zbiory:

R — zbiór rezystorów liniowych

E — zbiór źródeł napięcia występujących wewnątrz (m_1+m_2) — wrotnika

J — zbiór źródeł prądu występujących wewnątrz (m_1+m_2) — wrotnika

A — zbiór wrotowych źródeł napięcia

B — zbiór wrotowych źródeł prądu.

Ponadto przez T i L oznaczymy odpowiednio zbiór elementów należących do drzewa i dopełnienia. Formułujemy równania fundamentalnych pętli i rozcięć, gdzie

$$Di = 0 \quad Bu = 0$$

gdzie

$$D = \left[\begin{array}{ccc|c} F_{11} & F_{12} & F_{13} & 100 \\ F_{21} & F_{22} & F_{23} & 010 \\ F_{31} & F_{32} & F_{33} & 001 \end{array} \right], \quad B = \left[\begin{array}{ccc|c} 100 & -F_{11}^T & -F_{21}^T & -F_{31}^T \\ 010 & -F_{12}^T & -F_{22}^T & -F_{32}^T \\ 001 & -F_{13}^T & -F_{23}^T & -F_{33}^T \end{array} \right]$$

oraz

$$i = [i_{JL}^T \ i_{BL}^T \ i_{RL}^T \ | \ i_{ET}^T \ i_{AT}^T \ i_{RT}^T]^T, \quad u = [u_{JL}^T \ u_{BL}^T \ u_{RL}^T \ | \ u_{ET}^T \ u_{AT}^T \ u_{RT}^T]^T,$$

a następnie równania elementowe

$$u_{RL} = R_L i_{RL} \quad u_{RT} = R_T i_{RT},$$

gdzie $R_L = \text{diag}(R_{L_1}, \dots, R_{L_m})$, $R_T = \text{diag}(R_{T_1}, \dots, R_{T_r})$ przy czym przez m oznaczono liczbę rezystorów liniowych zaliczanych do dopełnienia a przez r do drzewa.

Na podstawie powyższych zależności otrzymujemy po elementarnych przekształceniach równanie (1), gdzie

$$H = \left[\begin{array}{c|c} F_{23} R^{-1} F_{23}^T & F_{22} - F_{23} R^{-1} F_{33}^T R_T F_{32} \\ \hline -F_{22}^T + F_{32}^T G^{-1} F_{33} R_L^{-1} F_{23}^T & F_{32}^T G^{-1} F_{32} \end{array} \right] \quad (3)$$

oraz

$$S = \left[\begin{array}{c|c} -F_{23} R^{-1} F_{13}^T & -F_{21} + F_{23} R^{-1} F_{33}^T R_T F_{31} \\ \hline F_{12}^T - F_{32}^T G^{-1} F_{33} R_L^{-1} F_{13}^T & -F_{23}^T G^{-1} F_{31} \end{array} \right] \begin{bmatrix} u_{ET} \\ i_{JL} \end{bmatrix} \quad (4)$$

przy czym

$$R = R_L + F_{33}^T R_T F_{33} \quad (5)$$

$$G = R_T^{-1} + F_{33} R_L^{-1} F_{33}^T = G_T + F_{33} G_L F_{33}^T \quad (6)$$

gdzie $G_T = R_T^{-1}$, $G_L = R_L^{-1}$.

Wzory (3) i (4) umożliwiają sformułowanie równania (1). Dla złożonych obwodów najbardziej czasochłonnym ogniwem tego procesu jest wyznaczenie macierzy R^{-1} i G^{-1} występujących zarówno we wzorze (3) jak i (4).

W celu nadania interpretacji macierzy R i G utworzymy obwód $\mathcal{N}_R$ zawierając w obwodzie $\mathcal{N}$ wszystkie źródła napięcia i usuwając z tego obwodu wszystkie źródła prądu. Łatwo zauważyć, że macierz R jest macierzą impedancyjną fundamentalnych pętli obwodu $\mathcal{N}_R$. Inwersję tej macierzy R^{-1} nazwiemy macierzą admitancyjną fundamentalnych pętli (w skrócie MAFP). Podobnie G oznacza macierz admitancyjną fundamentalnych rozcięć obwodu $\mathcal{N}_R$. Jej inwersję G^{-1} nazwiemy macierzą impedancyjną fundamentalnych rozcięć (w skrócie MIFR).

W niniejszej pracy przedstawiono szybki i efektywny algorytm generacji macierzy R^{-1} i G^{-1} . Ze względu na przejrzystość tego algorytmu przyjęto, że gałęzie zbioru RT zostały ponumerowane rozpoczynając od tych, które nie tworzą w obwodzie $\mathcal{N}$ fundamentalnych rozcięć wyłącznie ze źródłami prądu. Elementy tego zbioru

wyznaczają fundamentalne rozcięcia, którym nadajemy tę samą numerację. Pozostałe elementy zbioru RT utworzą zbiór Ψ . Dla gałęzi zbioru RL stosujemy bieżącą numerację rozpatrując kolejne fundamentalne rozcięcia. Gałęzie ze zbioru RL tworzące w obwodzie $\mathcal{N}$ fundamentalne pętle tylko ze źródłami napięcia są numerowane jako ostatnie i zaliczone do zbioru ε .

Wykażemy, że macierz R jest nieosobliwa. W tym celu rozpatrzmy macierz $F_{33}^T R_T F_{33}$, która jest dodatnio półokreślona co można uzasadnić następująco. Z definicji iloczynu skalarnego wynika, że dla dowolnych macierzy A o wymiarach $r \times r$ oraz N o wymiarach $r \times m$ zachodzi

$$\langle N^T A N x, x \rangle = \langle A N x, N x \rangle$$

Stąd otrzymujemy $\langle F_{33}^T R_T F_{33} x, x \rangle = \langle R_T F_{33} x, F_{33} x \rangle = \langle R_T w, w \rangle$, gdzie $w = F_{33} x$. Ponieważ R_T jest macierzą diagonalną o dodatnich elementach znajdujących się na głównej przekątnej, więc dla każdego $w \in R^r$ zachodzi $\langle R_T w, w \rangle \geq 0$ co dowodzi dodatniej półokreśloności macierzy $F_{33}^T R_T F_{33}$. A zatem macierz $\langle F_{33}^T R_T F_{33} x, x \rangle = \langle R_T F_{33} x, F_{33} x \rangle \geq 0$ dla każdego wektora x o odpowiednim wymiarze. A zatem macierz $R = R_L + F_{33}^T R_T F_{33}$ jest dodatnio określona, czyli $\det R \neq 0$. Analogicznie dowodzi się nieosobliwości macierzy G .

3. METODA GENERACJI MACIERZY R^{-1} ORAZ G^{-1}

Wykażemy, że struktury macierzy R oraz G umożliwiają opracowanie efektywnej i łatwej do zalgorytmizowania metody generacji macierzy R^{-1} i G^{-1} co stanowi zasadnicze ogniwo procedury tworzenia opisu hybrydowego. Na wstępie ustalimy pewne związki pomiędzy R^{-1} oraz G^{-1} co umożliwi zastosowanie w procesie obliczeniowym różnych wariantów metody.

Skorzystamy ze wzoru macierzowego Woodbury'ego [2], [7—8]:

$$(A + P S^{-1} V)^{-1} = A^{-1} - A^{-1} P (S + V A^{-1} P)^{-1} V A^{-1}. \quad (7)$$

Stosując powyższą formułę do macierzy R (wzór (5)) otrzymujemy

$$R^{-1} = G_L - G_L F_{33}^T G^{-1} F_{33} G_L. \quad (8)$$

Podobnie, w przypadku macierzy G znajdujemy

$$G^{-1} = R_T - R_T F_{33} R^{-1} F_{33}^T R_T. \quad (9)$$

Zależności (8) i (9) pozwalają obliczać inwersję macierzy $R(G)$ gdy znana jest inwersja macierzy $G(R)$.

Oznaczamy kolejne wiersze macierzy F_{33} przez $\xi_1 = [\xi_{11} \ \xi_{12} \ \dots \ \xi_{1m}], \dots, \xi_r = [\xi_{r1} \ \xi_{r2} \ \dots \ \xi_{rm}]$. Wówczas macierz R można przedstawić w postaci:

$$R = R_L + \xi_1^T R_T \xi_1 + \dots + \xi_r^T R_T \xi_r \quad (10)$$

i do wyznaczenia macierzy R^{-1} zastosować r -krotnie wzór (7). W ten sposób znając

z k -tego kroku obliczeniowego macierz $R_k^{-1} = \left(R_L + \sum_{i=1}^k \xi_i^T R_{T_i} \xi_i \right)^{-1}$ obliczamy $R_{k+1}^{-1} = (R_k + \xi_{k+1}^T R_{T_{k+1}} \xi_{k+1})^{-1}$ podstawiając w (7): $A = R_k$, $P = \xi_{k+1}^T$, $S = R_{T_{k+1}}^{-1}$, $V = \xi_{k+1}$. Otrzymujemy

$$R_{k+1}^{-1} = R_k^{-1} - R_k^{-1} \xi_{k+1}^T (R_{T_{k+1}}^{-1} + \xi_{k+1} R_k^{-1} \xi_{k+1}^T)^{-1} \xi_{k+1} R_k^{-1}, \quad (11)$$

gdzie macierz $(R_{T_{k+1}}^{-1} + \xi_{k+1} R_k^{-1} \xi_{k+1}^T)$ ma wymiar 1×1 i jej odwrócenie wymaga wykonania jednej operacji. Należy ponadto zauważyć, że macierz ξ_{k+1} ma elementy równe 0, ± 1 co sprawia, że we wzorze (11) niektóre operacje mnożenia zastępuje się krótkimi operacjami dodawania. W ten sposób w kolejnych krokach obliczeniowych modyfikuje się macierze, których wymiar jest stały i wynosi $m \times m$.

Efektywność metody generacji macierzy R^{-1} można znacznie zwiększyć, jeżeli omawianą procedurę zastosuje się do pewnych zdekomponowanych obwodów sukcesywnie łączonych i powiększanych. Oznacza to wykonywanie wielu etapów na macierzach o mniejszych wymiarach, co ma istotny wpływ na skuteczność algorytmu.

W celu wyjaśnienia idei algorytmu rozpatrzmy obwód $\mathcal{N}_R$ i usuniemy z niego elementy zbioru Ψ . Podzielimy gałęzie dopełnienia obwodu $\mathcal{N}_R$ na dwa rozłączne zbiory $\mathcal{L}_1$ oraz $\mathcal{L}_2$. Drzewo tego obwodu podzielimy na trzy rozłączne zbiory $\mathcal{T}_1$, $\mathcal{T}_2$ i $\mathcal{T}_{12}$ ($\mathcal{T} = \mathcal{T}_1 \cup \mathcal{T}_2 \cup \mathcal{T}_{12}$) w ten sposób, aby gałęzie zbioru $\mathcal{T}_1$ tworzyły fundamentalne rozcięcia tylko z elementami zbioru $\mathcal{L}_1$, gałęzie zbioru $\mathcal{T}_2$ tylko z elementami zbioru $\mathcal{L}_2$ zaś gałęzie zbioru $\mathcal{T}_{12}$ zarówno z elementami zbioru $\mathcal{L}_1$ jak i $\mathcal{L}_2$. Zwierając chwilowo gałęzie zbioru $\mathcal{T}_{12}$ otrzymujemy obwód $\mathcal{M}_{12}$ złożony z dwóch podobowodów $\mathcal{M}_1$ oraz $\mathcal{M}_{12}$. Przy odpowiednim ponumerowaniu gałęzi macierz R_{12} obwodu $_{12}$ można na podstawie wzoru (10) przedstawić w postaci

$$R_{12} = \begin{bmatrix} R_{L_1} + \sum_{i=1}^r \beta_i^T R_{T_i} \beta_i & 0 \\ 0 & R_{L_2} + \sum_{i=r_1+1}^{r_1+r_2} \gamma_i^T R_{T_i} \gamma_i \end{bmatrix},$$

przy czym

$$R_L = \begin{bmatrix} R_{L_1} & 0 \\ 0 & R_{L_2} \end{bmatrix}, \quad \xi_i = [\beta_i \ 0] \text{ lub } \xi_i = [0 \ \gamma_i]$$

zaś r_1 (r_2) oznacza liczbę elementów zbioru $\mathcal{T}_1$ ($\mathcal{T}_2$).

W celu sformułowania macierzy R obwodu $\mathcal{N}_R$ należy uwzględnić następnie składniki wzoru (10) zawierające elementy zbioru $\mathcal{T}_{12}$. Przypisując tym elementom numery od $r_1 + r_2 + 1$ do r otrzymujemy

$$R = R_{12} + \sum_{i=r_1+r_2+1}^r \xi_i^T R_{T_i} \xi_i$$

Proces odwracania macierzy R można podzielić na trzy etapy. W pierwszym etapie wyznaczymy macierz

$$(R_{L_1} + \sum_{i=1}^{r_1} \beta_i^T R_{T_i} \beta_i)^{-1};$$

stosując r_1 razy iteracje wynikające ze wzoru Woodbury'ego (patrz wzór (11) z zamianą ξ na β). W drugim etapie oblicza się analogicznie macierz

$$(R_{L_1} + \sum_{i=r_1+1}^{r_1+r_2} \gamma_i^T R_{T_i} \gamma_i)^{-1}.$$

W ten sposób znajdujemy macierz

$$R_{12}^{-1} = \begin{bmatrix} (R_{L_1} + \sum_{i=1}^{r_1} \beta_i^T R_{T_i} \beta_i)^{-1} & 0 \\ 0 & (R_{L_2} + \sum_{i=r_1+1}^{r_1+r_2} \gamma_i^T R_{T_i} \gamma_i)^{-1} \end{bmatrix}$$

co umożliwi obliczenie inwersji macierzy R w wyniku wykonania iteracji (11) dla $k = r_1 + r_2, \dots, r$ oraz $R_{r_1+r_2} = R_{12}$.

Zgodnie z powyższym niektóre operacje wykonuje się na macierzach o mniejszych wymiarach $r_1 \times r_1$ oraz $r \times r_2$.

Powyższą koncepcję podziału elementów dopełnienia na dwa podzbiory $\mathcal{L}_1$ i $\mathcal{L}_2$ można uogólnić na większą liczbę podzbiorów i usystematyzować. W rezultacie zamiast podziału na dwa podzbiory dokonuje się podziału na wiele podzbiorów systematycznie powiększanych co wydatnie poprawia efektywność algorytmu.

Kluczową sprawą jest tu sposób podziału obwodu na podobowody i sterowanie procedurą. Szczegółowy opis procesu obliczeniowego zawiera podany niżej algorytm.

ALGORYTM GENERACJI MACIERZY R^{-1}

1. Z obwodu $\mathcal{N}_R$ usuwamy gałęzie zawierające elementy zbioru Ψ i zwieramy węzły, do których były dołączone oraz usuwamy gałęzie zawierające elementy zbioru ε . Tak przekształcony obwód oznaczamy przez $\mathcal{N}'_R$. Definiujemy zbiory: $\tilde{R}T = RT \setminus \Psi$, $\tilde{R}L = RL \setminus \varepsilon$ mające odpowiednio $\tilde{r}$ oraz $\tilde{m}$ elementów. Wyodrębniamy macierz F_{33} .

2. Rozpatrujemy obwód $\mathcal{N}'_R$ i tworzymy w nim zbiór A wszystkich fundamentalnych rozcięć. Wyznaczamy zbiór $\Phi \subset A$ rozłącznych fundamentalnych rozcięć oraz tworzymy zbiór $\Gamma = A \setminus \Phi$. Zwierając wszystkie gałęzie drzewa występujące w Γ tworzymy podobowody $\mathcal{M}_i$, dla których formułujemy odpowiednie podmacierze $R_{\mathcal{M}_i}^{-1}$. W tym celu najpierw znajdujemy macierze diagonalne odpowiadające gałęziom

dopełnienia a następnie stosując wzór Woodbury'ego uwzględniamy gałęzie drzewa występujące w podobwodach nie tworzących pętli jednogałęziowej.

3. Sprawdzamy, czy w zbiorze Γ istnieje taki element, którego gałęzie dopełnienia zawierają się wśród gałęzi dopełnienia któregoś z elementów zbioru Φ . Jeżeli odpowiedź jest pozytywna to powiększamy odpowiedni podobwód o gałąź drzewa stosując wzór Woodbury'ego usuwając jednocześnie ten element ze zbioru Γ . Czynność powtarzamy aż do uzyskania odpowiedzi negatywnej.

4. Sprawdzamy, czy w każdym podobwodzie elementy dopełnienia mają kolejne rosnące indeksy. Jeżeli warunek ten nie jest spełniony, to podobwód uzupełniamy o odpowiednie elementy (podobwody) zbioru RL i rozszerzamy jego MAFP powiększając jej główną przekątną o elementy typu $R_{L_j}^{-1}$ tak, aby indeksy elementów tej przekątnej tworzyły ciąg rosnący. Przyjmujemy takie uporządkowanie podobwodów, aby ich gałęzie dopełnienia rozpatrywane w każdym podobwodzie wg rosnącej numeracji utworzyły łącznie zbiór gałęzi o numerach $1, 2, 3, \dots, \tilde{m}$.

5. Tworzymy nowe podobwody złożone z dwóch kolejnych i wyznaczamy ich MAFP o postaci blokowo-diagonalnej. Następnie otrzymane podobwody powiększamy o takie gałęzie zbioru RT, które wyznaczają fundamentalne rozcięcia o dopełnieniu mieszczącym się w dopełnieniu rozpatrywanego podobwodu. MAFP powiększonego podobwodu znajduje się stosując wzór Woodbury'ego.

6. Krok 5 powtarzamy aż do wyznaczenia MAFP $(R')^{-1}$ obwodu $\mathcal{N}'_R$.

7. Formułujemy macierz R^{-1} (blokowo-diagonalną) obwodu $\mathcal{N}_R$ złożoną z bloków $(R')^{-1}$ oraz Δ^{-1} , gdzie Δ jest macierzą diagonalną utworzoną z elementów zbioru ε .

W przypadku macierzy G (wzór (6)) postępujemy podobnie. Oznaczamy kolejne kolumny macierzy F_{33} przez $\eta_1 = [\xi_{11} \xi_{21} \dots \xi_{r1}]^T, \dots, \eta_m = [\xi_{1m} \xi_{2m} \dots \xi_{rm}]^T$ i przedstawiamy macierz G w postaci

$$G = G_T + \eta_1 G_{L_1} \eta_1^T + \dots + \eta_m G_{L_m} \eta_m^T. \quad (12)$$

Zależność (12) umożliwia obliczenie macierzy G^{-1} w wyniku r -krotnego zastosowania wzoru Woodbury'ego (7). W ten sposób znając z k -tego kroku obliczeniowego macierz

$$G_k^{-1} = (G_T + \sum_{i=1}^k \eta_i G_{L_i} \eta_i^T)^{-1} \text{ otrzymujemy} \quad (13)$$

$$G_{k+1}^{-1} = G_k^{-1} - G_k^{-1} \eta_{k+1} (G_{L_{k+1}}^{-1} + \eta_{k+1}^T G_k^{-1} \eta_{k+1})^{-1} \eta_{k+1}^T G_k^{-1}.$$

W celu zwiększenia efektywności algorytmu stosujemy opisaną procedurę numeryczną, do pewnych obwodów zdekomponowanych w sposób opisany niżej.

ALGORYTM GENERACJI MACIERZY G^{-1}

1. Wykonujemy punkt 1) algorytmu generacji macierzy R^{-1} .

2. W obwodzie $\mathcal{N}'_R$ tworzymy zbiór Θ wszystkich fundamentalnych pętli. Wyznaczamy zbiór $\pi \subset \Theta$ rozłącznych fundamentalnych pętli oraz tworzymy zbiór $\Omega = \Theta \setminus \pi$. Usuwając wszystkie gałęzie dopełnienia występujące w Ω tworzymy

podobwody κ_i , dla których formułujemy odpowiednie MIFR $G_{\kappa_i}^{-1}$. W tym celu najpierw znajdujemy macierze diagonalne odpowiadające gałęziom drzewa a następnie stosując wzór Woodbury'ego uwzględniamy gałęzie dopełnienia.

3. Sprawdzamy czy w zbiorze Ω istnieje taki element, którego gałęzie drzewa znajdują się wśród gałęzi drzewa któregoś z elementów zbioru Θ . Jeżeli odpowiedź jest pozytywna, to powiększamy odpowiedni podobwód o gałąź dopełnienia stosując wzór Woodbury'ego jednocześnie usuwając ten element ze zbioru Ω . Czynność tę powtarzamy aż do uzyskania odpowiedzi negatywnej.

4. Sprawdzamy czy w każdym podobwodzie elementy drzewa mają kolejne rosnące indeksy. Jeżeli warunek ten nie jest spełniony, to podobwód uzupełniamy o odpowiednie elementy zbioru RT i rozszerzamy jego MIFR powiększając jej główną przekątną o elementy typu G_T^{-1} . Przyjmujemy takie uporządkowanie podobwodów, aby ich gałęzie drzewa rozpatrywane w każdym podobwodzie wg rosnącej numeracji utworzyły łącznie zbiór gałęzi o numerach $1, 2, \dots, \tilde{n}$.

5. Tworzymy nowe podobwody złożone z dwóch kolejnych i wyznaczamy ich MIFR o postaci blokowo-diagonalnej. Następnie otrzymane podobwody powiększamy o takie gałęzie zbioru RL , które wyznaczają fundamentalne pętle o drzewie mieszczącym się w drzewie rozpatrywanego podobwodu. Macierz powiększonego podobwodu znajduje się stosując wzór Woodbury'ego.

6. Krok 5 powtarzamy aż do wyznaczenia MIFR $(G')^{-1}$ obwodu $\mathcal{N}'_R$.

7. Formułujemy macierz G^{-1} obwodu $\mathcal{N}_R$ złożoną z bloków $(G')^{-1}$ oraz δ^{-1} , gdzie δ jest macierzą diagonalną utworzoną z konduktancji elementów zbioru Ψ . W celu zrealizowania powyższych algorytmów należy wygenerować macierz F_{33} będącą podmacierzą macierzy fundamentalnych rozcięć. Analizując macierz F_{33} wyznacza się odpowiednie struktury topologiczne. Dalej wykonuje się obliczenia macierzowe. Fakt, że wiele działań jest realizowanych na macierzach zero-jedynkowych pozwala zastąpić długie operacje mnożenia krótkimi operacjami dodawania.

4. DYSKUSJA I PRZYKŁADY LICZBOWE

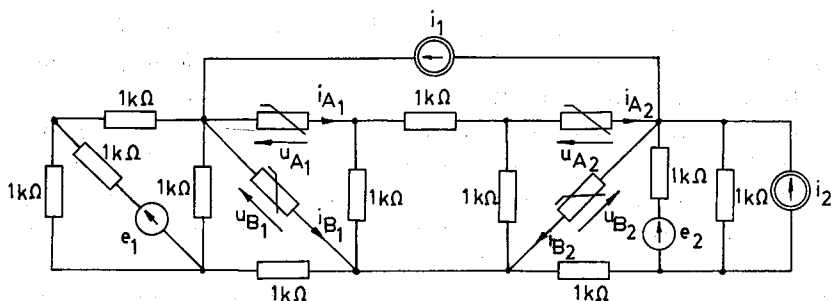
Rozpatrzmy obwód przedstawiony na rys. 3 zawierający dwa rezystory nieliniowe uzależnione napięciowo:

$$\begin{aligned} i_{A1} &= f_1(u_{A1}) \\ i_{A2} &= f_2(u_{A2}) \end{aligned}$$

oraz dwa rezystory nieliniowe uzależnione prądowo:

$$\begin{aligned} u_{B1} &= f_3(i_{B1}) \\ u_{B2} &= f_4(i_{B2}). \end{aligned}$$

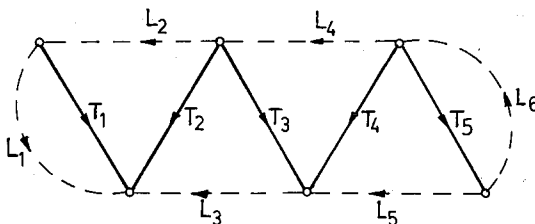
W celu sformułowania równania hybrydowego (2) należy wyznaczyć macierz H oraz wektor s korzystając ze wzorów (3) i (4). Występujące w tych wzorach macierze R^{-1} i G^{-1} obliczymy stosując omówiony wyżej algorytm.



Rys. 3. Układ zawierający rezystory nieliniowe

Na rys. 4 pokazano graf obwodu $\mathcal{N}'_R = \mathcal{N}_R$ utworzonego z gałęzi drzewa $RT = \{T_1, T_2, T_3, T_4, T_5\}$ oraz dopełnienia $RL = \{L_1, L_2, L_3, L_4, L_5, L_6\}$. Macierze F_{33} i F_{33}^T są następujące:

$$F_{33} = \begin{matrix} & L_1 & L_2 & L_3 & L_4 & L_5 & L_6 \\ \begin{matrix} T_1 \\ T_2 \\ T_3 \\ T_4 \\ T_5 \end{matrix} & \begin{bmatrix} 1 & -1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & -1 & -1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 & -1 & -1 \end{bmatrix} \end{matrix} \quad F_{33}^T = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ -1 & 1 & 0 & 0 & 0 \\ 0 & 1 & -1 & 0 & 0 \\ 0 & 0 & -1 & 1 & 0 \\ 0 & 0 & 0 & 1 & -1 \\ 0 & 0 & 0 & 0 & -1 \end{bmatrix}$$

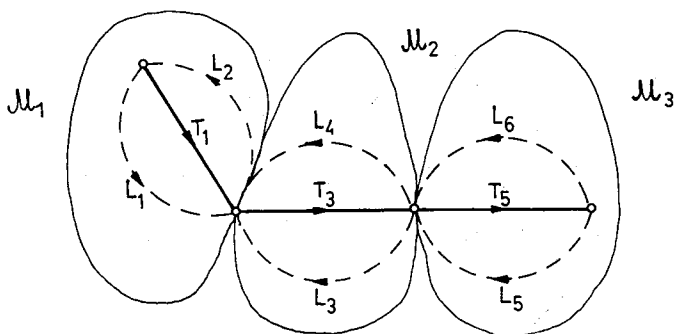


Rys. 4. Graf obwodu z rys. 3

ZASTOSOWANIE ALGORYTMU GENERACJI MACIERZY R^{-1}

Tworzymy zbiór Λ fundamentalnych rozcięć $\Lambda = \{T_1 L_1 L_2, T_2 L_2 L_3, T_3 L_3 L_4, T_4 L_4 L_5, T_5 L_5 L_6\}$ po czym wyodrębniamy zbiory $\Phi = \{T_1 L_1 L_2, T_3 L_3 L_4, T_5 L_5 L_6\}$ oraz $\Gamma = \Lambda / \Phi = \{T_2 L_2 L_3, T_4 L_4 L_5\}$. Budujemy podobwoły $\mathcal{M}_1, \mathcal{M}_2, \mathcal{M}_3$ (rys. 5) i wyznaczamy ich macierze $R_{\mathcal{M}_i}^{-1} (i = 1, 2, 3)$. Otrzymujemy

$$R_{\mathcal{M}_1}^{-1} = \left(\begin{bmatrix} R_{L_1} & \\ & R_{L_2} \end{bmatrix} + \begin{bmatrix} 1 \\ -1 \end{bmatrix} R_{T_1} [1 \ -1] \right)^{-1} = \begin{bmatrix} 2/3 & 1/3 \\ 1/3 & 2/3 \end{bmatrix}$$

Rys. 5. Podobwody $\mathcal{M}_1, \mathcal{M}_2, \mathcal{M}_3$

i podobnie $R_{\mathcal{M}_2}^{-1} = \begin{bmatrix} 2/3 & -1/3 \\ -1/3 & 2/3 \end{bmatrix}$, $R_{\mathcal{M}_3}^{-1} = \begin{bmatrix} 2/3 & -1/3 \\ -1/3 & 2/3 \end{bmatrix}$

Tworzymy nowy podobwód $\mathcal{M}_4$ (rys. 6) utworzony z podobwodów $\mathcal{M}_1$ i $\mathcal{M}_2$ oraz powiększony o gałąź T_2 po czym wyznaczamy macierz

$$R_{\mathcal{M}_4}^{-1} = \left(\begin{bmatrix} R_{\mathcal{M}_1} & \\ & R_{\mathcal{M}_2} \end{bmatrix} + \begin{bmatrix} 0 \\ 1 \\ 1 \\ 0 \end{bmatrix} R_{T_2} [0 \ 1 \ 1 \ 0] \right)^{-1} = \begin{bmatrix} 13/21 & 5/21 & -2/21 & 1/21 \\ 5/21 & 10/21 & -4/21 & 2/21 \\ -2/21 & -4/21 & 10/21 & -5/21 \\ 1/21 & 2/21 & -5/21 & 13/21 \end{bmatrix}.$$

Następnie rozpatrujemy podobwód złożony z $\mathcal{M}_4$ i $\mathcal{M}_3$ i powiększamy go o gałąź T_4 . Dochodzimy w ten sposób do obwodu $\mathcal{N}'_R = \mathcal{N}_R$, którego macierz R^{-1} wynosi:

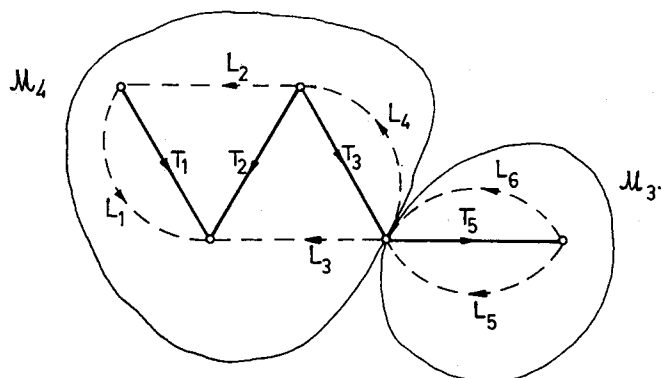
$$R^{-1} = \left(\begin{bmatrix} R_{\mathcal{M}_4} & \\ & R_{\mathcal{M}_3} \end{bmatrix} + [0 \ 0 \ 0 \ 1 \ 1 \ 0]^T R_{T_4} [0 \ 0 \ 0 \ 1 \ 1 \ 0] \right)^{-1} =$$

$$= \begin{bmatrix} 0,6181 & 0,2361 & -0,0905 & 0,0353 & -0,0138 & 0,0069 \\ 0,2361 & 0,4723 & -0,1807 & 0,0686 & -0,0277 & 0,0138 \\ -0,0905 & -0,1807 & 0,4516 & -0,1743 & 0,0694 & -0,0346 \\ 0,0353 & 0,0686 & -0,1743 & 0,4526 & -0,1791 & 0,0899 \\ -0,0138 & -0,0277 & 0,0694 & -0,1791 & 0,4722 & -0,2361 \\ 0,0069 & 0,0138 & -0,0346 & 0,0899 & -0,2361 & 0,6180 \end{bmatrix}.$$

Stosując opisaną wyżej procedurę należy wykonać 90 operacji mnożenia (bez uwzględnienia symetrii). Gdyby obliczać macierz R^{-1} metodą eliminacji Gaussa należałoby wykonać około 216 operacji. W przypadku uwzględnienia symetrii uzyskuje się 2,7-krotne zmniejszenie liczby operacji mnożenia.

ZASTOSOWANIE ALGORYTMU GENERACJI MACIERZY G^{-1}

Tworzymy zbiór Θ fundamentalnych pętli $\Theta = \{L_1 T_1, L_2 T_1 T_2, L_3 T_2 T_3, L_4 T_3 T_4, L_5 T_4 T_5, L_6 T_5\}$ po czym wyodrębniamy zbiory $\Pi = \{L_1 T_1, L_3 T_2 T_3, L_5 T_4 T_5\}$ oraz $\Omega = \Theta \setminus \Pi = \{L_2 T_1 T_2, L_4 T_3 T_4, L_6 T_5\}$. Budujemy podobwody $\kappa_1, \kappa_2, \kappa_3$ (rys. 7) i wyznaczamy ich macierze $G_{\kappa_i}^{-1}$ ($i = 1, 2, 3$). Otrzymujemy:

Rys. 6. Podobwody $\mathcal{M}_3, \mathcal{M}_4$

$$G_{\kappa_1}^{-1} = (G_{T_1} + 1 G_{L_1} 1)^{-1} = 1/2,$$

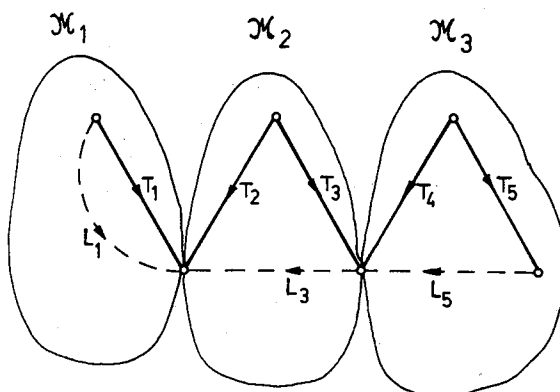
$$G_{\kappa_2}^{-1} = \left(\begin{bmatrix} G_{T_2} & \\ & G_{T_3} \end{bmatrix} + \begin{bmatrix} 1 \\ -1 \end{bmatrix} G_{L_3} [1 \ -1] \right)^{-1} = \begin{bmatrix} 2/3 & 1/3 \\ 1/3 & 2/3 \end{bmatrix}$$

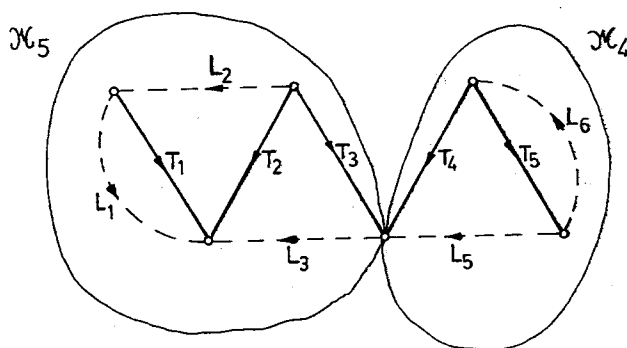
$$G_{\kappa_3}^{-1} = \left(\begin{bmatrix} G_{T_4} & \\ & G_{T_5} \end{bmatrix} + \begin{bmatrix} 1 \\ -1 \end{bmatrix} G_{L_5} [1 \ -1] \right)^{-1} = \begin{bmatrix} 2/3 & 1/3 \\ 1/3 & 2/3 \end{bmatrix}.$$

W zbiorze Ω znajdujemy element $L_6 T_5$, którego gałąź drzewa T_5 znajduje się wśród gałęzi drzewa elementu $L_5 T_4 T_5$ zbioru Θ . Zgodnie z punktem 3 algorytmu powiększamy podobwód κ_3 o gałąź L_6 otrzymując podobwód κ_4 , którego macierz $G_{\kappa_4}^{-1}$ wynosi

$$G_{\kappa_3}^{-1} = \left(G_{\kappa_3} + \begin{bmatrix} 0 \\ -1 \end{bmatrix} G_{L_6} [0 \ -1] \right)^{-1} = \begin{bmatrix} 3/5 & 1/5 \\ 1/5 & 2/5 \end{bmatrix}$$

Tworzymy nowy podobwód κ_5 złożony z κ_1 i κ_2 i powiększamy go o gałąź L_2 (rys. 8). Wyznaczamy macierz $G_{\kappa_6}^{-1}$ tego podobwodu:

Rys. 7. Podobwody $\kappa_1, \kappa_2, \kappa_3$

Rys. 8. Podobwody κ_4 , κ_5

$$G_{\kappa_5}^{-1} = \left(\begin{bmatrix} G_{\kappa_1} & \\ & G_{\kappa_2} \end{bmatrix} + \begin{bmatrix} -1 \\ 1 \\ 0 \end{bmatrix} G_{L_2} [-1 \ 1 \ 0] \right)^{-1} = \begin{bmatrix} 5/13 & 2/13 & 1/13 \\ 2/13 & 6/13 & 3/13 \\ 1/13 & 3/13 & 8/13 \end{bmatrix}$$

Następnie rozpatrujemy podobwód złożony z κ_5 i κ_4 i powiększamy go o gałąź L_4 . Dochodzimy w ten sposób do obwodu $\mathcal{N}'_R = \mathcal{N}_R$, którego macierz G^{-1} wynosi:

$$G^{-1} = \left(\begin{bmatrix} G_{\kappa_5} & \\ & G_{\kappa_4} \end{bmatrix} + [0 \ 0 \ -1 \ 1 \ 0]^T G_{L_4} [0 \ 0 \ -1 \ 1 \ 0] \right)^{-1} =$$

$$= \begin{bmatrix} 0,3820 & 0,1458 & 0,0556 & 0,0208 & 0,0069 \\ 0,1458 & 0,4375 & 0,1666 & 0,0625 & 0,0208 \\ 0,0556 & 0,1666 & 0,4444 & 0,1666 & 0,0555 \\ 0,0208 & 0,0625 & 0,1666 & 0,4375 & 0,1459 \\ 0,0069 & 0,0208 & 0,0555 & 0,1459 & 0,3830 \end{bmatrix}.$$

Stosując opisaną procedurę należy wykonać 71 operacji mnożenia bez uwzględnienia symetrii, podczas gdy użycie metody eliminacji Gaussa wymagałoby wykonania około 125 operacji mnożenia. Przyjęto tu, że dane są konduktancje wszystkich rezystorów liniowych. Przy uwzględnieniu symetrii zmniejszenie liczby operacji mnożenia jest około 2-krotne.

Podsumowując stwierdzamy, że obliczanie macierzy R^{-1} oraz G^{-1} przy wykorzystaniu zaproponowanych w pracy algorytmów wymaga wykonania łącznie 161 operacji mnożenia, bez uwzględnienia symetrii, podczas gdy zastosowanie metody eliminacji Gaussa prowadzi do około 341 operacji mnożenia. Uzyskano więc ponad dwukrotne zmniejszenie liczby długich operacji. Jeszcze bardziej korzystny efekt uzyskuje się uwzględniając symetrię.

Na podstawie wzorów (8) i (9) stwierdzamy, że istnieją dwie inne możliwości obliczenia obu macierzy R^{-1} i G^{-1} , a mianowicie:

- obliczenia R^{-1} za pomocą podanego w pracy algorytmu oraz wyznaczenie G^{-1} ze wzoru (9);

- b) obliczenie G^{-1} za pomocą podanego w pracy algorytmu oraz wyznaczenie R^{-1} ze wzoru (8).

W rozpatrywanym przypadku wersja a) prowadzi do łącznej liczby operacji mnożenia wynoszącej 140 (bez uwzględnienia symetrii), podczas gdy wersja b) wymaga wykonania 143 operacji mnożenia. Tak więc obie wersje a) i b) są w przybliżeniu jednakowo efektywne i wykazują przewagę nad zastosowaniem odpowiednich algorytmów do oddzielnego obliczenia R^{-1} i G^{-1} . Możliwość określenia a priori liczby operacji mnożenia dla każdego spośród trzech dyskutowanych wyżej wariantów pozwala każdorazowo dobierać optymalny schemat obliczeniowy.

Na podstawie wyznaczonych macierzy R^{-1} i G^{-1} oraz korzystając ze wzorów (3) i (4) sformułowano dla badanego obwodu następujące równanie hybrydowe:

$$\begin{bmatrix} f_1(u_{A1}) \\ f_2(u_{A2}) \\ f_3(i_{B1}) \\ f_4(i_{B2}) \end{bmatrix} + \begin{bmatrix} -0,4514 & -0,0694 & -0,7222 & 0,1041 \\ -0,0694 & -0,4722 & -0,1111 & 0,7083 \\ 0,7222 & 0,1111 & -0,4444 & -0,1666 \\ -0,1041 & -0,7083 & -0,1666 & -0,4375 \end{bmatrix} \begin{bmatrix} u_{A1} \\ u_{A2} \\ i_{B1} \\ i_{B2} \end{bmatrix} = \begin{bmatrix} 0,0902 & -0,0347 & 0,8264 & -0,0346 \\ 0,0139 & -0,2361 & 0,8166 & -0,2361 \\ 0,0554 & 0,0555 & 0,2778 & 0,0555 \\ 0,0209 & 0,1457 & -0,2709 & 0,1457 \end{bmatrix} \begin{bmatrix} e_1 \\ e_2 \\ i_1 \\ i_2 \end{bmatrix}.$$

Niżej wykazemy, że wyniki prezentowane w niniejszej pracy mogą być bezpośrednio wykorzystane do formułowania równania Sandberga-Willsona opisującego układy diodowo-tranzystorowe. W tym celu rozpatrzmy układ DC utworzony z tranzystorów, diod, rezystorów liniowych o dodatnich rezystancjach oraz niezależnych źródeł napięcia i prądu. Tranzystory i diody wyodrębnimy na zewnątrz obwodu oraz zastąpimy modelami Ebersa-Molla. Otrzymamy w ten sposób liniowy n -wrotnik zawierający rezystory i źródła niezależne obciążony diodami lub kombinacjami dioda-źródło sterowane. Elementy obciążające n -wrotnik można opisać równaniem

$$i = Tf(u), \quad (14)$$

gdzie T jest macierzą blokowo-diagonalną związaną z opisem Ebersa-Molla tranzystorów i diod [5–6], zaś $f(u) = [f_1(u_1) \dots f_n(u_n)]^T: R^n \rightarrow R^n$ przy czym $f_i(u_i) = K_i(e^{\lambda_i u_i} - 1)$ oznacza charakterystykę i -tej diody występującej samodzielnie lub modelu tranzystora. Wynika stąd wniosek, że w celu sformułowania równania układu diodowo-tranzystorowego należy opisać n -wrotnik równaniem

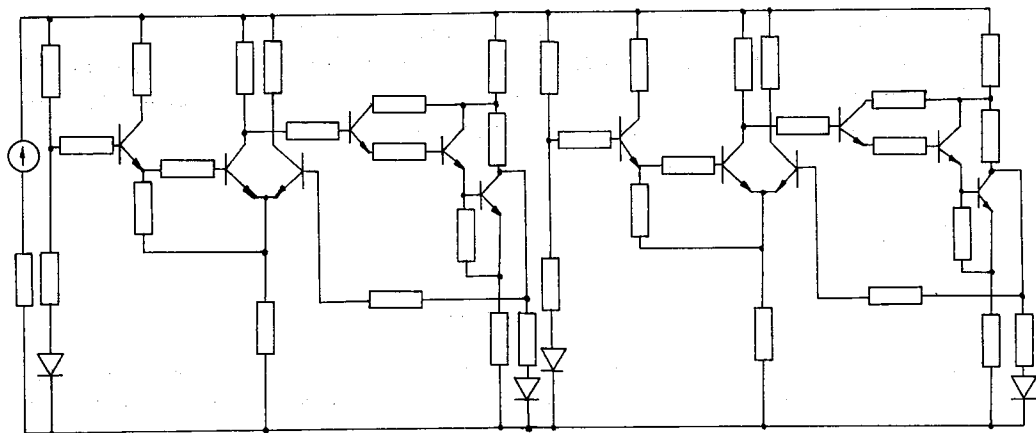
$$i = -Cu + b, \quad (15)$$

gdzie C jest szczególnym przypadkiem macierzy hybrydowej i nosi nazwę zwarciowej macierzy admitancyjnej. Na podstawie (3) i (4) znajdujemy

$$C = F_{23}R^{-1}F_{23}^T \quad (16)$$

$$b = -F_{23}R^{-1}F_{13}^T u_{ET} + (-F_{21} + F_{23}R^{-1}F_{23}^T R_T F_{31}) i_{JL}, \quad (17)$$

gdzie R jest macierzą określoną wzorem (5). Uwzględniając (14) i (15) dochodzimy do znanego równania Sandberga-Willsona



Rys. 9. Układ diodowo-tranzystorowy

$$Tf(u) + Cu = b \quad (18)$$

odgrywającego bardzo ważną rolę zarówno w analizie jakościowej jak i ilościowej [4 – 6]. Sformułowanie tego równania wymaga wyznaczenia macierzy C oraz wektora b zawierających inwersję macierzy R . W złożonych układach diodowo-tranzystorowych obliczenie macierzy R^{-1} jest bardzo czasochłonnym ogniwem procesu formułowania równania (18). Omówiony w rozdziale 3 niniejszej pracy algorytm generacji macierzy R^{-1} pozwala znacznie usprawnić procedurę tworzenia równania (18).

Zastosowanie tego algorytmu do wyznaczenia macierzy R^{-1} układu przedstawionego na rys. 9 wymagało wykonania, bez uwzględnienia symetrii, około 6-krotnie mniej operacji mnożenia niż w przypadku metody eliminacji Gaussa. Przy uwzględnieniu symetrii zmniejszenie jest około 7-krotne.

WNIOSKI

Przedstawiony w pracy algorytm generacji macierzy R^{-1} i G^{-1} stanowiący główne ogniwo formułowania równań hybrydowych obwodów nieliniowych jest efektywny, przejrzysty i wymaga wykonania stosunkowo małej liczby długich operacji. Fakt, że wiele działań jest realizowanych na macierzach strukturalnych pozwala zastąpić długie operacje mnożenia krótkimi operacjami dodawania. Algorytm wyznaczenia macierzy R^{-1} oraz G^{-1} wykorzystuje w nieco odmienny sposób rzadkości macierzy. Proces obliczania R^{-1} oraz G^{-1} odbywa się bez potrzeby jawnego formułowania macierzy R i G . Liczba wykonanych operacji mnożenia jest znacznie mniejsza niż w metodzie eliminacji Gaussa. Opracowanie kilku wariantów metody umożliwia dobór procedury optymalnej. Naturalną cechą algorytmu jest łatwość modyfikowania wyników w przypadku zmian pewnych parametrów. Metoda może być bezpośrednio użyta do formułowania równań Sandberga-Willsona opisujących DC obwody diodowo-tran-

zystorowe. W rezultacie uzyskuje się znaczne usprawnienie procesu tworzenia opisu matematycznego szerokiej klasy obwodów nieliniowych. Potencjalną wadą metody jest możliwość kumulacji błędów zaokrągleń ale eksperymenty numeryczne wskazują, że dokładność otrzymanych wyników nie budzi zastrzeżeń.

Pracę wykonano w ramach CPBP 02.14

BIBLIOGRAFIA

1. L. O. Chua and P. M. Lin: *Komputerowa analiza układów elektronicznych*. Warszawa: WNT, 1981
2. L. O. Chua and R. L. P. Ying: *Finding all solutions of piecewise-linear circuits*. Int. J. Cir. Theor. Appl. 1982, 10, p. 201—229
3. A. S. Householder: *The theory of matrices in numerical analysis*. Blaisdell Publishing Company, New York, 1965
4. I. W. Sandberg and A. N. Willson: *Some theorems on properties of DC equations of nonlinear networks*. Bell Syst. Tech. J., 1969, 48 p. 1—34
5. M. Tadeusiewicz: *On the solvability and comutation of DC transistor networks*. Int. J. Cir. Theor. Appl., 1981, 9, p. 251—267
6. M. Tadeusiewicz: *An algorithm for finding bounds on the location of the all solutions of piecewise-linear transistor circuits*. Proc. ECCTD'87, Paris, 1987, September 1—4, p. 735—740
7. M. Tadeusiewicz & al.: *An algorithm for generation of nodal impedance matrix for networks containing mutually coupled inductors*. Int. J. El. Pow. En. Syst., 1988, 10 p. 36—40
8. M. Tadeusiewicz, J. Bek: *Komputerowe formułowanie równań układów diodowo-tranzystorowych*. XIKKTOiUE Teksty referatów, Łódź-Rytro 1988 p. 162—167

M. TADEUSIEWICZ, J. BEK

A METHOD FOR FORMULATION OF HYBRID EQUATIONS FOR NONLINEAR RESISTIVE CIRCUITS

Summary

A method for formulation of hybrid equations describing DC nonlinear circuits is developed in the paper. It has been shown that finding some matrices inversions is the main part of the procedure. The matrices have stuctures which enable us work out an effective method expressed in terms of diacoptic and Woodbury's formula. In this way many steps are performed on matrices having small sizes associated with some subnetworks succesively enlarged. It is shown that the algorithm can be strightforwardly by used for formulation of the Sandberg-Willson equation describing DC diode-transistor circuits. Numerical examples are demonstrated confirming effectiveness of the algorithm.

Detektory fazoczułe w pomiarach składowych immitancji dwójników

ANATOL JAWOREK

*Instytut Maszyn Przepływowych,
Polska Akademia Nauk w Gdańsku*

Otrzymano 1990.03.10

Autoryzowano do druku 1990.05.15

Przedstawiono zasadę pomiaru składowej rzeczywistej i urojonej immitancji dwójników biernych, będących w szczególności parametrycznymi przetwornikami pomiarowymi, za pomocą trzech metod detekcji fazoczułej. Pomiar polega na całkowaniu odpowiedzi dwójnika w stanie ustalonym na pobudzenie harmoniczne w ściśle określonym przedziale kąta fazowego. Przeprowadzono analizę dokładności pomiaru składowej rzeczywistej i urojonej dwójników biernych. Wyznaczono błędy pomiarowe spowodowane przez niedokładne określenie kątów fazowych przedziałów całkowania oraz przez składowe harmoniczne zawarte w sygnale mierzonym.

1. WSTĘP

Detektor fazoczuły jest układem, na którego wyjściu otrzymuje się sygnał stałonapięciowy proporcjonalny do amplitudy detekowanego sygnału zmiennonapięciowego i jego przesunięcia fazowego względem sygnału odniesienia. W pracy przedstawiono metodę pomiaru składowych rzeczywistej i urojonej impedancji lub admitancji dwuelementowych, liniowych dwójników biernych za pomocą detektorów fazoczułych. Potrzeba rozróżnienia składowych immitancji występuje w przypadku diagnostyki dwójników biernych lub w pomiarach sygnałów generowanych przez parametryczne przetworniki pomiarowe. Rzeczywisty przetwornik pomiarowy posiada składową rzeczywistą i urojoną immitancji, przy czym często tylko jedna z nich w sposób jednoznaczny zależy od mierzonej wielkości fizycznej i niesie użyteczną informację o wartości wielkości mierzonej. Występuje wówczas potrzeba wyeliminowania składowej pasożytniczej. Ma to miejsce np. w przetwornikach indukcyjnych, w których mierzonym zmianom indukcyjności towarzyszyć mogą zmiany rezystancji spowodowane stratami w rdzeniu i uzwojeniu albo w przetwornikach pojemnościowych o dużej upływności [1] lub rezystancyjnych o niepomijalnej susceptancji [2],

gdy z powodu polaryzacji elektrod zanurzonych w cieczy konieczne jest zasilanie elektrod prądem przemiennym.

Metody mostkowe, polegające na kompensowaniu poszczególnych składowych immitancji, ze względu na długotrwały proces równoważenia, lub metody techniczne polegające na wyznaczeniu modułu i kąta przesunięcia fazowego, z powodu konieczności przeliczania wyników pomiarów, nie umożliwiają pomiarów dynamicznych zadanej wielkości fizycznej w czasie rzeczywistym [3—5]. Detektory fazoczułe umożliwiają pomiar zmian wartości wielkości mierzonej w czasie nie przekraczającym kilku okresów sygnału zasilającego przetwornik pomiarowy, co pozwala na dostatecznie wierne odtworzenie przebiegu czasowego. W niektórych specjalnych rozwiązaniach czas pomiaru może być skrócony do jednego lub nawet do pół okresu sygnału pobudzającego.

Diagnostyka poszczególnych elementów zamontowanych w gotowych układach elektronicznych i zbocznikowanych przez inne elementy występujące w układzie nie byłaby możliwa bez zastosowania detekcji fazoczułej [6].

Wyróżnić można trzy typy integracyjnych detektorów fazoczułych do bezpośredniego pomiaru składowych immitancji, w których pomiar polega na całkowaniu sygnału odpowiedzi dwójnika na pobudzenie harmoniczne w ściśle określonym dla danej składowej przedziale kąta fazowego:

1. Detektor fazoczuły z całkowaniem sygnału odpowiedzi w stałym przedziale, wyznaczonym przez fazę sygnału odniesienia i wynoszącym połowę okresu sygnału odniesienia [7—9].

2. Detektor fazoczuły z całkowaniem sygnału odpowiedzi w zmiennym przedziale, zależnym od fazy sygnału odniesienia (pobudzenia dwójnika) i fazy odpowiadającej wartości maksymalnej sygnału odpowiedzi badanego dwójnika [10—12].

3. Detektor fazoczuły z próbkowaniem sygnału odpowiedzi (sample-and-hold), w którym momenty próbkowania wyznaczane są przez sygnał odniesienia [13—15].

Pierwotnie detekcja fazoczuła została zastosowana przez Sigdella [7] do wyznaczania pojemności przetworników pomiarowych. Detektor fazoczuły całkował mierzony sygnał jednopółkowo w przedziale od 0 do π . Orsini [13] rozwinął tę metodę na przypadek pomiaru obu składowych immitancji dwójnika równoległego złożonego z pojemności i konduktancji lub dwójnika szeregowego złożonego z indukcyjności i rezystancji, stosując detektor fazoczuły z próbkowaniem. Stosowane są również bardziej złożone rozwiązania. Hoja [19] zastosował całkowanie dwukrotne do wyznaczania parametrów dwójników trzelementowych. Maeda [20] przez zastosowanie całkowania wielokrotnego wyeliminował wpływ parametrów linii łączącej badany dwójnik z układem pomiarowym.

Pomimo licznych zastosowań detekcji fazoczułej w aparaturze naukowej i przyrządach pomiarowych [16—22], brak jest opracowań systematyzujących zagadnienie pomiaru immitancji za pomocą detektorów fazoczułych, a tylko nieliczne prace poświęcone są zagadnieniu dokładności przetwarzania sygnałów za ich pomocą [13, 16]. W szczególności nie są znane z literatury rozważania na temat wpływu zmian

przedziałów całkowania lub zniekształceń mierzonego sygnału na wynik pomiaru składowych immitancji dla poszczególnych typów detektorów fazoczułych.

W pracy podjęto trzy grupy zagadnień:

1. Przedstawiono zasadę detekcji fazoczułej dla trzech typów detektorów fazoczułych z całkowaniem, stosowanych do pomiarów składowych immitancji dwójników biernych, zwłaszcza będących przetwornikami pomiarowymi.

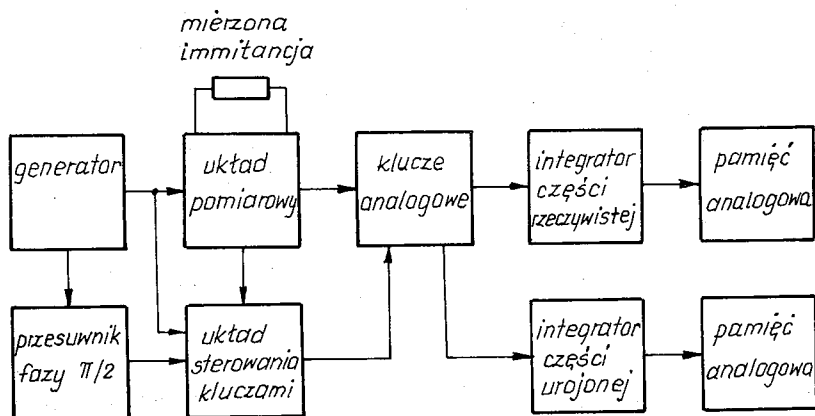
2. Przeprowadzono teoretyczne oszacowania błędów pomiaru składowych immitancji spowodowanych niedokładnym określeniem kątów fazowych przedziałów całkowania.

3. Obliczono błędy pomiarowe spowodowane przez składowe harmoniczne zawarte w sygnale mierzonym.

2. UKŁADY POMIAROWE DWÓJNIKÓW BIERNYCH

Niezależnie od wielu możliwych praktycznych realizacji detekcji fazoczułej z przeznaczeniem do pomiaru składowych immitancji dwuelementowych dwójników biernych, podstawowe operacje wykonywane na sygnale odniesienia i odpowiedzi pozostają niezmienione. Schemat blokowy układu zawierającego detektor fazoczuły i mierzący składowe immitancji dwójników biernych przedstawiono na rys. 1.

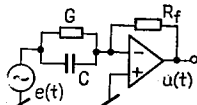
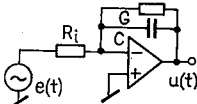
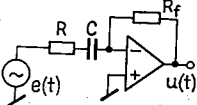
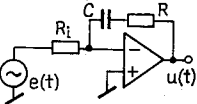
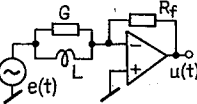
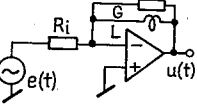
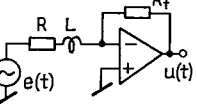
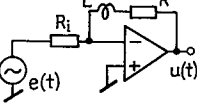
W pracy rozważono podstawowe układy detektorów fazoczułych do pomiaru składowych immitancji dwójników biernych. Badany dwójnik załączony jest w obwodzie wzmacniacza operacyjnego w sposób przedstawiony w tabeli 1. Dwójnik równoległy $G \parallel C$ lub $G \parallel L$ załącza się na wejściu wzmacniacza operacyjnego a w pętli sprzężenia zwrotnego znajduje się rezystancja wzorcowa R_f . Dwójnik szeregowy $R + L$ lub $R + C$ załącza się w pętli sprzężenia zwrotnego wzmacniacza operacyjnego a na wejściu wzmacniacza umieszcza się rezystancję wzorcową R_i .



Rys. 1. Schemat blokowy układu do pomiaru składowych immitancji dwójników za pomocą detekcji fazoczułej

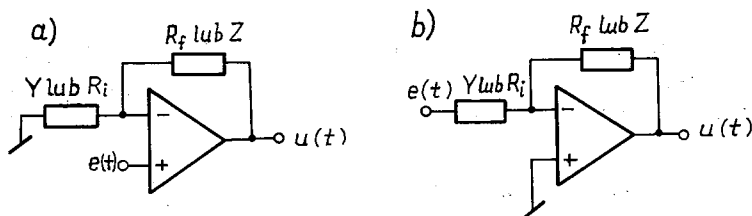
Tabela 1

Podstawowe obwody pomiarowe dwójników biernych

Nr Ozn.	Układ pomiarowy	Odpowiedź czasowa Transformata	Moduł immitancji Przesun. fazowe
1 $G \parallel C$		$u(t) = -ER_f Y \sin(\omega t + \varphi)$ $\frac{\hat{u}}{\hat{e}} = R_f(G + j\omega C)$	$ Y = \sqrt{G^2 + \omega^2 C^2}$ $\varphi = \operatorname{atn} \frac{\omega C}{G}$
2 —		$u(t) = \frac{-E}{R_i Y }\sin(\omega t - \varphi)$ $\frac{\hat{u}}{\hat{e}} = \frac{1}{R_i} \frac{G - j\omega C}{ Y ^2}$	$ Y = \sqrt{G^2 + \omega^2 C^2}$ $\varphi = \operatorname{atn} \frac{\omega C}{G}$
3 —		$u(t) = \frac{-ER_f}{ Z }\sin(\omega t + \varphi)$ $\frac{\hat{u}}{\hat{e}} = R_f \frac{R + j/(\omega C)}{ Z ^2}$	$ Z = \sqrt{R^2 + 1/(\omega C)^2}$ $\varphi = \operatorname{atn} \frac{1}{R\omega C}$
4 $R + C$		$u(t) = \frac{-E}{R_i} Z \sin(\omega t - \varphi)$ $\frac{\hat{u}}{\hat{e}} = \frac{1}{R_i}(R - j/(\omega C))$	$ Z = \sqrt{R^2 + 1/(\omega C)^2}$ $\varphi = \operatorname{atn} \frac{1}{R\omega C}$
5 $G \parallel L$		$u(t) = -ER_f Y \sin(\omega t - \varphi)$ $\frac{\hat{u}}{\hat{e}} = R_f(G - j/(\omega L))$	$ Y = \sqrt{G^2 + 1/(\omega L)^2}$ $\varphi = \operatorname{atn} \frac{1}{G\omega L}$
6 —		$u(t) = \frac{-E}{R_i Y }\sin(\omega t + \varphi)$ $\frac{\hat{u}}{\hat{e}} = \frac{1}{R_i} \frac{G + j/(\omega L)}{ Y ^2}$	$ Y = \sqrt{G^2 + 1/(\omega L)^2}$ $\varphi' = \operatorname{atn} \frac{1}{G\omega L}$
7 —		$u(t) = \frac{-ER_f}{ Z }\sin(\omega t - \varphi)$ $\frac{\hat{u}}{\hat{e}} = R_f \frac{R - j\omega L}{ Z ^2}$	$ Z = \sqrt{R^2 + \omega^2 L^2}$ $\varphi = \operatorname{atn} \frac{\omega L}{R}$
8 $R + L$		$u(t) = \frac{-E}{R_i} Z \sin(\omega t + \varphi)$ $\frac{\hat{u}}{\hat{e}} = \frac{1}{R_i}(R + j\omega L)$	$ Z = \sqrt{R^2 + \omega^2 L^2}$ $\varphi = \operatorname{atn} \frac{\omega L}{R}$

Źródło pobudzające układ pomiarowy może być załączone do wejścia nieinwersyjnego wzmacniacza operacyjnego, wówczas mierzona immitancja jest załączona między wejściem inwersyjnym a ziemią lub w pętli sprzężenia zwrotnego wzmacniacza (rys. 2). Źródło pobudzające może być również załączone do wejścia inwersyjnego poprzez mierzoną immitancję lub rezystancję wzorcową, tak jak podano w tabeli 1.

Istnieje więc szesnaście możliwych konfiguracji pomiarowych, ale tylko cztery z nich (1, 4, 5, 8 z tabeli 1), zapewniają proporcjonalność sygnału wyjściowego ze wzmacniacza do mierzonej składowej immitancji. Konfiguracje ze źródłem podłączonym do wejścia nieinwersyjnego stosuje się wówczas, gdy moduł mierzonej immitancji ma bardzo małą wartość (mniejszą od ok. 10Ω), lub gdy jeden z zacisków dwójnika (przetwornika) powinien być uziemiony.



Rys. 2. Sposób podłączenia źródła pobudzającego do układu pomiarowego składowych dwójników

Oprócz metod opisanych w pracy istnieje również sposób polegający na mierzeniu napięcia i prądu na elemencie badanym [6], co umożliwia pomiar parametrów dwójników znajdujących się w dowolnym miejscu sieci, oddzielnym od źródeł sygnału wieloma innymi elementami czynnymi i biernymi. Sygnał napięciowy jest wówczas sygnałem odniesienia.

Możliwe jest również zastąpienie rezystancji wzorcowych R_f lub R_i przez pojemność wzorcową [23]. Przedziały całkowania ulegają wówczas zamianie w stosunku do podanych w tabeli 1, tzn. składową rzeczywistą mierzy się całkując odpowiedź w przedziałach podanych dla składowej urojonej i odwrotnie. Chociaż prezentowane w tym artykule zależności odnoszą się do układów pomiarowych i błędów przetwarzania detektora jednopółkowego to jednak mogą być one rozszerzone z niewielkimi modyfikacjami na przypadek detektorów dwupółkowych.

3. SYGNAŁY WYJŚCIOWE UKŁADÓW POMIAROWYCH

3.1. SYGNAŁY WYJŚCIOWE ZE WZMACNIACZA POMIAROWEGO

Ze względu na proporcjonalność napięcia wyjściowego do wartości wielkości mierzonej, rozważone zostaną tylko układy 1, 4, 5 i 8 z tabeli 1, przy czym największe znaczenie praktyczne w zastosowaniu do parametrycznych przetworników pomiarowych wielkości nieelektrycznych mają układy 1 i 8.

Zakłada się, że operacyjny wzmacniacz pomiarowy jest idealny tzn. posiada nieskończone wzmocnienie napięciowe i nieograniczone pasmo przenoszonych częstotliwości, oraz, że pozostałe bloki toru pomiarowego nie wnoszą błędów.

Przy pobudzeniu każdego z rozważanych układów sygnałem harmonicznym o postaci:

$$e(t) = E \sin(\omega t), \quad (1)$$

przy czym: E — amplituda, ω — pulsacja generatora zasilającego, na wyjściu poszczególnych układów pomiarowych otrzymuje się napięcia zestawione w tabeli 1. $|Y|$, $|Z|$ są modułami admitancji i impedancji dwójnika, φ — kątem przesunięcia fazowego sygnału odpowiedzi względem sygnału odniesienia. Sygnały te są całkowane w odpowiednich przedziałach właściwych danej metodzie detekcji fazowej i podane zostały w tabeli 2.

Moduł wzmocnienia obwodów pomiarowych z tabeli 1 dla prądu zmiennego w układzie z zamkniętą pętlą sprzężenia zwrotnego dla dwójników równoległych $G||C$ i $G||L$ zdefiniowany jest wyrażeniem:

$$K = R_f |Y| \quad (2)$$

a dla dwójników szeregowych $R+C$ i $R+L$:

$$K = \frac{|Z|}{R_i} \quad (3)$$

Tabela 2

Przedziały całkowania w każdym okresie sygnału odniesienia dla trzech metod detekcji fazoczułej

Metoda	Dwójnik	Składowa immitancji	
		rzeczywista	urojona
stały przedział całkowania	wszystkie	$0; \pi$	$-\frac{\pi}{2}; \frac{\pi}{2}$
zmienny przedział całkowania	$G C, R+L$	$0; \varphi_m = \frac{\pi}{2} - \varphi$	$\varphi_m = \frac{\pi}{2} - \varphi; \frac{\pi}{2}$
	$R+C, G L$	$\varphi_m = \frac{\pi}{2} + \varphi; \pi$	$\frac{\pi}{2}; \varphi_m = \frac{\pi}{2} + \varphi$
próbkowanie (sample-and-hold)	wszystkie	$\delta\left(\frac{\pi}{2}\right)$	$\delta(0)$

W dalszych rozważaniach założono również, że moduł funkcji przejścia układu całkującego równy jest jedności.

3.2. DETEKTOR FAZOCZUŁY Z CAŁKOWANIEM W STAŁYM PRZEDZIALE

Przedział całkowania w tej metodzie detekcji fazoczułej zależy tylko od mierzonej składowej immitancji i jest stały dla danej składowej tzn. nie zależy od konfiguracji dwójnika ani od przesunięcia fazowego wnoszonego przez dwójnik.

Przy pomiarze składowej rzeczywistej przedział całkowania zawiera się w granicach $(0 \div \pi) + 2n\pi$, $n = 0, 1, 2, \dots$. Wartości napięć będących wynikiem całkowania wynoszą dla poszczególnych dwójników:

dla dwójnika $G\|C$ lub $R+L$:

$$U_{Re} = -KE \int_0^{\pi} \sin(\omega t + \varphi) d\omega t = -2KE \cos(\varphi) \quad (4)$$

dla dwójnika $R+C$ lub $G\|L$:

$$U_{Re} = -KE \int_0^{\pi} \sin(\omega t - \varphi) d\omega t = -2KE \cos(\varphi). \quad (5)$$

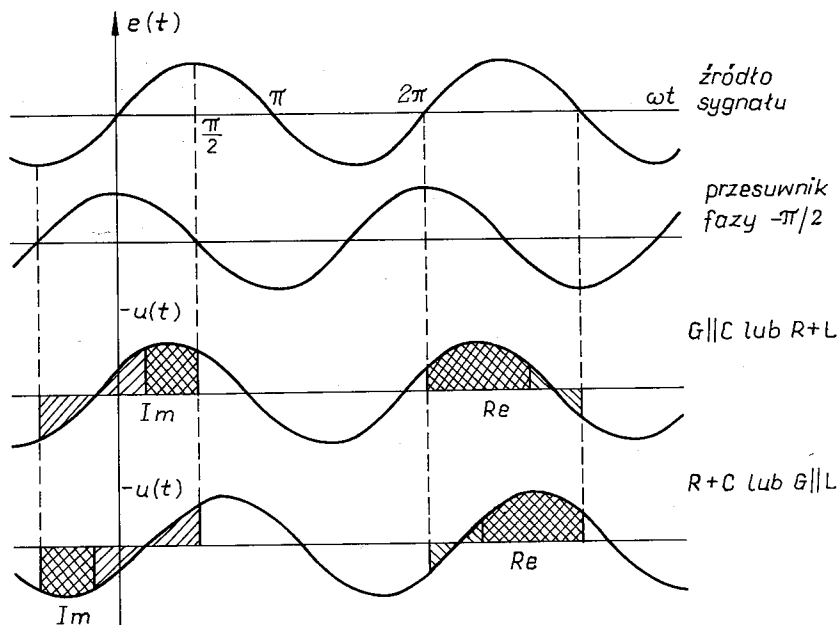
Przy pomiarze składowej urojonej całkowanie sygnału odpowiedzi odbywa się w granicach $(-\pi/2 \div +\pi/2) + 2n\pi$, $n = 0, 1, 2, \dots$, a wyniki całkowania wynoszą: dla dwójnika $G\|C$ lub $R+L$:

$$U_{Im} = -KE \int_{-\pi/2}^{+\pi/2} \sin(\omega t + \varphi) d\omega t = -2KE \sin(\varphi) \quad (6)$$

dla dwójnika $R+C$ lub $G\|L$:

$$U_{Im} = -KE \int_{-\pi/2}^{+\pi/2} \sin(\omega t - \varphi) d\omega t = 2KE \sin(\varphi). \quad (7)$$

Zaletą stałego przedziału całkowania jest dwukrotnie większe napięcie wyjściowe na wyjściu integratora w porównaniu z układem ze zmiennym przedziałem całkowania. Tłumione są również parzyste harmoniczne sygnału pochodzącego ze wzmacniacza pomiarowego, co jest korzystne w przypadku nieliniowości badanego



Rys. 3. Zasada pomiaru składowych immitancji dwójników biernych za pomocą detektora fazoczułego z całkowaniem w stałym przedziale

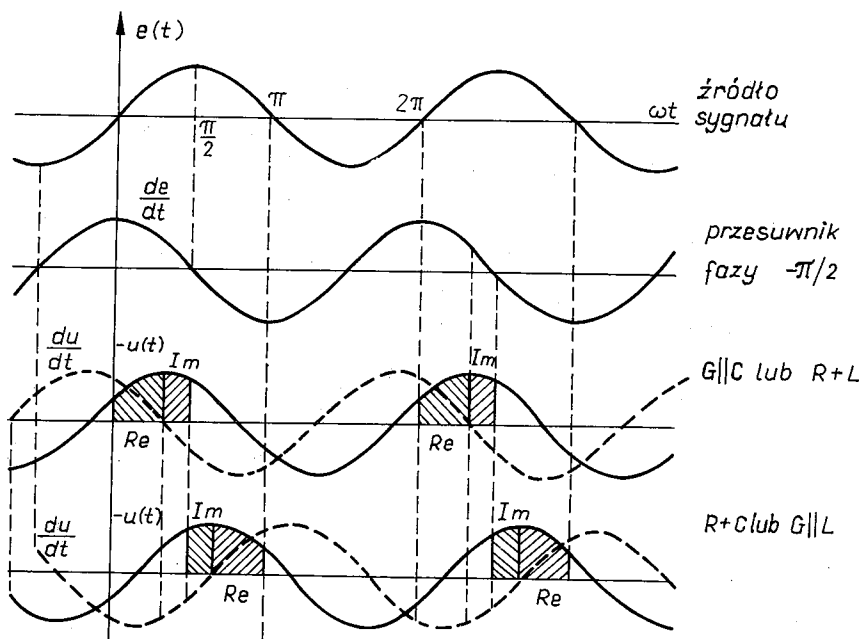
dwójnika. Korzystne jest również to, że przedział całkowania nie jest wyznaczany przez sygnał odpowiedzi, co w przypadku zniekształceń i zakłóceń (szumów) wprowadzanych przez dwójnik może prowadzić do błędnego wyznaczenia przedziału całkowania a w następstwie do błędów pomiarowych.

Przebieg sygnału odpowiedzi układu pomiarowego przedstawiono na rys. 3, z zaznaczonymi przedziałami całkowania. Podwójnie zakreskowane są te pola pod krzywą, które nie redukują się w wyniku całkowania.

3.3. DETEKTOR FAZOCZUŁY Z CAŁKOWANIEM W ZMIENNYM PRZEDZIALE

W detektorach fazoczulych z całkowaniem w zmiennym przedziale, wartości kątów fazowych określających przedziały całkowania wyznaczane są przez punkty charakterystyczne (maksima, minima, przejścia przez zero) zarówno sygnału pobudzającego (zasilającego) jak i sygnału odpowiedzi dwójnika. Dlatego przedział całkowania zmienia się ze zmianą parametrów badanego dwójnika. Zasadę detekcji fazoczulej z zaznaczonymi przedziałami całkowania przedstawiono na rys. 4.

Przy pomiarze konduktancji dwójnika równoległego $G \parallel C$ lub rezystancji dwójnika szeregowego $R + L$ przedział całkowania zawiera się w granicach $(0 \div \pi/2 - \varphi) + 2n\pi$, $n = 0, 1, 2, \dots$. Wartość bezwzględna napięcia uzyskanego na wyjściu układu całkującego wynosi:



Rys. 4. Zasada pomiaru składowych immitancji dwójników biernych za pomocą detektora fazoczulego z całkowaniem w zmiennym przedziale

$$U_{Re} = -KE \int_0^{\pi/2 - \varphi} \sin(\omega t + \varphi) d\omega t = -KE \cos(\varphi). \quad (8)$$

Jest więc proporcjonalna do składowej rzeczywistej admitancji $G = |Y| \cos(\varphi)$ lub impedancji $R = |Z| \cos(\varphi)$.

Pomiar susceptancji dwójnika równoległego $G \parallel C$ oraz reaktancji dwójnika szeregowego $R + L$ polega na scałkowaniu odpowiedzi dwójnika w przedziale $(\pi/2 - \varphi \div \pi/2) + 2n\pi$, $n = 0, 1, 2, \dots$. W wyniku scałkowania otrzymuje się napięcie, którego wartość wynosi:

$$U_{Im} = -KE \int_{\pi/2 - \varphi}^{\pi/2} \sin(\omega t + \varphi) d\omega t = -KE \sin(\varphi), \quad (9)$$

czyli jest proporcjonalna do składowej urojonej admitancji $\omega C = |Y| \sin(\varphi)$ lub impedancji $\omega L = |Z| \sin(\varphi)$.

W przypadku pary dwójników $G \parallel C$ i $R + L$ kąt przesunięcia fazowego przyjmuje wartości ujemne, natomiast dla pozostałych dwóch rozważanych dwójników wartości dodatnie i dlatego wybiera się inne przedziały całkowania.

Pomiaru rezystancji dwójnika szeregowego $R + C$ lub konduktancji dwójnika równoległego $G \parallel L$ dokonuje się przez scałkowanie napięcia odpowiedzi w przedziałach kąta fazowego $(\pi/2 + \varphi \div \pi) + 2n\pi$, $n = 0, 1, 2, \dots$, otrzymując w wyniku napięcie:

$$U_{Re} = -KE \int_{\pi/2 + \varphi}^{\pi} \sin(\omega t - \varphi) d\omega t = -KE \cos(\varphi), \quad (10)$$

które jest proporcjonalne do składowej rzeczywistej impedancji $R = |Z| \cos(\varphi)$, lub admitancji $G = |Y| \cos(\varphi)$.

Przy pomiarach reaktancji dwójnika szeregowego $R + C$, lub susceptancji dwójnika równoległego $G \parallel L$ wybiera się przedział całkowania w granicach $(\pi/2 \div \pi/2 + \varphi) + 2n\pi$:

$$U_{Im} = -KE \int_{\pi/2}^{\pi/2 + \varphi} \sin(\omega t - \varphi) d\omega t = -KE \sin(\varphi). \quad (11)$$

Otrzymana w wyniku całkowania wartość napięcia jest proporcjonalna do składowej urojonej badanej impedancji $1/\omega C = -|Z| \sin(\varphi)$, lub admitancji $1/\omega L = -|Y| \sin(\varphi)$.

3.4. DETEKTOR FAZOCZUŁY Z PRÓBKOWANIEM SYGNAŁU ODPOWIEDZI

Detekcja fazoczuła z próbkowaniem sygnału odpowiedzi (*sample-and-hold*) polega na pobraniu próbki wartości chwilowej sygnału odpowiedzi w momentach czasu odpowiadających kątom fazowym sygnału odniesienia $0 + 2n\pi$ przy pomiarze składowej urojonej lub $\pi/2 + 2n\pi$ podczas pomiaru składowej rzeczywistej. Kąty fazowe próbkowania nie zależą od konfiguracji dwójnika badanego.

Detekcja fazoczuła z próbkowaniem należy również do metod całkowych. Próbkowanie sygnału odpowiedzi można bowiem uznać za całkowanie sygnału

w nieskończenie krótkim odcinku czasu. Równoważne jest to pomnożeniu sygnału całkowanego (funkcji podcałkowej) przez funkcję Diraca $\delta(\omega t)$.

Dla składowej rzeczywistej immitancji poszczególnych dwójników otrzymuje się następujące wartości sygnałów wyjściowych:

dla dwójnika $G \parallel C$ lub $R + L$:

$$U_{Re} = -KE \int_0^{\pi} \sin(\omega t + \varphi) \delta\left[\frac{\pi}{2}\right] d\omega t = -KE \cos(\varphi) \quad (12)$$

dla dwójnika $R + C$ lub $G \parallel L$:

$$U_{Re} = -KE \int_0^{\pi} \sin(\omega t - \varphi) \delta\left[\frac{\pi}{2}\right] d\omega t = -KE \cos(\varphi). \quad (13)$$

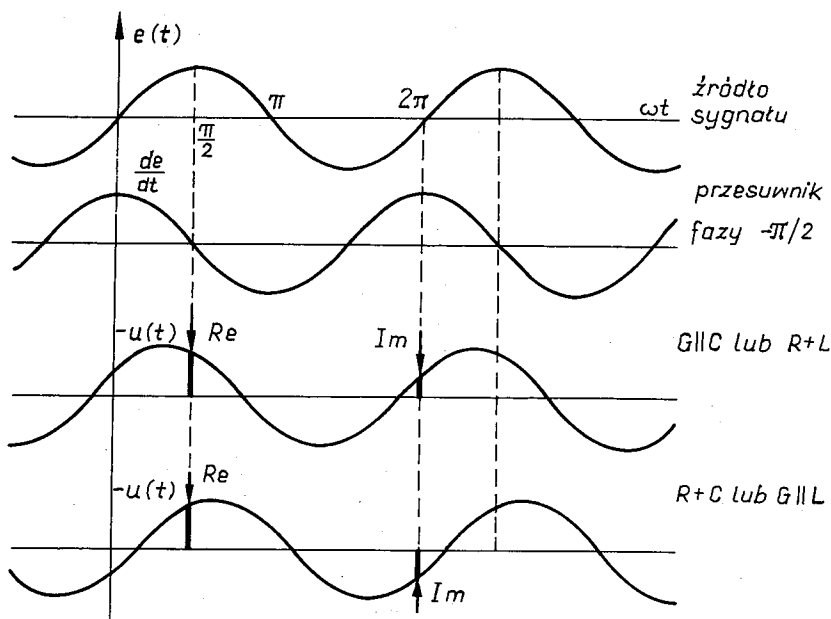
Dla składowej urojonej immitancji poszczególnych dwójników otrzymuje się następujące wartości sygnałów wyjściowych:

dla dwójnika $G \parallel C$ lub $R + L$:

$$U_{Im} = -KE \int_{-\pi/2}^{+\pi/2} \sin(\omega t + \varphi) \delta(0) d\omega t = -KE \sin(\varphi) \quad (14)$$

dla dwójnika $R + C$ lub $G \parallel L$:

$$U_{Im} = -KE \int_{-\pi/2}^{+\pi/2} \sin(\omega t - \varphi) \delta(0) d\omega t = KE \sin(\varphi). \quad (15)$$



Rys. 5. Zasada pomiaru składowych immitancji dwójników biernych za pomocą detektora fazoczułego z próbkowaniem

Momenty próbkowania i odpowiadające im wartości mierzonych wielkości zaznaczono strzałkami na krzywych odpowiedzi dwójników na rys. 5.

4. BŁĄD POMIARU SPOWODOWANY PRZEZ NIEDOKŁADNE WYZNACZENIE PRZEDZIAŁU CAŁKOWANIA

Niedokładność określenia kątów fazowych wyznaczających przedziały całkowania lub kąta fazowego próbkowania sygnału odpowiedzi może wynikać z zakłóceń sygnału odniesienia lub odpowiedzi, zawartości wyższych harmonicznych lub napięć i prądów niezerównoważenia wzmacniaczy operacyjnych w torze pomiarowym, a także dryfu poziomu napięcia odniesienia w układach komparatorów.

Przesunięcie granic przedziałów całkowania od wartości nominalnej φ_d do wartości $\varphi_d + \Delta\varphi_d$ (dolna granica przedziału całkowania) i od wartości φ_g do $\varphi_g + \Delta\varphi_g$ (górna granica przedziału całkowania) powoduje przyrost napięcia na wyjściach integratorów. Błąd względny całkowania, odniesiony do wartości wielkości mierzonej w przypadku bezbłędneho wyznaczenia przedziałów całkowania wyraża się stosunkiem:

$$\delta_\varphi = \frac{\int_{\varphi_d}^{\varphi_g} \sin(\omega t \pm \varphi) d\omega t - \int_{\varphi_d + \Delta\varphi_d}^{\varphi_g + \Delta\varphi_g} \sin(\omega t \pm \varphi) d\omega t}{\int_{\varphi_d}^{\varphi_g} \sin(\omega t \pm \varphi) d\omega t}. \quad (16)$$

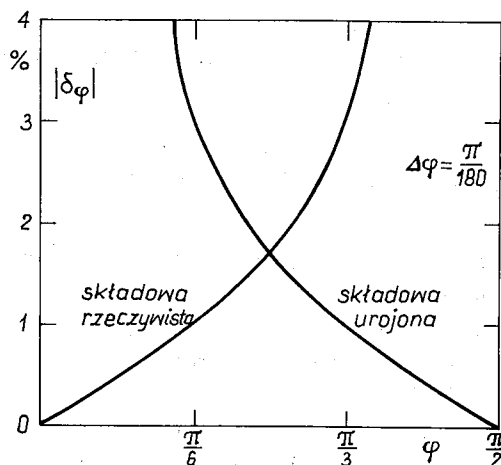
W kolejnych podrozdziałach podano wyrażenia na błąd względny pomiaru powstały na skutek niedokładności wyznaczenia kątów fazowych początku i końca przedziału całkowania. Podano również zależności przybliżone, w których założono jednakową wartość bezwzględną i znak błędu z jakim wyznaczono kąty fazowe górnej i dolnej granicy całkowania $\Delta\varphi_g = \Delta\varphi_d = \Delta\varphi$. Przyjęto również, że dla małych kątów fazowych $\Delta\varphi \leq \pi/180^\circ$ zachodzą równości $\sin(\Delta\varphi) \cong \Delta\varphi$ i $\cos(\Delta\varphi) \cong 1$.

4.1. DETEKTOR FAZOCZUŁY Z CAŁKOWANIEM W STAŁYM PRZEDZIALE

Wyrażenia na błąd pomiaru obu składowych immitancji przy pomiarach tą metodą są następujące:

dla składowej rzeczywistej dwójników $G \parallel C$ i $R + L$:

$$\begin{aligned} \delta_{\varphi Re} &= 1 - \frac{1}{2\cos(\varphi)} \int_{0 + \Delta\varphi_d}^{\pi + \Delta\varphi_g} \sin(\omega t + \varphi) d\omega t \\ &= 1 - \frac{1}{2} [\cos(\Delta\varphi_d) + \cos(\Delta\varphi_g) - \operatorname{tg}(\varphi) [\sin(\Delta\varphi_d) + \sin(\Delta\varphi_g)]] \\ &\cong \Delta\varphi \operatorname{tg}(\varphi), \end{aligned} \quad (17)$$



Rys. 6. Wartość bezwzględna względnego błędzi pomiaru spowodowanego przez niedokładne określenie przedziałów całkowania w pomiarach za pomocą detektora fazoczułego z całkowaniem w stałym przedziale i z próbkowaniem

dla składowej urojonej dwójników $G \parallel C$ i $R + L$:

$$\begin{aligned} \delta_{\varphi Im} &= 1 - \frac{1}{2\sin(\varphi)} \int_{-\frac{\pi}{2} + \Delta\varphi_d}^{\frac{\pi}{2} + \Delta\varphi_g} \sin(\omega t + \varphi) d\omega t \\ &= 1 - \frac{1}{2} [\cos(\Delta\varphi_d) + \cos(\Delta\varphi_g) - \operatorname{ctg}(\varphi) [\sin(\Delta\varphi_d) + \sin(\Delta\varphi_g)]] \\ &\cong -\Delta\varphi \operatorname{ctg}(\varphi). \end{aligned} \quad (18)$$

Pomiar składowej rzeczywistej dwójników szeregowego $R + C$ i równoległego $G \parallel L$ obarczony jest błędem wynoszącym:

$$\begin{aligned} \delta_{\varphi Re} &= 1 - \frac{1}{2\cos(\varphi)} \int_{0 + \Delta\varphi_d}^{\pi + \Delta\varphi_g} \sin(\omega t - \varphi) d\omega t \\ &= 1 - \frac{1}{2} [\cos(\Delta\varphi_d) + \cos(\Delta\varphi_g) + \operatorname{tg}(\varphi) [\sin(\Delta\varphi_d) + \sin(\Delta\varphi_g)]] \\ &\cong -\Delta\varphi \operatorname{tg}(\varphi), \end{aligned} \quad (19)$$

a pomiar składowej urojonej błędem:

$$\delta_{\varphi Im} = 1 - \frac{1}{2\sin(\varphi)} \int_{-\frac{\pi}{2} + \Delta\varphi_d}^{\frac{\pi}{2} + \Delta\varphi_g} \sin(\omega t - \varphi) d\omega t$$

$$= 1 - \frac{1}{2} [\cos(\Delta\varphi_d) + \cos(\Delta\varphi_g) - \operatorname{ctg}(\varphi) [\sin(\Delta\varphi_d) + \sin(\Delta\varphi_g)]]$$

$$\cong \Delta\varphi \operatorname{ctg}(\varphi). \quad (20)$$

Moduły wartości błędów pomiaru składowej rzeczywistej $|\delta_{\varphi Re}|$ określonych wyrażeniami przybliżonymi są jednakowe dla wszystkich czterech rodzajów dwójników. Podobnie jest w przypadku błędów pomiaru składowej urojonej $|\delta_{\varphi Im}|$. Wadą detektora fazoczułego z całkowaniem w stałym przedziale jest silna wrażliwość wyniku pomiaru na błąd wyznaczenia przedziału całkowania przy małym udziale mierzonej składowej w badanej immitancji tzn. przy wartościach kąta fazowego φ dążących do $\pi/2$ w przypadku pomiaru składowej rzeczywistej oraz dążących do 0 w przypadku pomiaru składowej urojonej. Wykresy modułów błędów (17)–(20) przedstawiono na rys. 6. dla wartości $\Delta\varphi = \pi/180$.

4.2. DETEKTOR FAZOCZUŁY Z CAŁKOWANIEM W ZMIENNYM PRZEDZIALE

Dla dwójnika równoległego $G \parallel C$ i dwójnika szeregowego $R + L$ wyrażenia na błąd pomiaru spowodowany przez niedokładne wyznaczenie kątów fazowych przedziałów całkowania są jednakowe i wynoszą dla składowej rzeczywistej:

$$\delta_{\varphi Re} = 1 - \frac{1}{\cos(\varphi)} \int_{0 + \Delta\varphi_d}^{\frac{\pi}{2} - \varphi + \Delta\varphi_g} \sin(\omega t + \varphi) d\omega t = 1 - \cos(\Delta\varphi_d) + \operatorname{tg}(\varphi) \sin(\Delta\varphi_d) - \frac{\sin(\Delta\varphi_g)}{\cos(\varphi)}$$

$$\cong \Delta\varphi \frac{1 - \sin(\varphi)}{\cos(\varphi)} \quad (21)$$

i dla składowej urojonej:

$$\delta_{\varphi Im} = 1 - \frac{1}{\sin(\varphi)} \int_{\frac{\pi}{2} - \varphi + \Delta\varphi_d}^{\frac{\pi}{2} + \Delta\varphi_g} \sin(\omega t + \varphi) d\omega t = 1 - \cos(\Delta\varphi_g) + \operatorname{ctg}(\varphi) \sin(\Delta\varphi_g) - \frac{\sin(\Delta\varphi_d)}{\sin(\varphi)}$$

$$\cong \Delta\varphi \frac{1 - \cos(\varphi)}{\sin(\varphi)}. \quad (22)$$

Podobne zależności otrzymuje się dla dwójników $R + C$ i $G \parallel L$, przy pomiarze składowej rzeczywistej immitancji:

$$\delta_{\varphi Re} = 1 - \frac{1}{\cos(\varphi)} \int_{\frac{\pi}{2} + \varphi + \Delta\varphi_d}^{\pi + \Delta\varphi_g} \sin(\omega t - \varphi) d\omega t = 1 - \cos(\Delta\varphi_g) - \operatorname{tg}(\varphi) \sin(\Delta\varphi_g) + \frac{\sin(\Delta\varphi_d)}{\cos(\varphi)}$$

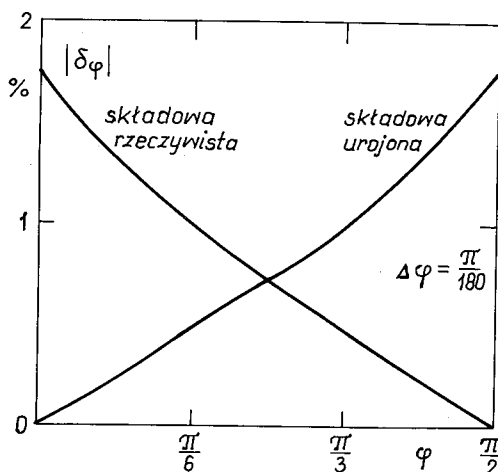
$$\cong \Delta\varphi \frac{1 - \sin(\varphi)}{\cos(\varphi)} \quad (23)$$

i przy pomiarze składowej urojonej:

$$\delta_{\varphi Im} = 1 - \frac{1}{\sin(\varphi)} \int_{\frac{\pi}{2} + \Delta\varphi_d}^{\frac{\pi}{2} + \varphi + \Delta\varphi_g} \sin(\omega t - \varphi) d\omega t = 1 - \cos(\Delta\varphi_d) + \operatorname{ctg}(\varphi) \sin(\Delta\varphi_d) - \frac{\sin(\Delta\varphi_g)}{\sin(\varphi)}$$

$$\cong -\Delta\varphi \frac{1 - \cos(\varphi)}{\sin(\varphi)} \quad (24)$$

Błąd względny pomiaru składowych immitancji spowodowany przez niedokładne określenie przedziałów całkowania w detektorze fazoczułym z całkowaniem w zmiennym przedziale przedstawiono na rys. 7. Wynik pomiaru tą metodą jest najmniej wrażliwy ze wszystkich trzech rozważanych metod detekcji fazoczułej, na błędne określenie przedziałów całkowania, ponadto w całym zakresie zmian kąta przesunięcia fazowego φ błąd pomiaru ma wartość ograniczoną.



Rys. 7. Wartość bezwzględna względnego błędu pomiaru spowodowanego przez niedokładne określenie przedziałów całkowania w pomiarach za pomocą detektora fazoczułego z całkowaniem w zmiennym przedziale

4.3. DETEKTOR FAZOCZUŁY Z PRÓBKOWANIEM

W przypadku metody detekcji fazoczułej polegającej na próbkowaniu sygnału odpowiedzi, niedokładne określenie kąta fazowego próbkowania powoduje pobranie próbki sygnału w niewłaściwym momencie czasu tzn. dla kątów fazowych $\varphi + \Delta\varphi$.

Błąd pomiaru składowej rzeczywistej immitancji dla wszystkich dwójników wynosi:

$$\delta_{\varphi Re} = 1 - \cos(\Delta\varphi) + \operatorname{tg}(\varphi) \sin(\Delta\varphi) \cong \Delta\varphi \operatorname{tg}(\varphi). \quad (25)$$

Błąd pomiaru składowej urojonej immitancji dwójników $G \parallel C$ i $R + L$ dany jest wyrażeniem:

$$\delta_{\varphi Im} = 1 - \cos(\Delta\varphi) - \operatorname{ctg}(\varphi) \sin(\Delta\varphi) \cong -\Delta\varphi \operatorname{ctg}(\varphi) \quad (26)$$

a dwójników $R + C$ i $G \parallel L$:

$$\delta_{\varphi Im} = 1 - \cos(\Delta\varphi) + \operatorname{ctg}(\varphi) \sin(\Delta\varphi) \cong \Delta\varphi \operatorname{ctg}(\varphi). \quad (27)$$

Zależności przybliżone opisujące błąd pomiaru spowodowany niedokładnym określeniem kąta fazowego próbkowania są identyczne jak w przypadku detekcji fazoczułej z całkowaniem w stałym przedziale, a wykresy błędów przedstawione zostały na rys. 6. Również w przypadku tej metody, błąd pomiaru rośnie nieograniczenie jeśli udział mierzonej składowej immitancji jest dużo mniejszy od składowej ortogonalnej.

5. BŁĄD POMIARU SPOWODOWANY SKŁADOWYMI HARMONICZNYMI

W rozdziale określono udział składowych harmoniczných zawartych w sygnale wyjściowym ze wzmacniacza pomiarowego w wyniku pomiaru składowych immitancji.

Błąd względny pomiaru spowodowany obecnością w mierzonym sygnale i -tej harmoniczných definiowany jest jako stosunek wartości całki z danej harmoniczných obliczonej w odpowiednim dla mierzonej składowej immitancji przedziale całkowania $\langle \varphi_d, \varphi_g \rangle$ do wartości całki ze składowej podstawowej obliczonej w tym samym przedziale:

$$\delta_i = \frac{A_i \int_{\varphi_d}^{\varphi_g} \sin(i\omega t + \varphi_i) d\omega t}{A_1 \int_{\varphi_d}^{\varphi_g} \sin(\omega t \pm \varphi) d\omega t}. \quad (28)$$

Rozważania zostaną ograniczone do drugiej i trzeciej harmoniczných, które mają największy wpływ na wynik pomiaru ze względu na ich amplitudę na ogół większą od pozostałych harmoniczných. A zatem w wyrażeniu (28) $i = 2; 3$. Wyrażenie w mianowniku jest napięciem na wyjściu układu całkującego jednego z układów detektorów, odpowiadającym mierzonej składowej immitancji badanego dwójnika przy braku harmoniczných.

Założono następujące postacie drugiej i trzeciej harmoniczných sygnału w torze pomiarowym:

$$u_2(\omega t) = A_2 \sin(2\omega t + \varphi_2) \quad u_3(\omega t) = A_3 \sin(3\omega t + \varphi_3) \quad (29), (30)$$

gdzie A_2, A_3 — amplitudy, φ_2, φ_3 — fazy odpowiednio drugiej i trzeciej harmoniczných.

Uwzględniając, że próbkowanie sygnału odpowiada δ — całkowaniu, dla metody detekcji fazoczułej z próbkowaniem błąd pomiaru można zapisać jako stosunek wartości i -tej harmoniczných w momencie próbkowania do wartości chwilowej pierwszej harmoniczných:

$$\delta_i = \frac{A_i \sin(\varphi_p + \varphi_i)}{A_1 \sin(\varphi_p \pm \varphi)}, \quad (31)$$

przy czym φ_p jest kątem fazowym, dla którego następuje pobranie próbki sygnału odpowiedzi (0 lub $\pi/2$).

W przypadku detekcji dwupołówkowej, o ile w układzie nie przewidziano uśredniania napięcia wyjściowego w czasie dłuższym od okresu sygnału odniesienia, wyrażenia na tłumienie składowych harmoniczných są identyczne jak w przypadku detekcji jednapołówkowej, z tą różnicą, że w każdym półokresie następuje zmiana znaku wyniku całkowania każdej harmoniczných. Jeśli scałkowany sygnał jest uśredniany w okresie równym wielokrotności składowej podstawowej, następuje redukcja wyników całkowania o przeciwnych znakach.

5.1. DETEKTOR FAZOCZUŁY Z CAŁKOWANIEM W STAŁYM PRZEDZIALE

W detektorze fazoczulym z całkowaniem w stałym przedziale, wynik całkowania drugiej harmoniczných wynosi zero zarówno przy pomiarze składowej rzeczywistej jak i urojonej immitancji. Jest to podstawowa zaleta tej metody detekcji fazoczulęj, stanowiąca o jej przewadze nad pozostałymi dwiema metodami w sytuacji, gdy detekowany sygnał jest zniekształcony lub zakłócony. To samo odnosi się do wszystkich parzystych harmoniczných.

Dla trzeciej harmoniczných błędy pomiarowe składowych immitancji wyrażają się następującymi zależnościami:

przy pomiarze składowej rzeczywistej:

$$\delta_{3Re} = \frac{A_3}{3EK} \frac{\cos(\varphi_3)}{\cos(\varphi)} \quad (32)$$

przy pomiarze składowej urojonej:

$$\delta_{3Im} = \frac{-A_3}{3EK} \frac{\sin(\varphi_3)}{\sin(\varphi)}. \quad (33)$$

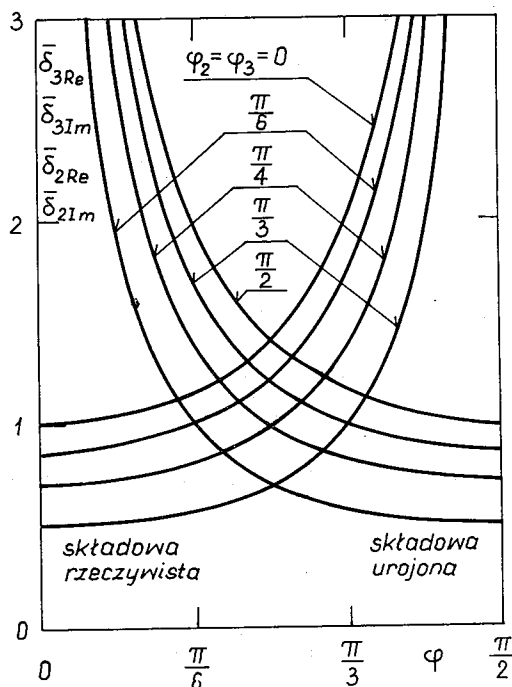
Wykresy błędu pomiarowego w formie znormalizowanej zgodnie z definicją:

$$\bar{\delta}_{3Re} = \left| \delta_{3Re} \frac{3EK}{A_3} \right| = \left| \frac{\cos(\varphi_3)}{\cos(\varphi)} \right|, \quad (34)$$

$$\bar{\delta}_{3Im} = \left| \delta_{3Im} \frac{3EK}{A_3} \right| = \left| \frac{\sin(\varphi_3)}{\sin(\varphi)} \right|, \quad (35)$$

przedstawiono na rys. 8. Błąd pomiaru jest również zależny od fazy trzeciej harmoniczných. Wartość liczbowa błędu względnego można oszacować z zależności (34) i (35) lub rys. 8 po podstawieniu konkretnych wartości napięć i wzmacnienia.

W detektorze fazoczulym z całkowaniem w stałym przedziale, wynik całkowania parzystych harmoniczných wynosi zero zarówno dla pomiaru składowej rzeczywistej jak i urojonej immitancji. W pewnym zakresie zmienności kąta fazowego φ ,



Rys. 8. Błąd względny pomiaru składowych immitancji spowodowany składowymi harmonicznymi w pomiarach za pomocą detektora fazoczułego z całkowaniem w stałym przedziale i z próbkowaniem

skuteczniejsze jest również tłumienie trzeciej harmonicznej, w porównaniu z pozostałymi metodami detekcji.

5.2. DETEKTOR FAZOCZUŁY Z CAŁKOWANIEM W ZMIENNYM PRZEDZIALE

Wyznaczone z równania (28) błędy pomiaru składowych immitancji za pomocą detektora fazoczułego z całkowaniem w zmiennym przedziale, przy detekcji jedno- i dwupółkowej wynoszą dla składowej rzeczywistej dwójników $G \parallel C$ i $R + L$:

$$\delta_{2Re} = \frac{A_2}{2EK} \frac{\cos(2\varphi - \varphi_2) + \cos(\varphi_2)}{\cos(\varphi)}, \quad (36)$$

$$\delta_{3Re} = \frac{A_3}{3EK} \frac{\sin(3\varphi - \varphi_3) + \cos(\varphi_3)}{\cos(\varphi)}, \quad (37)$$

i składowej urojonej:

$$\delta_{2Im} = \frac{A_2}{2EK} \frac{-\cos(2\varphi + \varphi_2) + \cos(\varphi_2)}{\sin(\varphi)}, \quad (38)$$

$$\delta_{3Im} = \frac{A_3}{3EK} \frac{\sin(3\varphi + \varphi_3) - \sin(\varphi_3)}{\sin(\varphi)} \quad (39)$$

Błędy pomiaru składowej rzeczywistej dwójników $R + C$ i $G \parallel L$ wynoszą:

$$\delta_{2Re} = \frac{A_2}{2EK} \frac{-\cos(2\varphi + \varphi_2) - \cos(\varphi_2)}{\cos(\varphi)}, \quad (40)$$

$$\delta_{3Re} = \frac{A_3}{3EK} \frac{\sin(3\varphi + \varphi_3) + \cos(\varphi_3)}{\cos(\varphi)} \quad (41)$$

i składowej urojonej:

$$\delta_{2Im} = \frac{A_2}{2EK} \frac{\cos(2\varphi + \varphi_2) - \cos(\varphi_2)}{\sin(\varphi)}, \quad (42)$$

$$\delta_{3Im} = \frac{A_3}{3EK} \frac{-\sin(3\varphi + \varphi_3) + \sin(\varphi_3)}{\sin(\varphi)}. \quad (43)$$

Wyrażenia (36)–(43) przy założeniu $\varphi_2 = \varphi_3 = 0$, znormalizowane zgodnie z definicjami (34) i (35) oraz

$$\bar{\delta}_{2Re} = \left| \delta_{2Re} \frac{2EK}{A_2} \right| \quad (44)$$

$$\bar{\delta}_{2Im} = \left| \delta_{2Im} \frac{2EK}{A_2} \right|, \quad (45)$$

dla drugiej harmonicznej wynoszą:

$$\bar{\delta}_{2Re} = 2\cos(\varphi), \quad (46)$$

$$\bar{\delta}_{3Re} = \frac{\sin(3\varphi) + 1}{\cos(\varphi)}, \quad (47)$$

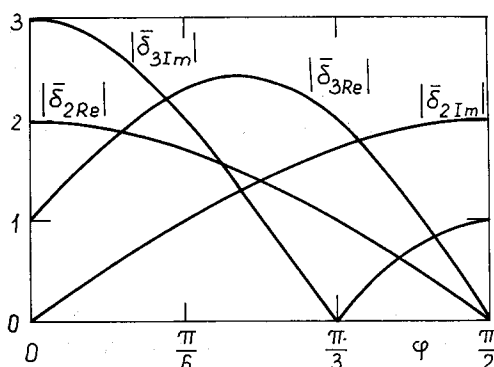
$$\bar{\delta}_{2Im} = 2\sin(\varphi), \quad (48)$$

$$\bar{\delta}_{3Im} = \left| \frac{\sin(3\varphi)}{\sin(\varphi)} \right|, \quad (49)$$

Przykładowe wykresy błędów (46)–(49) dla $\varphi_2 = \varphi_3 = 0$, w zależności od kąta przesunięcia fazowego φ przedstawiono na rys. 9. Wykresy te ilustrują charakter zmienności wartości błędu w postaci znormalizowanej. Wartości liczbowe błędów otrzymuje się po podstawieniu konkretnych wartości napięć i immitancji. Dla innych wartości kątów fazowych φ_2, φ_3 charakter zmienności jest taki sam, z tym, że wykres błędu jest przesunięty wzdłuż osi odciętych.

Tłumienie i -tej składowej harmonicznej po przejściu przez detektor fazoczuły uzależnione jest od fazy danej harmonicznej oraz od kąta przesunięcia fazowego składowej podstawowej sygnału odpowiedzi. Wartość średnia drugiej harmonicznej jest przez detektor fazoczuły z całkowaniem w zmiennym przedziale tłumiona co najmniej dwukrotnie, a wartość średnia trzeciej harmonicznej co najmniej trzykrotnie w stosunku do wartości średniej z modułu.

Tłumienie składowych harmonicznych w tej metodzie detekcji fazoczułej jest mniej skuteczne niż w przypadku detekcji fazoczułej z całkowaniem w stałym przedziale, lecz lepsze niż przez detektor fazoczuły z próbkowaniem. W przypadku detekcji dwupołkowej, nie następuje całkowite wytłumienie drugiej harmonicznej, jak w przypadku pozostałych metod detekcji po uśrednieniu wyniku całkowania w czasie dłuższym od okresu sygnału podstawowego. Zaletą tego detektora jest jednak to, że wartość błędu jest zawsze ograniczona dla każdej wartości kąta fazowego φ z przedziału od 0 do $\pi/2$, i dla obu składowych harmonicznych. Może on więc być stosowany w przypadkach, gdy udział mierzonej składowej immitancji jest mały w porównaniu ze składową ortogonalną.



Rys. 9. Błąd względny pomiaru składowych immitancji spowodowany składowymi harmonicznymi w pomiarach za pomocą detektora fazoczułego z całkowaniem w zmiennym przedziale

5.3. DETEKTOR FAZOCZUŁY Z PRÓBKOWANIEM

Błędy wnoszone przez wyższe harmoniczne w przypadku detekcji fazoczułej z próbkowaniem określone są ogólnym wyrażeniem (31).

Dla składowej rzeczywistej wszystkich typów rozważanych dwójników błędy wynoszą:

$$\delta_{2Re} = \frac{-A_2 \sin(\varphi_2)}{EK \cos(\varphi)}, \quad (50)$$

$$\delta_{3Re} = \frac{-A_3 \cos(\varphi_3)}{EK \cos(\varphi)}, \quad (51)$$

a dla składowej urojonej immitancji:

$$\delta_{2Im} = \frac{A_2 \sin(\varphi_2)}{EK \sin(\varphi)}, \quad (52)$$

$$\delta_{3Im} = \frac{A_3 \sin(\varphi_3)}{EK \sin(\varphi)}. \quad (53)$$

Charakter zależności błędu od kąta przesunięcia fazowego φ dla obu harmonicznych jest podobny jak w przypadku detekcji fazoczułej z całkowaniem w stałym przedziale (por. rys. 8).

Wpływ danej harmonicznej na wynik próbkowania zależy od wartości chwilowej danej harmonicznej w momencie pobrania próbki sygnału, a więc także od amplitudy i kąta przesunięcia fazowego i -tej harmonicznej. Wartość chwilowa i -tej harmonicznej znormalizowana względem wartości średniej z modułu, w zależności od kąta fazowego φ_i jest jednakowa dla wszystkich dwójników. Względny błąd pomiaru dąży do nieskończoności przy małym udziale mierzonej składowej względem składowej ortogonalnej.

Tłumienie składowych harmonicznych w przypadku detekcji fazoczułej dwupółkowej jest takie samo, z zastrzeżeniem, że w każdym półokresie sygnału podstawowego następuje zmiana znaku wyniku próbkowania drugiej harmonicznej. W przypadku dodatkowego uśredniania w czasie dłuższym od jednego okresu następuje redukcja drugiej harmonicznej, czyli wyeliminowanie jej wpływu na wynik pomiaru.

ZAKOŃCZENIE

Przedstawiono zasadę pomiaru składowych rzeczywistej i urojonej immitancji dwuelementowych, liniowych dwójników biernych za pomocą trzech typów detektorów fazoczułych. Metody te są szczególnie przydatne w pomiarach dynamicznych za pomocą parametrycznych przetworników pomiarowych typu RLC, przy porównywalnym udziale składowych ortogonalnych w mierzonej immitancji. Podano podstawowe zależności określające napięcie wyjściowe z detektora fazoczułego dla czterech typów dwójników, reprezentujących przetwornik pomiarowy. Określone zostały wartości błędów spowodowane przez niedokładne określenie przedziałów całkowania w poszczególnych metodach detekcji fazoczułej. Z przeprowadzonych rozważań wynika, że najmniej wrażliwy na błąd określenia przedziału całkowania jest pomiar za pomocą detektora fazoczułego z całkowaniem w zmiennym przedziale.

Określone zostały również wartości błędów spowodowane przez składowe harmoniczne zawarte w mierzonym sygnale. Z przytoczonych obliczeń wynika, że najmniej wrażliwy na składowe harmoniczne jest pomiar za pomocą detektora fazoczułego z całkowaniem w zmiennym przedziale, a błąd ma zawsze wartość ograniczoną. Najskuteczniejsze tłumienie drugiej harmonicznej jest w detektorze fazoczułym z całkowaniem w stałym przedziale.

BIBLIOGRAFIA

1. A. L. Stott, R. G. Green, K. Seraji: *Comparison of the use of internal and external electrodes for the measurement of the capacitance and conductance of fluid in pipes*. J. Phys. E: Sci. Instrum. 1985, vol. 18, 587—92
2. B. Durkiewicz: *Metoda admitancyjna pomiaru zawartości faz w przepływie dwufazowym*. Sympozjum Wymiany Ciepła i Masy. Warszawa-Jabłonna, 1986

3. A. Szulce: *Mostki elektryczne pomiarowe*. Warszawa 1977
4. A.C.J. Beerens: *Messgeräte und Messmethoden in der Elektronik*. Hamburg, 1971
5. B. Szadkowski: *Synteza metod pomiaru immitancji*. ZN Politechniki Śląskiej. Elektryka 1984, Nr 93
6. J. Hoja: *Metoda pomiarowa parametrów elementów RLC wykorzystująca sygnały proporcjonalne do prądu i napięcia na elemencie mierzonym*. Rozprawy Elektrotechniczne. 1984, vol. 30 Nr 3, 643—58
7. J.E. Sigdell: *A principle for capacitance measurement, suitable for linear evaluation of capacitance transducers*. IEEE Trans. Instrum. Meas. 1972, vol. 21, No 1, 60—4
8. R.V. Jones, J. C. S. Richards: *The design and some applications of sensitive capacitance micrometers*. J. Phys. E: Sci. Instrum. 1973, vol. 6, No 7, 589—600
9. H. Fisher: *Lock-in amplifiers obtain measurements even if there is more noise than signal*. Laser Focus. Nov. 1977, 82—8
10. A. Jaworek: *Pomiary składowych impedancji dwójników*. PAK 1987, vol. 33, Nr 2, 38—9
11. A. Jaworek: *Wyznaczanie składowych impedancji parametrycznych przetworników pomiarowych za pomocą detekcji fazoczułej*. VII Krajowa Konferencja Metrologów. Wrocław, 1986
12. A. Jaworek: *Sposób i układ do pomiaru składowych impedancji lub admitancji dwójników*. Patent P-149187
13. L.Q. Orsini: *Conversion of immittance parameters to dc voltage*. IEEE Trans. Instrum. Meas. 1973, vol. 21, No 3, 196—8
14. A. Jaworek: *Sposób szybkiego pomiaru składowych immitancji dwójników metodą próbkowania odpowiedzi*. PAK 1987, vol. 33, Nr 5, 102—3
15. A. Jaworek: *Sposób i układ do pomiaru składowych rzeczywistej i urojonej immitancji dwójników*. Patent P-149188
16. J. Hoja: *Optymalizacja parametrów detektora fazoczułego przy niekorzystnym stosunku składowych ortogonalnych sygnału pomiarowego*. XVII Międzyuczelniana Konferencja Metrologów, Poznań, Wrzesień 1984
17. G.B. Clayton: *Experiments with operational amplifiers*. Pt. 11: An op-amp used as a phase sensitive detector. Wireless World, July 1973, 355—6
18. C.T. Morse: *A computer controlled apparatus for measuring ac properties of materials over the frequency range 10^{-5} to 10^5 Hz*. J. Phys. E: Sci. Instrum. 1974, vol. 7, 657—62
19. J. Hoja: *Bezpośrednia metoda pomiaru elementów składowych dwójników wieloelementowych o znanej strukturze*. Praca doktorska. Politechnika Gdańska, Gdańsk, 1979
20. K. Maeda: *An Automatic Precision 1 MHz Digital LCR Meter*. Hewlett-Packard J., March 1974, 2—9
21. M. Ortenberg von, W. Staguhn, F. Böbel, S. Takeyama, T. Sakakibara, N. Miura: *Phase-sensitive detection of magnetoresistivity by the four-probe technique in pulsed high magnetic fields*. J. Phys. E: Sci. Instrum. 1989, vol. 22, No 6, 369—62
22. S.M. Huang, A.L. Stott, R.G. Green, M.S. Beck: *Electronic transducers for industrial measurement of low value capacitances*. J. Phys. E: Sci. Instrum. 1988, vol. 21, No 3, 242—50
23. J.C.R. Richards: *Some aspects of transducer immittance measurement*. J. Phys. E: Sci. Instrum. 1982, vol. 15, 1251—56

A. JAWOREK

PHASE-SENSITIVE DETECTORS IN MEASUREMENTS OF IMMITTANCE COMPONENTS OF A TWO-TERMINAL

Summary

The principles of measurement of real and imaginary parts of a complex immittance of a passive two-terminal by means of three types of phase-sensitive detectors have been presented. The two-terminal

may in particular be a measuring transducer. The measurement is carried out by integration in an appropriate time intervals of steady-state time-response of a two-terminal to a harmonic excitation voltage. The accuracy of measurement by these methods has been considered. The measurement errors due to inaccuracy in phase-angles of integration as well as the harmonic distortions have been theoretically determined.

Złożoność układów cyfrowych VLSI

JANUSZ STOKŁOSA

Ośrodek Informatyki, Politechnika Poznańska

Otrzymano 1989.9.5.

Autoryzowano do druku 1989.12.30

W pracy przedstawiono główne nurty problematyki złożonościowej w odniesieniu do projektowania układów cyfrowych VLSI. Wprowadzono pojęcia elementarne dotyczące złożoności obliczeniowej. Omówiono modele obliczeń, w szczególności równoległych. Przedstawiono możliwości paralelizowalności algorytmów i etapy projektowania algorytmów równoległych. Przedyskutowano złożoność komunikacyjną i S-komunikacyjną oraz jej związki ze złożonością układu VLSI rozumianą jako AT^2 , przy czym A jest powierzchnią układu scalonego, a T czasem przetwarzania w tym układzie danych wejściowych. Zwrócono uwagę na wybrane zagadnienia układów kombinacyjnych, w szczególności na możliwości realizacji układu VLSI za pomocą układów kombinacyjnych i ich złożoność, także w kontekście badania hierarchii boolowskiej.

1. WSTĘP

Na różnych etapach powstawania układu cyfrowego bardzo wielkiej skali scalenia (VLSI) projektant układu ma do czynienia z problematyką złożonościową. Obejmuje ona złożoność algorytmu realizowanego sprzętowo, złożoność samego układu wyłożonego na płycie krzemowej oraz złożoność algorytmów wykorzystywanych w trakcie projektowania. Co do realizowanych algorytmów i algorytmów ułatwiających projektowanie, to istotna jest tu przede wszystkim ich złożoność czasowa. Za miarę złożoności układu scalonego wykonanego na płycie krzemowej przyjmuje się iloczyn AT^2 , gdzie A jest powierzchnią układu, natomiast T czasem przetwarzania przez ten układ danych wejściowych — przy założeniu stałego wymiaru charakterystycznego λ .

Praca składa się z sześciu rozdziałów. W rozdziale 2 dokonano wprowadzenia w problematykę złożonościową. Zdefiniowano pojęcie problemu, algorytmu, złożoności obliczeniowej algorytmu, przedstawiono uproszczoną klasyfikację problemów. Trzeci rozdział poświęcono modelom obliczeń. Omówiono maszynę o dostępie

Pracę wykonano w ramach CPBP 02.17 w czasie pobytu autora w Laboratoire de Génie Informatique, Institut IMAG, Université Joseph Fourier w Grenoble.

swobodnym i jej wersję równoległą, maszynę PRAM. Zwrócono uwagę na inne modele, bliższe rzeczywistym architekturom komputerów. Obliczenia równoległe są przedmiotem rozdziału 4, w którym szczególnie skoncentrowano się na omówieniu możliwości dekompozycji równoległej algorytmów oraz wskazaniu etapów projektowania algorytmów równoległych. Rozdział 5 dotyczy złożoności układów VLSI i jej związków ze złożonością komunikacyjną i S -komunikacyjną. W rozdziale 6 przedstawiono wybrane zagadnienia dotyczące układów kombinacyjnych i hierarchii boolowskiej problemów, co ma istotne znaczenie z punktu widzenia możliwości realizacji algorytmów za pomocą układów kombinacyjnych. Uwagi podsumowujące kończą artykuł.

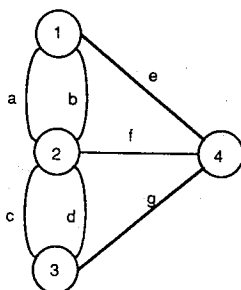
Problematykę tę zaprezentowano biorąc pod uwagę aktualny (wiosna 1989) stan wiedzy w tej dziedzinie. Ma ona jeszcze w znacznym stopniu charakter teoretyczny, niemniej niektóre wyniki już dzisiaj mają istotne znaczenie w praktyce projektowania układów cyfrowych VLSI [8, 12, 13, 15, 21, 30, 32, 43, 45].

2. WPROWADZENIE W PROBLEMATYKĘ ZŁOŻONOŚCIOWĄ

Rozdział ten jest poświęcony wprowadzeniu elementarnych pojęć dotyczących problemów, algorytmów i złożoności [15, 31, 39, 40].

Pod pojęciem *problemu* rozumie się pytanie ogólne, zwykle zależne od pewnej liczby parametrów, na które należy dać odpowiedź. Problem jest zadany wówczas, jeśli przedstawimy (i) ogólny opis jego parametrów oraz (ii) zdanie stwierdzające czym powinna charakteryzować się odpowiedź, czyli *rozwiązanie problemu*. *Przypadek szczególny* problemu uzyskuje się wówczas, gdy ustali się wartości wszystkich jego parametrów.

Jako przykład weźmy problem Eulera, w którym pytamy czy w grafie nieskierowanym o n wierzchołkach istnieje ścieżka przechodząca przez wszystkie krawędzie dokładnie jeden raz. Przypadkiem szczególnym tego problemu jest znane powszechnie zagadnienie mostów na Pregole w Królewcu: czy można wyznaczyć drogę spaceru po mieście tak, aby odwiedzić każdą z czterech części 1, 2, 3, 4 miasta i przejść przez każdy z mostów a, b, c, d, e, f, g dokładnie jeden raz? (rys. 1). W przypadku



Rys. 1. Problem mostów królewieckich

problemu mostów królewieckich odpowiedź jest negatywna, gdyż z twierdzenia Eulera [10] wynika, że przez wszystkie krawędzie grafu można przejść w sposób ciągły, przechodząc przy tym każdą krawędź dokładnie jeden raz, wtedy i tylko wtedy, gdy graf jest spójny i albo nie ma wierzchołków o nieparzystym stopniu, albo ma dokładnie dwa takie wierzchołki.

Problem Eulera należy do klasy *problemów decyzyjnych* charakteryzujących się tym, że ich rozwiązaniem może być albo odpowiedź *tak*, albo odpowiedź *nie*. Wiele problemów, w tym optymalizacyjne, można sformułować w postaci problemów decyzyjnych.

Zadany problem można opisać na wiele sposobów. Będziemy zakładać w dalszym ciągu, że z każdym problemem decyzyjnym jest związane ustalone kodowanie, oznaczające tutaj odwzorowanie poszczególnych przypadków problemu w słowa w pewnym alfabecie skończonym A . *Rozmiarem* (wejścia) danego przypadku problemu nazywamy długość słowa kodującego ten przypadek, czyli liczbę liter tego słowa. Mówimy, że rozpatrujemy problem o rozmiarze n , jeśli wszystkie jego przypadki szczególne mają rozmiar co najwyżej n .

Powracając do problemu Eulera, przedstawiony na rys. 1 przypadek szczególny tego problemu można zakodować w alfabecie $A = \{0, 1, \dots, 9, a, b, \dots, z, :, ;, /, -\}$ np. w postaci słowa: 1:2:3:4;a:b/-/e/c:d/f/g, w którym przed średnikiem wylistowano wierzchołki grafu, a za nim przedstawiono krawędzie grafu przyjmując porządek leksykograficzny. Jeśli przyjąć taką zasadę kodowania, to rozmiar rozpatrywanego przypadku wynosi 24.

Można także stosować, i stosuje się, inne miary problemów. W przypadkach dotyczących grafów bardziej naturalną jest miara w postaci liczby wierzchołków lub krawędzi, w problemie mnożenia macierzy kwadratowych — stopień macierzy, w problemie wartościowania wyrażeń boolowskich — liczba zmiennych itd.

Algorytm jest procedurą służącą do rozwiązywania problemów. Mówimy, że algorytm rozwiązuje problem, jeśli może być zastosowany do każdego przypadku szczególnego tego problemu i zawsze jest zagwarantowane uzyskanie wyniku. Wszystkie problemy z jakimi mamy do czynienia można podzielić na dwie rozłączne klasy:

- problemy nierozstrzygalne, czyli takie, do rozwiązania których nie istnieją algorytmy, choć pewne przypadki szczególne mogą być rozwiązane, np. przez zgadnięcie,

- problemy algorytmiczne, tzn. takie, które można rozwiązać za pomocą algorytmów.

W dalszym ciągu artykułu skoncentrujemy się wyłącznie na problemach algorytmicznych. Możliwość rozwiązania problemu za pomocą algorytmu, to jeszcze zbyt mało. Należy zapytać o koszty z tym związane. Najczęściej stosowanymi miarami złożoności algorytmów są: złożoność czasowa i złożoność pamięciowa, przy czym ta ostatnia nie będzie istotna dla naszych specyficznych rozważań.

Złożoność czasowa algorytmu jest to czas potrzebny do wykonania algorytmu wyrażony w funkcji rozmiaru problemu. W praktyce złożonościowej istotne są duże

rozmiary problemów, zatem w celu ich scharakteryzowania posługujemy się oszacowaniami asymptotycznymi.

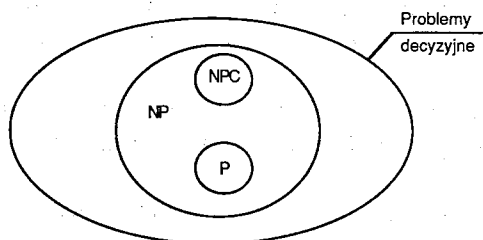
Niech $f, g: \mathcal{N}_0 \rightarrow \mathcal{N}_0$ będą funkcjami, przy czym $\mathcal{N}_0$ oznacza zbiór liczb naturalnych z zerem. Mówimy, że $f = O(g)$, jeśli istnieje stała rzeczywista $c > 0$ i $m \in \mathcal{N}_0$ takie, że $f(n) \leq c \cdot g(n)$ dla wszystkich $n \geq m$. Oszacowanie $f = O(g)$ oznacza określenie ograniczenia dolnego g na wartość funkcji f . Przykładowo, jeśli $f(n) = 5n^2 + 6n - 4$, to $f(n) = O(n^2)$.

Piszemy $f = \Omega(g)$, jeśli $g = O(f)$ — w ten sposób określamy ograniczenie górne na wartości funkcji g . Jeśli $f = O(g)$ i $f = \Omega(g)$, to mówimy, że f jest dokładnie rzędu g i piszemy $f = \Theta(g)$.

Odwołując się jeszcze raz do praktyki obliczeniowej, będziemy rozpatrywać — chyba, że wyraźnie zaznaczono inaczej — złożoność pesymistyczną, tj. złożoność dla najgorszego przypadku, choć niekiedy bardziej istotna może okazać się wartość średnia (oczekiwana) złożoności.

Algorytm o złożoności czasowej wielomianowej, lub w skrócie algorytm *wielomianowy*, to taki algorytm, którego złożoność czasowa jest $O(p(n))$, przy czym $p(n)$ jest wielomianem, a n oznacza rozmiar problemu. Jeśli powyższy warunek nie jest spełniony, to mówimy o *algorytmie wykładniczym*.

Powracając do problemu Eulera, to z przytoczonego twierdzenia wynika, że aby określić rząd złożoności obliczeniowej należy zbadać stopień każdego wierzchołka sprawdzając przy tym, czy każdy z n wierzchołków ma połączenie czy nie z pozostałymi $n-1$ wierzchołkami, przy czym wystarczy wykonać tylko $n(n-1)/2$ takich sprawdzeń. Jeśli przyjąć, że na jedno sprawdzenie zużywa się jedną jednostkę czasu, to złożoność czasowa problemu Eulera $T(n) = O(n^2)$. Jest to zatem problem wielomianowy. Do klasy problemów wielomianowych należą też np. problemy o złożoności $O(n \log n)$, $O(n^3)$ itd., a do klasy problemów wykładniczych np. problemy o złożoności $O(2^n)$ czy $O(n!)$, chociaż z matematycznego punktu widzenia $f(n) = n!$ nie jest funkcją wykładniczą.



Rys. 2. Uproszczona hierarchia problemów decyzyjnych stosowana w zagadnieniach złożonościowych

Należy rozróżnić pomiędzy złożonością algorytmu a złożonością problemu, gdyż bardzo często dany problem można rozwiązać za pomocą więcej niż jednego algorytmu. Uproszczoną hierarchię problemów decyzyjnych stosowaną w zagadnieniach złożonościowych (będącą fragmentem tzw. hierarchii wielomianowej problemów [41] (por. także [3, 15, 19, 38, 49]) przedstawiono na rys. 2. Klasa P, to

problemy o złożoności wielomianowej. Klasa NP zawiera wszystkie problemy z klasy P, czyli $P \subseteq NP$, a także problemy, których status jest mniej pewny: takie problemy, że znane są tylko algorytmy wykładnicze służące do ich rozwiązania. W teorii złożoności obliczeniowej przyjmuje się powszechnie hipotezę, że $P \neq NP$. Klasa NPC problemów NP — zupełnych jest podklasą problemów w pewnym sensie najtrudniejszych obliczeniowo w klasie NP. Znajdują się w niej takie problemy, jak, przykładowo, problem spełnialności wyrażeń boolowskich, problem komiwojażera, upakowania, dopasowania trójwymiarowego, kliki itd. [15, 29, 31]. Gdyby udało się udowodnić, że choć jeden problem z klasy NPC można rozwiązać za pomocą algorytmu wielomianowego, to $P = NP$, co wydaje się mało prawdopodobne.

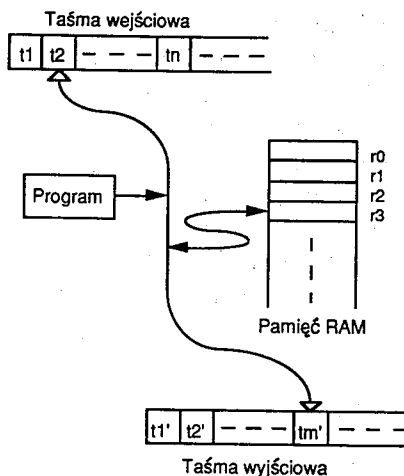
3. MODELOWANIE OBLICZEŃ

Teoria złożoności obliczeniowej za punkt wyjścia bierze modele formalne obliczeń [29, 31]. Innymi słowy, rzeczywiste procesy obliczeniowe są modelowane za pomocą takich narzędzi jak maszyna Turinga, maszyna o dostępie swobodnym, maszyna RASP i wiele innych. Okazuje się, że modele te są zasadne w tym sensie, że uzyskane za ich pomocą wyniki są potwierdzane w praktyce obliczeniowej. Ponadto każdy z nich można zasymulować za pomocą innego kosztem wielomianowej ilości czasu co oznacza, że definicje formalne klasy P i klasy NP są niezależne od przyjętego modelu obliczeń. W celu pełniejszego określenia klas P i NP posłużymy się modelem w postaci maszyny o dostępie swobodnym (w skrócie: maszyna RAM).

Maszyna RAM jest idealizowanym komputerem składającym się z (i) pamięci o dostępie swobodnym (RAM), (ii) pamięci, w której jest przechowywany program pracy maszyny, (iii) taśmy wejściowej, (iv) taśmy wyjściowej (rys. 2). Pamięć RAM składa się z ciągu rejestrów $r_0, r_1, r_2, \dots$. Każdy z nich może przechowywać dowolną liczbę naturalną lub dowolną literę słowa wejściowego. Nie nakłada się żadnego ograniczenia na liczbę rejestrów. Wszystkie obliczenia odbywają się w rejestrze r_0 zwanym akumulatorem (rys. 3).

Dane do maszyny są wprowadzane z potencjalnie nieskończonej taśmy wejściowej. Jest ona podzielona na komórki, w których umieszcza się litery słowa wejściowego po jednym w każdej komórce (w szczególności mogą to być dowolne liczby naturalne). Z taśmy wejściowej można wyłącznie odczytywać dane. Jeśli litera słowa wejściowego została przeczytana, to głowica przesuwą się o jedną komórkę w prawo.

Taśma wyjściowa jest tego samego kształtu i jest obsługiwana na podobnej zasadzie jak taśma wejściowa z tym, że służy wyłącznie do zapisywania wyników. Program jest sekwencją arbitralnie zaetykietowanych instrukcji, wśród których wyróżnia się zwykle instrukcje wejścia/wyjścia, współpracy akumulatora z pamięcią RAM, arytmetyczne i logiczne, skoków oraz instrukcję zatrzymania. Taki model służy do realizacji obliczeń wartości funkcji i do rozpoznawania (akceptowania) języków formalnych. Język formalny L w pewnym alfabecie skończonym A jest to podzbiór zbioru A^* wszystkich słów skończonych w tym alfabecie łącznie ze słowem pustym, $L \subseteq A^*$.



Rys. 3. Schemat maszyny RAM

Z tego co zostało powiedziane wyżej, wynika, że każdemu problemowi decyzyjnemu można przyporządkować pewien język formalny. Jeśli słowo wejściowe $x \in L$ reprezentuje przypadek szczególny problemu, to fakt akceptacji tego słowa oznacza rozwiązanie reprezentowanego przez x przypadku szczególnego badanego problemu.

Będziemy przyjmować, że na wykonanie każdej instrukcji maszyna zużywa jedną jednostkę czasu oraz, że złożoność czasowa $T(n)$ problemu rozwiązywanego za pomocą maszyny RAM jest to maksymalna wartość, brana po wszystkich przypadkach szczególnych o rozmiarze n , z sumy czasów potrzebnych do wykonania wszystkich instrukcji programu. Należy podkreślić, że maszyna RAM jest modelem obliczeń szeregowych. Do modelowania równoległych procesów obliczeniowych można wykorzystać maszynę PRAM składającą się z pewnej liczby, powiedzmy H , równolegle pracujących maszyn RAM (procesorów) dzielących wspólną pamięć. Złożoność czasowa maszyny PRAM jest to maksimum ze złożoności czasowych branych po wszystkich maszynach RAM.

Maszyna PRAM jest więc modelem systemu wieloprocesorowego dzielącego wspólną pamięć. Przypadki szczególne takich systemów są modelowane za pomocą różnych odmian tej maszyny [11, 14, 23, 47]. Jednakże system wieloprocesorowy dzielący wspólną pamięć nie wyczerpuje wszystkich, w szczególności rzeczywistych, sytuacji obliczeniowych. W [22] wyróżnia się synchroniczne systemy równoległe w postaci układów rozproszonych (*distributed*), komputerów łańcuchowych (*pipeline*) i matrycowych (*array*). Do systemów asynchronicznych zalicza się tam komputery sterowane przepływem danych (*dataflow*). Jako szczególnie przydatne do realizacji w strukturach VLSI wymienia się układy systoliczne [27, 30, 45]. W zakresie badania związków istniejących pomiędzy różnymi równoległymi systemami obliczeniowymi z punktu widzenia złożoności obliczeniowej uzyskano pewne znaczące rezultaty (por. np. [24, 28, 32, 44, 46]), tym niemniej w celu otrzymania pełnej charakterystyki konieczne są dalsze badania.

4. OBLICZENIA RÓWNOLEGŁE

Niech para uporządkowana (T, H) będzie miarą złożoności algorytmu równoległego, gdzie T jest czasem obliczeń (złożonością czasową), a H liczbą procesorów. Zarówno T jak i H zależą od rozmiaru problemu.

Algorytm równoległy o złożoności (T_0, H_0) można przekształcić w rodzinę algorytmów o złożoności (T, H) , przy czym $TH = T_0H_0$ dla $H \leq H_0$ [14]. Wynika to stąd, że jeśli dysponujemy tylko H procesorami, to każdy z nich musi symulować H_0/H procesorów, zatem spowolnienie wynosi H_0/H , czyli $T = T_0H_0/H$. Bezpośrednio z powyższego wynika wniosek, że algorytm równoległy o złożoności (T, H) można przekształcić w algorytm szeregowy o złożoności czasowej TH . (Wystarczy przyjąć $H = 1$).

Dla $k \geq 1$ klasę NC^k definiujemy jako klasę języków (odp. problemów) akceptowanych (odp. rozwiązywanych) przez maszynę PRAM w czasie $T(n) = O(\log^k)$ przy wykorzystaniu $H(n) = n^k$ procesorów [36]. Przyjmujemy ponadto, że

$$NC = \bigcup_{k \geq 1} NC^k.$$

Z powyżej zapisanego wniosku o algorytmach równoległych wynika, że $NC \subseteq P$. Nie wiadomo dotąd, czy inkluzja ta jest właściwa. Pewne fakty zdają się wskazywać, że $NC \neq P$ [1, 36, 47]. Oznaczałoby to, że nie wszystkie problemy z klasy P są paralelizowalne w powyższym sensie. Udowodniono, że spora liczba problemów należy do NC . Wśród nich [9, 26]: cztery operacje arytmetyczne $(+, -, *, /)$ wykonywane na liczbach binarnych, operacje macierzowe (mnożenie, odwracanie, obliczanie rzędu, obliczanie wartości wyznacznika), wyznaczanie największych wspólnych dzielników, wielomianów, sortowanie, badanie spójności grafu, wyznaczanie minimalnego drzewa rozpinającego w grafie.

Dodajmy, że klasa NC jest niezależna od przyjętego w definicji modelu obliczeń. W szczególności jest to ta sama klasa, gdy w charakterze modelu obliczeń równoległych wziąć układ kombinacyjny (w tym przypadku procesorami są bramki logiczne).

Pojęcie równoległości wprowadzone w [47] różni się od rozpatrywanego w przypadku klasy NC . Nie wymaga się tutaj, aby czas obliczeń był ograniczony do $O(\log^k)$. Konsekwencją tego może być, iż pewne problemy spoza NC mogą mieć rozwiązania równoległe. Inną korzyścią wynikającą z braku wymagania tak dużego przyspieszenia jest, że liczba procesorów koniecznych w implementacji równoległej algorytmu, zwykle nierozsądnie duża (powiedzmy n, n^2, n^5 itp.) może być zredukowana.

W [47] zdefiniowano dwie klasy problemów. Klasa PC zawiera te problemy należące do P , dla których można uzyskać przyspieszenie o więcej niż współczynnik stały, jeśli stosować równoległość obliczeń. Klasa PC^* jest klasą problemów silniejszych, takich że przyspieszenie jest proporcjonalne do liczby H procesorów.

Jest oczywiste, że $PC^* \subseteq PC$ i $NC \subseteq PC$. Pytanie czy $PC^* \subset PC \subset P$ jest problemem otwartym. Ponadto łatwo zauważyć, że nierówność $NC \neq P$ implikuje $NC \neq PC$.

Klasy NC i PC nie wyczerpują wszystkich znanych dotąd klasyfikacji problemów z punktu widzenia równoległości obliczeń. Istnieją także inne oparte na modelach obliczeń w większym stopniu oddalonych od rzeczywistych architektur komputerów. Jednakże we wszystkich tych modelach z różnych powodów nie uwzględnia się długości przewodów łączących poszczególne procesory, co ma także wpływ na czas przetwarzania. Średnia długość przewodów w funkcji liczby n procesorów dla architektur w postaci grafu zupełnego i kostki wielowymiarowej (*hypercube*) w przypadku, gdy schemat połączeń jest zanurzony w siatce planarnej (dla układu VLSI wykonanego technologią planarną) wynosi, odpowiednio, $\Theta(n^2)$ i $\Theta(n/\log n)$ [1]. Przewody te wnoszą opóźnienie istotnie wpływające na jakość algorytmu.

W procesie projektowania obliczeń równoległych, w szczególności w układach VLSI, można wyróżnić cztery etapy: (i) wybór algorytmu (w postaci acyklicznego grafu skierowanego wskazującego obliczenia elementarne i ich współzależność [8]), (ii) wybór architektury wieloprocessorowej, (iii) uszeregowanie zadań, dzięki któremu algorytm jest wykonywany poprawnie, (iv) określenie złożoności algorytmu. W [33] zaproponowano, aby za miarę złożoności przyjąć głębokość acyklicznego grafu skierowanego (tzn. długość maksymalnej ścieżki wejściowo-wyjściowej grafu), co przy założeniu, że układ jest synchroniczny odpowiada złożoności czasowej tego układu. Zaproponowano także, aby uwzględniać opóźnienie komunikacyjne τ , tj. czas pomiędzy momentem wytworzenia informacji, rozumianej w szerokim sensie, a momentem wykorzystania tej informacji. Opóźnienie komunikacyjne, będące funkcją liczby n procesorów zmienia się w zależności od architektury od n (pierścien), poprzez $\sqrt{n}$ (siatka) do $\log n$ (kostka wielowymiarowa). W przypadku komputerów będących aktualnie w eksploatacji τ zawarte jest w granicach od jednostek do milionów elementarnych kroków [33]. Opóźnienie komunikacyjne jest zatem istotnie zależne od architektury komputera.

Złożoność obliczeniowa algorytmu zależy więc zarówno od architektury systemu (w szczególności od opóźnienia), jak i od uszeregowania elementarnych zadań jakie muszą być wykonane, aby osiągnąć poprawny wynik końcowy.

5. ZŁOŻONOŚĆ KOMUNIKACYJNA PROBLEMÓW A ZŁOŻONOŚĆ UKŁADÓW VLSI

Zagadnienia komunikacji odgrywają istotną rolę w bezpośrednim określaniu złożoności układów VLSI.

Przyjmijmy, że język $L \subseteq \{0,1\}^*$ ma być rozpoznany przez dwa odległe komputery, lub inaczej, że problem ma być rozwiązany za pomocą dwu komputerów. Każdy komputer otrzymuje połowę bitów wejściowych i obliczenie przebiega z zastosowaniem tego samego protokołu komunikacyjnego pomiędzy dwoma komputerami (wiele interesujących języków nie może być rozpoznanych z pominięciem komunikacji). Minimalna liczba bitów jaką należy wymienić w celu rozpoznania zbioru słów zerojedynkowych o parzystej długości, czyli języka $L \cap \{0,1\}^{2n} (n = 1, 2, 3, \dots)$, mini-

malizowana ze względu na wszystkie podziały bitów wejściowych na dwie równe części i rozważana jako funkcja n , jest nazywana *złożonością komunikacyjną języka L* (dla prostoty rozważamy tylko te słowa, które mają długość parzystą). Analogicznie definiuje się złożoność komunikacyjną problemu decyzyjnego.

Taką miarę złożoności języków (a tym samym problemów) zaproponowano w [36]. Motywacją do badania miary złożoności komunikacyjnej jest to, że jest ona ściśle związana z ograniczeniem dolnym na AT^2 układu scalonego rozpoznającego język L (czyli rozwiązującego problem modelowany przez L). W celu określenia złożoności komunikacyjnej przyjmijmy kilka dodatkowych definicji.

Protokołem o $2n$ wejściach nazywamy parę $D_n = (\pi, \Phi)$, gdzie

(1) π jest podziałem zbioru $\{1, 2, \dots, 2n\}$ na dwa różne zbiory S_I i S_{II} . Odpowiada to podziałowi słowa wejściowego na dwie równe części, każda przeznaczona dla innego komputera.

(2) $\Phi: \{0, 1\}^n \times \{0, 1, \$\}^* \rightarrow \{0, 1\}^* \cup \{\text{accept}, \text{reject}\}$ jest funkcją.

Intuicyjnie pierwszy argument funkcji Φ jest częścią lokalną wejścia, a drugi reprezentuje historię (wprowadzanie poprzednich wiadomości), przy czym $\$$ jest separatorem wiadomości.

Wartość funkcji Φ określa następną wiadomość. Dla danego słowa $c \in \{0, 1, \$\}^*$ funkcja Φ ma własność braku przedrostka, tzn., że dla żadnych dwóch $y, y' \in \{0, 1\}^n$ nie zachodzi warunek, że $\Phi(y, c)$ jest właściwym przedrostkiem $\Phi(y', c)$. Ta własność zapewnia, że przy wymianie wiadomości nie jest potrzebny żaden dodatkowy symbol końca transmisji.

Obliczeniem protokołu D_n ze względu na wejście $x \in \{0, 1\}^{2n}$ jest słowo $c = c_1 \$ c_2 \$ \dots \$ c_k \$ c_{k+1}$, gdzie $k \geq 0$, $c_i \in \{0, 1\}^*$ dla $i = 0, 1, \dots, k$, $c_{k+1} \in \{\text{accept}, \text{reject}\}$, takie że dla każdej liczby całkowitej j , $0 \leq j < k$, mamy:

(a) jeśli j jest nieparzyste, to $c_{j+1} = \Phi(x_I, c_1 \$ c_2 \$ \dots \$ c_j)$, gdzie x_I jest wejściem x ograniczonym do zbioru S_I ,

(b) jeśli j jest parzyste, to $c_{j+1} = \Phi(x_{II}, c_1 \$ c_2 \$ \dots \$ c_j)$, gdzie x_{II} jest wejściem x ograniczonym do zbioru S_{II} .

Innymi słowy, współpraca komputerów przebiega w ten sposób, że kolejna wysyłana wiadomość zależy od lokalnego wejścia i wszystkich poprzednich wymian wiadomości. Oczywiście ten proces jest całkowicie deterministyczny. *Długością* obliczenia c jest całkowita długość wiadomości w c z pominięciem symboli separujących $\$$ i końcowych *accept/reject*.

Niech $L \subseteq \{0, 1\}^*$ będzie językiem, a $\Delta = \langle D_n \rangle$ ciągiem protokołów. Mówimy, że Δ rozpoznaje L , jeśli dla każdego n i dla każdego $x \in \{0, 1\}^{2n}$ obliczenie D_n ze względu na wejście x jest zawsze skończone i kończy się symbolem *accept* wtedy i tylko wtedy, gdy $x \in L$.

Niech f będzie funkcją ze zbioru liczb naturalnych w zbiór liczb naturalnych. Mówimy, że język L jest *rozpoznawalny z komunikacją f* , co zapisujemy w postaci $L \in \text{COMM}[f]$, jeśli istnieje ciąg protokołów $\Delta = \langle D_n \rangle$ taki, że dla wszystkich n i $x \in \{0, 1\}^{2n}$ obliczenie D_n ze względu na x ma długość co najwyżej $f(n)$.

Dla każdego języka $L = \{0, 1\}^*$ można zdefiniować nieskończony ciąg funkcji

boolowskich $\langle h_n(L) \rangle$ (gdzie dla każdego n , $h_n(L): \{0, 1\}^n \rightarrow \{0, 1\}$) taki, że

$$x = x_1 x_2 \dots x_n \in L \cap \{0, 1\}^n \Leftrightarrow h_n(x_1, x_2, \dots, x_n) = 1.$$

Złożonością komunikacyjną $C(\pi, \Phi)$ protokołu $D_n = (\pi, \Phi)$ nazywamy maksimum, brane po długościach, ze wszystkich obliczeń D_n . Przyjmijmy ponadto, że

$$C(h) = \min \{C(\pi, \Phi) \mid \text{protokół } (\pi, \Phi) \text{ oblicza funkcję boolowską } h\}.$$

W charakterze przykładu weźmy wyniki uzyskane w [16]. Niech G będzie grafem o n wierzchołkach, a ponadto niech $f_1(n)$ oznacza funkcję boolowską odpowiadającą problemowi spójności grafu, $f_2(n)$ funkcję boolowską odpowiadającą problemowi s - t -spójności grafu (pytamy, czy pomiędzy wierzchołkami s i t istnieje ścieżka), a $f_3(n)$ funkcję boolowską odpowiadającą problemowi dwudzielności grafu. Wówczas

$$C(f_1) = C(f_2) = C(f_3) = \Theta(n \log n).$$

Thompson [43] wykazał, że dla dowolnej funkcji boolowskiej h

$$AT^2 = \Omega(C^2(h)).$$

Wynik ten jest bezpośrednim łącznikiem pomiędzy zagadnieniami złożoności komunikacyjnej a złożonością układów VLSI wykonywanych na płycie krzemowej.

Hromkovič [21] zauważył, że określanie złożoności zgodnie z [36] charakteryzuje się dwiema wadami: (i) w pewnych przypadkach uzyskuje się „nierozsadne” ograniczenie dolne na AT^2 (istnieją problemy bardzo trudne, które mogą być rozwiązane ze złożonością komunikacyjną zerową), (ii) dla pewnych specyficznych problemów trudno jest udowodnić nietrywialne ograniczenie dolne na złożoność komunikacyjną. Złożoność S -komunikacyjna usuwa obydwie wymienione wady oryginalnej złożoności komunikacyjnej.

Niech A będzie problemem ze zbiorem zmiennych wejściowych X i niech $Y \subseteq X$. Złożoność komunikacyjną zgodną z Y określamy jako minimum po wszystkich podziałach zbioru X dzielących Y na dwie równe części. Oznacza to, że podział zmiennych wejściowych zbioru $X \setminus Y$ jest dowolny (np. wszystkie zmienne ze zbioru $X \setminus Y$ mogą być przekazane do jednego komputera). Jeśli porównać oryginalną koncepcję teorii VLSI u Thompsona [43], to nie ma wątpliwości, że dla każdego podzbioru Y zbioru X kwadrat złożoności komunikacyjnej zgodnej z Y daje ograniczenie dolne na AT^2 .

Niech $\mathcal{N} = \mathcal{N}_0 \setminus \{0\}$, a $|Z|$ niech będzie liczbą elementów zbioru Z . Ponadto niech $h: \{0, 1\}^n \rightarrow \{0, 1\}$ będzie funkcją boolowską n zmiennych $x_1, x_2, \dots, x_n$ i niech $X = \{x_1, x_2, \dots, x_n\}$.

Protokół o n wejściach z X jest uporządkowaną trójką $D_n = (Y, \pi, \Phi)$, gdzie:

(1) π jest podziałem zbioru $X = \{x_1, x_2, \dots, x_n\}$ na dwa różne zbiory S_I i S_{II} takie, że $|S_I| = m$, $|S_{II}| = k$ dla pewnych $m, k \in \mathcal{N}$ $m + k = n$; S_I zawiera r zmiennych ze zbioru Y , a S_{II} zawiera s zmiennych z Y . (Intuicyjnie rzecz biorąc S_I i S_{II} zawierają zmienne wejściowe, odpowiednio, pierwszego i drugiego komputera),

(2) $\Phi: \{0, 1\}^m \times \{0, 1, \$\}^* \cup \{0, 1\}^k \times \{0, 1, \$\}^* \rightarrow \{0, 1\}^+ \cup \{\text{accept, reject}\}$ jest funkcją

z własnością braku przedrostka, przy czym $\{0, 1\}^+$ jest zbiorem wszystkich skończonych ciągów zerojedynkowych bez ciągu pustego.

Obliczenie protokołu D_n i długość obliczenia są zdefiniowane analogicznie jak w przypadku rozpatrywanym powyżej. W taki sam sposób określamy też protokół obliczający funkcję boolowską.

Złożoność komunikacyjna protokołu $D_n = (Y, \pi, \Phi)$, oznaczana jako $SC(Y, \pi, \Phi)$, jest równa maksimum z długości wszystkich obliczeń tego protokołu.

Niech

$$SC(h, Y) = \min \{SC(Y, \pi, \Phi) \mid \text{protokół } (Y, \pi, \Phi) \text{ oblicza } h\}.$$

Wówczas złożoność S-komunikacyjną funkcji boolowskiej h określamy jako

$$SC(h) = \max \{SC(h, Y) \mid Y \subseteq X\}.$$

Oczywiście pomiędzy złożonością komunikacyjną określoną w [36] a złożonością S-komunikacyjną zachodzi zależność

$$C(h) = SC(h, X).$$

Dla funkcji $f: \mathcal{N}_0 \rightarrow \mathcal{N}_0$, $SCOMM[f]$ jest klasą problemów (odp. języków), które mogą być rozwiązane (odp. rozpoznane) za pomocą ciągu protokołów ze złożonością S-komunikacyjną f . Bezpośrednio z przytoczonych powyżej określeń wynika, że dla każdej funkcji boolowskiej h jest spełniony warunek

$$C(h) \leq SC(h),$$

co oznacza, że dla każdej funkcji $f: \mathcal{N}_0 \rightarrow \mathcal{N}_0$

$$SCOMM[f] \subseteq COMM[f].$$

Szczegółową charakterystykę złożoności S-komunikacyjnej zawiera praca [21]. Zatrzymajmy się jeszcze przez chwilę nad dwoma ważnymi wynikami tam zawartymi.

Sformułujmy następujący problem decyzyjny: czy dane wyrażenie boolowskie F ma niezerową złożoność S-komunikacyjną? Okazuje się, że jest to problem NP-zupełny.

Drugi wynik dotyczy złożoności układu VLSI obliczającego wartości funkcji boolowskiej. Niech $m \leq n$, $m, n \in \mathcal{N}$ a ponadto niech f będzie funkcją boolowską n zmiennych, której wartości zależą dokładnie od wartości m zmiennych. Wówczas dla każdego układu VLSI obliczającego f

$$AT^2 \geq (m/2)(\log_2 m - 2\log_2(\log_2 m)).$$

Omówione powyżej protokoły są protokołami deterministycznymi. Oprócz nich ważne są też protokoły deterministyczne o k nawrotach, protokoły niedeterministyczne i probabilistyczne [17, 21, 36]. Uogólnienia deterministycznych protokołów komunikacyjnych na sieci rozproszone komputerów dokonano w [44]. Hierarchie złożoności komunikacyjnej badano w [20]. W [12] przedstawiono możliwości zastosowania protokołów probabilistycznych w obliczeniach VLSI.

Wszystkie wyniki dotyczące miary złożoności układu VLSI w postaci iloczynu AT^2 uzyskano przy założeniu stałego wymiaru charakterystycznego λ układu. W pracy [7]

zaprezentowano inne podejście. Za miarę złożoności przyjęto T/λ^2 przy stałej powierzchni układu A .

Uzasadnieniem takiego podejścia jest fakt, że w ostatnich latach obserwuje się wyraźne zmniejszenie wymiaru charakterystycznego λ , a ponadto ze względu na możliwości technologii waflowej (WSI) powierzchnia A wydaje się być mniej istotna. Wykazano także zasadę, w myśl której przejście od modelu przy stałym λ do modelu przy stałym A dokonuje się poprzez zastąpienie A przez $1/\lambda^2$ oraz T przez T/λ .

6. UKŁADY KOMBINACYJNE I HIERARCHIA BOOLOWSKA

W tym rozdziale będziemy rozpatrywać układy kombinacyjne (bez sprzężeń zwrotnych) o n wejściach i m wyjściach. Układ kombinacyjny dobrze nadaje się do modelowania procesów obliczeniowych. Wspomniano już o tym przy omawianiu klasy NC problemów. (Związki z innymi modelami badano m.in. w pracach [34, 37, 42]). Możliwości realizacji technicznych są oczywiste.

Rozmiarem s układu nazywamy liczbę jego bramek, a głębokością d układu — maksymalną liczbę bramek braną po wszystkich ścieżkach wejściowo-wyjściowych. Jeśli przyjąć, że każda bramka przetwarzając sygnał wnosi jednostkowe opóźnienie, to (zauważono to już w rozdziale 4) głębokość układu określa jego złożoność czasową.

Dla danej funkcji boolowskiej $h: \{0, 1\}^n \rightarrow \{0, 1\}$ rozmiar układu obliczającego wartości tej funkcji definiujemy jako

$$S(h) = \min \{s \mid \text{istnieje układ kombinacyjny o jednym wyjściu obliczający wartości funkcji } h\},$$

a głębokość takiego układu jako

$$D(h) = \min \{d \mid \text{istnieje układ kombinacyjny o jednym wyjściu obliczający wartości funkcji } h\}.$$

Analogicznie określa się rozmiar układu (m wyjściowego) dla rodziny funkcji boolowskich.

W 1958 r. Łupanov wykazał (por. np. [18, 38]), że dla każdej n argumentowej funkcji boolowskiej h

$$S(h) \leq (1 + O(1/\sqrt{n})) \cdot 2^n/n.$$

W [35] udowodniono, że dla każdej funkcji boolowskiej h

$$D(h) \geq 0,25 \cdot S(h) \cdot \log_2 S(h) - O(S(h)).$$

Oczywiście w przypadku projektowania układu VLSI chodzi o to, aby uzyskać układ kombinacyjny o minimalnej głębokości, a tym samym charakteryzujący się dużymi możliwościami obliczeń równoległych [8].

Definicje rozmiaru rozszerzymy teraz na języki w alfabecie $\{0, 1\}$. Niech $L \subseteq \{0, 1\}^*$. Dla $n \geq 0$ niech $\chi_{n,L}: \{0, 1\}^n \rightarrow \{0, 1\}$ będzie funkcją taką, że

$$\chi_{n,L}(x) = \begin{cases} 1, & \text{jeśli } x \in L \\ 0, & \text{jeśli } x \notin L. \end{cases}$$

Wówczas rozmiar układu akceptującego język $L \subseteq \{0, 1\}^*$ określamy następująco:

$$S_L(n) = S(\chi_n, L).$$

Mówimy, że język L ma rozmiar wielomianowy, jeśli istnieje wielomian p taki, że $S_L(n) \leq p(n)$.

Każdy język w klasie P ma rozmiar wielomianowy [38], co oznacza, że każdy problem z klasy P można rozwiązać za pomocą „niewielkiego” układu kombinacyjnego, tj. układu o rozmiarze ograniczonym przez wielomian.

Niech $\langle h_i \rangle$ będzie rodziną funkcji boolowskich. Przez $P/poly$ oznaczamy zbiór rodzin $\langle h_i \rangle$ zawierających wszystkie układy o rozmiarach wielomianowych. Badaniom klasy $P/poly$ poświęcono wiele uwagi (por. np. [4, 19, 25, 38]) próbując znaleźć jej powiązania z hierarchią wielomianową problemów. Jednakże związki np. z klasą NP nie są do końca wyjaśnione.

Biorąc za punkt wyjścia klasę NP definiuje się hierarchię boolowską [48]. Niech $L \subseteq \Sigma^*$ będzie językiem w alfabecie Σ . Przez $\bar{L}$ oznaczamy dopełnienie języka L , czyli $\bar{L} = \Sigma^* \setminus L$. Każdy poziom hierarchii boolowskiej składa się ze zbiorów reprezentowanych przez pewną ustaloną strukturę operatorów boolowskich na zbiorach NP :

$$NP(0) = P,$$

$$NP(1) = NP,$$

$$NP(2) = \{L_1 \cap \bar{L}_2 \mid L_1, L_2 \in NP\},$$

$$NP(3) = \{(L_1 \cap \bar{L}_2) \cup L_3 \mid L_1, L_2, L_3 \in NP\},$$

$$NP(4) = \{((L_1 \cap \bar{L}_2) \cup L_3) \cap \bar{L}_4 \mid L_1, L_2, L_3, L_4 \in NP\},$$

$$\vdots$$

Hierarchię boolowską definiujemy jako

$$BH = \bigcup_{k \geq 1} NP(k).$$

Hierarchia BH ma następującą interpretację fizyczną: zbiory w tej hierarchii reprezentują problemy rozwiązywalne sprzętowo, przy czym reprezentacja ta polega na domknięciu klasy NP ze względu na operatory boolowskie.

Dokonyamy teraz krótkiej algebraicznej charakterystyki klas $NP(n)$ za pomocą trzech operatorów. Niech L_i , $i = 1, 2, \dots, k$, będą językami. Wówczas operatory C , D , E określamy następująco:

— operator C

$$C(L_1) = L_1,$$

a dla $k > 1$

$$C(L_1, \dots, L_k) = \begin{cases} C(L_1, \dots, L_{k-1}) \cup L_k & \text{dla } k \text{ parzystego} \\ C(L_1, \dots, L_{k-1}) \cap \overline{L_k} & \text{dla } k \text{ nieparzystego} \end{cases}$$

— operator D

$$D(L_1) = L_1,$$

a dla $k > 1$

$$D(L_1, \dots, L_k) = L_1 \setminus D(L_2, \dots, L_k),$$

— operator E

$$E(L_1) = L_1,$$

$$E(L_1, L_2) = L_1 \setminus L_2,$$

a dla $k > 2$

$$E(L_1, \dots, L_k) = \begin{cases} E(L_1, \dots, L_{k-1}) \cup L_k & \text{dla } k \text{ nieparzystego} \\ E(L_1, \dots, L_{k-1}) \cup (L_{k-1} \setminus L_k) & \text{dla } k \text{ parzystego} \end{cases}$$

W [48] wykazano, że

$$\begin{aligned} \text{NP}(n) &= \{C(L_1, \dots, L_n) \mid L_i \in \text{NP}, i = 1, 2, \dots, n\} = \\ &= \{C(L_1, \dots, L_n) \mid L_1 \supseteq \dots \supseteq L_n, L_i \in \text{NP}, i = 1, 2, \dots, n\} = \\ &= \{D(L_1, \dots, L_n) \mid L_i \in \text{NP}, i = 1, 2, \dots, n\} = \\ &= \{D(L_1, \dots, L_n) \mid L_1 \supseteq \dots \supseteq L_n, L_i \in \text{NP}, i = 1, 2, \dots, n\} = \\ &= \{E(L_1, \dots, L_n) \mid L_i \in \text{NP}, i = 1, 2, \dots, n\} = \\ &= \{E(L_1, \dots, L_n) \mid L_1 \supseteq \dots \supseteq L_n, L_i \in \text{NP}, i = 1, 2, \dots, n\}. \end{aligned}$$

Powyższe postaci normalne są analogiczne do postaci normalnych Hausdorffa, w szczególności do znanych z teorii układów kombinacyjnych postaci normalnych wyrażeń boolowskich.

Wyniki uzyskane dotąd w zakresie badania hierarchii boolowskiej (obszerne ich zestawienie zawierają prace [5, 6]) przedstawiają w nowym świetle zagadnienia złożonościowe, w szczególności ze względu na możliwości rozwiązywania problemów za pomocą układów kombinacyjnych.

PODSUMOWANIE

W odniesieniu do układów cyfrowych VLSI problematyka złożonościowa zaczęła się rozwijać przed około dziesięć laty. W niniejszej pracy naszkicowano główne jej nurty: klasyfikację problemów ze szczególnym uwzględnieniem algorytmów równoległych rozwiązujących zadane problemy, złożoność komunikacyjną i jej związki ze złożonością układów VLSI, zagadnienia rozwiązywania problemów za pomocą układów kombinacyjnych w powiązaniu z hierarchią boolowską. Mają one już dzisiaj bardzo istotne znaczenie w praktyce projektowania układów VLSI, zarówno wykorzystującej metody tradycyjne, wspomaganie komputerowe, jak i prowadzonej automatycznie za pomocą kompilatorów krzemowych. Można tu wymienić projektowanie algorytmów, projektowanie logiczne, rozmieszczanie elementów i prowadzenie połączeń między nimi, zagadnienia testowania [8, 13, 45].

Wiele zagadnień złożonościowych istotnych w projektowaniu układów VLSI wymaga dalszych, pogłębionych badań. Odnosi się to szczególnie do projektowania algorytmów równoległych, klasyfikacji problemów ze względu na możliwości wykonywania obliczeń równoległych oraz hierarchii boolowskiej.

BIBLIOGRAFIA

1. A. Bertoni, M. Goldwurm, G. Mauri, N. Sabadini: *Parallel algorithms and the classification of problems*. Becker J. D., Eisele I. (eds.), WOPLOT 86 — Parallel Processing: Logic, Organization, and Technology. LNCS 253, Springer, Berlin, 1987, 206—226
2. P. Boas van Emde: *The second machine class 2, an encyclopedic view on the parallel computation thesis*. Mirkowska G., Rasiowa H. (eds.), Mathematical Problems in Computation Theory. Banach Center Publications, vol. 24, PWN, Warszawa, 1988, 235—256
3. R. V. Book: *Towards a theory of relativizations: positive relativizations*. Brandenburg F. J., Vidal-Naquet G., Wirsig M. (eds.), STACS 87 — Proceedings. LNCS 247, Springer, Berlin, 1987
4. R. B. Boppana, J. C. Lagarias: *One-way functions and circuit complexity*. Selman A. L. (ed.), Structure in Complexity Theory, LNCS 223, Springer, Berlin, 1986, 51—65
5. J. Y. Cai, T. Gundermann, J. Hartmanis, L. A. Hemachandra, V. Sewelson, K. Wagner, G. Wechsung: *The boolean hierarchy I: Structural properties*. SIAM Journal on Computing, 1988, vol. 17, No 6, 1232—1252
6. J. Y. Cai, T. Gundermann, J. Hartmanis, L. A. Hemachandra, V. Sewelson, K. Wagner, G. Wechsung: *The boolean hierarchy II: Applications*. SIAM Journal on Computing, 1989 vol. 18, No. 1, 95—111
7. H. C. Card, G. Gulak, R. D. McLoad, W. Pries: (λ, T) complexity measures for VLSI computations in constant chip area. IEEE Transactions on Computers, 1987, vol. C-36, No. 1, 112—117
8. K. Chmiel, J. Stokłosa: *Algorytmy kompilacji krzemowej*. Studia z Automatyki, t. XV, PWN, Warszawa-Poznań, 1989, 5—46
9. S. A. Cook: *An overview of computational complexity*. Communications of the ACM, 1983, vol. 26, No. 6, 400—408
10. N. Deo: *Teoria grafów i jej zastosowania w technice i informatyce*. PWN, Warszawa, 1980.
11. F. E. Fich, P. Rade, A. Wigderson: *Relations between concurrent-write models of parallel computation*. SIAM Journal on Computing, 1988, vol. 17, No. 3, 606—627
12. M. Fürer: *Universal hashing in VLSI*. Reif J. H. (ed.), VLSI Algorithms and Architectures. LNCS 319, Springer, Berlin, 1988, 312—318
13. D. D. Gajski: *Silicon Compilation*. Addison-Wesley, Reading, MA, 1988
14. Z. Galil: *Optimal parallel algorithms*. Bertolazzi P., Luccio F. (eds.), VSLI: Algorithms and Architectures. Elsevier, New York, 1985, 3—11
15. M. R. Garey, D. S. Johnson: *Computers and Intractability. A Guide to the Theory of NP-Completeness*. Freeman, San Francisco, 1979
16. A. Hajnal, W. Maass, G. Turán: *On the communication complexity of graph properties*. Proceedings of the 20th Annual ACM Symposium on Theory of Computing. ACM Press, New York, 1988, 196—191
17. B. Halstenberg, R. Reischuk: *On different modes of communication*. Proceedings of the 20th Annual ACM Symposium on Theory of Computing, ACM Press, New York, 1988, 162—172
18. M. A. Harrison: *Wstęp do teorii sieci przełączających i teorii automatów*. Warszawa, PWN, 1973
19. L. A. Hemachandra: *Counting in structural complexity theory* Ph. D. Thesis, TR-87-840, Department of Computer Science, Cornell University, Ithaca, NY, 1987
20. J. Hromkovič: *Communication complexity hierarchy*. Theoretical Computer Science, 1986, vol. 48, 109—115
21. J. Hromkovič: *The advantages of a new approach to defining the communication complexity for VLSI*. Theoretical Computer Science, 57, 1988, 97—111

22. K. Hwang, F.A. Briggs: *Computer Architecture and Parallel Processing*. McGraw-Hill, New York, 1984 (wyd. II — 1986).
23. I. Immerman: *Expressibility and parallel complexity*. SIAM Journal on Computing, 1989, vol. 18, No. 3, 625—638
24. A.R. Karlin, E. Upfal: *Parallel hashing: an efficient implementation of shared memory*. Journal of the ACM, 1988, vol. 35, No. 4, 876—892
25. R.M. Karp, R. Lipton: *Some connections between nonuniform and uniform complexity classes*. Proceedings of the 12th Annual ACM Symposium on Theory of Computing, ACM Press, New York, 1980, 302—309
26. C.P. Kruskal, L. Rudolph, M. Snir: *Efficient parallel algorithms for graph problems*. Proceedings of the International Conference on Parallel Processing, Computer Society Press, Washington DC, 1986, 869—876
27. H.T. Kung: *The structure of parallel algorithms*. Advances in Computers, vol. 19. Academic Press, New York, 1980, 65—112
28. C.E. Leiserson, B.M. Maggs: *Communication-efficient parallel graph algorithms*. Proceedings of the International Conference on Parallel Processing. Computer Society Press, Washington, DC, 1986, 861—868
29. M. Machtey, P. Young: *An Introduction to the General Theory of Algorithms*. Elsevier, New York, 1978
30. C. Mead, L. Conway: *Introduction to VLSI Systems*. Addison-Wesley, Reading, MA, 1980
31. B. Mikołajczak, J. Stokłosa: *Złożoność obliczeniowa algorytmów*. Wydawnictwo Politechniki Poznańskiej, Poznań, 1984 (wyd. II — 1986)
32. A. Moitra, S.S. Iyengar: *Discussion of parallel algorithms*. TR-86-759, Department of Computer Science, Cornell University, Ithaca, New York, 1986
33. C.H. Papadimitriou, M. Yannakakis: *Towards an architecture-independent analysis of parallel algorithms*. Proceedings of the 20th Annual ACM Symposium on Theory of Computing, ACM Press, New York, 1988, 510—513
34. I. Parberry: *An improved simulation of space and reversal bounded deterministic Turing machines by width and depth bounded uniform circuits*. Information Processing Letters, 1987, vol. 24, 363—367
35. M.S. Paterson, L.G. Valiant: *Circuit size is nonlinear in depth*. Theoretical Computer Science, 1976, vol. 2, 397—400
36. N. Pippenger: *On simultaneous resource bounds*. Proceedings of the 20th IEEE FOCS, IEEE, New York, 1979, 307—311
37. N. Pippenger, M.J. Fischer: *Relations among complexity measures*. Journal of the ACM, 1979, vol. 26, No. 2, 361—381
38. U. Schöning: *Complexity and Structure*. LNCS 211, Springer, Berlin, 1986
39. J. Stokłosa: *Automaty i języki*. Mikołajczak B. (red.), Algebraiczna i strukturalna teoria automatów. PWN, Warszawa-Łódź, 1984, 22—47
40. J. Stokłosa: *Zadania ze złożoności obliczeniowej algorytmów*. Wydawnictwo Politechniki Poznańskiej, Poznań, 1989
41. L.J. Stockmeyer: *The polynomial-time hierarchy*. Theoretical Computer Science, 1977, vol. 3, 1—22
42. L. Stockmeyer, U. Vishkin: *Simulation of parallel random machines by circuits*. SIAM Journal on Computing, 1984, vol. 13, No. 2, 409—422
43. C.D. Thompson: *Area-time complexity for VLSI*. Proceedings of the 11th Annual ACM Symposium on Theory of Computing. ACM Press, New York, 1979, 81—88; także w: Journal of Computer and System Sciences (w druku)
44. P. Tiwari: *Lower bounds on communication complexity in distributed computer networks*. Journal of the ACM, 1987, vol. 34, No. 4, 921—936
45. D.J. Ullman: *Computational Aspects of VLSI*. Computer Science Press, Rockville, MA, 1984
46. E. Upfal, A. Wigderson: *How to share memory in distributed systems*. Journal of the ACM, vol. 34, No. 1, 1987, 116—127

47. J. S. Vitter, R. A. Simons: *New class of parallel complexity: a study of unification and other complete problems for P*. IEEE Transactions on Computers, 1986, vol. C-35, No. 5, 403—418
48. G. Wechsung: *On the boolean closure of NP*. Budach L. (ed.), Fundamentals of Computation Theory. LNCS 199, Springer, Berlin, 1985, 485—493
49. C. Wrathal: *Complete sets and the polynomial-time hierarchy*. Theoretical Computer Science, 1977, vol. 3, 23—33

J. STOKŁOSA

COMPLEXITY OF VLSI DIGITAL SYSTEMS

Summary

In the paper the main trends in the design of VLSI digital systems from the complexity point of view are presented. The basic notions dealing with computational complexity are introduced. Possibility of parallelization of algorithms and stages in the design of parallel algorithms are presented. Communication and S-communication complexities are discussed and relations with the AT^2 complexity of VLSI systems are outlined. Some problems of boolean circuits in the context of the boolean hierarchy are examined.

Optoelektroniczny przetwornik względnej różnicy częstotliwości

JANUSZ MINDYKOWSKI

*Instytut Elektroenergetyki Okrętowej
Wyższa Szkoła Morska w Gdyni*

Otrzymano 1990.05.16

Autoryzowano do druku 1990.06.12

W artykule przedstawiono rozwiązanie nowego optoelektronicznego przetwornika względnej różnicy częstotliwości, szczególnie przydatnego dla potrzeb synchronizacji generatora z pracującą siecią niesztynną, np. w warunkach okrętowych. Przedstawiono podstawy teoretyczne rozważanego przetwornika, wyprowadzono jego charakterystykę przetwarzania, to znaczy zależność sygnału wyjściowego od względnej różnicy częstotliwości sieci i synchronizowanego z nią generatora. Przeprowadzono analizę teoretyczną stałej przetwornika $\mathcal{K}$ oraz błędu nieliniowości $\delta_{\mathcal{K}}$ w zależności od częstotliwości sprowadzonej η i względnej różnicy dwóch częstotliwości σ . Oszacowano główne właściwości metrologiczne. Sformułowano warunki konstrukcyjne dla osiągnięcia określonych parametrów technicznych.

1. WSTĘP

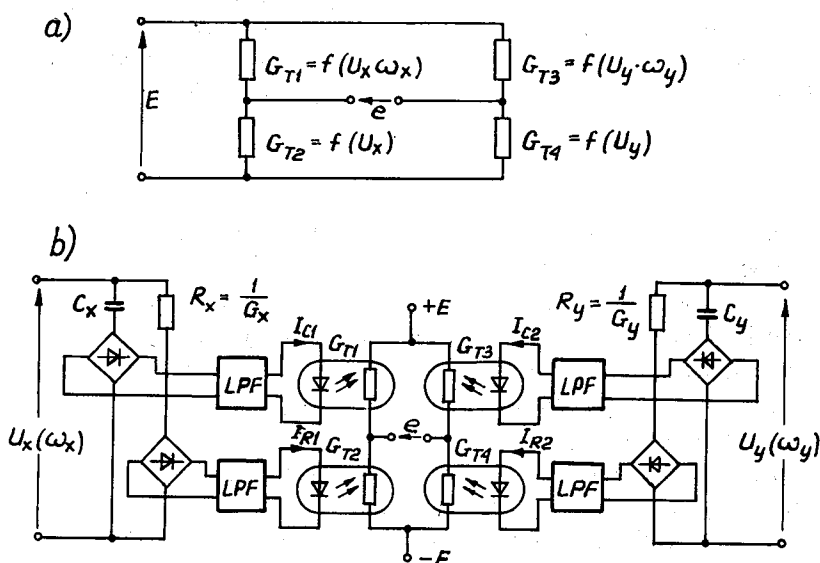
W „Rozprawach Elektrotechnicznych” 4/1989 [6] przedstawiono zagadnienie pomiaru odchyłki częstotliwości, rozumianej jako różnica między wartością częstotliwości mierzonej i jej wartością powinna. Podano koncepcje stosowanych przetworników optoelektronicznych, prezentowanych na tle znanych rozwiązań. Podkreślono, że ważną informacją dla użytkownika może być albo odchyłka częstotliwości od wartości znamionowej (czyli stanowiącej cechę konstrukcyjną danego urządzenia), albo względna różnica dwóch częstotliwości. Pierwsze z zagadnień wiąże się z oceną stanu pracy urządzeń technicznych, natomiast drugie dotyczy synchronizacji generatora z pracującą siecią niesztynną. Rozwiązania omawiane w pracach [2, 6] dotyczą przypadku, gdy częstotliwość znamionowa ma wartość stałą. Przedstawiony tam transoptorowy przetwornik odchyłki częstotliwości od wartości znamionowej nie jest optymalnym przyrządem dla przeprowadzenia prądniczy z siecią okrętową. Sieć ta wykazuje bowiem istotne odchylenia częstotliwości od jej wartości znamionowej, a o przebiegu synchronizacji decyduje rzeczywista wartość względnej różnicy częstotliwości sieci i generatora. Nasuwa to myśl opracowania nowego typu przetwornika [3, 4, 7] którego sygnał wyjściowy byłby właśnie miarą tej różnicy.

Niniejsza praca prezentuje wstępną koncepcję takiego przetwornika, jego właściwości metrologiczne i wymagania konstrukcyjne.

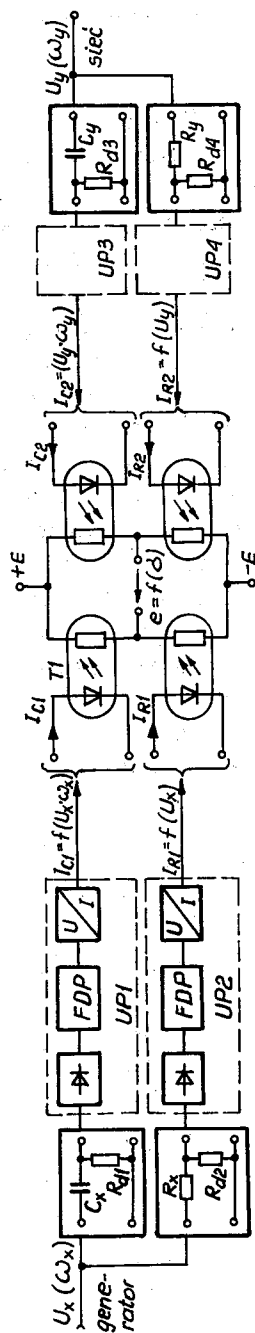
2. TEORETYCZNE PODSTAWY DZIAŁANIA TORÓW POMIAROWYCH NOWEGO PRZETWORNIKA WZGLĘDNEJ RÓŻNICY CZĘSTOTLIWOŚCI

Sygnal wyjściowy przetwornika powinien być miarą względnej różnicy między częstotliwością f_x napięcia U_x wytwarzanego przez synchronizowaną prądnicę — a częstotliwością f_y napięcia U_y sieci nieszytywnej. Przetwornik ma mieć charakterystykę zbliżoną do proporcjonalnej. W dalszych rozważaniach będą brane pod uwagę układy analogowe, wykorzystujące znane transoptorowe elementy przetwarzające typu „ $I/\Delta G$ ” [5].

Koncepcja transoptorowego przetwornika względnej różnicy częstotliwości przedstawiona w publikacjach [4, 7] bazuje na wykorzystaniu elementów tego typu zasilanych poprzez układy prostownicze i filtry dolnoprzepustowe. Transoptory są sterowane wyprostowanymi prądami, z których I_{C1} oraz I_{C2} mają charakter pojemnościowy, zaś I_{R1} oraz I_{R2} — rezystancyjny; ilustruje to rys. 1.



Rys. 1. Koncepcja transoptorowego przetwornika względnej różnicy częstotliwości, a) idea rozwiązania, b) schemat układu; G_{T1} , G_{T2} , G_{T3} , G_{T4} — konduktancje odpowiednich transoptorów, e — sygnał wyjściowy przetwornika, σ — względna różnica częstotliwości f_x i f_y , FDP — filtry dolnoprzepustowe. R_x , R_y , C_x , C_y — rezystancje i pojemności wejściowe w gałęziach prądu odpowiednio I_{R1} , I_{R2} , I_{C1} , I_{C2} , $U_x(\omega_x)$ — napięcie synchronizowanego generatora, $U_y(\omega_y)$ — napięcie sieci okretowej, E — stabilizowane napięcie pomocnicze zasilające mostek



Rys. 2. Zmodyfikowana koncepcja przetwornika względnej różnicy częstotliwości, T1, T2, T3, T4 — transistory, e — sygnał wyjściowy przetwornika, σ — względna różnica częstotliwości ω_x i ω_y , $U_x(\omega_x)$ — napięcie synchronizowanego generatora, $U_y(\omega_y)$ — napięcie sieci, $C_x, C_y, R_{d1}, R_{d2}, R_{d3}, R_{d4}$ — elementy kształtujące impedancję wejściową układów UZ; UZ1, UZ2, UZ3, UZ4 — układy zasilające odpowiednich transoptorów obejmujące prostowniki, filtry dolnoprzepustowe i przetworniki typu „napięcie na prąd”

Analiza teoretyczna przeprowadzona w pracach [4, 7] wskazuje, iż możliwe jest uzyskanie praktycznie liniowego przetwornika względnej różnicy częstotliwości. Uzyskuje się to pod warunkiem, że przy opisanym sposobie zasilania charakterystyki przetwarzania dla transoptorów są praktycznie liniowe. Muszą tu być spełnione pewne wymagania konstrukcyjne.

Wstępne badania [6] udowodniły, że dla osiągnięcia stawianych celów konieczne jest zasilanie transoptorów poprzez bardziej rozbudowane układy, zawierające prostownik liniowy, aktywny filtr dolnoprzepustowy oraz przetwornik typu „napięcie na prąd”. Opis takich układów, ich parametry i wyniki badań doświadczalnych przedstawiono we wspomnianej pracy. Przeprowadzone doświadczenia potwierdziły zasadniczo liniowy charakter zależności $G_T = f(I_s)$, gdzie G_T jest konduktancją wyjściową transoptora, a I_s prądem wyjściowym wyżej opisanego układu zasilającego. Badania te wykazały również, iż konieczne jest zmniejszenie impedencji wejściowej układu przetwarzającego w torze zasilania transoptorów, przy zachowaniu możliwie małego poboru mocy z obwodu badanego. Doprowadziło to do zastosowania członów wejściowych, kształtujących charakter prądu danej gałęzi (jako pojemnościowy lub rezystancyjny), odpowiednio czwórników typu $C-R_d$ lub $G-R_d$. Rys. 2 przedstawia schemat zmodyfikowanej koncepcji transoptorowego przetwornika względnej różnicy częstotliwości, gdzie do zasilania transoptorów wykorzystano układy UZ zawierające prostowniki liniowe, aktywne filtry dolnoprzepustowe i przetworniki typu „napięcie na prąd”, zamiast rozważanych wstępnie mostków Greatza z filtrami pasywnymi.

Szczegółowe rozważania na temat układów zasilających o charakterze pojemnościowym względnie rezystancyjnym podane są w pracy [6].

Zgodnie z przedstawionymi tam równaniami (27) i (30), prądy sterujące transoptory T1, T2, T3 i T4 można opisać zależnościami:

$$\left. \begin{aligned} I_{C1} &= \frac{K_{UZ1} U_x \omega_x R_{d1} C_x}{\sqrt{(\omega_x R_{d1} C_x)^2 + 1}} \\ I_{R1} &= K_{UZ2} \frac{q_x}{q_x + 1} U_x \\ I_{C2} &= \frac{K_{UZ3} U_y \omega_y R_{d3} C_y}{\sqrt{(\omega_y R_{d3} C_y)^2 + 1}} \\ I_{R2} &= K_{UZ4} \frac{q_y}{q_y + 1} U_y \end{aligned} \right\} \quad (1)$$

Znaczenie symboli we wzorach (1) jest następujące:

$K_{UZ1}, K_{UZ2}, K_{UZ3}, K_{UZ4}$ — współczynniki wzmocnienia odpowiednich układów zasilających

U_x — napięcie synchronizowanego generatora o pulsacji ω_x

U_y — napięcie sieci okrętowej o pulsacji ω_y

- C_x, R_x, R_{d1}, R_{d2} — elementy wyznaczające odpowiedni charakter i wartości prądów sterujących transoptory T1 i T2
 C_y, R_y, R_{d3}, R_{d4} — elementy wyznaczające odpowiedni charakter i wartości prądów sterujących transoptory T3 i T4
 q_x, q_y — liczby bezwymiarowe, odpowiednio równe:

$$q_x = R_{d2} G_x$$

$$q_y = R_{d4} G_y$$

przy czym $G_x = 1/R_x$ i $G_y = 1/R_y$ są admitancjami wejściowymi torów o charakterze rezystancyjnym.

Przyjmuje się, że spełnione są warunki konstrukcyjne:

$$\left. \begin{aligned} \omega_x R_{d1} C_x &\ll 1 \\ \omega_y R_{d3} C_y &\ll 1 \\ q_x &\ll 1 \\ q_y &\ll 1 \end{aligned} \right\} \quad (2)$$

Wówczas układ równań (1) przyjmuje przybliżoną postać:

$$\left. \begin{aligned} I_{C1} &\cong K_{UZ1} U_x R_{d1} C_x \omega_x \\ I_{R1} &\cong K_{UZ2} q_x U_x \\ I_{C2} &\cong K_{UZ3} U_y R_{d3} C_y \omega_y \\ I_{R2} &\cong K_{UZ4} q_y U_y \end{aligned} \right\} \quad (1a)$$

Uwzględniając ogólną charakterystykę transoptora $G_T = AI_s$, [5], dostaje się

$$\left. \begin{aligned} G_{T1} &= (A_1^* R_{d1} C_x) U_x \omega_x \\ G_{T2} &= (A_2^* q_x) U_x \\ G_{T3} &= (A_3^* R_{d3} C_y) U_y \omega_y \\ G_{T4} &= (A_4^* q_y) U_y \end{aligned} \right\} \quad (3)$$

gdzie symbole

$$A_1^* = A_1 K_{UZ1}$$

$$A_2^* = A_2 K_{UZ2}$$

$$A_3^* = A_3 K_{UZ3}$$

$$A_4^* = A_4 K_{UZ4}$$

wyrażają stałe konstrukcyjne transoptorów wraz z odnośnymi układami zasilającymi.

3. CHARAKTERYSTYKA PRZETWARZANIA PRZETWORNIKA ZMODYFIKOWANEGO

Napięcie na wyjściowej przekątnej mostka jest określone równaniem:

$$\frac{e}{E} = \frac{G_{T1}}{G_{T1} + G_{T2}} - \frac{G_{T3}}{G_{T3} + G_{T4}} \quad (4)$$

gdzie:

E — stabilizowane napięcie pomocnicze zasilające mostek.

Po uwzględnieniu zależności (3), napięcie wyjściowe przetwornika względnej różnicy częstotliwości opisuje wyrażenie:

$$\frac{e}{E} = \frac{A_4^* A_1^* R_{d1} C_x \omega_x q_y - A_3^* A_2^* R_{d3} C_y \omega_y q_x}{(A_1^* R_{d1} C_x \omega_x + A_2^* q_x)(A_3^* R_{d3} C_y \omega_y + A_4^* q_y)}. \quad (5)$$

W momencie zrównania się obu częstotliwości mostek powinien być w równowadze. Wobec tego dla $\omega_x = \omega_y$ ma być $e = 0$. Postulat ten jest spełniony, jeżeli zachodzi związek:

$$A_1^* A_4^* R_{d1} C_x q_y = A_2^* A_3^* R_{d3} C_y q_x$$

co daje warunek konstrukcyjny:

$$A_4^* q_y = \frac{A_2^* A_3^* R_{d3} C_y q_x}{A_1^* R_{d1} C_x}. \quad (6)$$

Po podstawieniu (6) do równania (5) dostaje się:

$$\frac{e}{E} = \frac{\omega_x - \omega_y}{\left(\frac{q_x}{C_y R_{d3}} + \frac{A_1^* R_{d1} C_x}{A_2^* R_{d3} C_y} \omega_x \right) \left(\frac{A_2^* R_{d3} C_y}{A_1^* R_{d1} C_x} + \frac{C_y R_{d3}}{q_x} \omega_y \right)}. \quad (7)$$

Wprowadza się teraz oznaczenia skrótowe:

$$\frac{e}{E} = \varepsilon \quad (8)$$

$$\frac{A_1^* R_{d1} C_x}{A_2^* R_{d3} C_y} = B_0 \quad (9)$$

oraz

$$\frac{q_x}{C_y R_{d3}} = \omega_0. \quad (10)$$

Pulsację prądnicy ω_x wiąże się z pulsacją sieci ω_y za pomocą zależności:

$$\omega_x = \omega_y (1 + \sigma) \quad (11)$$

gdzie σ oznacza względną różnicę obu częstotliwości. W ten sposób dochodzi się do wzoru:

$$\varepsilon = \frac{1}{1 + \frac{B_0 \omega_y}{\omega_0} \left(1 + \frac{\omega_0}{B_0 \omega_y} \right) + \sigma}. \quad (12)$$

Nazwą „częstotliwości sprowadzonej” określa się ułamek:

$$\frac{B_0 \omega_y}{\omega_0} = \eta \quad (13)$$

co pozwala zapisać równanie (12) w postaci dogodnej do analizy:

$$\varepsilon = \frac{1}{1 + \eta} \frac{\sigma}{\left(1 + \frac{1}{\eta}\right) + \sigma}. \quad (14)$$

Należy tu wyraźnie podkreślić, iż zmiany częstotliwości sieci nieszytywnej ω_y , wywołują zmiany wartości częstotliwości sprowadzonej, czyli $\eta = var$. Łącząc równania (9), (10) i (13) można wyznaczyć wartość znamionową częstotliwości sprowadzonej η_n

$$\eta_n = \frac{A_1^* R_{d1} C_x}{A_2^* R_{d2} G_x} \omega_{yn} \quad (15)$$

Dla danej znamionowej częstotliwości sieci okrętowej ω_{yn} dobór odpowiedniej wartości η_n może odbywać się poprzez regulację rezystancji R_{d1} i/lub R_{d2} , regulację pojemności C_x i/lub admitancji G_x oraz zmianę stosunku A_1/A_2 na drodze zmiany odległości „ a ” transoptorów T_1 i/lub T_2 .

Należy tu jednak pamiętać o konsekwencjach wynikających z warunku równowagi (6). Techniczna realizowalność wielkości η_n związana jest z żadaną wartością e_{max} sygnału wyjściowego, gdyż zasilające napięcie pomocnicze E ograniczone jest obciążalnością elementów mostka oraz uzyskaniem liniowości charakterystyk transoptorów $G_T = f(I_s)$. Napięcie pomocnicze E może być uzyskane z dodatkowego wtórnego uzwojenia jednego z transformatorów wejściowych (obniżających napięcie przetwarzane do poziomu wymaganego dla zastosowanej konstrukcji torów pomiarowych.)

4. LINIOWOŚĆ I STAŁA UKŁADU

Omawiany układ mógłby być w przybliżeniu liniowym przetwornikiem względnej różnicy częstotliwości, o ile przedział spotykanych wartości η będzie należycie dobrany. Należy więc wybrać $\eta_n = B_0 \cdot \omega_{yn}/\omega_0$ tak, by wpływ zmian ω_y na ε był minimalny. Celem bliższego naświetlenia tego zagadnienia wykorzystuje się równanie (14). Po przekształceniu otrzymuje się

$$\frac{\varepsilon}{\sigma} = \frac{1}{\kappa} \quad (16)$$

gdzie wyrażenie

$$\kappa = (1 + \eta) \left[\left(1 + \frac{1}{\eta}\right) + \sigma \right] \quad (16a)$$

oznacza stałą rozważanego układu.

Częstotliwość ω_y podlega zmianom w pewnym przedziale wartości

$$\omega_{ym} \leq \omega_y \leq \omega_{yM}. \quad (17)$$

Zgodnie z równaniem (13), przy $B_0 = \text{const.}$ i $\omega_0 = \text{const.}$ wywołuje to odpowiednie zmiany częstotliwości sprowadzonej

$$\eta_m \leq \eta \leq \eta_M \quad (18)$$

gdzie η_m i η_M są odpowiednio minimalną i maksymalną wartością tej wielkości. Przyjmuje się, że można je wyrazić jako:

$$\eta_m = (1 - d) \quad (19)$$

$$\eta_M = (1 + d) \quad (20)$$

przy czym $0 < d < 1$.

Wartościom η_m i η_M zgodnie z równaniem (16a), odpowiadają odpowiednio stałe przetwornika κ_m i κ_M .

Dla przedziału η określonego nierównością (18) zachodzi więc zmienność κ taka, że pojawia się wartość średnia

$$\kappa_S = \frac{1}{2}(\kappa_m + \kappa_M). \quad (21)$$

W rozpatrywanym przedziale występuje więc wpływ wielkości η na wartość κ , dający się scharakteryzować za pomocą błędu względnego:

$$\delta_\kappa = \frac{\kappa_M - \kappa_m}{\kappa_M + \kappa_m} 100\%. \quad (22)$$

Liniowość i stała rozważanego układu wiążą się bezpośrednio z wielkościami δ_κ i κ . Zarówno błąd δ_κ , jak i stała przetwornika κ są funkcjami dwóch zmiennych: d oraz σ . Związków $\delta_\kappa = f(d, \sigma)$ oraz $\kappa = f(d, \sigma)$ poszukiwano w oparciu o analizę numeryczną. W wyniku obliczeń wartości δ_κ jako funkcji σ przy parametrze d wyznaczono wartość σ tak, by dla danej pary (d, σ) otrzymać założoną wartość błędu δ_κ . Obliczenia wykonano dla $-1\% \leq \delta_\kappa \leq +1\%$ co $0,1\%$ i $0,05 \leq d \leq 0,35$ co $0,05$. Wyniki takiej procedury przedstawiono w tablicy 1.

Okazało się, że zależność $\kappa_S = f(d, \sigma)$ może być opisana równaniem:

$$\kappa_S = A_\kappa + B_\kappa \cdot d + C_\kappa \cdot \sigma + D_\kappa \cdot \sigma \cdot d. \quad (23)$$

Równanie (23) badano metodą analizy regresyjnej i otrzymano wartości współczynników:

$$A_\kappa = 3,9665; \quad B_\kappa = 0,4283; \quad C_\kappa = 1,9913 \quad \text{ i } \quad D_\kappa = 0,1770$$

przy współczynniku korelacji $R^2 = 0,9996$ dla przedziału $0,05 \leq d \leq 0,35$. Następnie wyznaczono błędy przybliżenia δ_p [%] (tablica 2), rozumiane jako różnice między wartościami κ_S wg tablicy 1 i ich przybliżeniami wg równania (23). Z analizy przeprowadzonej dla wszystkich punktów tablicy 1 wynika, że $\delta_{pm} = -0,53\%$ (przy $d = 0,35$, $\sigma = +0,0202$) oraz $\delta_{pM} = +0,39\%$ (przy $d = 0,20$, $\sigma = -0,2708$). Nadto ustalono, że $\delta_p \cong 0\%$ (przy $d = 0,10$; $\sigma = 0,0941$ oraz $d = 0,30$; $\sigma = 0,0847$). Dalej przyjęto, że wartość błędu δ_κ [%] daje się wyrazić zależnością:

$$\delta_\kappa = A_\delta + B_\delta \cdot \sigma + C_\delta \cdot \sigma^2. \quad (24)$$