

W oparciu o równanie (1) oraz równanie (3) można napisać

$$\frac{H_c}{c_{pK}} \frac{\partial E}{\partial t} = E \frac{\partial T_K}{\partial t} = \frac{H_c D_K}{c_{pK} \rho_K} \nabla^2 E, \quad (4)$$

które może być przepisane jako

$$E \frac{\partial T_K}{\partial t} = \frac{H_c}{c_{pK}} \left(\frac{\partial E}{\partial t} - \frac{D_K}{\rho_K} \nabla^2 E \right) \quad (5)$$

oraz dodatkowo

$$\frac{H_c}{c_{pK}} \frac{\partial E}{\partial t} = \partial T_K. \quad (6)$$

Powracając do rozważań rozdziału 2^o.2.1. zagadnienie transportu entalpii w strumieniu dyfuzji roztwór – kryształ posiada rozwiązanie:

$$T_K(K, t) = T_{K0}(K, t) + \frac{H_c}{c_{pK}} \left[\ln \left| \int_0^t \int_{\Omega_K} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^E e^{-\lambda_{n,m,k} \frac{D_K}{\rho_K} C(t-\tau)} \frac{8}{abc} \cdot W(x, \xi; y, \eta; z, \beta) \right\} |E_0| e^{S\tau} d\xi d\eta d\beta d\tau - \ln |E_0| \right] \right] \quad (7)$$

Dla rozwiązania (7) obowiązuje funkcja Greena dana równaniem (12) przy dyfuzyjnym transporcie stężenia roztwór – kryształ.

2^o.2.3. Zmiany temperatury od przewodnictwa cieplnego roztwór – kryształ (zmienny punkt $Q(\xi, \eta, \beta, \tau)$).

Punktem wyjścia rozważań jest postać jednorodna równania przewodnictwa cieplnego roztwór – kryształ ($\Pi_{3.3}^{hp}$) przepisana poniżej

$$E \frac{\partial T_K}{\partial t} - \left(- \frac{\lambda_K}{c_{pK} \rho_K} \right) (1-E) \nabla^2 T_K = 0 \quad (1)$$

Dla równania (1) wprowadzamy nową postać funkcji temperatury kryształów

$$T_K(x, y, z, t) = \bar{V}(x, y, z, t) \bar{T}_K(t) \quad (2)$$

i z równania (1) otrzymujemy następujący układ równań

$$\frac{\partial^2 \bar{V}}{\partial x^2} + \frac{\partial^2 \bar{V}}{\partial y^2} + \frac{\partial^2 \bar{V}}{\partial z^2} + \lambda \bar{V} = 0 \quad (3)$$

oraz

$$\frac{d\bar{T}_K}{dt} = -\lambda \frac{\lambda_K}{c_{pK} \rho_K} \left(1 - \frac{1}{E} \right) \bar{T}_K. \quad (4)$$

Na bazie TWIERDZENIA TW3 z Części I [1] warunki brzegowe na brzegu $F_{K(a,b,c)t}$ elementu objętość–czas $\bar{\Omega}_{K(a,b,c)t}$ są równe zero — patrz rozdział 1^o.2.1.

W konsekwencji takiego podejścia uzyskujemy:

— wartości własne

$$\lambda_{n,m,k} = \left(\frac{n\pi}{a}\right)^2 + \left(\frac{m\pi}{b}\right)^2 + \left(\frac{k\pi}{c}\right)^2. \quad (5)$$

— funkcje własne dla punktu $K(x, y, z, t)$ o stałych współrzędnych

$$\bar{V}_{n,m,k}(x, y, z) = \sqrt{\frac{8}{abc}} W(x, y, z) \quad (6)$$

— funkcje własne dla punktu $Q(\xi, \eta, \beta, \tau)$ o zmiennych współrzędnych

$$\bar{V}_{n,m,k}(\xi, \eta, \beta) = \sqrt{\frac{8}{abc}} W(\xi, \eta, \beta). \quad (7)$$

Po scałkowaniu równania (4) otrzymuje się

$$\bar{T}_K(t) = D_{n,m,k}^{T_K} e^{-\lambda_{n,m,k} \frac{\lambda_K}{c_{pK} \rho_K} [J_K(t)]} \quad (8)$$

gdzie

$$[J_K(t)] = \int \left[1 - \frac{1}{|E_0| e^{5t}} \right] dt.$$

We wzorze (8) stała całkowania ma postać

$$D_{n,m,k}^{T_K} = \iiint_{\Omega_K} T_{K0} \sqrt{\frac{8}{abc}} W(\xi, \eta, \beta) d\xi d\eta d\beta. \quad (9)$$

Funkcja źródłowa dla przewodnictwa cieplnego roztwór–kryształ ze wzoru (2^o.1.2) o postaci

$$T_K^k(t) = T_{K0} + \frac{H_{GK}}{c_{pK}} St. \quad (10)$$

Całość powyższych rozważań prowadzi do pełnego rozwiązania zagadnienia przewodnictwa cieplnego roztwór–kryształ

$$T_K(K, t) = \int_0^t \iiint_{\Omega_K} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^{T_K} e^{-\lambda_{n,m,k} \frac{\lambda_K}{c_{pK} \rho_K} [J_K(t-\tau)]} \frac{8}{abc} W(x, \xi, y, \eta, z, \beta) \right\} \cdot$$

$$\cdot \left(T_{K0} + \frac{H_{GK}}{c_{pK}} S\tau \right) d\xi d\eta d\beta d\tau. \quad (11)$$

Funkcja Greena dla rozwiązania (11) ma postać

$$G_{Ph}^T(K, t; Q, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^T e^{-\lambda_{n,m,k} \frac{\lambda_k}{c_{k,k} \rho_k} [U_k(t-\tau)]} \frac{8}{abc} W(x, \xi; y, \eta; z, \beta) \right\} \quad (12)$$

i odpowiada stałemu zjawiskowemu obrazowi punktów [11–12], [15]

$$\begin{array}{ccc} K(x, y, z, t) & \longleftrightarrow & Q(\xi, \eta, \beta, \tau) \\ \in \bar{\Omega}_{K(a,b,c)t} & \text{Stały wzajemnościowy obraz.} & \in \Omega_{Rt} \end{array} \quad (13)$$

2^o.3. Zjawiskowe rozwiązania dla pola wirowego

2^o.3.1. Pole wektora prędkości roztwór – kryształ (zmienny punkt $Q'(\xi', \eta', \beta', \tau)$).

Zaczynamy nasze rozważania od układu (Π_1^H) mającego postać

$$\frac{DE}{Dt} = + E \operatorname{div} \mathbf{v}_K \quad (1)$$

$$E \frac{DT_K}{Dt} = + E T_K \operatorname{div} \mathbf{v}_K \quad (2)$$

$$\rho_K E \frac{D\mathbf{v}_K}{Dt} = \eta_w E \nabla^2 \mathbf{v}_K + \rho_K E \mathbf{v}_K \operatorname{div} \mathbf{v}_K + \frac{\eta_w}{3} E \operatorname{grad} \operatorname{div} \mathbf{v}_K \quad (3)$$

przy warunku pływania

$$\rho_K E \frac{D\mathbf{v}_K}{Dt} = C(1-E) \frac{D\mathbf{v}}{Dt} \quad (4)$$

gdzie

$$\frac{D}{Dt} = \frac{\partial}{\partial t} - \mathbf{v}_K \operatorname{grad} \quad [13]. [2-8]$$

Teraz równanie (1) mnożymy obustronnie przez element „ $\rho_K \mathbf{v}_K$ ” i stosujemy je następnie w równaniu (3) otrzymując rezultat

$$\rho_K E \frac{D\mathbf{v}_K}{Dt} - \mathbf{v}_K \rho_K \frac{DE}{Dt} = \eta_w E \nabla^2 \mathbf{v}_K + \frac{\eta_w}{3} E \operatorname{grad} \operatorname{div} \mathbf{v}_K. \quad (5)$$

Założenie o stałym obrazie wzajemnościowym w stanie statycznym dla elementu $\bar{\Omega}_{K(a,b,c)}$ pozwala określić równanie (5) w postaci układu równań, do których tyczy się Uwaga U10:

Uwaga U10. Treść Uwagi U6 zachowuje ważność dla rozważań tego rozdziału. Otrzymujemy więc

$$\eta_w E \nabla^2 \mathbf{v}_K + \frac{\eta_w}{3} E \operatorname{grad} \operatorname{div} \mathbf{v}_K = 0 \quad (6)$$

oraz

$$\rho_K E \frac{Dv_K}{Dt} - v_K \rho_K \frac{DE}{Dt} = 0. \quad (7)$$

Z zastosowania reguły matematycznej [14], [16]

$$\text{rot} \text{rot} v_K = \text{grad} \text{div} v_K - \nabla^2 v_K \quad (8)$$

w równaniu (6) otrzymujemy

$$\left(\eta_w + \frac{\eta_w}{3} \right) E \text{grad} \text{div} v_K - \eta_w E \text{rot} \text{rot} v_K = 0 \quad (9)$$

kontynuując rozważania równaniu (9) nadajemy postać

$$\bar{\nabla}^2 v_K = \hat{\alpha} \text{grad} \text{div} v_K - \hat{\beta} \text{rot} \text{rot} v_K, \quad (10)$$

gdzie

$$\hat{\alpha} = \left(\eta_w + \frac{\eta_w}{3} \right) E \quad \text{oraz} \quad \hat{\beta} = \eta_w E.$$

Dla stanu statycznego, przy $t=0$ każdy punkt posiadający współrzędne $K(x, y, z, t)=K(x, y, z, t=0)=K(x, y, z, 0)$ posiada pole wektora prędkości $v_{K0}(x, y, z, 0)=v_{K0}$ zaś rozkład statyczny pola wektora prędkości kryształów dany jest równaniem

$$v_{K0}(K, 0) = v_{K0}(x, y, z, 0) = \iiint_{\Omega_K} \Gamma_{GK}^W(K, Q') v_{K0}(Q') d\xi' d\eta' d\beta' d\tau'. \quad (11)$$

Uwaga U11. Zawartość Uwagi U7 odnosi się także do zagadnień wzrostu kryształów, a więc punktu $K(x, y, z, t)$ oraz pola wektora prędkości „ v_{K0} ”. Statyczny tensor Greena dla elementu $\bar{\Omega}_{K(a,b,c)}$ zawiera dwa przypadki [16–17], [2–8]:

— przypadek izotropowy, Dodatek 3

$$\Gamma_{GK}^W(K, Q') = \Gamma_{GK}(K, Q') \quad (12)$$

— przypadek anizotropowy, Dodatek 4

$$\Gamma_{GK}^W(K, Q') = \Gamma_{GKA}(K, Q'). \quad (13)$$

Następną część naszych rozważań stanowi dynamiczna strona pola wektora prędkości kryształów „ v_K ” związana ze współrzędną czasową „ t ” dla elementu $\bar{\Omega}_{K(a,b,c)t}$

$$\rho_K E \frac{Dv_K}{Dt} - v_K \rho_K \frac{DE}{Dt} = 0. \quad (14)$$

Równanie (14) całkujemy obecnie otrzymując

$$\ln|N^s| + \ln|v_K| = \ln|E| \quad (15)$$

i następnie po przekształceniach

$$v_K(t) = \frac{E^k(t)}{N_1^s}, \quad (16)$$

gdzie: $E^k(t)$ — funkcja źródłowa dla ułamka fazy krystalicznej w punkcie $K(x, y, z, t)$ dla $x=a, y=b, z=c$ zgodnie z równaniem (2^o.1.1)

$$E^k(t) = |E_0|e^{St} \quad (17)$$

N_1^s — stała całkowania otrzymywana z równania (11) zawierającego związek $N_1^s = v_K(K, 0)$ dla stanu statycznego.

Do rozważań wprowadzamy teraz funkcję źródłową pola wektora prędkości kryształów

$$v_K^k(t) = v_{K0} + \left\{ 1 + \frac{C^z(t)[1-E^k(t)]}{\rho_K E^k(t)} \right\} g t \quad (18)$$

i w ten sposób możemy sformułować kompletną postać rozwiązania zagadnienia rozkładu pola wektora prędkości roztwór – kryształ

$$v_{K(1)}(K, t) = \int_0^t \iiint_{\Omega_K} \left\{ \frac{1}{N_1^s} |E_0| e^{S(t-\tau)} \Gamma_{GK}^W(K, Q') \right\} \left\langle v_{K0} + \left\{ 1 + \frac{C^z(\tau)[1-E^k(\tau)]}{\rho_K E^k(\tau)} \right\} g \tau \right\rangle d\xi' \cdot d\eta' d\beta' d\tau \quad (19)$$

Uwaga U14. Patrz Uwaga U8.

Przyjmując istnienie równoważnej reakcji punktowej wzrostu kryształów w punkcie $K(x, y, z, t)$ otrzymujemy dodatkowe zależności

$$v = v_K \quad (20)$$

oraz

$$\rho_K E = C(1-E) \quad (21)$$

które modyfikują dynamiczną część równania tematowego w postaci (14) do:

$$C(1-E) \frac{Dv}{Dt} - v_K \rho_K \frac{DE}{Dt} = 0. \quad (22)$$

Równanie (22) całkujemy konsekwentnie do poprzednich rozważań otrzymując

$$v_K(t) = \frac{-v^z(t)C^z(t)}{\rho_K \ln|1-|E_0|e^{St}|N_1^s|} \quad (23)$$

gdzie:

N_1^s — stała całkowania otrzymana z równania (11), przy $N_1^s = v_K(K, 0)$,

$v^z(t)$ — funkcja źródłowa pola wektora prędkości roztworu, równanie (1^o.1.3),

$C^z(t)$ — funkcja źródłowa stężenia dla przesyconego roztworu, równanie (1^o.1.1).

Zależności (20) oraz (21) modyfikują funkcję źródłową pola wektora prędkości kryształów (2^o.1.3) do postaci

$$\mathbf{v}_k^k(t) = \mathbf{v}_{k0} + \left\{ \frac{\rho_K \mathbf{E}^k(t)}{C^z(t)[1 - \mathbf{E}^k(t)]} + 1 \right\} \mathbf{g}t. \quad (24)$$

Zmodyfikowana przez równania (20–24) postać pola wektora prędkości roz-
twór – kryształ może być zapisana następująco

$$\mathbf{v}_{K(2)}(K, t) = \iiint\limits_0^t \iiint\limits_{\Omega_K} \left\{ \frac{-v^z(t-\tau)}{\rho_K \ln|1 - |E_0|e^{St(t-\tau)}|N_1^s|} \Gamma_{GK}^W(K, Q') \right\} - \left\langle \mathbf{v}_{k0} + \left\{ \frac{\rho_K \mathbf{E}^k(\tau)}{C^z(\tau)[1 - \mathbf{E}^k(\tau)]} + 1 \right\} \mathbf{g}\tau \right\rangle \cdot d\xi' d\eta' d\beta' d\tau \quad (25)$$

Uwaga U15. Patrz Uwaga U8.

Rozwiązaniu (19) odpowiada następująca funkcja Greena

$$\mathbf{G}_{Ph(2)}^v(K, t; Q', \tau) = \left\{ \frac{1}{N_1^s} |E_0| e^{St(t-\tau)} \Gamma_{GK}^W(K, Q') \right\} \quad (26)$$

zaś rozwiązaniu (25) przyporządkowana jest następująca funkcja Greena

$$\mathbf{G}_{Ph(2)}^v(K, t; Q', \tau) = \left\{ \frac{-v^z(t-\tau) C^z(t-\tau)}{\rho_K \ln|1 - |E_0|e^{St(t-\tau)}|N_1^s|} \Gamma_{GK}^W(K, Q') \right\}. \quad (27)$$

Oba równania (26) oraz (27) dotyczą tego samego zagadnienia stałości zjawiskowego obrazu dla punktów [11–12], [15]

$$\begin{array}{ccc} K(x, y, z, t) & \longleftrightarrow & Q'(\xi', \eta', \beta', \tau) \\ \in \bar{\Omega}_{K(a,b,c)t} & \text{Stały wzajemnościowy obraz.} & \in \Omega_{Rt} \end{array} \quad (28)$$

3^o. Rozwiązanie równania różniczkowego cząstkowego dla procesu zarodkowania wtórnego

Rozważania zaczynamy od przypomnienia w tym miejscu równania różniczkowego cząstkowego dla procesu zarodkowania wtórnego

$$n\rho_K K_V L v_K^2 + \frac{B}{2\rho_K} \frac{\partial^2 L}{\partial t^2} = h \left[\nabla^2 L + \frac{\bar{k}}{\bar{k}-2} \text{grad div } L \right]. \quad (1)$$

Funkcją źródłową dla równania (1) jest początkowa wartość wektora rozmiaru liniowego kryształu przy jego zderzeniach:

$$L_0 = \text{const (warunek nieodkształcalności dla zderzeń)}. \quad (2)$$

W związku ze stałym obrazem wzajemnościowym dla elementu $\bar{\Omega}_{K(a,b,c)}$ równanie (1) przedstawiamy jako układ równań:

$$h \left[\nabla^2 \mathbf{L} + \frac{\bar{k}}{\bar{k}-2} \text{graddiv} \mathbf{L} \right] = 0 \quad (3)$$

oraz

$$n\rho_K K_V \mathbf{L} v_K^2 + \frac{B}{2\rho_K} \frac{\partial^2 \mathbf{L}}{\partial t^2} = 0. \quad (4)$$

Stosując matematyczną regułę [14], [16]

$$\text{rotrot} \mathbf{L} = \text{graddiv} \mathbf{L} - \nabla^2 \mathbf{L} \quad (5)$$

do równania (3), równanie to przekształca się do postaci

$$h \left(1 + \frac{\bar{k}}{\bar{k}-2} \right) \text{graddiv} \mathbf{L} - h \text{rotrot} \mathbf{L} = 0. \quad (6)$$

Równanie (6) można przedstawić jako

$$\hat{\nabla}^2 \mathbf{L} = 0 = \tilde{\alpha} \text{graddiv} \mathbf{L} - \tilde{\beta} \text{rotrot} \mathbf{L} \quad (7)$$

gdzie

$$\tilde{\alpha} = h \left(1 + \frac{\bar{k}}{\bar{k}-2} \right) \quad \text{oraz} \quad \tilde{\beta} = h.$$

Dla elementu $\bar{\Omega}_{K(a,b,c)}$ w stanie statycznym i o anizotropowych własnościach kryształu wewnątrz warstwy fluidalnej lub adsorbcyjnej w Dodatku 5 wyprowadzono anizotropowy tensor Greena $\Gamma_{GKA}^W(K, Q')$ [2–4], [6], [8].

Rozważmy teraz zagadnienie dynamiczne oparte na równaniu (4), które konsekwentnie doprowadzamy do postaci

$$\frac{d^2 \mathbf{L}}{dt^2} + \frac{2n\rho_K^2 K_V v_K^2}{B} \mathbf{L} = 0. \quad (8)$$

Rzeczywiste rozwiązanie równania (8) ma postać

$$\mathbf{L}(t) = \mathbf{L}_0 \sin \sqrt{\frac{2n\rho_K^2 K_V v_K^2}{B}} t. \quad (9)$$

Do pełnego rozwiązania równania procesu zarodkowania wtórnego przypominamy funkcję źródłową w postaci

$$\mathbf{L}^k(t) = \mathbf{L}_0 \quad (10)$$

i w konsekwencji otrzymujemy

$$\mathbf{L}(K, t) = \int_0^t \iiint_{\Omega_K} \left\{ \mathbf{L}_0 \Gamma_{GKA}^W(K, Q') \sin \left[\sqrt{\frac{2n\rho_K^2 K_V v_K^2}{B}} (t - \tau) \right] \right\} \mathbf{L}^k(t) d\xi' d\eta' d\beta' d\tau. \quad (11)$$

Rozwiązanie (11) poprzez funkcję $f(L)$ stanowi zakłócenie kinetyki V_{GK} reakcji wzrostu kryształów i jako takie mieści się w obrazie wzajemnościowym ułamka fazy krystalicznej — źródło masy, procesu wzrostu kryształów — wzory (2^o.2.1.1-3). Funkcja Greena dla procesu zarodkowania wtórnego ma postać [11–12], [15]

$$G_{Ph}^L(K, t; Q', \tau) = \left\{ L_0 \Gamma_{GKA}^W(K, Q') \sin \left[\sqrt{\frac{2n\rho_K^2 K_V v_K^2}{B}} (t - \tau) \right] \right\} \quad (12)$$

4^o. Kompletne rozwiązanie dla procesu wzrostu kryształów oraz procesu zarodkowania wtórnego jako zakłócenia kinetyki wzrostu kryształów

W oparciu o zawartość rozdziałów A^o.2^o.1. — A^o.2^o.3.1. — dla wzrostu kryształów oraz A^o.3.^o. — dla zarodkowania wtórnego jako zakłócenia kinetyki wzrostu kryształów kompletne rozwiązanie tematowego zagadnienia ma postać [2–8]:

Kompletna postać
rozwiązania

Funkcje źródłowe w punkcie
K(x, y, z, t)

$$\begin{bmatrix} \text{Ułamek objętościowy} \\ \text{fazy krystalicznej} \\ \\ \text{Temperatura kryształów} \\ \\ \text{Pole wektora prędkości} \\ \text{kryształów} \end{bmatrix} = \begin{bmatrix} \text{Przyrost ułamka objętościowego fazy} \\ \text{krystalicznej (formuła A}^{\circ}\text{.2}^{\circ}\text{.1.1) oraz} \\ \text{zakłócenie (formuła A}^{\circ}\text{.3}^{\circ}\text{.12)} \\ \\ \text{Przyrost temperatury kryształów} \\ \text{(formuła A}^{\circ}\text{.2}^{\circ}\text{.1.2) oraz zakłócenie} \\ \text{(formuła A}^{\circ}\text{.3}^{\circ}\text{.11)} \\ \\ \text{Przyrost pola wektora prędkości} \\ \text{kryształów (formuła A}^{\circ}\text{.2}^{\circ}\text{.1.3) oraz} \\ \text{zakłócenie (formuła A}^{\circ}\text{.3}^{\circ}\text{.11)} \end{bmatrix} +$$

Dyfuzja
Q → K

Przewodniczywo cieplne
Q → K

$$+ \begin{bmatrix} \text{Przyrost ułamka fazy krystalicznej} \\ \text{(formuła A}^{\circ}\text{.2}^{\circ}\text{.2.1.11)} \\ \\ \text{Przyrost temperatury fazy krystalicznej} \\ \text{(formuła A}^{\circ}\text{.2}^{\circ}\text{.2.2.7)} \end{bmatrix} + \begin{bmatrix} * \\ \\ \text{Przyrost temperatury fazy} \\ \text{krystalicznej (formuła A}^{\circ}\text{.2}^{\circ}\text{.2.3.11)} \\ * \end{bmatrix} +$$

Pole wektora prędkości
Q' → K

$$\begin{bmatrix} \text{Ciągłość ułamka fazy krystalicznej} \\ \text{(formuła A}^{\circ}\text{.2}^{\circ}\text{.3.1.1)} \\ \text{Ciągłość temperatury fazy krystalicznej} \\ \text{(formuła A}^{\circ}\text{.2}^{\circ}\text{.3.1.2)} \\ \text{Przyrost pola wektora prędkości fazy krystalicznej} \\ \text{(formuła A}^{\circ}\text{.2}^{\circ}\text{.3.1.19) lub (formuła A}^{\circ}\text{.2}^{\circ}\text{.3.1.25)} \end{bmatrix} \quad (1)$$

* — współrzędne wektora stanu, w których nie ma związków między zjawiskami fizycznymi.

B⁰. ISTNIENIE WARUNKÓW POCZĄTKOWYCH I BRZEGOWYCH

1⁰. Zagadnienie brzegowe dla procesów składowych, procesu ciągłej krystalizacji masowej

Dla ustalenia postulatów istnienia warunków brzegowych dla procesów składowych, procesu ciągłej krystalizacji masowej niezbędne jest wprowadzenie następujących założeń:

Założenie ZG1. Brzeg obszaru przestrzeni robocza — czas Ω_{Rt} jest dany ciągłą powierzchnią F_{Rt} , regularną, zorientowaną zewnętrżnie wektorem $\mathbf{n}_R^{(z)}$, którego cyrkulacja zależy od zadań sterowania całością procesu ciągłej krystalizacji masowej.

Założenie ZG2. Kształt powierzchni F_{Rt} przyjmuje się jako prostopadłościenny, chociaż może być walcowy — współrzędne walcowe lub kuliste — współrzędne kuliste.

Założenie ZG3. Na brzegu F_{Rt} określony jest brzegowy wektor stanu o współrzędnych będących funkcjami ciągłymi: $S(t)$ — dla masy, $U(t)$ — dla energii cieplnej, $W(t)$ — dla pędu, dla procesów składowych krystalizacji masowej, danych układami równań (I) oraz (II) z równaniem (III.1).

Założenie ZG4. Wymiary powierzchni F_{Rt} są stałe.

Założenie ZG5. Układy równań (I) oraz (II) z równaniem (III.1) są transformowalne do odpowiednich układów równań źródłowych (I^b) oraz (II^b) z równaniem (III^b.1) na stałym geometrycznie brzegu F_{Rt} .

Założenie ZG6. Funkcje źródłowe na brzegu będziemy nazywać „FUNKCJAMI BRZEGOWYMI”.

Założenie ZG7. Zachodzą relacje: punktowi $Z(x, y, z, t)$ odpowiada punkt $R(x, y, z, t) \in F_{(a,b,c)t}^r \in F_{Rt}$ zaś dla punktu $K(x, y, z, t)$ odpowiadającym jest punkt $P(x, y, z, t) \in F_{(a,b,c)t}^p \in F_{Rt}$.

1⁰.1. Proces zarodkowania pierwotnego

1⁰.1.1. Zagadnienia brzegowe zjawisk fizykalnych roztwór — brzeg

W oparciu o TWIERDZENIE TW2 udowodnione w Części I [1] tej pracy, na bazie założeń ZG1–ZG7, dla reakcji zarodkowania pierwotnego, układ równań różniczkowych cząstkowych konstytutywnych stanu (I) w otoczeniu $x=a$, $y=b$, $z=c$ punktu $R(x, y, z, t) \in F_{(a,b,c)t}^r \in F_{Rt}$ sprowadza się do brzegowego układu równań różniczkowych zwyczajnych względem czasu

$$\begin{array}{ll}
 \frac{dC}{dt} = \pm S_c(t) & (1) \\
 (I^b) \quad C \frac{dT}{dt} = \pm \frac{1}{c_p} U_T(t) & (2) \\
 C \frac{dv}{dt} = \pm W_v(t) & (3)
 \end{array}
 \int_0^t$$

Funkcje brzegowe: $C^b(t)$, $T^b(t)$, $v^b(t)$.

Uwaga U1. Przy całkowaniu układu (I^b) nie obowiązuje ważność Uwagi U1.

Warunki początkowe na brzegu: C_{0b}, T_{0b}, v_{0b} .

Dla stanu ustalonego warunków początkowych obowiązuje ponadto układ równości:

$$C_0 = C_{0b} \quad (4)$$

$$(I^{b'}) \quad T_0 = T_{0b} \quad (5)$$

$$v_0 = v_{0b} \quad (6).$$

1^o.2. Proces wzrostu kryształów

1^o.2.1. Zagadnienia brzegowe zjawisk fizycznych kryształów — brzeg

Udowodnione w Części I [1] TWIERDZENIE TW2 oraz założenia ZG1 — ZG7 pozwalają na sformułowanie zagadnienia brzegowego dla reakcji wzrostu kryształów, w oparciu o sprowadzenie układu równań różniczkowych cząstkowych konstytutywnych stanu (II) w otoczeniu $x=a, y=b, z=c$ punktu $P(x, y, z, t) \in F_{(a,b,c)t} \in F_{Rt}$ do układu równań zwyczajnych względem czasu

$$\begin{aligned} \frac{dE}{dt} &= \pm \frac{1}{\rho_K} S_E(t) & (1) & \text{Funkcje brzegowe: } E^b(t), T_K^b(t), v_K^b(t). \\ (II^b) \quad E \frac{dT_K}{dt} &= \pm \frac{1}{c_{pK} \rho_K} U_{T_K}(t) & (2) & \text{Uwaga U1. Przy całkowaniu układu} \\ E \frac{dv_K}{dt} &= \pm \frac{1}{\rho_K} W_{v_K}(t) & (3) & \text{(II}^b\text{) nie obowiązuje ważność Uwagi} \\ & & & \text{U2.} \end{aligned}$$

Warunki początkowe na brzegu: E_{0b}, T_{K0b}, v_{K0b} .

Dla układu (II^b) następujące warunki są spełnione:

$$E_0 = E_{0b} \quad (4)$$

$$(II^{b'}) \quad T_{K0} = T_{K0b} \quad (5)$$

$$v_{K0} = v_{K0b} \quad (6).$$

1^o.3. Proces zarodkowania wtórnego

zachodzi tutaj: $L_0 = \text{const.}$ (1) oraz $L_0 = L_{0b}$ (2)

2^o. Funkcje brzegowe

2^o.1. Proces zarodkowania pierwotnego

Układ równań różniczkowych zwyczajnych (I^b) jest całkowany w przedziale czasu $[0, t]$ dla otrzymania odpowiednich funkcji brzegowych:

— z równania (I^b.1) dla stężenia roztworu

$$C^b(t) = C_{0b} \pm \int_0^t S_C(t) dt \quad (1)$$

— z równania (I^b.2) dla temperatury roztworu

$$T^b(t) = T_{0b} \pm \int_0^t U_{TC}(t) dt, \quad \text{dla } U_{TC}(t) = \frac{U_T(t)}{c_p C^b(t)} \quad (2)$$

— dla równania (I^b.3) dla pola wektora prędkości roztworu

$$\mathbf{v}^b(t) = \mathbf{v}_{0b} \pm \int_0^t \mathbf{W}_{vC}(t) dt, \quad \text{dla } \mathbf{W}_{vC}(t) = \frac{\mathbf{W}_v(t)}{C^b(t)}. \quad (3)$$

Funkcje brzegowe (1), (2) oraz (3) generują odpowiednie zjawiska fizyczne dla pól potencjalnych: dyfuzja stężenia i entalpii, przewodnictwo cieplne oraz dla pola wirowego: pole wektora prędkości w układzie roztwór – roztwór, z powierzchni F_{Rt} do wnętrza Ω_{Rt} . Brzgowy wektor stanu działający w punkcie $R(x, y, z, t)$ według współrzędnych $[C^b(t), T^b(t), \mathbf{v}^b(t)]$ uwzględnia, że

$$R(x, y, z, t) \in F'_{(a,b,c)t} \in F_{Rt}. \quad (4)$$

2^o.2. Proces wzrostu kryształów

Całkujemy układ równań różniczkowych zwyczajnych (II^b) w przedziale czasu $[0, t)$ dla otrzymania funkcji brzegowych:

— z równania (II^b.1) dla ułamka fazy krystalicznej

$$E^b(t) = E_{0b} \pm \int_0^t S_{EG}(t) dt, \quad \text{dla } S_{EG}(t) = \frac{S_E(t)}{\rho_K} \quad (1)$$

— z równania (II^b.1) dla temperatury kryształów

$$T_K^b(t) = T_{K0b} \pm \int_0^t U_{T_{KE}}(t) dt, \quad \text{dla } U_{T_{KE}}(t) = \frac{U_{T_K}(t)}{c_{pK} \rho_K E^b(t)} \quad (2)$$

— z równania (II^b.3) dla pola wektora prędkości kryształów

$$\mathbf{v}_K^b(t) = \mathbf{v}_{K0b} \pm \int_0^t \mathbf{W}_{v_{KE}}(t) dt, \quad \text{dla } \mathbf{W}_{v_{KE}}(t) = \frac{\mathbf{W}_{v_K}(t)}{\rho_K E^b(t)}. \quad (3)$$

Funkcje brzegowe (1), (2), a także (3) generują odpowiednie zjawiska fizyczne z powierzchni F_{Rt} w kierunku wnętrza Ω_{Rt} . Są to zjawiska dla pól potencjalnych: dyfuzja stężenia i entalpii, przewodnictwo cieplne, a także pola wirowego: pole wektora prędkości — wszystkie w reakcji roztwór – kryształ. Funkcje brzegowe $[E^b(t), T_K^b(t), \mathbf{v}_K^b(t)]$ tworzą wektor stanu działający w punkcie $P(x, y, z, t)$ podlegającym relacji:

$$P(x, y, z, t) \in F^p_{(a,b,c)t} \in F_{Rt}. \quad (4)$$

2⁰.3. Zarodkowanie wtórne

Funkcja brzegowa zarodkowania wtórnego ma postać:

$$L^b(t) = L_0 \quad (1)$$

przy czym

$$L_0 = L_{0b} \quad (2)$$

3⁰. Rozwiązania zjawisk fizykalnych przy istnieniu warunków początkowych i brzegowych dla procesu zarodkowania pierwotnego3⁰.1. Rozwiązania zjawisk fizykalnych dla pól potencjalnych3⁰.1.1. Dyfuzyjny transport stężenia roztwór—roztwór (zmienny punkt $Q(\xi, \eta, \beta, \tau)$).

Funkcja Greena dla dyfuzyjnego transportu stężenia roztwór—roztwór dla warunków początkowych ma postać daną równaniem ($A^0.1^0.2.1.19$)

$$G_{Ph}^C(Z, t; Q, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^C e^{-\lambda_{n,m,k} D_i(t-\tau)} \frac{8}{abc} W(x, \xi; y, \eta; z, \beta) \right\} \quad (1)$$

przy czym stała całkowania jest dana równaniem w postaci ($A^0.1^0.2.1.16$)

$$D_{n,m,k}^C = \iiint_{\Omega_r} C_0 \sqrt{\frac{8}{abc}} W(\xi, \eta, \beta) d\xi d\eta d\beta. \quad (2)$$

W oparciu o założenia ZG2, p prstopadłościennym kształcie powierzchni F_{Rt} , funkcje Greena dla poszczególnych powierzchni mają postacie związane z punktem $R(x, y, z, t)$

— dla powierzchni „xy”

$$G_{Ph}^{Cxy}(R^{xy}, t; Q^{\xi\eta}, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1}}^{\infty} D_{n,m}^{Cxy} e^{-\lambda_{n,m} D_i(t-\tau)} \frac{4}{ab} W(x, \xi; y, \eta) \right\} \quad (3)$$

przy stałej całkowania

$$D_{n,m}^{Cxy} = \iint_{F_{Rxy}} C_{0b} \sqrt{\frac{4}{ab}} W(x, \xi; y, \eta) d\xi d\eta \quad (4)$$

— dla powierzchni „yz”

$$G_{Ph}^{Cyz}(R^{yz}, t; Q^{\xi\tau}, \tau) = \left\{ \sum_{\substack{m=1 \\ k=1}}^{\infty} D_{m,k}^{Cyz} e^{-\lambda_{m,k} D_i(t-\tau)} \frac{4}{bc} W(y, \eta; z, \beta) \right\} \quad (5)$$

przy stałej całkowania

$$D_{m,k}^{Cyz} = \iint_{F_{Ryz}} C_{0b} \sqrt{\frac{4}{bc}} W(y, \eta; z, \beta) d\eta d\beta \quad (6)$$

— dla powierzchni „xz”

$$G_{Ph}^{C_{xz}}(R^{xz}, t; Q^{\xi\beta}, \tau) = \left\{ \sum_{n=1}^{\infty} D_{n,k}^{C_{xz}} e^{-\lambda_{n,k} D_c(t-\tau)} \frac{4}{ac} W(x, \xi; z, \beta) \right\} \quad (7)$$

przy stałej całkowania

$$D_{n,k}^{C_{xz}} = \iint_{F_{R_{xz}}} C_{0b} \sqrt{\frac{4}{ac}} W(x, \xi; z, \beta) d\xi d\beta. \quad (8)$$

Funkcja Greena dla całej prostopadłościenniej powierzchni F_{R_t} ma postać

$$G_{Ph}^{CS}(R, t; Q, \tau) = 2[G_{Ph}^{C_{xy}}(R^{xy}, t; Q^{\xi\eta}, \tau) + G_{Ph}^{C_{yz}}(R^{yz}, t; Q^{\eta\beta}, \tau) + \\ + G_{Ph}^{C_{xz}}(R^{xz}, t; Q^{\xi\beta}, \tau)]. \quad (9)$$

Dla postawienia pełnego rozwiązania zagadnienia dyfuzji stężenia roztwór—roztwór, przypominamy sobie:

— funkcję źródłową dla warunków początkowych objętość—czas, z równania (A^{0.1^o.1.1})

$$C^z(t) = C_0 - \int_0^t V_B G dt, \quad (10)$$

— funkcję brzegową dla warunków brzegowych powierzchnia—czas, z równania (B^{0.2^o.1.1})

$$C^b(t) = C_{0b} \pm \int_0^t S_C(t) dt. \quad (11)$$

W związku z całością powyższych rozważań pełne rozwiązanie zagadnienia transportu stężenia w strumieniu dyfuzji roztwór—roztwór w obecności warunków początkowych i brzegowych jest dane wzorem

$$C(Z, R, t) = \int_0^t \iiint_{\Omega_x} G_{Ph}^C(Z, t; Q, \tau) (C_0 - \int_0^{\tau} V_B G dt) d\xi d\eta d\beta d\tau + \\ + \int_0^t \iint_{F_x} G_{Ph}^{CS}(R, t; Q, \tau) [C_{0b} \pm \int_0^{\tau} S_C(t) dt] d\xi d\eta d\beta d\tau. \quad (12)$$

3^o.1.2. Dyfuzyjny transport entalpii roztwór—roztwór (zmienny punkt $Q(\xi, \eta, \beta, \tau)$).

Na podstawie rozważań rozdziału B^{0.3^o.1.1}. można sformułować rozwiązanie zagadnienia dyfuzyjnego transportu entalpii w układzie roztwór—roztwór

$$\begin{aligned}
T(Z, R, t) = T_0(Z, t) + \frac{H_c}{c_p} \left[\ln \left| \int_0^t \iiint_{\Omega_x} G_{Ph}^C(Z, t; Q, \tau) (C_0 - \int_0^{\tau} V_B G dt) d\xi d\eta d\beta d\tau \right| - \right. \\
\left. - \ln |C_0| \right] + T_{0b}(R, t) + \frac{H_c}{c_p} \left[\ln \left| \int_0^t \iint_{F_x} G_{Ph}^{CS}(R, t; Q, \tau) [C_{0b} \pm \right. \right. \\
\left. \left. \int_0^{\tau} S_c(t) dt \right] d\xi d\eta d\beta d\tau \right| - \ln |C_{0b}| \right]. \quad (1)
\end{aligned}$$

We wzorze (1) zjawiskowe funkcje Greena są określane następująco:

- zagadnienia warunków początkowych — wzór (B^{0.3}0.1.1.1),
- zagadnienia warunków brzegowych — wzór (B^{0.3}0.1.1.9).

3^{0.1.3}. Zmiany temperatury od przewodnictwa cieplnego roztwór — roztwór (zmien-
ny punkt Q(ξ, η, β, τ)).

Funkcja Greena dla zmian temperatury od przewodnictwa cieplnego roz-
twór — roztwór przy istnieniu warunków początkowych jest dana wzorem
(A^{0.1}0.2.3.12)

$$G_{Ph}^T(Z, t; Q, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^T e^{-\lambda_{n,m,k} \frac{\lambda_z}{c_p} [J(t-\tau)]} \frac{8}{abc} W(x, \xi; y, \eta; z, \beta) \right\} \quad (1)$$

dla stałej całkowania podanej wzorem (A^{0.1}0.2.3.9)

$$D_{n,m,k}^T \iiint_{\Omega_x} T_0 \sqrt{\frac{8}{abc}} W(\xi, \eta, \beta) d\xi d\eta d\beta \quad (2)$$

oraz

$$[J(t)] = \int \frac{dt}{C_0 - \int_0^t V_B G dt}. \quad (3)$$

Dla zagadnień powierzchniowych wprowadzamy wspólną funkcję

$$[J^u(t)] = \int \frac{dt}{C_{0b} \pm \int_0^t S_c(t) dt} \quad (4)$$

i odpowiednio dla poszczególnych powierzchni mamy składowe ogólnej powierzch-
niowej funkcji Greena:

— dla powierzchni „xz”

$$G_{Ph}^{Txy}(R^{xy}, t; Q^{\xi\eta}, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1}}^{\infty} D_{n,m}^{Txy} e^{-\lambda_{n,m} \frac{\lambda_z}{c_p} [J^u(t-\tau)]} \frac{4}{ab} W(x, \xi; y, \eta) \right\} \quad (5)$$

przy stałej całkowania

$$D_{n,m}^{T_{xy}} = \iint_{F_{xy}} T_{0b} \sqrt{\frac{4}{ab}} W(\xi, \eta) d\xi d\eta. \quad (6)$$

— dla powierzchni „yz”

$$G_{Ph}^{T_{yz}}(R^{yz}, t; Q^{\eta\xi}, \tau) = \left\{ \sum_{m=1}^{\infty} D_{m,k}^{T_{yz}} e^{-\lambda_{n,k} \frac{\lambda_z}{c_p} [y^y(t-\tau)]} \frac{4}{bc} W(y, \eta; z, \beta) \right\} \quad (7)$$

przy stałej całkowania

$$D_{m,k}^{T_{yz}} = \iint_{F_{yz}} T_{0b} \sqrt{\frac{4}{bc}} W(\eta, \beta) d\eta d\beta. \quad (8)$$

— dla powierzchni „xz”

$$G_{Ph}^{T_{xz}}(R^{xz}, t; Q^{\xi\beta}, \tau) = \left\{ \sum_{n=1}^{\infty} D_{n,k}^{T_{xz}} e^{-\lambda_{n,k} \frac{\lambda_x}{c_p} [x^x(t-\tau)]} \frac{4}{ac} W(x, \xi; z, \beta) \right\} \quad (9)$$

przy stałej całkowania

$$D_{n,k}^{T_{xz}} = \iint_{F_{xz}} T_{0b} \sqrt{\frac{4}{ac}} W(\xi, \beta) d\xi d\beta. \quad (10)$$

Pełna powierzchniowa funkcja Greena ma postać:

$$G_{Ph}^{TS}(R, t; Q, \tau) = 2[G_{Ph}^{T_{xy}}(R^{xy}, t; Q^{\xi\eta}, \tau) + G_{Ph}^{T_{yz}}(R^{yz}, t; Q^{\eta\beta}, \tau) + \\ + G_{Ph}^{T_{xz}}(R^{xz}, t; Q^{\xi\beta}, \tau)] \quad (11)$$

Wprowadza się teraz:

— funkcję źródłową dla warunków początkowych z równania (A^{0.10.1.2}), w punkcie Z(x, y, z, t)

$$T^z(t) = T_0 - \frac{H_c}{c_p} \ln \left| \frac{C^z(t)}{C_0} \right| \quad (12)$$

oraz

— funkcję brzegową dla warunków brzegowych z równania (B^{0.20.1.2}), w punkcie R(x, y, z, t)

$$T^b(t) = T_{0b} \pm \int_0^t U_{TC}(t) dt \quad (13)$$

na sformułowanie pełnego rozwiązania zagadnienia zmian temperatury od zjawiska przewodnictwa cieplnego w obecności warunków początkowych i brzegowych

$$T(Z, R, t) = \int_0^t \iiint_{\Omega_R} G_{Ph}^T(Z, t; Q, \tau) \left[T_0 - \frac{H_B}{c_p} \ln \left| \frac{C^z(\tau)}{C_0} \right| \right] d\xi d\eta d\beta d\tau +$$

$$+ \int_0^t \iiint_{F_R} G_{Ph}^{TS}(R, t; Q, \tau) \left[T_{0b} \pm \int_0^t U_{TC}(t) dt \right] d\xi d\eta d\beta d\tau. \quad (14)$$

3^o.2. Zjawiskowe rozwiązania dla pola wirowego

3^o.2.1. Pole wektora prędkości roztwór—roztwór (zmienny punkt $Q'(\xi', \eta', \beta', \tau)$)

Punktem wyjścia do naszych rozważań jest funkcja Greena dla pola wektora prędkości krystalizującego roztworu w postaci ($A^0.1^0.3.1.19$) mająca postać:

$$G_{Ph}^V(Z, t; Q', \tau) = \left\{ \frac{1}{N_1} \left[C_0 - \int_0^{t \rightarrow (t-\tau)} v_B G dt \right] \Gamma_G^W(Z, Q') \right\} \quad (1)$$

gdzie stała:

$$N_1 = \iiint_{\Omega_R} \Gamma_G^W(Z, Q') v_0(Q') d\xi' d\eta' d\beta' \quad (2)$$

Tensor Greena zagadnienia tematowego spełnia:

$$\Gamma_G^W(Z, Q') = \Gamma_G(Z, Q') \quad \text{dla } Z \equiv R \in F'_{(a,b,c)} \in F_R \quad (3)$$

— przypadek anizotropowy, patrz Dodatek 2

$$\Gamma_G^W(Z, Q') = \Gamma_{GA}(Z, Q') \quad \text{dla } Z \equiv R \in F'_{(a,b,c)} \in F_R. \quad (4)$$

Funkcja Greena (1) posiada następujące składowe powierzchniowe

— dla powierzchni „xy”

$$G_{Ph}^{vxy}(R^{xy}, t; Q'^{\xi'\eta'}, \tau) = \left\{ \frac{1}{N_1^{xy}} \left[C_{0b} \pm \int_0^{t \rightarrow (t-\tau)} S_C(t) dt \right] \Gamma_G^{Wxy}(R, Q'^{\xi'\eta'}) \right\} \quad (5)$$

gdzie

$$N_1^{xy} = \iint_{F_{Rxy}} \Gamma_G^{Wxy}(R^{xy}, Q'^{\xi'\eta'}) v_0^{xy}(Q'^{\xi'\eta'}) d\xi' d\eta' \quad (6)$$

— dla powierzchni „yz”

$$G_{Ph}^{vyz}(R^{yz}, t; Q'^{\eta'\beta'}, \tau) = \left\{ \frac{1}{N_1^{yz}} \left[C_{0b} \pm \int_0^{t \rightarrow (t-\tau)} S_C(t) dt \right] \Gamma_G^{Wyz}(R^{yz}, Q'^{\eta'\beta'}) \right\} \quad (7)$$

gdzie

$$N_1^{yz} = \iint_{F_{xyz}} \Gamma_G^{Wyz}(R^{yz}, Q'^{\eta'\beta'}) v_0^{yz}(Q'^{\eta'\beta'}) d\eta' d\beta' \quad (8)$$

— dla powierzchni „xz”

$$G_{Ph}^{vzx}(R^{xz}, t; Q'^{\xi'\beta'}, \tau) = \left\{ \frac{1}{N_1^{xz}} \left[C_{0b} \pm \int_0^{t \rightarrow (t-\tau)} S_c(t) dt \right] \Gamma_G^{Wxz}(R^{xz}, Q'^{\xi'\beta'}) \right\} \quad (9)$$

gdzie

$$N_1^{xz} = \iint_{F_{xz}} \Gamma_G^{Wxz}(R^{xz}, Q'^{\xi'\beta'}) v_0^{xz}(Q'^{\xi'\beta'}) d\xi' d\beta'. \quad (10)$$

W oparciu o powyższe rozważania pełna postać funkcji Greena dla zagadnień powierzchniowych powierzchnia — czas jest:

$$G_{Ph}^{vS}(R, t; Q', \tau) = 2[G_{Ph}^{vy}(R^{xy}, t; Q'^{\xi'\eta'}, \tau) + G_{Ph}^{vz}(R^{yz}, t; Q'^{\eta'\beta'}, \tau) + G_{Ph}^{vx}(R^{xz}, t; Q'^{\xi'\beta'}, \tau)] \quad (11)$$

Kontynuując dalsze rozważania przypomnijmy sobie:

— funkcję źródłową pola wektora prędkości roztwór — roztwór dla warunków początkowych w punkcie $Z(x, y, z, t)$, wzór $A^0.1^0.1.3$

$$v^z(t) = v_0 + gt \quad (12)$$

— funkcję brzegową pola wektora prędkości roztwór brzeg dla warunków brzegowych w punkcie $K(x, y, z, t)$, wzór $B^0.2^0.1.3$

$$v^b(t) = v_0 \pm \int_0^t W_{vc}(t) dt \quad (13)$$

tworzące wraz z odpowiadającymi im funkcjami Greena pełne rozwiązanie dla pola wektora prędkości roztworu w obecności warunków początkowych i brzegowych

$$v(Z, R, t) = \int_0^t \iiint_{\Omega_z} G_h^v(Z, t; Q', \tau) (v_0 + g\tau) d\xi' d\eta' d\xi' d\tau + \int_0^t \iint_{F_x} G_h^v(R, t; Q', \tau) \cdot [v_{0b} \pm \int_0^{t \rightarrow \tau} W_{vc}(t) dt] d\xi' d\eta' d\xi' d\tau \quad (14)$$

oraz dodatkowo obowiązują zależności

$$\frac{DC}{Dt} = -C \operatorname{div} v \quad (15)$$

$$C \frac{DT}{Dt} = -CT \text{div}. \quad (16)$$

3⁰.3. Kompletne rozwiązanie zarodkowania pierwotnego przy istnieniu warunków początkowych i brzegowych

Procedura kompletnego rozwiązania dla procesu zarodkowania pierwotnego jest podobna jak przy obecności tylko warunków początkowych w rozdziale A⁰.1⁰.4. i bazuje na wektorach stanu o współrzędnych: masa, energia i pęd, a uwzględnienie warunków brzegowych i pamięci brzegowej (proces także zachodzi na brzegu) daje:

Kompletna postać rozwiązania	=	Funkcje źródłowe w punkcie Z(x, y, z, t)	+	
$\begin{bmatrix} \text{Stężenie} \\ \text{Temperatura} \\ \text{Pole wektora prędkości} \end{bmatrix}$		$\begin{bmatrix} \text{Obniżenie stężenia} \\ \text{(formuła A}^0.1^0.1.1) \\ \text{Obniżenie temperatury} \\ \text{(formuła A}^0.1^0.1.2) \\ \text{Przyrost pola wektora prędkości} \\ \text{(formuła A}^0.1^0.1.3) \end{bmatrix}$	+	
Funkcje brzegowe w punkcie R(x, y, z, t) (Istnieje pamięć brzegowa)		Dyfuzja Z ← Q → R		
$\begin{bmatrix} \text{Zmiany stężenia} \\ \text{(formuła B}^0.2^0.1.1) \\ \text{Zmiany temperatury} \\ \text{(formuła B}^0.2^0.1.2) \\ \text{Zmiany pola wektora prędkości} \\ \text{(formuła B}^0.2^0.1.3) \end{bmatrix}$	+	$\begin{bmatrix} \text{Zmiany stężenia} \\ \text{(formuła B}^0.3^0.1.1.12) \\ \text{Zmiany temperatury} \\ \text{(formuła B}^0.3^0.1.2.1) \\ * \end{bmatrix}$	+	
Przewodnictwo cieplne Z ← Q → R		Pole wektora prędkości Z ← Q' → R		
$\begin{bmatrix} * \\ \text{Zmiany temperatury} \\ \text{(formuła B}^0.3^0.1.3.14) \\ * \end{bmatrix}$	+	$\begin{bmatrix} \text{Ciągłość stężenia} \\ \text{(formuła B}^0.3^0.2.1.15) \\ \text{Ciągłość temperatury} \\ \text{(formuła B}^0.3^0.2.1.16) \\ \text{Zmiany pola wektora prędkości} \\ \text{(formuła B}^0.3^0.2.1.14) \end{bmatrix}$		(1)

* — współrzędne wektora stanu, w których nie ma związków, między zjawiskami fizycznymi.

4⁰. Rozwiązanie zjawisk fizycznych przy istnieniu warunków początkowych i brzegowych dla procesu wzrostu kryształów.

4⁰.1. Rozwiązania zjawisk fizycznych dla pól potencjalnych.

4⁰.1.1. Dyfuzyjny transport stężenia roztwór—kryształ (zmienny punkt Q(ξ, η, β, τ)).

Funkcja Greena dla dyfuzyjnego transportu stężenia roztwór—kryształ dla warunków początkowych jest dana równaniem (A⁰.2⁰.2.1.12)

$$G_{Ph}^E(K, t; Q, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^E e^{-\lambda_{n,m,k} \frac{D_E}{\rho_E} C(t-\tau)} \frac{8}{abc} W(x, \xi; y, \eta; z, \beta) \right\} \quad (1)$$

przy stałej całkowania danej wzorem (A⁰.2⁰.2.1.9)

$$D_{n,m,k}^E = \iiint_{\Omega_E} E_0 \sqrt{\frac{8}{abc}} W(\xi, \eta, \beta) d\xi d\eta d\beta. \quad (2)$$

Założenie ZG2, o prostopadłościennym kształcie powierzchni F_{Rt} pozwala na otrzymanie funkcji Greena dla poszczególnych powwwierzchni związanych z punktem $P(x, y, z, t)$

— dla powierzchni „xy”

$$G_{Ph}^{Exy}(P^{xy}, t; Q^{\xi\eta}, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1}}^{\infty} D_{n,m}^{Exy} e^{-\lambda_{n,m} \frac{D_E}{\rho_E} C(t-\tau)} \frac{4}{ab} W(x, \xi; y, \eta) \right\} \quad (3)$$

przy stałej całkowania

$$D_{n,m}^{Exy} = \iint_{F_{Rxy}} E_{0b} \sqrt{\frac{4}{ab}} W(x, \xi; y, \eta) d\xi d\eta \quad (4)$$

— dla powierzchni „yz”

$$G_{Ph}^{Eyz}(P^{yz}, t; Q^{\eta\xi}, \tau) = \left\{ \sum_{\substack{m=1 \\ k=1}}^{\infty} D_{m,k}^{Eyz} e^{-\lambda_{m,k} \frac{D_E}{\rho_E} C(t-\tau)} \frac{4}{abc} W(y, \eta; z, \beta) \right\} \quad (5)$$

przy stałej całkowania

$$D_{m,k}^{Eyz} = \iint_{F_{Ryz}} E_{0b} \sqrt{\frac{4}{bc}} W(y, \eta; z, \beta) d\eta d\beta \quad (6)$$

— dla powierzchni „xz”

$$G_{Ph}^{Exz}(P^{xz}, t; Q^{\xi\beta}, \tau) = \left\{ \sum_{\substack{n=1 \\ k=1}}^{\infty} D_{n,k}^E e^{-\lambda_{n,k} \frac{D_E}{\rho_E} C(t-\tau)} \frac{4}{ac} W(x, \xi; z, \beta) \right\} \quad (7)$$

przy stałej całkowania

$$D_{n,k}^{Exz} = \iint_{F_{Rxz}} E_{0b} \sqrt{\frac{4}{ac}} W(\xi, \beta) d\xi d\beta. \quad (8)$$

Funkcja Greena dla całej powierzchni F_{Rt} ma postać

$$G_{Ph}^{ES}(P, t; Q, \tau) = 2[G_{Ph}^{Exy}(P^{xy}, t; Q^{\xi\eta}, \tau) + G_{Ph}^{Eyz}(P^{yz}, t; Q^{\eta\beta}, \tau) + G_{Ph}^{Exz}(P^{xz}, t; Q^{\xi\beta}, \tau)]. \quad (9)$$

Przypomnijmy sobie teraz dla pełnego rozwiązania zagadnienia dyfuzji stężenia roztwór – kryształ:

— funkcję źródłową w punkcie $K(x, y, z, t)$ dla warunków początkowych z równania (A^{0.2.0.1.1}), w zagadnieniach objętość – czas

$$E^k(t) = |E_0|e^{S^t} \quad (10)$$

— funkcję brzegową w punkcie $P(x, y, z, t)$ dla warunków brzegowych z równania (B^{0.2.2.1}), w zagadnieniach powierzchnia – czas

$$E^b(t) = E_{0b} \pm \int_0^t S_{EG}(t) dt. \quad (11)$$

Tak więc w oparciu o całość powyższych rozważań można postawić pełne rozwiązanie zagadnienia transportu stężenia w strumieniu dyfuzji roztwór – kryształ w obecności warunków początkowych i brzegowych

$$E(K, P, t) = \int_0^t \iiint_{\Omega_K} G_{Ph}^E(K, t; Q, \tau) |E_0| e^{S^t} d\xi d\eta d\beta d\tau + \int_0^t \iint_{F_K} G_{Ph}^{ES}(P, t; Q, \tau) [E_{0b} \pm \int_0^{\tau} S_{EG}(t) dt] d\xi d\eta d\beta d\tau. \quad (12)$$

4^{0.1.2}. Dyfuzyjny transport entalpii roztwór – kryształ (zmienny punkt $Q(\xi, \eta, \beta, \tau)$).

Rozważania rozdziału B^{0.4.0.1.1}. dają możliwość postawienia rozwiązania dla zagadnienia dyfuzyjnego transportu entalpii w układzie roztwór – kryształ przy istnieniu warunków początkowych i brzegowych

$$T_K(K, P, t) = T_{K0}(K, t) + \frac{H_c}{c_{pK}} \left[\ln \left| \int_0^t \iiint_{\Omega_K} G_{Ph}^E(K, t; Q, \tau) |E_0| e^{S^t} d\xi d\eta d\beta d\tau \right| - \ln |E_0| \right] + \\ + T_{K0b}(P, t) + \left[\ln \left| \iint_{F_K} G_{Ph}^{ES}(P, t; Q, \tau) \left[E_{0b} \mp \int_0^{\tau} S_{EG}(t) dt \right] d\xi d\eta d\beta d\tau \right| - \right. \\ \left. - \ln |E_{0b}| \right]. \quad (1)$$

Zjawiskowe funkcje Greena dla wzoru (1) są podane wzorami:

- dla zagadnień warunków początkowych — wzór (B^{0.4}^{0.1.1.1})
- dla zagadnień warunków brzegowych — wzór (B^{0.4}^{0.1.1.9}),

4^{0.1.3}. Zmiany temperatury od przewodnictwa cieplnego roztwór–kryształ (zmien- ny punkt Q(ξ, η, β, τ)).

Funkcja Greena dla zjawiska przewodnictwa cieplnego roztwór–kryształ przy istnieniu warunków początkowych dana jest postacią wzoru (A^{0.2}^{0.2.3.12})

$$G_{Ph}^{T_K}(K, t; Q, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k}^{T_K} e^{-\lambda_{n,m,k} \frac{\lambda_K}{c_{pK} \rho_K} [J_K(t-\tau)]} \frac{8}{abc} W(x, \xi; y, \eta; z, \beta) \right\} \quad (1)$$

przy stałej całkowania, która ze wzoru (A^{0.2}^{0.2.3.9}) może być przepisana jako

$$D_{n,m,k}^{T_K} = \iiint_{\Omega_K} T_{K0} \sqrt{\frac{8}{abc}} W(\xi, \eta, \beta) d\xi d\eta d\beta \quad (2)$$

oraz

$$[J_K(t)] = \int \left[1 - \frac{1}{|E_0| e^{St}} \right] dt. \quad (3)$$

Przy opracowywaniu zagadnień powierzchniowych wprowadza się wspólną funkcję

$$[J_K^u(t)] = \int \left[1 - \frac{dt}{E_{0b} \mp \int_0^t S_{EG}(t) dt} \right] dt. \quad (4)$$

Składowe powierzchniowej funkcji Greena mają postacie:

— dla powierzchni „xy”

$$G_{Ph}^{T_K^{xy}}(P^{xy}, t; Q, \xi, \eta, \tau) = \left\{ \sum_{\substack{n=1 \\ m=1}}^{\infty} D_{n,m}^{T_K^{xy}} e^{-\lambda_{n,m} \frac{\lambda_K}{c_{pK} \rho_K} [J_K^u(t-\tau)]} \frac{4}{ab} W(x, \xi; y, \eta) \right\} \quad (5)$$

przy stałej całkowania

$$D_{n,m}^{T_K^{xy}} = \iint_{F_{Kxy}} T_{K0b} \sqrt{\frac{4}{ab}} W(\xi, \eta) d\xi d\eta \quad (6)$$

— dla powierzchni „yz”

$$G_{Ph}^{T_K^{yz}}(P^{yz}, t; Q, \eta, \beta, \tau) = \left\{ \sum_{\substack{m=1 \\ k=1}}^{\infty} D_{m,k}^{T_K^{yz}} e^{-\lambda_{m,k} \frac{\lambda_K}{c_{pK} \rho_K} [J_K^u(t-\tau)]} \frac{4}{bc} W(y, \eta; z, \beta) \right\} \quad (1)$$

przy stałej całkowania

$$D_{m,k}^{T_{\kappa}yz} = \iint_{F_{xyz}} T_{K0b} \sqrt{\frac{4}{bc}} W(\eta, \beta) d\eta d\beta \quad (2)$$

— dla powierzchni „xz”

$$G_{Ph}^{T_{\kappa}xz}(P^{xz}, t; Q, \xi^{\beta}, \tau) = \left\{ \sum_{n=1}^{\infty} D_{n,k}^{T_{\kappa}xz} e^{-\lambda_{n,k} \frac{\lambda_{\kappa}}{c_{pK} \rho_{\kappa}} [U_K^y(t-\tau)]} \frac{4}{ac} W(x, \xi, z, \beta) \right\} \quad (3)$$

przy stałej całkowania

$$D_{n,k}^{T_{\kappa}xz} = \iint_{F_{Kxz}} T_{K0b} \sqrt{\frac{4}{ac}} W(\xi, \beta) d\xi d\beta \quad (4)$$

Pełna funkcja Greena powierzchni F_{Rxx}

$$G_{Ph}^{T_{\kappa}S}(P, t; Q, \tau) = 2[G_{Ph}^{T_{\kappa}xy}(P^{xy}, t; Q^{\xi\eta}, \tau) + G_{Ph}^{T_{\kappa}yz}(P^{yz}, t; Q^{\eta\beta}, \tau) + G_{Ph}^{T_{\kappa}xz}(P^{xz}, t; Q^{\xi\beta}, \tau)]. \quad (5)$$

Wprowadzamy teraz:

— funkcję źródłową dla warunków początkowych z równania (A⁰.2⁰.1.2), w punkcie $K(x, y, z, t)$

$$T_K^k(t) = T_{K0} + \frac{H_{GK}}{c_{pK}} St \quad (6)$$

— funkcję brzegową dla warunków brzegowych z równania (B⁰.2⁰.2.2), w punkcie $P(x, y, z, t)$

$$T_K^b(t) = T_{K0b} \pm \int_0^t U_{T_{KE}}(t) dt. \quad (7)$$

Rozważania powyższe dają możliwość przedstawienia pełnego rozwiązania zagadnienia przewodnictwa cieplnego w układzie roztwór—kryształ w obecności warunków początkowych i brzegowych.

$$T_K(K, P, t) = \int_0^t \iiint_{\Omega_{\kappa}} G_{Ph}^{T_{\kappa}}(K, t; Q, \tau) \left(T_{K0} + \frac{H_{GK}}{c_{pK}} S\tau \right) d\xi d\eta d\beta d\tau + \\ + \int_0^t \iint_{F_K} G_{Ph}^{T_{\kappa}S}(P, t; Q, \tau) \left[T_{K0b} \mp \int_0^{\tau} U_{T_{KE}}(t) dt \right] d\xi d\eta d\beta d\tau. \quad (8)$$

4^o.2. Zjawiskowe rozwiązania dla pola wirowego.4^o.2.1. Pole wektora prędkości roztwór—kryształ (zmienny punkt $Q'(\xi', \eta', \beta', \tau)$).

Punktem wyjściowym do rozważań są zjawiskowe funkcje Greena dla wektora prędkości roztwór—kryształ przy obecności warunków początkowych:

— z równania (A⁰.2^o.3.1.26) funkcja Greena ma postać

$$G_{Ph(1)}^{v_K}(K, t; Q', \tau) = \left\{ E^k(t-\tau) \frac{1}{N_1^s} \Gamma_{GK}^W(K, Q') \right\} \quad (1)$$

— z równania (A⁰.2^o.3.1.27) funkcja Greena ma postać

$$G_{Ph(2)}^{v_K}(K, t; Q', \tau) = \left\{ \frac{-v^z(t-\tau)C^z(t-\tau)}{\rho_K \ln|1 - E^k(t-\tau)|N_1^s} \Gamma_{GK}^W(K, Q') \right\}. \quad (2)$$

Do rozważań bierzemy element przestrzenny wspólny dla równań (1) oraz (2) jako statyczny

$$G_{Ph}^{v_K}(K, Q') = \left\{ \frac{1}{N_1^s} \Gamma_{GK}^W(K, Q') \right\} \quad (3)$$

gdzie stała

$$N_1^s = \iiint_{\Omega_r} \Gamma_{GK}^W(K, Q') v_{K0}(Q') d\xi' d\eta' d\beta' \quad (4)$$

oraz odpowiadające tym zagadnieniom przypadki statycznego tensora Greena są:

— przypadek izotropowy, patrz Dodatek 3

$$\Gamma_{GK}^W(K, Q') = \Gamma_{GK}(K, Q') \quad \text{dla } K \equiv P \in F_{(a,b,c)}^P \in F_R \quad (5)$$

— przypadek anizotropowy, patrz Dodatek 4

$$\Gamma_{GK}^W(K, Q') = \Gamma_{GKA}(K, Q') \quad \text{dla } K \equiv P \in F_{(a,b,c)}^P \in F_R. \quad (6)$$

Funkcja Greena dla przestrzennych zagadnień (3) posiada następujące składowe powierzchniowe:

$$G_{Ph}^{v_K^{xy}}(P^{xy}; Q'^{\xi'\eta'}) = \left\{ \frac{1}{N_1^{xy}} \Gamma_{GK}^{Wxy}(P^{xy}; Q'^{\xi'\eta'}) \right\} \quad (7)$$

gdzie stała

$$N_1^{xy} = \iint_{F_{Rxy}} \Gamma_{GK}^{Wxy}(P^{xy}, Q'^{\xi'\eta'}) v_{K0}(Q'^{\xi'\eta'}) d\xi' d\eta' \quad (8)$$

— dla powierzchni „yz”

$$G_{Ph}^{v_{k,yz}}(P^{yz}; Q'^{\eta'\beta'}) = \left\{ \frac{1}{N_1^{syx}} \Gamma_{GK}^{W_{yz}}(P^{yz}; Q'^{\eta'\beta'}) \right\} \quad (9)$$

gdzie stała

$$N_1^{syx} = \iint_{F_{xy}} \Gamma_{GK}^{W_{yz}}(P^{yz}, Q'^{\eta'\beta'}) v_{K0}(Q'^{\eta'\beta'}) d\eta' d\beta' \quad (10)$$

— dla powierzchni „xz”

$$G_{Ph}^{v_{k,xz}}(P^{xz}; Q'^{\xi'\beta'}) = \left\{ \frac{1}{N_1^{xxz}} \Gamma_{GK}^{W_{xz}}(P^{xz}; Q'^{\xi'\beta'}) \right\} \quad (11)$$

gdzie stała

$$N_1^{xxz} = \iint_{F_{xz}} \Gamma_{GK}^{W_{xz}}(P^{xz}, Q'^{\xi'\beta'}) v_{K0}(Q'^{\xi'\beta'}) d\xi' d\beta'. \quad (12)$$

Pełna funkcja Greena dla zagadnień dotyczących tylko powierzchni może być zapisana jako

$$G_{Ph}^{v_{kS}}(P; Q') = 2 \left[G_{Ph}^{v_{kxy}}(P^{xy}; Q'^{\xi'\eta'}) + G_{Ph}^{v_{k,yz}}(P^{yz}; Q'^{\eta'\beta'}) + \right. \\ \left. + G_{Ph}^{v_{k,xz}}(P^{xz}; Q'^{\xi'\beta'}) \right]. \quad (13)$$

Z równania (1), widać, że część czasowa funkcji Greena dla zagadnień powierzchniowych ma postać

$$G_{Ph(1)}^{v_k}(t; \tau) = \{E^b(t - \tau)\}, \quad (14)$$

zaś z równań (13) i (14) iloczyn pozwala otrzymać

$$G_{Ph(1)}^{v_{kS}}(P, t; Q', \tau) = G_{Ph}^{v_{kS}}(P; Q') G_{Ph}^{v_k}(t; \tau) \quad (15)$$

Dla dalszych rozważań przypominamy sobie teraz:

— funkcję źródłową pola wektora prędkości roztwór—kryształ dla warunków początkowych w punkcie $K(x, y, z, t)$, wzór ($A^0.2^0.1.3$)

$$v_K^k(t) = v_{K0} + \left\{ 1 + \frac{C^z(t)[1 - E^k(t)]}{\rho_K E^k(t)} \right\} g^t \quad (16)$$

— funkcję brzegową pola wektora prędkości kryształ—brzeg dla warunków brzegowych w punkcie $P(x, y, z, t)$ z równania ($B^0.2^0.2.3$)

$$v_{K_1}^b(t) = v_{K0b} \pm \int_0^t W_{v_k E}(t) dt \quad (17)$$

otrzymując na podstawie całosci rozważań pierwszy wariant pola wektora prędkości kryształów w obecności warunków początkowych i brzegowych — bez warunku pływania

$$v_{K(1)}(K, P, t) = \int_0^t \iiint_{\Omega_K} G_{Ph(1)}^{v_K}(K, t; Q', \tau) \left\langle v_{K0} + \left\{ 1 + \frac{C^z(\tau)[1 - E^k(\tau)]}{\rho_K E^k(\tau)} \right\} g\tau \right\rangle d\xi' d\eta' d\beta' d\tau +$$

$$\int_0^t \iiint_{F_K} G_{Ph(1)}^{v_K S}(P, t; Q', \tau) \left[v_{K0b} \mp \int_0^{\tau} W_{v_K E}(t) dt \right] d\xi' d\eta' d\beta' d\tau. \quad (18)$$

Uwaga U16. Patrz Uwaga U8.

Drugi wariant wykorzystuje warunek pływania kryształów wewnątrz roztworu. Dla tego przypadku część czasowa funkcji Greena ma postać:

$$G_{Ph(2)}^{v_K}(t; \tau) \left\{ \frac{-v^b(t-\tau)C^b(t-\tau)}{\rho_K \ln|1 - E^b(t-\tau)|N_1^b} \right\} \quad \text{oraz} \quad N_1^b = v_K^b(K, 0) \quad (19)$$

Z równań (13) oraz (19) otrzymujemy pełną powierzchniową funkcję Greena:

$$G_{Ph(2)}^{v_K S}(P, t; Q', \tau) = G_{Ph}^{v_K S}(P; Q') G_{Ph(2)}^{v_K}(t; \tau). \quad (20)$$

Dla drugiego wariantu zagadnienia tematowego warunek pływania modyfikuje:

— funkcję źródłową pola wektora prędkości kryształów dla warunków początkowych w punkcie $K(x, y, z, t)$ do postaci $(A^0.2^0.3.1.24)$

$$v_{K_1}^k(t) = v_{K0} + \left\{ \frac{\rho_K E^k(t)}{C^z(t)[1 - E^k(t)]} + 1 \right\} g t \quad (21)$$

— funkcję brzegową pola wektora prędkości kryształów dla warunków brzegowych w punkcie $P(x, y, z, t)$ daną wzorem (17) do postaci

$$v_{K_1}^b(t) = v_{K0b} \mp \int_0^t W_{v_K CE}(t) dt \quad (22)$$

gdzie

$$W_{v_K CE} = \frac{W_{v_K}(t)}{C^b(t)[1 - E^b(t)]}$$

Dla drugiego wariantu pole wektora prędkości kryształów wewnątrz przesyconego roztworu z uwzględnieniem warunków początkowych i brzegowych daje zapisać w postaci wzoru

$$\begin{aligned}
 v_{K(2)}(K, P, t) = & \iiint\limits_0^t \iiint\limits_{\Omega_K} G_{Ph(2)}^{v_K}(K, t; Q', \tau) \left\langle v_{K0} + \left\{ \frac{\rho_K E^k(\tau)}{C^z(\tau)[1 - E^k(\tau)]} + 1 \right\} g\tau \right\rangle d\xi' d\eta' d\beta' d\tau + \\
 & + \iiint\limits_0^t G_{Ph(2)}^{v_K S}(P, t; Q', \tau) \left[v_{K0b} \mp \int_0^{\tau} W_{v_{KE}}(t) dt \right] d\xi' d\eta' d\beta' d\tau. \quad (23)
 \end{aligned}$$

dla:

$$\frac{DE}{dt} = E \operatorname{div} v_K \quad (24) \quad \text{oraz} \quad E \frac{DT_K}{Dt} = ET_K \operatorname{div} v_K \quad (25)$$

Uwaga U17. Patrz Uwaga U8.

5⁰. Rozwiązanie równania różniczkowego cząstkowego dla procesu zarodkowania wtórnego.

Funkcja Greena dla procesu zarodkowania wtórnego przy obecności warunków początkowych dana jest wzorem (A⁰.3⁰.12)

$$G_{Ph}^L(K, t; Q', \tau) = \left\{ L_0 \Gamma_{GKA}^W(K, Q') \sin \left[\sqrt{\frac{2n\rho_K^2 K_V v_K^2}{B}} (t - \tau) \right] \right\}, \quad (1)$$

przy czym stała L_0 jest początkowym wymiarem wektora rozmiaru liniowego kryształu.

Zgodnie z założeniem ZG2, o prostopadłościennym kształcie powierzchni F_{Rt} , funkcje Greena dla poszczególnych powierzchni mają postacie związane z punktem $P(x, y, z, t)$:

— dla powierzchni „xy”

$$G_{Ph}^{Lxy}(P^{xy}, t; Q'^{\xi'\eta'}, \tau) = \left\{ L_0^{xy} \Gamma_{GKA}^{Wxy}(P^{xy}, Q'^{\xi'\eta'}) \sin \left[\sqrt{\frac{2n^{xy}\rho_K^2 K_V^{xy} v_K^{xy2}}{B^{xy}}} (t - \tau) \right] \right\} \quad (2)$$

— dla powierzchni „yz”

$$G_{Ph}^{Lyz}(P^{yz}, t; Q'^{\eta'\beta'}, \tau) = \left\{ L_0^{yz} \Gamma_{GKA}^{Wyz}(P^{yz}, Q'^{\eta'\beta'}) \sin \left[\sqrt{\frac{2n^{yz}\rho_K^2 K_V^{yz} v_K^{yz2}}{B^{yz}}} (t - \tau) \right] \right\} \quad (3)$$

— dla powierzchni „xz”

$$G_{Ph}^{Lxz}(P^{xz}, t; Q'^{\xi'\beta'}, \tau) = \left\{ L_0^{xz} \Gamma_{GKA}^{Wxz}(P^{xz}, Q'^{\xi'\beta'}) \sin \left[\sqrt{\frac{2n^{xz}\rho_K^2 K_V^{xz} v_K^{xz2}}{B^{xz}}} (t - \tau) \right] \right\} \quad (4)$$

we wzorach (2), (3) oraz (4) może być

$$v_K = v_{K(1)} \text{ wzór (B⁰.4⁰.2.1.18)} \quad \text{lub} \quad v_K = v_{K(2)} \text{ wzór (B⁰.4⁰.2.1.23)} \quad (5)$$

Pełna funkcja Greena dla zagadnień powierzchniowych ze wzorów (2), (3) oraz (4) może zapisana jako

$$G_{Ph}^{LS}(P, t; Q', \tau) = 2[G_{Ph}^{Lxy}(P^{xy}, t; Q'^{x'\eta'}, \tau) + G_{Ph}^{Lyz}(P^{yz}, t; Q'^{\eta'\beta'}, \tau) + G_{Ph}^{Lxz}(P^{xz}, t; Q'^{x'\beta'}, \tau)]. \quad (6)$$

Dla całości rozważań należy przypomnieć sobie, że:

— funkcja źródłowa dla zagadnienia zarodkowania wtórnego w obecności warunków początkowych, w punkcie $K(x, y, z, t)$ ma z równania (A⁰.3.10) postać równania

$$L^k(t) = L_0 \quad (7)$$

— funkcje brzegowe dla zagadnienia zarodkowania wtórnego przy istnieniu warunków brzegowych, ma w punkcie $P(x, y, z, t)$ na podstawie równania B⁰.2⁰.3.1) formułę

$$L^b(t) = L_{0b}. \quad (8)$$

Jako wynik powyższych rozważań dla warunków początkowych i brzegowych otrzymujemy pełne rozwiązanie dla procesu zarodkowania wtórnego w postaci

$$L(K, P, t) = \int_0^t \iiint_{\Omega_K} G_{Ph}^L(K, t; Q', \tau) L^K(t) d\xi' d\eta' d\beta' d\tau + \int_0^t \iint_{F_K} G_{Ph}^{LS}(P, t; Q', \tau) L^b(t) \cdot d\xi' d\eta' d\beta' d\tau. \quad (9)$$

Rozwiązanie (9) wchodzi poprzez funkcję $f(L)$ w zakłócenie kinetyki V_{GK} reakcji wzrostu kryształów.

6⁰. Kompletne rozwiązanie dla procesu wzrostu kryształów oraz procesu zarodkowania wtórnego jako zakłócenia kinetyki wzrostu kryształów dla warunków początkowych i brzegowych

Po uwzględnieniu warunków początkowych i brzegowych, a także pamięci brzegowej — proces zachodzi również na brzegu F_{Rt} , rozwiązanie zagadnienia tematowego dla procesu wzrostu kryształów zakłóconego w kinetyce procesem zarodkowania wtórnego ma postać:

Kompletna postać
rozwiązania

$$\begin{bmatrix} \text{Ułamek objętościowy} \\ \text{fazy krystalicznej} \\ \\ \text{Temperatura kryształów} \\ \\ \text{Pole wektora prędkości} \\ \text{kryształów} \end{bmatrix} =$$

Funkcje źródłowe w punkcie
 $K(x, y, z, t)$

$$\begin{bmatrix} \text{Przyrost ułamka objętościowego} \\ \text{fazy krystalicznej (formuła A}^0.2^0.1.1) \\ \text{oraz zakłócenie (formuła A}^0.3^0.11) \\ \\ \text{Przyrost temperatury kryształów} \\ \text{(formuła A}^0.2^0.1.2) \text{ oraz zakłócenie} \\ \text{(formuła A}^0.3^0.11) \\ \\ \text{Przyrost pola wektora prędkości} \\ \text{kryształów (formuła A}^0.2^0.1.3) \text{ oraz} \\ \text{(formuła A}^0.3^0.11) \end{bmatrix} +$$

Funkcje brzegowe w punkcie

 $P(x, y, z, t)$

(Istnieje pamięć brzegowa)

$$+ \left[\begin{array}{l} \text{Zmiany ułamka fazy krystalicznej} \\ \text{(formuła B}^0.2^0.2.2) \text{ oraz zakłócenie} \\ \text{(formuła B}^0.5^0.9) \\ \text{Zmiany temperatury kryształów} \\ \text{(formuła B}^0.2^0.2) \text{ oraz zakłócenie} \\ \text{(formuła B}^0.5^0.9) \\ \text{Zmiany pola wektora prędkości kryształów} \\ \text{(form. B}^0.2^0.2.3) \text{ oraz zakł. (form. B}^0.5^0.9) \end{array} \right] + \left[\begin{array}{l} \text{Zmiany ułamka fazy krystalicznej} \\ \text{(formuła B}^0.4^0.1.1.12) \\ \text{Zmiany temperatury kryształów} \\ \text{(formuła B}^0.4^0.1.2) \\ * \end{array} \right] +$$

Przewodnictwo cieplne

 $K \leftarrow Q \rightarrow P$

$$+ \left[\begin{array}{l} \text{Zmiany temperatury kryształów} \\ \text{(formuła B}^0.4^0.1.3.8) \\ * \end{array} \right] + \left[\begin{array}{l} \text{Ciągłość ułamka fazy krystalicznej} \\ \text{(formuła B}^0.4^0.2.1.24) \\ \text{Ciągłość temperatury kryształów} \\ \text{(formuła B}^0.4^0.2.1.25) \\ \text{Zmiany pola wektora prędkości} \\ \text{kryształów (formuła B}^0.4^0.2.1.18) \\ \text{lub (formuła B}^0.4^0.2.1.23) \end{array} \right]$$

Pole wektora prędkości

 $K \leftarrow Q' \rightarrow P$

* — współrzędne wektora stanu, w których nie ma związku między zjawiskami fizycznymi.

4. WNIOSKI KOŃCOWE DOTYCZĄCE ANALITYCZNEGO ROZWIĄZANIA KONSTYTUTYWNEGO ROZŁOŻONEGO PARAMETRYCZNIE MODELU DLA PROCESU CIĄGŁEJ KRYSTALIZACJI MASOWEJ

Całkowite rozwiązanie konstytutywnego rozłożonego parametrycznie modelu dla procesu ciągłej krystalizacji masowej obejmuje kompletne rozwiązania analityczne wszystkich jego procesów składowych:

— procesu zarodkowania pierwotnego, układ równań różniczkowych cząstkowych konstytutywnych stanu (I),

— procesu wzrostu kryształów, układ równań różniczkowych cząstkowych konstytutywnych stanu (II),

— procesu zarodkowania wtórnego, równanie różniczkowe cząstkowe wektora rozmiaru liniowego kryształu (III.1), dla następujących przypadków,

A⁰. — istnienia warunków początkowych przy rozwiązywaniach zjawisk fizycznych,

B⁰. — istnienia warunków początkowych i brzegowych przy rozwiązywaniach zjawisk fizycznych.

Rozwiązanie analityczne układów (I) oraz (II) z równaniem (III.1) było w pracy możliwe dzięki:

— wprowadzeniu wspólnej normy $N_e = 1$, a co za tym idzie tych samych wymiarów „ a , b , c ” dla czasu „ t ” w elementach $\bar{\Omega}_{(a,b,c)}$ oraz $\bar{\Omega}_{K(a,b,c)t}$, dla równań różniczkowych cząstkowych pól potencjalnych — dyfuzyjny transport stężenia i entalpii, przewodnictwo cieplne,

— konsekwentnie wprowadzonej zasady wzajemności dla izotropowych i anizotropowych niejennorodnych ośrodków z pamięciami przestrzennymi i czasowymi, STATYCZNEJ I DYNAMICZNEJ — dla pola wirowego — pole wektora prędkości [11–12],

— konstytutywnej dekompozycji układów (I) oraz (II) do; — pól potencjalnych i pola wirowego, każdy z układów,

— funkcji źródłowych i brzegowych, każdy z układów — funkcja źródłowa i funkcja brzegowa dla procesu zarodkowania wtórnego otrzymywane są z warunków technologicznych, $L_0 = L_{0b} = \text{const}$

— własności zjawiskowych funkcji Greena, sploty odpowiednich funkcji Greena z funkcjami źródłowymi i brzegowymi,

— sumowanie efektów wektorów stanu: [funkcji źródłowych] + + [ewentualnie funkcji brzegowych] + [zjawisk pól potencjalnych i wirowych przy (dla pamięci brzegu)

obecności warunków początkowych i ewentualnie brzegowych]. Pamięć na brzegu i $P \in F_{Rt}$ daje istnienie funkcji brzegowych odpowiednio dla punktów $R \in F_{Rt}$ w tym sumowaniu. Zjawiskowe rozwiązania dla pól potencjalnych i wirowych tworzą także funkcje Greena o następujących współrzędnych

— dla $(x, y, z, t; \xi^R, \eta^R, \beta^R, \tau; t)$ — w każdym zjawiskowym rozwiązaniu,

— punkty: $A; Q^R; A_t$ oznaczają: $A(x, y, z, t) \equiv Z(x, y, z, t) \equiv K(x, y, z, t)$, $Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q(\xi, \eta, \beta, \tau) \equiv Q'(\xi', \eta', \beta', \tau)$ zaś punkt $A_t \equiv Z(x=a, y=b, z=c) \equiv K(x=a, y=b, z=c, t)$ — co odpowiada dynamicznej zasadzie wzajemności [11–12].

Otrzymane analityczne rozwiązania dla procesów składowych procesu ciągłej krystalizacji masowej [18–19], [9]:

— dają przebieg zmiennych stanu procesów składowych od poszczególnych zjawisk fizykalnych,

— pozwalają skonstruować kompletną informację dla celów sterowania zjawiskami fizykalnymi procesów składowych poprzez odpowiednie warunki brzegowe z punktu widzenia wydajności i jakości krystalizacji,

— pozwalają na poszukiwanie optymalnego kształtu komory krystalizatora z punktu widzenia przebiegu procesów składowych, procesu ciągłej krystalizacji masowej,

— pozwalają na identyfikację zjawisk fizykalnych i ich współczynników,

— pokazują więzi i wpływy pomiędzy poszczególnymi zjawiskami procesów składowych ciągłej krystalizacji masowej,

— posiadają praktycznie nieograniczoną dokładność związaną z ilością wyrazów szeregów rozwiązań analitycznych wziętych do obliczeń,

— bazują na minimalnej ilości bardzo prostych pomiarów podstawowych zmiennych stanu w stanie statycznym, prostych pomiarach geometrycznych i prostych badaniach kinetycznych.

Oryginalność rozwiązań analitycznych dotyczy ich elementów:

- stałość współczynników fizykalnych — dowód Część I [1],
- konstytutywna dekompozycja układów (I) oraz (II) dla otrzymania funkcji źródłowych i brzegowych,
- rozkład jednorodnych postaci układów (I^h) oraz (II^h) na równania pól potencjalnych i wirowych,
- przyjęcie normy funkcji własnych $N_e=1$ dla unormowania amplitud wszystkich funkcji własnych, przy analitycznych rozwiązaniach zjawisk fizykalnych pól potencjalnych,
- przyjęcie stałości obrazu wzajemnościowego dla otrzymania rozwiązań części statycznej i dynamicznej odpowiednio do pól wirowych procesów zarodkowania pierwotnego i wzrostu kryształów a także dla procesu zarodkowania wtórnego,
- użycie funkcji Greena zjawisk fizykalnych dla warunków początkowych do konstrukcji zjawiskowych funkcji Greena dla warunków brzegowych tych zjawisk fizykalnych,
- konstrukcja kompletnych rozwiązań układów równań (I) oraz (II) z równaniem (III.1) na bazie sumowania odpowiednich wektorów stanu dla warunków początkowych i brzegowych,
- uwzględnienie pamięci na brzegu F_{Rt} poprzez wprowadzenie do konstrukcji kompletnych rozwiązań wektorów stanu funkcji brzegowych: w punkcie $R(x, y, z, t) \in F_{Rt}$ dla układu równań różniczkowych cząstkowych konstytutywnych stanu (I) — proces zarodkowania pierwotnego oraz w punkcie $P(x, y, z, t) \in F_{Rt}$ dla układu równań różniczkowych cząstkowych konstytutywnych stanu (II) z równaniem (III.1) — proces wzrostu kryształów zakłócony w kinetyce procesem zarodkowania wtórnego. Uwzględnienie pamięci na brzegu F_{Rt} odpowiada zagadnieniu zachodzenia danego procesu składowego procesu ciągłej krystalizacji masowej na brzegu krystalizatora,
- dowolność kształtu przestrzeni roboczej krystalizatora Ω_{Rt} .

5. DODATKI

DODATEK I

Statyczny izotropowy tensor Greena dla elementu $\bar{\Omega}_{(a,b,c)}$ przesyconego roztworu [16–17], [2–8]

Tensor statycznych deformacji elementu roztworu $\bar{\Omega}_{(a,b,c)}$ może być zapisany w postaci wzoru

$$\varepsilon_{iklm} = \alpha \delta_{ik} \delta_{lm} + \beta (\delta_{il} \delta_{km} + \delta_{im} \delta_{kl}). \quad (1)$$

Do tensora (1) przyporządkowany jest statyczny tensor Greena

$$\Gamma_G(Z, Q') = \frac{1}{8\pi} \left[\frac{1}{\alpha} \Gamma_\alpha(Z, Q') + \frac{1}{\beta} \Gamma_\beta(Z, Q') \right] \quad (2)$$

Składowa tensora $\Gamma_\alpha(Z, Q')$ odpowiada operacji [graddivv] i ma postać szczegółową

$$\Gamma_\alpha(Z, Q') = \frac{1}{|r|} I - \frac{r \otimes r}{|r|^3} \quad (3)$$

gdzie: I — macierz diagonalna, $r = r_Z - r_Q$, — promień działania sił wiskozyjnych w przesyconym roztworze [m].

Składowa tensora $\Gamma_\beta(Z, Q')$ odpowiadająca operacji [—rotrotv]

$$\Gamma_\beta(Z, Q') = \frac{1}{|r|} I + \frac{r \otimes r}{|r|^3} \quad (4)$$

Wprowadzając postacie współczynników α oraz β

$$\alpha = a + b \text{ oraz } \beta = a \text{ przy } a = \eta_w \text{ oraz } b = \frac{1}{3} \eta_w \quad (5)$$

otateczna postać tensora Greena (2) jest następująca

$$\Gamma_G(Z, Q') = \frac{1}{32\pi\eta_w} \left[7 \frac{1}{|r|} I + \frac{r \otimes r}{|r|^3} \right] \quad (6)$$

dla η_w — współczynnik lepkości dynamicznej $\left[\frac{Ns}{m^2} \right]$.

DODATEK 2.

Statyczny anizotropowy tensor Greena dla elementu $\bar{\Omega}_{(a,b,c)}$ przesyconego roztworu [2–8]

Metodą małego parametru można wyprowadzić anizotropowy tensor Greena dla elementu $\bar{\Omega}_{(a,b,c)}$ traktowanego jako miejscowa anizotropia. Formuła szeregu tego anizotropowego tensora Greena jest przedstawiona poniżej

$$\Gamma_{GA}(Z, Q') = \Gamma_k \ln(Z, Q') = \Gamma_0^p(Z, Q') * \varepsilon_{pq}(Z, Q') \Gamma_{k-1}^{qn}(Z, Q') \quad (1)$$

dla $k=1(1)w$ oraz

$$\Gamma_0^p(Z, Q') = \Gamma_G(Z, Q') \quad (2)$$

w — liczba mówiąca o potędze parametru.

DODATEK 3

Statyczny izotropowy tensor Greena dla kryształu wewnątrz warstwy fluidalnej lub adsorbcyjnej — rzadki przypadek [16–17], [2–8]

Biorąc pod uwagę, że dla kryształu może zająć

$$\hat{\alpha} = \hat{a} + \hat{b} \text{ oraz } \hat{\beta} = \hat{a} \text{ przy } \hat{a} = \eta_w E \text{ oraz } \hat{b} = \frac{1}{3} \eta_w E \quad (1)$$

W oparciu o Dodatek 1 można napisać izotropowy tensor Greena jako

$$\Gamma_{GK}(K, Q') = \frac{1}{32\pi\eta_w E} \left[7 \frac{1}{|\mathbf{r}_k|} I + \frac{\mathbf{r}_k \otimes \mathbf{r}_k}{|\mathbf{r}_k|^3} \right] \quad (2)$$

gdzie: I — macierz diagonalna, η_w — współczynnik lepkości dynamicznej $\left[\frac{Ns}{m^2} \right]$,
 $E = E(x, y, z)$ — rozkład statyczny ułamka fazy krystalicznej, $\mathbf{r}_k = \mathbf{r}_K - \mathbf{r}_Q$,
 — promień działania sił wiskozyjnych [m].

DODATEK 4

Statyczny anizotropowy tensor Greena dla kryształu wewnątrz warstwy fluidalnej lub adsorbcyjnej [2–8]

W oparciu o Dodatek 2 oraz Dodatek 3 można napisać tematowy tensor Greena w następującej formie

$$\Gamma_{GKA}(K, Q') = \Gamma_{ln}(K, Q') = \Gamma_{lp}(K, Q') * \hat{\varepsilon}_{pq}(K, Q') \Gamma_{k-1}^{qn}(K, Q') \quad (1)$$

dla $k = 1(1)w$ oraz

$$\Gamma_{lp}(K, Q') = \Gamma_{GK}(K, Q') \quad (2)$$

w — liczba mówiąca o potędze małego parametru. Stałe do części izotropowej tensora (1) w postaci (2) bierzemy z Dodatku 3.

DODATEK 5

Anizotropowy tensor Greena dla zagadnień stycznych procesu zarodkowania wtórnego [2–8]

Wprowadzamy współczynniki materiałowe

$$\tilde{\alpha} = \tilde{a} + \tilde{b} \text{ oraz } \tilde{\beta} = \tilde{a} \text{ przy } \tilde{a} = h \text{ oraz } \tilde{b} = \frac{h\bar{k}}{\bar{k}-2}. \quad (1)$$

W oparciu o Dodatek 4 można postawić anizotropowy tensor Greena dla procesu zarodkowania wtórnego

$$\Gamma_{GKA}^W(K, Q') = \Gamma_k^W(K, Q') = \Gamma_0^W(K, Q') * \overset{\Delta}{\varepsilon}_{pq}(K, Q') \cdot \Gamma_{k-1}^{qn}(K, Q') \quad (2)$$

przy $k=1(1)w$, oraz

$$\Gamma_0^W(K, Q') = \frac{1}{8\pi} \left[\frac{1}{\tilde{a} + \tilde{b}} \Gamma_\alpha(K, Q') + \frac{1}{\tilde{a}} \Gamma_\beta(K, Q') \right] \quad (3)$$

gdzie tensory $\Gamma_\alpha(K, Q')$ oraz $\Gamma_\beta(K, Q')$ są definiowane w Dodatku 1, dla warunku $\mathbf{r} = \mathbf{r}_k$.

Pole deformacji $\bar{\Omega}_{K(a,b,c)}$ dla stanu statycznego przy rozważaniu procesu wtórnego jest

$$\tilde{\varepsilon}_{iklm} = \tilde{\alpha} \delta_{ik} \delta_{lm} + \tilde{\beta} (\delta_{il} \delta_{km} + \delta_{im} \delta_{kl}). \quad (4)$$

SYMBOLIKA

C	— stężenie krystalizującego roztworu	kg/m ³
C_s	— stężenie nasycenia dla krystalizującego roztworu	kg/m ³
C_0	— stężenie początkowe dla krystalizującego roztworu, proces zaczyna się dla $C_0 = C_s$	kg/m ³
V_B	— intensywność reakcji zarodkowania pierwotnego w lokalnym elemencie $\Omega_{(a,b,c)t}$ [m, m, m, s]	mol/s
G	— własna masa molowa dla reakcji zarodkowania pierwotnego	kg/(mol m ³)
H_B	— entalpia własna reakcji zarodkowania pierwotnego	J/kg
T	— temperatura krystalizującego roztworu	°K
T_0	— temperatura początkowa krystalizującego roztworu	°K
D_c	— współczynnik dyfuzji roztwór—roztwór	m ² /s
λ_c	— współczynnik przewodnictwa cieplnego roztwór—roztwór	W/(m°K)
H_c	— entalpia własna stężenia krystalizującego roztworu	J/kg
c_p	— ciepło właściwe krystalizującego roztworu	J/(kg°K)
$\mathbf{v}$	— pole wektora prędkości krystalizującego roztworu	m/s
$\mathbf{v}_0$	— początkowe pole wektora prędkości krystalizującego roztworu (od ciągłego mieszania)	m/s
η_w	— współczynnik lepkości dynamicznej	Ns/m ²
g	— stała grawitacyjna	m/s ²
$f(L)$	— funkcja liniowego rozmiaru kryształu w kinetyce wzrostu kryształów	m ²
E	— ułamek objętościowy fazy krystalicznej wewnątrz przesyconego roztworu [objętość kryształu/ $\bar{\Omega}_{K(a,b,c)t}$]	
$(1-E)$	— ułamek objętościowy roztworu	
E_0	— początkowa wartość ułamka objętościowego fazy krystalicznej wewnątrz przesyconego roztworu	

ρ_K	— gęstość właściwa kryształu	kg/m ³
L	— wektor liniowy rozmiar kryształu	m
L_0	— wartość początkowa wektora liniowego rozmiaru kryształu	m
V_{GK}	— intensywność reakcji wzrostu kryształów w lokalnym elemencie $\bar{\Omega}_{K(a,b,c)t}$ [m, m, m, s]	mol/s
G_K	— własna masa molowa dla reakcji wzrostu kryształów	kg/(mol m) ³
H_{GK}	— entalpia własna reakcji wzrostu kryształów	J/kg
T_K	— temperatura kryształów	°K
T_{K0}	— temperatura początkowa kryształów	°K
D_K	— współczynnik dyfuzji roztwór — kryształ	m ² /s
λ_K	— współczynnik przewodnictwa cieplnego roztwór — kryształ	W/(m°K)
c_{pK}	— ciepło właściwe rosnących kryształów	J/(kg°K)
v_K	— pole wektora prędkości rosnących kryształów	m/s
v_{K0}	— początkowe pole wektora prędkości rosnących kryształów (od ciągłego mieszania)	m/s
n	— statystyczna liczba zderzeń kryształów na powierzchnię kryształu	liczba/m ²
$K_{\bar{v}}$	— objętościowy współczynnik kształtu kryształu	
B	— funkcja tarcia kryształów	[kg/m ³] ²
$\bar{k}$	— stała Poissona dla kryształów	
h	— moduł ściskania kryształu	N/m ²
$\bar{\Omega}_{(a,b,c)t}$	— lokalnie wybrany, rozłożony parametrycznie element objętość — czas dla reakcji zarodkowania pierwotnego $x=a, y=b, z=c$ dla czasu „t” [m, m, m, s], zjawiskowo domknięty	m ³ — s
$F_{(a,b,c)t}$	— zewnętrznie zorientowana powierzchnia dla elementu $\bar{\Omega}_{(a,b,c)t}$	m ² — s
$n^{(z)}$	— wektor orientacyjny, zewnętrzny, normalny do powierzchni $F_{(a,b,c)t}$ mający cyrkulację decydującą o znakach zjawisk fizycznych dla pól potencjalnych i wirowych przy sumowaniu ich wpływów na określone współrzędne wektorów stanu w bilansach masy, energii i pędu	
$\bar{\Omega}_{K(a,b,c)t}$	— lokalnie wybrany, rozłożony parametrycznie element objętość — czas dla reakcji wzrostu kryształów zakłóconej w kinetyce procesem zarodkowania wtórnego, mający wymiary $x=a, y=b, z=c$ [dla czasu „t” [m, m, m, s], zjawiskowo domknięty	m ³ — s
$F_{K(a,b,c)t}$	— zewnętrznie zorientowana powierzchnia dla elementu $\bar{\Omega}_{K(a,b,c)t}$	m ² — s
$n_K^{(z)}$	— wektor orientacyjny, zewnętrzny, normalny do powierzchni $F_{K(a,b,c)t}$ mający cyrkulację decydującą o znakach zjawisk fizycznych dla pól potencjalnych i wirowych	

przy sumowaniu ich wpływów na określone współrzędne wektorów stanu w bilansach masy, energii i pędu. Cyrkulacja wektora $\mathbf{n}^{(z)}$ jest przeciwna do cirkulacji wektora $\mathbf{n}_K^{(z)}$.

- Ω_{Rt} — przestrzeń robocza — czas dla krystalizatorów $m^3 - s$
 F_{Rt} — zewnętrznie zorientowana powierzchnia dla Ω_{Rt} $m^2 - s$
 $\mathbf{n}_R^{(z)}$ — wektor orientacyjny, zewnętrzny, normalny do powierzchni F_{Rt} , którego cirkulacja zależy od zagadnień sterowania na brzegu daną powierzchnią F_{Rt} , będący podstawą do określenia relacji między cirkulacjami:

$$\mathbf{n}^{(z)} \xleftarrow{(\tilde{A})} \mathbf{n}_R^{(z)} \xrightarrow{(\tilde{B})} \mathbf{n}_K^{(z)}$$

- $(\tilde{A})$ — ważną rolę odgrywa proces zarodkowania pierwotnego, znaki „ $\pm$ ” w zagadnieniach funkcji brzegowych,
 $(\tilde{B})$ — ważne są procesy wzrostu kryształów zakłócony w kinetyce procesem zarodkowania wtórnego, znaki „ $\pm$ ” w zagadnieniach funkcji brzegowych.

Wybór wariantu $(\tilde{A})$ lub $(\tilde{B})$ zależy od wymaganych własności produktu procesu ciągłej krystalizacji masowej.

Uwaga U14. Elementy $\bar{\Omega}_{(a,b,c)t}$ oraz $\bar{\Omega}_{K(a,b,c)t}$ są zjawiskowo domknięte tzn. wpływy poszczególnych zjawisk na współrzędne wektorów stanu odpowiadające tym elementom mogą być rozpatrywane w sposób oddzielny.

BIBLIOGRAFIA

1. W. Niemiec: *Równania konstytutywne ośrodka ciągłego dla procesu ciągłej krystalizacji masowej*. Kw. Elektr. i Telem. 1993, t.39, z.2 cz. I
2. W. Niemiec: *Model matematyczny o parametrach rozłożonych procesu ciągłej krystalizacji masowej do celów sterowania adaptacyjnego*. Praca doktorska. Politechnika Śląska w Gliwicach 1979
3. W. Niemiec: *A mathematical model of the distributed parameters of the continuous mass crystallization process for the adaptive control. Part II. The complete information for the phenomenal distributed parameter control and adaptive features of the continuous mass crystallization process*. Poznańskie Towarzystwo Przyjaciół Nauk, Wydział Nauk Technicznych, Prace Komisji Automatyki i Informatyki, Tom XV — 1989, s. 119 — 162
4. W. Niemiec: *On the constitutive theory of modelling and information for distributed parameter control of the continuous mass crystallization process: Part II: Complete information for phenomenal distributed parameter control of continuous mass crystallization process*. Third Int'l Conference on Liquid Metal Engineering and Technology in Energy Production, Oxford, England, 9 — 13 April, 1984, Paper No. 192
5. W. Niemiec: *On modelling and information construction for control of distributed parameter chemical processes in fluid phase*. 3^o IFAC Symposium "Control of Distributed Parameter Systems", Toulouse, France, 29 VI — 2 VII, 1982, Session 23.
6. 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 proces-*

- ses in gas, fluid and solid phase. 7th Miami Int'l Conference on Alternative Energy Sources, Miami Beach, Florida, USA, Session "Hydrocarbons/Energy Transfer", pp. 589–598
7. W. Niemiec: *On modelling information construction and optimization for distributed parameter control of multicomponent chemical processes in fluid— and gasphase. Part II. Information construction and optimization for distributed parameter control of multicomponent chemical processes in fluid— and gasphase.* 21st Annual Meeting of the Society of Engineering Science, Virginia Polytechnic Institute and University, Blacksburg, USA, 15–17 October, 1984, 16AM6–9
 8. W. Niemiec: *A mathematical dynamical model of distributed parameters of the continuous mass crystallization process. Part II. The solution of partial differential constitutive state equations for the continuous mass crystallization process.* 4th SMIRT Conference, Paris, France, 7–11 August, 1984, Session/paper L 1/7
 9. W. Niemiec: *Struktura sterowania adaptacyjnego procesów przemysłowych.* Zeszyty Naukowe Politechniki Śląskiej. „Automatyka” Zeszyt 51, 1981
 10. C.A. Coulson, A. Jeffrey: *Fale modele matematyczne.* WNT, Warszawa 1982
 11. W. Nowacki: *Dynamiczne zagadnienia termosprężystości.* PWN, Warszawa 1966
 12. W. Nowacki, Z. Olesiak: *Termodyfuzja w ciałach stałych.* PWN, Warszawa 1992
 13. B. Średniawa: *Hydrodynamika i teoria sprężystości.* PWN, Warszawa 1977
 14. T. Trajdos: *Matematyka dla inżynierów.* PWN, Warszawa 1974
 15. A.N. Tichonow, A.A. Samarski: *Równania fizyki matematycznej.* PWN, Warszawa 1963
 16. W.I. Smirnow: *Matematyka wyższa.* PWN, Warszawa 1960
 17. W. Kupradze: *Wybrane zagadnienia teorii sprężystości i termosprężystości.* Wyd. PAN, Wrocław—Warszawa—Kraków 1970
 18. Wł. Findeisen: *Wielopoziomowe układy sterowania.* WNT, Warszawa 1974
 19. S. Węgrzyn: *Podstawy automatyki.* PWN, Warszawa 1974

W. NIEMIEC

THE CONSTRUCTION OF CONSTITUTIVE EQUATIONS OF THE CONTINUOUS MEDIUM FOR THE CONTINUOUS MASS CRYSTALLIZATION PROCESS

PART (II). ANALYTICAL SOLUTION OF CONSTITUTIVE EQUATIONS OF THE CONTINUOUS MEDIUM FOR THE CONTINUOUS MASS CRYSTALLIZATION PROCESS DOCTORAL DISSERTATION – PART II

S u m m a r y

Deduced in Part I of this article original constitutive distributed parameter mathematical model for complete continuous mass crystallization process and containing all its composite processes:

- the primary nucleation process,
- the crystal growth process,
- the secondary nucleation process,

has been solved in an original and analytical way. Two cases of this analytical solution by:

- A° — the existence of the initial conditions,
- B° — the existence of the initial and boundary conditions,

for the solutions of physical phenomena of the potential and rotational fields in relations solution—solution and solution—crystal and the mechanical energies of the crystal have been considered. The wide scope of applications of the analytical solutions for the separate processes and whole continuous mass crystallization process was discussed.

Metoda identyfikacji parametrycznej z wygładzaniem danych początkowych

ANNA TRZMIELAK – STANISŁAWSKA

Instytut Informatyki, Uniwersytet Wrocławski

JANUSZ HALAWA

Instytut Cybernetyki Technicznej, Politechnika Wrocławska

Otrzymano 1992.02.26

Autoryzowano do druku 1992.04.29

W pracy omówiono zaproponowaną przez autorów metodę optymalizacji parametrycznej z wygładzaniem danych początkowych. Do wygładzania wykorzystano estymator jądrowy *cross-validation* (krzyżowej walidacji). Zaproponowaną metodę zastosowano do wyznaczania transmisji uproszczonych. Przedstawiono przykłady stosowania metody.

W pracy przedstawiono wyniki zastosowania metody optymalizacji parametrycznej [1] z wygładzaniem danych pomiarowych za pomocą estymatora jądrowego *cross-validation* [3]. W metodzie, zaproponowanej przez autorów, przyjmuje się, że dane odpowiedzi obiektu mierzone w dyskretnych punktach czasu obciążone są błędami losowymi. Identyfikowane są współczynniki transmitancji. Przyjmuje się transmitancje możliwie niskich rzędów. Wykorzystuje się je do projektowania suboptymalnych układów sterowania oraz budowy szybko działających układów automatycznego sterowania (stąd poszukiwanie możliwie niskiego rzędu transmitancji), a także w badaniach symulacyjnych. Proponowaną metodę można również zastosować do doboru nastaw regulatorów [2].

W metodzie przyjmuje się, że odpowiedź identyfikowanego obiektu obciążona jest błędem losowym o rozkładzie normalnym. Zakłada się, że odpowiedź w mierzonych punktach t_k jest dana wzorem:

$$\bar{y}(t_k) = y_l(t_k) + e_k,$$

gdzie e_k jest błędem losowym o rozkładzie normalnym, zaś l jest rzędem przyjętego modelu. Mając dane t_k i $\bar{y}_l(t_k)$ ($k = 1, 2, \dots, M$; $l = 1, 2, \dots$) oblicza się estymującą funkcję $y_l(t)$ wykorzystując estymator *cross-validation CV*.

Estymator ma postać:

$$y_i(t) = \frac{\sum_{i=1}^M \bar{y}_i(t_i) K((t-t_i)/h)}{\sum_{i=1}^M K((t-t_i)/h)} \quad i \bar{y}_i(t_i), \quad (1)$$

gdzie funkcja $K(t)$, zwana jądrem, jest gęstością pewnego rozkładu, a h nazywa się szerokością okna.

W metodzie przyjęto

$$K(t) = \begin{cases} 0 & \text{dla } t \leq -1 \\ 15(1-t(2-t))/16 & \text{dla } -1 \leq t \leq 1 \\ 0 & \text{dla } t \geq 1, \end{cases}$$

natomiast szerokość okna h jest wybierana metodą *cross validation*. W tej metodzie, dla każdego $i = 1, 2, \dots, M$ i dla naturalnego l konstruowany jest estymator

$$m_{li}(t) = \frac{\sum_{j=1, j \neq i}^M \bar{y}_l(t_j) K(t-t_j)/h}{\sum_{j=1, j \neq i}^M K(t-t_j)/h},$$

czyli estymator jądrowy zbudowany na punktach

$$(t_1, \bar{y}_l(t_1)), \dots, (t_{i-1}, \bar{y}_l(t_{i-1})), (t_{i+1}, \bar{y}_l(t_{i+1})), \dots, (t_M, \bar{y}_l(t_M)).$$

Szerokość okna h minimalizuje sumę kwadratów

$$\sum_{i=1}^M (\bar{y}_l(t_i) - m_{li}(t_i, h))^2,$$

w przedziale

$$\left[\frac{1}{M}, 3 \max_i \min_{j \neq i} |t_i - t_j| \right].$$

Estymator ten wygładza dane pomiarowe i pozwala wyznaczyć wartości funkcji $y_i(t)$ w dowolnym punkcie $t_i \in [0, t_{kon}]$, gdzie t_{kon} przyjmuje maksymalną wartość argumentu z pomiarów. Jest to bardzo istotne przy małej ilości punktów pomiarowych. Mając dane wartości w dowolnych punktach możemy wyznaczyć odpowiedź obiektu.

Zakłada się, że odpowiedź obiektu opisana jest funkcją postaci:

$$F(t, \alpha_1, \alpha_2, \dots, \alpha_l) = \sum_{i=1}^k g_i(\alpha_1, \alpha_2, \dots, \alpha_l) f_i(t, \alpha_1, \alpha_2, \dots, \alpha_l), \quad (2)$$

gdzie α_i są nieznanymi współczynnikami, l jest rzędem przyjętego modelu, funkcje g_i są dowolnymi funkcjami zależnymi tylko od α_i , funkcje f_i są dowolnymi funkcjami ciągłymi, których postać i ilość jest zależna od rozwiązania równania różniczkowego opisującego rozpatrywany model. Nieznane współczynniki α_i minimalizują sumę

$$\sum_{j=1}^N [y_i(t_j) - F(t_j, \alpha_1, \alpha_2, \dots, \alpha_l)]^2, \quad (3)$$

gdzie $N > l$, $t_j \in [0, t_{kon}]$, $y_i(t_j)$ są wartościami wyznaczonego estymatora jądrowego w punktach t_j . Chcąc wyznaczyć nieznanne współczynniki korzystamy z warunku koniecznego na istnienie minimum, który w tym przypadku ma postać:

$$\sum_{j=1}^N 2\{y_j - F(t_j, \alpha_1, \alpha_2, \dots, \alpha_l)\} \frac{\partial F(t_j, \alpha_1, \dots, \alpha_l)}{\partial \alpha_i} = 0 \quad (4)$$

dla $i = 1, 2, \dots, l$.

Warunek (4) prowadzi do nieliniowego układu równań. Układ ten rozwiązujemy korzystając ze zmodyfikowanej procedury Taylora–Spatha. Problem doboru wartości początkowych przybliżeń α_i ($i = 1, 2, \dots, l$) nie jest tutaj rozpatrywany.

Prezentowaną metodę zastosujemy do modelu, który jest opisany transmitancją drugiego rzędu. Dla sprawdzenia dokładności metody proces będziemy symulować.

Zakłada się, że model jest opisany transmitancją postaci:

$$G_2(s) = \frac{y_2(s)}{x(s)} = \frac{1}{(s+a)(s+b)}. \quad (5)$$

Odpowiedź obiektu opisanego transmisją (5) na skokową funkcję wymuszającą jest następującą funkcją:

$$y_2(t) = F(t, a, b) = k \left(\frac{1}{ab} + \frac{1}{a(a-b)} e^{-at} + \frac{1}{b(b-a)} e^{-bt} \right) \quad (6)$$

W tym przykładzie przyjmujemy $a = 2$ i $b = 5$, wówczas transmitancja (5) przyjmuje postać:

$$G_2(s) = \frac{y_2(s)}{x(s)} = \frac{1}{s^2 + 7s + 10},$$

a odpowiedź (6) postać:

$$y_2(t) = 0.1 - (1/6)e^{-2t} + (1/15)e^{-5t},$$

Funkcję $y_2(t)$ rozpatruje się w przedziale $[0, t_{kon}]$. Wartość t_{kon} wyznacza się z zależności:

t	$y_2(t)$ dokładne	$\bar{y}_2(t)$ bez wygładzania	$\bar{y}_2(t)$ z wygładzaniem $h = 0.766$ $p = 0.1$	$\bar{y}_2(t)$ $h = 0.643$ $p = 0.2$
0.0	0.000	0.000	0.000	0.000
0.2	0.013	-0.003	0.013	0.013
0.4	0.034	0.019	0.034	0.034
0.6	0.053	0.046	0.052	0.052
0.8	0.068	0.055	0.066	0.066
1.0	0.078	0.072	0.076	0.076
1.2	0.085	0.082	0.083	0.083
1.4	0.090	0.093	0.088	0.088
1.6	0.093	0.089	0.091	0.091
1.8	0.095	0.093	0.093	0.093
2.0	0.097	0.094	0.095	0.094
2.2	0.098	0.096	0.096	0.095
2.4	0.099	0.096	0.097	0.096
2.6	0.099	0.097	0.097	0.097
2.8	0.099	0.096	0.97	0.097
3.0	0.100	0.097	0.097	0.097
3.2	0.100	0.100	0.098	0.097
3.4	0.100	0.097	0.098	0.097
3.6	0.100	0.105	0.098	0.097
3.8	0.100	0.101	0.098	0.097
4.0	0.100	0.100	0.098	0.097

$$|y_c - y(t_{kon})| \leq \varepsilon, \quad \text{gdzie } y_c = \lim_{t \rightarrow \infty} y(t).$$

gdzie ε jest dostatecznie małą liczbą nieujemną.

Zaburzamy $y_2(t)$ błędem losowym należącym do przedziału $[-p, p]$, który ma rozkład normalny o odchyleniu standardowym $\sigma = 0.05$, gdzie p jest zadaną wielkością. Przeprowadzono badania dla różnych wartości p , które fizycznie odzwierciedlają błędy wynikające z pomiarów (dokładność przyrządów, dokładność odczytu itp.).

A więc, zakłada się, że dane są wartości pomierzone funkcji $\bar{y}_2(t)$ w 10 punktach ($M = 10$) należących do przedziału $[0, 4]$. Mając dane $(t_i, \bar{y}_2(t_i))$ ($i = 1, 2, \dots, 10$) można zbudować estymator jądrowy $y_1(t)$, a następnie obliczyć nieznane współczynniki a i b , które minimalizują

$$\sum_{j=1}^N \left(y_1(t_j) - k \left(\frac{1}{ab} + \frac{1}{a(a-b)} e^{-at_j} + \frac{1}{b(b-a)} e^{-bt_j} \right) \right)^2.$$

Dokładność uzyskanych wyników zależy od ilości punktów pomiarowych (M), od poziomu błędu jakim obarczone są wielkości mierzone (p) oraz od N . Przeprowadzono badania dla różnych M , p i N . I tak dla $p = 0.1$, $M = 10$ i $N = 20$, otrzymano $a = 1.98$, $b = 5.16$ transmitancję jest postaci

$$G_2(s) = \frac{y_2(s)}{x(s)} = \frac{1}{s^2 + 7.1400s + 10.2168}, \quad (7)$$

Natomiast $p = 0.2$, $M = 10$ i $N = 20$ otrzymano $a = 1.98$, $b = 5.18$ i transmitancję (5) ma postać:

$$G_2(s) = \frac{y_2(s)}{x(s)} = \frac{1}{s^2 + 7.1600s + 10.2564}, \quad (8)$$

Wartości odpowiedzi układów opisanych transmitancjami (7) i (8) na skokową funkcję wymuszającą podane są w powyższej tabeli.

Wnioski

Proponowana metoda jest szczególnie przydatna do wyznaczania modeli uproszczonych w przypadku małych ilości zaburzonych danych pomiarowych. Punkty pośrednie wyliczone są za pomocą estymatora jądrowego. Użyty tutaj do estymacji jądrowej estymator cross-validation daje dobre wyniki.

BIBLIOGRAFIA

1. J.E. Dennis, R.B. Schnabel: *Numerical Methods for Unconstrained Optimization and Non-linear Equations*. Prentice-Hall, 1983
2. J. Halwa: *Metody wyznaczania transmitancji uproszczonych i ich zastosowanie w automatyce i elektroenergetyce*. Prace Naukowe Instytutu Cybernetyki Technicznej Politechniki Wrocławskiej (88), Seria: Monografie (21), Wrocław, 1991
3. G. Hart-Olejniczak, A.S. Kozek, A. Trzmielak-Stanisławska: *Mat-stat System*. Raport nr P-34, Instytut Informatyki Uniwersytetu Wrocławskiego

A. TRZMIELAK-STANISŁAWSKA, J. HALAWA

PARAMETRIC IDENTIFICATION METHOD WITH SMOOTHING THE INITIAL DATA

Summary

The paper discusses the method of parametric optimization with the smoothing of the initial data. For the smoothing purpose use is made of the cross-validation kernel estimation. The method proposed is used for determining simplified transfer functions. Presented in the paper are also examples of using the above method.

SPIS TREŚCI

T. A d a m s k i: Rozkład prawdopodobieństwa chwilowej niestaości momentu próbkowania w oscyloskopach cyfrowych ze stroboskopowym sposobem próbkowania	235
T. A d a m s k i: Bity efektywne konwersji A/C w obecności chwilowej niestaości momentu próbkowania	253
A. D o m a n s k a: Podstawowe kryteria wzajemnego dopasowania głównych elementów układu konwersji analogowo-cyfrowej w urządzeniach pomiarowych z komputerami osobistymi	265
J. W o j a s: Analiza struktury energetycznej i przejść międzypasmowych w arsenku galu	275
T. S t a r e c k i: Modelowanie rezonansowych komór fotoakustycznych w układzie Helmholtza za pomocą analogii akusto-elektrycznych	307
M. G o t f r y d: Praca układu impulsowego do zmniejszania zniekształceń prądu sieciowego ze zmienną częstotliwością kluczowania	313
W. N i e m i e c: Równanie konstytutywne ośrodka ciągłego dla procesu ciągłej krystalizacji masowej. Część I: Konstrukcja równań konstytutywnych ośrodka ciągłego dla procesu ciągłej krystalizacji masowej	329
W. N i e m i e c: Równanie konstytutywne ośrodka ciągłego dla procesu ciągłej krystalizacji masowej. Część II: Analityczne rozwiązanie równań konstytutywnych ośrodka ciągłego dla procesu ciągłej krystalizacji masowej	381
A. T r z m i e l a k - S t a n i s ł a w s k a, J. H a l a w a: Metoda identyfikacji parametrycznej z wygładzeniem danych początkowych	437

CONTENTS

T. A d a m s k i: Probability distribution of sampling time jitter in equivalent time digital oscilloscopes	235
T. A d a m s k i: Effective bits of analog to digital conversion and sampling time jitter	253
A. D o m a n s k a: The basic criteria for reciprocal accomodation of the main elements of the analog-to-digital conversion system in PC-measurement instruments	265
J. W o j a s: Analysis of energetic structure and interband transitions of GaAs	275
T. S t a r e c k i: Modelling of photoacoustic Helmholtz resonators by means of acousto-electrical analogies	307
M. G o t f r y d: The action of pulse circuit working with variable switching frequency used for distortions decreasing of main current	313
W. N i e m i e c: Constitutive equations of the continuous medium for the continuous mass crystallization process. Part I: The construction of constitutive equations of the continuous medium for the continuous mass crystallization process	329
W. N i e m i e c: Constitutive equations of the continuous medium for the continuous mass crystallization process. Part II: Analytical solution of constitutive equations of the continuous medium for continuous mass crystallization process	381
A. T r z m i e l a k - S t a n i s ł a w s k a, J. H a l a w a: Parametric identification method with smoothing the initial data	437

Subscription for external subscribers:

The promotional subscription price in 1994 is \$ 80 including postage for institutions. Subscriptions should be sent too the publisher, Polish Scientific Publishers PWN Ltd, Journal Division, Miodowa 10, 00-251 Warsaw, POLAND, fax (48) (22) 26 09 50, (48) (22) 26 71 63, with a copy by fast to Electronics and Telecommunications Quarterly 00-665 Warsaw, ul. Nowowiejska 15/19, p. 470. Subscription is occupted after showing a checque or transfer documents. Our bank account is as follows:

Bank account: PBK VIII O/Warszawa nr 370028-1052-139-112. At subscriber's request this journal will be air mailed at additional postage to 50% of gross price to European countries and 65% overseas.

Kwartalnik Elektroniki i Telekomunikacji

WARUNKI PRENUMERATY

Wpłaty na prenumeratę przyjmowane są na okresy kwartalne:

na teren kraju

- jednostki kolportażowe „RUCH” S.A. i urzędy pocztowe oddawcze, właściwe dla miejsca zamieszkania lub siedziby prenumeratora, oraz doręczyciele w miejscowościach, gdzie dostęp do urzędu jest utrudniony,
- od osób lub instytucji, zamieszkających lub mieszczących się w miejscowościach, w których nie ma jednostek kolportażowych „RUCH”, wpłaty należy wnieść do „RUCHU” S.A. Oddział Warszawa, 00-958 Warszawa, ul. Towarowa 28. Konto: PBK XIII Oddział Warszawa nr 370044-1195-139-11. „RUCH” S.A. zapewnia dostawę pod wskazanym adresem pocztą zwykłą w ramach opłaconej prenumeraty.

na zagranicę

- „RUCH” S.A. Oddział Warszawa, 00-958 Warszawa, ul. Towarowa 28. Konto: PBK XIII Oddział Warszawa nr 370044-1195-139-11. Dostawa odbywa się pocztą zwykłą w ramach opłaconej prenumeraty, z wyjątkiem zlecenia dostawy pocztą lotniczą, której koszt w pełni pokrywa zleceniodawca.

Prenumerata ze zleceniem dostawy za granicę jest o 100 % wyższa od krajowej.

Dostawa zamówionej prasy następuje:

- przez jednostki kolportażowe „RUCH” S.A. — w sposób uzgodniony z zamawiającym,
- prenumerata pocztowa — pod wskazanym adresem, w ramach opłaconej prenumeraty.

Terminy przyjmowania przez „RUCH” S.A. wpłat na prenumeratę krajową i zagraniczną oraz przez Poczta Polską (tylko prenumerata krajowa):

„RUCH” S.A.

- do 20 XI na I kw. roku następnego
- do 20 II na II kw.
- do 20 V na III kw.
- do 20 VIII na IV kw.

Poczta Polska

- do 25 XI na I kw. roku następnego
- do 25 II na II kw.
- do 25 V na III kw.
- do 25 VIII na IV kw.

Bieżące numery można nabyć w Księgarni Wydawnictwa Naukowego PWN, ul. Miodowa 10, 00-251 Warszawa. Również można je nabyć, a także zamówić (przesyłka za zaliczeniem pocztowym) we Wzorcowni Ośrodka Rozpowszechniania Wydawnictw Naukowych PAN, Pałac Kultury i Nauki, 00-901 Warszawa.

UWAGA!

Poczynając od 1993 r. prenumeratę naszego czasopisma prowadzi również Zakład Kolportażu Wydawnictwa SIGMA-NOT. Zamówienia na prenumeratę należy składać w terminach kwartalnych. Wpłata na II półr. wynosi 60 000 zł

Warunkiem przyjęcia prenumeraty jest wpłata na konto wartości zamówionych egzemplarzy. Wpłat na prenumeratę należy dokonywać na blankietach do wpłat na rachunki bankowe na konto:

Zakład Kolportażu Wydawnictwa SIGMA-NOT Sp. z o.o.
00-950 Warszawa nr PBK III O/Warszawa 370015-1573-139-11

Zakład Kolportażu oferuje również możliwość zamawiania prenumeraty ze zleceniem wysyłki za granicę. Jej cena jest dwukrotnie wyższa od ceny prenumeraty normalnej, a zlecający powinien podać dokładny adres odbiorcy za granicą. Dodatkowych informacji udziela:

Zakład Kolportażu Wydawnictwa SIGMA-NOT Sp. z o.o.
00-716 Warszawa, skr. poczt. 1004, ul. Bartycka 20.
Telefony: 40-30-86, 40-35-89, 40-00-21 w. 249, 295, 299

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

Nr Indeksu 363189

PL ISSN 0867-6747

KWARTALNIK ELEKTRONIKI I TELEKOMUNIKACJI

ELECTRONICS AND TELECOMMUNICATIONS QUARTERLY

TOM XXXIX — ZESZYT 3

WYDAWNICTWO NAUKOWE PWN
WARSZAWA 1994

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

KWARTALNIK
ELEKTRONIKI I TELEKOMUNIKACJI

ELECTRONICS AND
TELECOMMUNICATIONS
QUARTERLY

TOM XXXIX — ZESZYT 3

WYDAWNICTWO NAUKOWE PWN
WARSZAWA 1994

RADA REDAKCYJNA

Przewodniczący

prof. dr inż. ADAM SMOLIŃSKI
członek rzeczywisty PAN

Członkowie

prof. dr hab. inż. DANIEL JÓZEF BEM, prof. dr hab. inż. MICHAŁ BIAŁKO — czł. koresp. PAN,
prof. dr hab. inż. STEFAN HAHN — czł. koresp. PAN, prof. dr inż. ANDRZEJ HAŁAS, prof. dr inż.
ZDZISŁAW KACHLICKI, prof. dr hab. inż. BOHDAN MROZIEWICZ, prof. dr inż. JERZY
OSIOWSKI, prof. dr inż. WITOLD ROSIŃSKI — czł. rzecz. PAN, prof. dr hab. inż. STEFAN
WĘGRZYN — czł. rzecz. PAN, prof. dr hab. inż. WIESŁAW WOLIŃSKI — czł. koresp. PAN,
prof. dr inż. ANDRZEJ ZIELIŃSKI, prof. dr inż. MARIAN ZIENTALSKI

REDAKCJA

Redaktor Naczelny

prof. dr hab. inż. WIESŁAW WOLIŃSKI

Zastępca Redaktora Naczelnego

doc. dr inż. KRYSZTYN PLEWKO

Sekretarz Odpowiedzialny

mgr KRYSZYNA LELAKOWSKA

ADRES REDAKCJI

00-665 Warszawa, ul. Nowowiejska 15/19 Politechnika, pok. 470
Instytut Telekomunikacji, Gmach im. prof. JANUSZA GROSZKOWSKIEGO

Dyżury Redakcji: środy i piątki, godz. 14–16

tel. 628 89 81 21007-737

Telefony domowe: Redaktora Naczelnego: 12 17 65

Zast. Red. Naczelnego: 26 83 41

Sekretarza Odpowiedzialnego: 25 29 18

WYDAWNICTWO NAUKOWE PWN

Warszawa, ul. Miodowa 10

Ark. druk. 11,5.	Podpisano do druku w styczniu 1994 r.
Papier offsetowy kl. III 80 g. B-1	Druk ukończono w lutym 1994 r.
Skład: Wydawnictwo GP-BIT Warszawa ul. Marymoncka 34.	Druk i oprawa: Drukarnia Braci Grodzickich, Żabieniec, ul. Przelotowa 7.

Analysis of additive noise introduced by sampling time jitter in sampling frequency converters

TOMASZ ADAMSKI

Instytut Podstaw Elektroniki, Politechnika Warszawska

Received 1993.04.02

Authorized 1993.09.11

Sampling frequency converters are electronic systems aimed to change sampling rate. In principle they work in purely deterministic way but a kind of sampling time jitter (STJ) can occur in such systems. This STJ is a result of the specific unsynchronized way of sampling applied in sampling frequency converters and causes errors which can be considered as an additive noise. In the paper a short analysis of this additive noise is given. It is shown that the STJ has as a rule the uniform probability distribution which is the start point to obtain the additive noise parameters: the mean value and the correlation function. Obtained results allow also to assess the influence of the STJ on the accuracy of measurements and signal processing in sampling systems similar to frequency converters like digital oscilloscopes with clock controlled sampling moments.

1. INTRODUCTION

In electronic systems with sampling, like digital oscilloscopes, digital waveform recorders, data acquisition systems, sampling frequency converters [1] and many others we cannot neglect random unstability of sampling time. Because of inherent noise or chaos like phenomena real sampling points are randomly shifted on the time axis. This shifting is called the sampling time jitter (STJ). The STJ introduces errors to measurements or signal processing and limits significantly accuracy of electronic systems with sampling. The paper deals with the STJ caused in sampling frequency converters by the way of sampling or more precisely by chaotic phenomena present in sampling frequency converters. In the section 2 we introduce shortly simple mathematical model of sampling with the STJ. In two next sections we answer the question how the STJ changes statistical properties of the input signal.

2. MATHEMATICAL MODEL OF SAMPLING WITH TIME JITTER

Assume the analog input signal $f: \mathbb{R} \rightarrow \mathbb{R}$ of the given sampling system is a periodic, $(B(\mathbb{R}), B(\mathbb{R}))$ measurable function sampled in time points $t_1, t_2, \dots, t_M$ (in particular we can have $t_k = k \cdot \Delta t$). Let $T(t_k)$ be the random variable describing the random shift of the real sampling point $t_k + T(t_k)$ related to assumed, deterministic one t_k . Therefore we can describe the all sampling process by a random sequence $(f(t_k + T(t_k)))_{k=1}^{\infty}$. To assess errors introduced by the STJ we have first to assess the value $er(t) \stackrel{\text{def}}{=} f(t + T(t)) - f(t)$ for a fixed sampling point t (particularly for $t_1, t_2, \dots, t_M$). If the signal f is a $(B(\mathbb{R}), B(\mathbb{R}))$ measurable function (for instance f is continuous) then $er(t)$ is a random variable for each sampling point t but its probability distribution depends on f and is in most cases difficult to calculate exactly.

Denote by $\mathcal{K}(A, \omega_g)$ a set of real band limited (by ω_g) functions over $\mathbb{R}$ with values in the interval $[-A, A]$. The worst case function (i.e. most sensitive to the STJ) from $\mathcal{K}(A, \omega_g)$ is sinusoid $f(t) = A \cdot \sin(\omega_g t)$. As a result we have for $f \in \mathcal{K}(A, \omega_g)$

$$D^2(er(t)) = D^2(f(t + T(t))) \leq A^2 \omega_g^2 D^2(T(t)). \quad (2.1)$$

The above error assessment deals with arbitrary sampling system where the STJ is significant.

A deterministic signal $f: \mathbb{R} \rightarrow \mathbb{R}$ can be treated as a stochastic process $(Q(t))_{t \in \mathbb{R}}$ defined as $Q(t)(\omega) = f(t)$ for each $t \in \mathbb{R}$, $\omega \in \mathbb{R}$. Then, we assume in the sequel that the considered input signal is a random one $(Q(t))_{t \in \mathbb{R}}$ which is more general approach. For example in appendix it is shown (theorem A.1) that if the function $f: \mathbb{R} \rightarrow \mathbb{R}$ is periodic then under certain assumptions (not specially restrictive) f can be considered as stationary stochastic process in the strict sense.

To be formally correct we also assume in the sequel that the input signal is a real, stochastic process $(Q(t))_{t \in \mathbb{R}} \in \mathcal{F} \otimes B(\mathbb{R})$ measurable (where $\mathcal{F} = \sigma(Q(t); t \in \mathbb{R})$).

3. STOCHASTIC PROCESSES WITH TIME JITTER

Assume, we know probabilistic properties of the real measurable stochastic process $(Q(t))_{t \in \mathbb{R}}$. In this section we would like to answer the question: how the STJ described by the real stochastic process $(T(t))_{t \in \mathbb{R}}$ changes statistical properties of the process $(Q(T(t) + t))_{t \in \mathbb{R}}$ and of the error process $(e(t))_{t \in \mathbb{R}} \stackrel{\text{def}}{=} Q(t + T(t)) - Q(t)_{t \in \mathbb{R}}$ (i.e. additive noise caused by the STJ). Hence we would like to assess influence of the "continuous jitter" or in other words phase noise on the random signal $(Q(t))_{t \in \mathbb{R}}$. In the sequel we state properties of the stochastic processes $(Q(t + T(t)))_{t \in \mathbb{R}}$, $(e(t))_{t \in \mathbb{R}}$ and random sequences $(Q(k\Delta t + T(k\Delta t)))_{k \in \mathbb{Z}}$, $(e(k\Delta t))_{k \in \mathbb{Z}}$, related to uniform sampling of the random signal with the STJ.

In all theorems of this section we assume that the stochastic process $(Q(t))_{t \in \mathbb{R}}$ is a real $\mathcal{F} \otimes B(\mathbb{R})$ measurable stochastic process and σ -fields $\mathcal{F} = \sigma(Q(t); t \in \mathbb{R})$ and

$\sigma(T(t), t \in \mathbb{R})$ are independent, i.e. jitter is independent from the process $(Q(t))_{t \in \mathbb{R}}$. Under these assumptions using theorem A.4 from the appendix we can obtain n -dimensional distributions of the process $Q(t+T(t))_{t \in \mathbb{R}}$ and $(e(t))_{t \in \mathbb{R}}$ from n -dimensional probability distribution of the process $(Q(t))_{t \in \mathbb{R}}$ i $(T(t))_{t \in \mathbb{R}}$.

In the case of $(Q(t+T(t))_{t \in \mathbb{R}}$ using formula (3.9) and the theorem A.5 we obtain for every $n \in \mathbb{N}$, every $B \in \mathcal{B}(\mathbb{R}^n)$ and every vector $(t_1, \dots, t_n) \in \mathbb{R}^n$, $t_1 < t_2 < \dots < t_n$

$$\begin{aligned} P((Q(t_1+T(t_1)), Q(t_2+T(t_2)), \dots, Q(t_n+T(t_n))) \in B) = \\ = \int_{\mathbb{R}^n} P((Q(t_1+x_1), Q(t_2+x_2), \dots, Q(t_n+x_n)) \in B) P_D(x_1, x_2, \dots, x_n), \end{aligned} \quad (3.1)$$

where $D = (T(t_1), T(t_2), \dots, T(t_n))$. For $(e(t))_{t \in \mathbb{R}}$ the situation is a little more complicated. Denote by f a function given by the formula

$f: \mathbb{R}^{2n} \ni (x_1, y_1, x_2, y_2, \dots, x_n, y_n) \rightarrow (x_1 - y_1, x_2 - y_2, \dots, x_n - y_n) \in \mathbb{R}^n$ and by $\tilde{T}(t)$, a random variable defined for every $t \in \mathbb{R}$ by the formula $\tilde{T}(t)(\omega) \equiv t$ for every $\omega \in \Omega$. Using theorem A.4 we have

$$P((e(t_1), e(t_2), \dots, e(t_n)) \in B) = P((Q(t_1+T(t_1)), Q(t_1), Q(t_2+T(t_2)), Q(t_2), \dots, Q(t_n+T(t_n))) \in$$

$$\in f^{-1}(B)) = \int_{\mathbb{R}^{2n}} P((Q(x_1), Q(y_1), Q(x_2), Q(y_2), \dots, Q(x_n), Q(y_n)) \in$$

$$\in f^{-1}(B)) P_G(d(x_1, y_1, x_2, y_2, \dots, x_n, y_n)) =$$

where $G(t_1+T(t_1), \tilde{T}(t_1), t_2+T(t_2), \tilde{T}(t_2), \dots, t_n+T(t_n), \tilde{T}(t_n))$. The random variable $(\tilde{T}(t_1), \tilde{T}(t_2), \dots, \tilde{T}(t_n))$ has the one point probability distribution then it is independent from the random variable $(t_1+T(t_1), t_2+T(t_2), \dots, t_n+T(t_n))$. Therefore using the Fubini theorem we have:

$$= \int_{\mathbb{R}^n} P((Q(t_1+x_1), Q(t_1), Q(t_2+x_2), Q(t_2), \dots, Q(t_n+x_n), Q(t_n)) \in$$

$$\in f^{-1}(B)) P_D(d(x_1, x_2, \dots, x_n)),$$

where $D = (T(t_1), T(t_2), \dots, T(t_n))$.

Theorem 3.1 Assume $(Q(t))_{t \in \mathbb{R}}$ is a real, stationary in the strict sense, $\mathcal{F} \otimes \mathcal{B}(\mathbb{R})$ measurable stochastic process (where $\mathcal{F} = \sigma(Q(t); t \in \mathbb{R})$). If a real, stationary the strict sense stochastic process $(T(t))_{t \in \mathbb{R}}$ (describing sampling time jitter) and the stochastic process $(Q(t))_{t \in \mathbb{R}}$ are independent then the stochastic processes $(Q(t+T(t)))_{t \in \mathbb{R}}$ and $(e(t))_{t \in \mathbb{R}}$ are stationary in the strict sense. If additionally the stochastic process $(Q(t))_{t \in \mathbb{R}}$ is stationary in the wide sense then the stochastic processes $(Q(t+T(t)))_{t \in \mathbb{R}}$ and $(e(t))_{t \in \mathbb{R}}$ are stationary in the wide sense too.

Proof. At first we prove that the stochastic process $Q(t+T(t))_{t \in \mathbb{R}}$ is stationary in the strict sense. To do it we have to show that every $n \in \mathbb{N}$, every vector $(t_1, t_2, \dots, t_n) \in \mathbb{R}^n$, $t_1 < t_2 < \dots < t_n$, every $t \in \mathbb{R}$ and every $B \in \mathcal{B}(\mathbb{R}^n)$ we have

$$\begin{aligned} P((Q(t_1+T(t_1)), Q(t_2+T(t_2)), \dots, Q(t_n+T(t_n))) \in B) = \\ = P(Q(t_1+t+T(t_1+t)), Q(t_2+t+T(t_2+t)), \dots, Q(t_n+t+T(t_n+t))) \in B). \end{aligned} \quad (3.2)$$

Indeed, because σ -fields $\sigma(T(t); t \in \mathbb{R})$ i $\sigma(Q(t); t \in \mathbb{R})$ are independent, then using theorem A.4 and stationary property of process $(Q(t))_{t \in \mathbb{R}}$ i $(T(t))_{t \in \mathbb{R}}$ we obtain

$$\begin{aligned} P((Q(t_1+T(t_1)), Q(t_2+T(t_2)), \dots, Q(t_n+T(t_n))) \in B) = \\ = \int_{\mathbb{R}^n} P((Q(\tilde{t}_1), Q(\tilde{t}_2), \dots, Q(\tilde{t}_n)) \in B) P_K(d(\tilde{t}_1, \tilde{t}_2, \dots, \tilde{t}_n)) = \end{aligned}$$

where $K = (t_1+T(t_1), t_2+T(t_2), \dots, t_n+T(t_n))$.

$$= \int_{\mathbb{R}^n} P((Q(t_1+\tilde{t}_1), Q(t_2+\tilde{t}_2), \dots, Q(t_n+\tilde{t}_n)) \in B) P_D(\tilde{t}_1, \tilde{t}_2, \dots, \tilde{t}_n) =$$

where $D = (T(t_1), T(t_2), \dots, T(t_n))$.

$$= \int_{\mathbb{R}^n} P((Q(t_1+t+\tilde{t}_1), Q(t_2+t+\tilde{t}_2), \dots, Q(t_n+t+\tilde{t}_n)) \in B) P_L(d(\tilde{t}_1, \tilde{t}_2, \dots, \tilde{t}_n)) =$$

where $L = (T(t_1+t), T(t_2+t), \dots, T(t_n+t))$

$$= \int_{\mathbb{R}^n} P((Q(\tilde{t}_1), Q(\tilde{t}_2), \dots, Q(\tilde{t}_n)) \in B) P_W(d(\tilde{t}_1, \tilde{t}_2, \dots, \tilde{t}_n)) =$$

where $W = (t_1+t+T(t_1+t), t_2+t+T(t_2+t), \dots, t_n+t+T(t_n+t))$

$$= P((Q(t_1+t+T(t_1+t)), Q(t_2+t+T(t_2+t)), \dots, Q(t_n+t+T(t_n+t))) \in B).$$

In the case of the error process $e(t) \stackrel{\text{df}}{=} Q(t+T(t)) - Q(t)$ our reasoning is similar. We have to prove that for every $n \in \mathbb{N}$, every $B \in (\mathbb{R}^n)$, every $t \in \mathbb{R}$ and every vector $(t_1, t_2, \dots, t_n) \in \mathbb{R}^n$ such that $t_1 < t_2 < \dots < t_n$ the following equality holds

$$\begin{aligned} P((Q(t_1+T(t_1)) - Q(t_1), Q(t_2+T(t_2)) - Q(t_2), \dots, Q(t_n+T(t_n)) - Q(t_n)) \in B) = \\ = P((Q(t_1+t+T(t_1+t)) - Q(t_1+t), Q(t_2+t+T(t_2+t)) - Q(t_2+t), \dots, \\ Q(t_n+t+T(t_n+t)) - Q(t_n+t)) \in B). \end{aligned} \quad (3.3)$$

Let f denotes a function defined by the formula:

$f: \mathbb{R}^{2n} \ni (x_1, x_2, \dots, x_{2n}) \rightarrow (x_1 - x_2, x_3 - x_4, \dots, x_{2n-1} - x_{2n}) \in \mathbb{R}^n$ and $\tilde{T}(t)$ for $t \in \mathbb{R}$ are random variables given as $\tilde{T}(t, \omega) \equiv t$ for each $\omega \in \Omega$ then the condition (3.3) can be written in the following way:

$$\begin{aligned} & P((Q(t_1 + T(t_1)), Q(t_1), Q(t_2 + T(t_2)), Q(t_2), \dots, Q(t_n + T(t_n)), Q(t_n)) \in f^{-1}(B)) = \\ & = P(Q(t_1 + t + T(t_1 + t)), Q(t_1 + t), \dots, Q(t_n + t + T(t_n + t)), Q(t_n + t)) \in f^{-1}(B)). \end{aligned} \quad (3.4)$$

Hence using theorem A.4 and stationary properties of the stochastic process $(Q(t))_{t \in R}$ and $(T(t))_{t \in R}$ we obtain.

$$\begin{aligned} & P((Q(t_1 + T(t_1)), Q(t_1), Q(t_2 + T(t_2)), Q(t_2), \dots, Q(t_n + T(t_n)), Q(t_n)) \in f^{-1}(B)) = \\ & = P((Q(t_1 + T(t_1)), Q(\tilde{T}(t_1)), Q(t_2 + T(t_2)), Q(\tilde{T}(t_2)), \dots, Q(t_n + T(t_n)), Q(\tilde{T}(t_n))) \in \\ & \in f^{-1}(B)) = \end{aligned}$$

$$= \int_{R^{2n}} P((Q(x_1), Q(y_1), Q(x_2), Q(y_2), \dots, Q(x_n), Q(y_n)) \in$$

$$\in f^{-1}(B)) P_H(d(x_1, y_1, x_1, y_2, \dots, x_n, y_n)) =$$

$$\text{where } H = (t_1 + T(t_1), \tilde{T}(t_1), t_2 + T(t_2), \tilde{T}(t_2), \dots, t_n + T(t_n), \tilde{T}(t_n))$$

$$= \int_{R^{2n}} P((Q(t_1 + x_1), Q(y_1), Q(t_2 + x_2), Q(y_2), \dots, Q(t_n + x_n), Q(y_n)) \in$$

$$\in f^{-1}(B)) P_M(d(x_1, y_1, x_2, y_2, \dots, x_n, y_n)) =$$

$$\text{where } M = (T(t_1), \tilde{T}(t_1), \dots, T(t_n), \tilde{T}(t_n))$$

$$= \int_{R^{2n}} P((Q(t + t_1 + x_1), Q(t + y_1), \dots, Q(t + t_n + x_n), Q(t + y_n)) \in$$

$$\in f^{-1}(B)) P_M(d(x_1, y_1, \dots, x_n, y_n)) =$$

$$= \int_{R^{2n}} P((Q(t + t_1 + x_1), Q(t + y_1), \dots, Q(t + t_n + x_n), Q(t + y_n)) \in$$

$$f^{-1}(B)) P_S(d(x_1, y_1, \dots, x_n, y_n)) =$$

$$\text{where } S = (T(t + t_1), \tilde{T}(t_1), \dots, T(t + t_n), \tilde{T}(t_n))$$

$$= P((Q(t_1 + t + T(t_1 + t)), Q(t + t_1), \dots, Q(t + t_n + T(t_n + t)), Q(t + t_n)) \in f^{-1}(B))$$

Hence the formula (3.4) is satisfied which gives as result (3.3) Q.E.D.

Remark. In the similar way we can prove the following modification of the theorem 3.1. If the stochastic process $(Q(t))_{t \in R}$ satisfies assumptions of the theorem 3.1 and the stochastic process $(T(k\Delta t))_{k \in Z}$ is a stationary sequence of the real random

variables independent from $(Q(t))_{t \in R}$ then the stochastic processes $(e(k\Delta t))_{k \in Z}$ and $(Q(k\Delta t + T(k\Delta t)))_{k \in Z}$ are stationary too.

From the theorem 3.1 theorems from the appendix we obtain also the following theorem:

Theorem 3.2 If $(Q(t))_{t \in R}$ is a real $\mathcal{F} \otimes B(R)$ measurable stationary process (where $\mathcal{F} = \sigma(Q(t); t \in R)$), $(T(t))_{t \in R}$ an arbitrary real process describing jitter, and processes $(Q(t))_{t \in R}$, $(T(t))_{t \in R}$ are independent then the STJ do not change 1-dimensional probability distributions i.e. for every $t \in R$ probability distributions of the random variables $Q(t)$ and $Q(t+T(t))$ are equal. Hence in particular if $Q(0) \in L^1(\Omega, \mathcal{M}, P)$ then $E(Q(t)) = E(Q(t+T(t)))$ for every $t \in R$ and similarly if $Q(0) \in L^2(\Omega, \mathcal{M}, P)$ then $D^2(Q(t)) = D^2(Q(t+T(t)))$ for every $t \in R$.

Remark. If a real stochastic process $(Q(t))_{t \in R}$ is continuous and the process describing jitter $(T(t))_{t \in R}$ is continuous then processes $(Q(T(t)))_{t \in R}$, $(Q(T(t)+t))_{t \in R}$ and $e(t)_{t \in R}$ are continuous too. It follows immediately from definition of the continuous process.

Theorem 3.3 Assume $(Q(t))_{t \in R}$ is a real $\mathcal{F} \otimes B(R)$ measurable (where $\mathcal{F} = \sigma(Q(t); t \in R)$) stochastic process of the second order with the mean value $m(t)$ and the autocorrelation function $R(t, s)$. If the stochastic process $(Q(t))_{t \in R}$ and process describing jitter $(T(t))_{t \in R}$ are independent and the following condition (3.5) is fulfilled

$$m(t+T(t)), R(t+T(t), t+T(t)) \in L^1(\Omega, \mathcal{M}, P) \text{ for every } t \in R \quad (3.5)$$

then the stochastic processes $(Q(t+T(t)))_{t \in R}$ and $(e(t))_{t \in R}$ are second order processes too and the mean value and the correlation function are given by the formulae

$$m_1(t) \stackrel{\text{df}}{=} E(Q(t+T(t))) = E(m(t+T(t))) \quad (3.6)$$

$$R_1(t_1, t_2) \stackrel{\text{df}}{=} E(Q(t_1+T(t_1))Q(t_2+T(t_2))) = E(R(t_1+T(t_1), t_2+T(t_2))) \quad (3.7)$$

$$m_2(t) \stackrel{\text{df}}{=} (E(t) = m_1(t) - m(t)) \quad (3.8)$$

$$R_2(t_1, t_2) \stackrel{\text{df}}{=} E(e(t_1)e(t_2)) = R(t_1, t_2) + R_1(t_1, t_2) - E(R(t_1, t_2 + T(t_2))) - E(R(t_1 + T(t_1), t_2)) \quad (3.9)$$

Proof. of the above theorem is based on the following lemma.

Lemma. If $(X(t))_{t \in T}$ is a real measurable stochastic process defined for a measurable space of parameters $(T, \mathcal{B})$ and for every $t \in T$ there exists $E(X(t))$ then the function $T \ni t \rightarrow E(X(t)) \in R$ is $(B, B(R))$ measurable.

The simple proof of this theorem will be omitted.

Corollary 3.4 If a stochastic process $(Q(t))_{t \in R}$ is a real measurable stochastic process of the second order and $(T(t))_{t \in R}$ is an arbitrary real stochastic process then

1) functions $R \ni t \rightarrow m(t) \stackrel{\text{df}}{=} E(Q(t)) \in R$ and $R^2 \ni (t_1, t_2) \rightarrow R(t_1, t_2) \stackrel{\text{df}}{=} E(Q(t_1)Q(t_2)) \in R$ are respectively $(B(R), B(R))$ and $(B(R^2), B(R))$ measurable.

2) $m(t_1 + T(t_1))$, $R(t_1, t_2)$, $R(t_1 + T(t_1), t_2 + T(t_2))$ are random variables.

Proof. The fact that $(m \cdot)$ is a measurable function and $m(t_1 + T(t_1))$ is a random variable follows immediately from the lemma. The stochastic process $(Q(t))_{t \in \mathbb{R}}$ is measurable, then processes $(Q(t_1))_{(t_1, t_2) \in \mathbb{R}^2}$ and $(Q(t_2))_{(t_1, t_2) \in \mathbb{R}^2}$ are measurable too. Hence a process $(Q(t_1), Q(t_2))_{(t_1, t_2) \in \mathbb{R}^2}$ is measurable. Because multiplication of real numbers is a continuous function from $\mathbb{R}^2$ to $\mathbb{R}$ then it is also a Borel function and the process $(Q(t_1)Q(t_2))_{(t_1, t_2) \in \mathbb{R}^2}$ is measurable. Therefore from the lemma we obtain that the function $\mathbb{R}^2 \ni (t_1, t_2) \rightarrow R(t_1, t_2) \stackrel{\text{df}}{=} E(Q(t_1)Q(t_2)) \in \mathbb{R}$ is $(B(\mathbb{R}^2), B(\mathbb{R}))$ measurable. Point 2) is a direct consequence of the point 1) Q.E.D.

Proof of the theorem 3.3. Because $R(t+T(t), t+T(t)) \in L^1(\Omega, \mathfrak{M}, P)$ from the theorem on change of variables (theorem A.5 from the appendix) we obtain for every fixed $t \in \mathbb{R}$

$$\int_{\Omega} R(t+T(t), t+T(t)) dP = \int_{\mathbb{R}} R(t+t_1, t+t_1) P_{T(t)}(dt_1) = \int_{\mathbb{R}} \left(\int_{\Omega} Q^2(t+t_1) dP \right) P_{T(t)}(dt_1)$$

The integrated function $Q^2(t+t_1)$ is nonnegative and $\mathcal{F} \otimes B(\mathbb{R})$ measurable as a function of the variable (ω, t_1) where $\mathcal{F} = \sigma(Q(t); t \in \mathbb{R})$. Hence from the Tonelli theorem we obtain, that the function $(\omega, t_1) \rightarrow Q^2(t+t_1)(\omega)$ is an element of $L^1(\Omega \times \mathbb{R}, \mathcal{F} \otimes B(\mathbb{R}), P \otimes P_{T(t)})$. Because σ -fields $\mathcal{F}$ and $\sigma(T(t))$ are independent, the measure $P \otimes P_{T(t)}$ on $\mathcal{F} \otimes B(\mathbb{R})$ is induced by the measurable function $(\text{id}, T(t)): \Omega \ni \omega \rightarrow (\omega, T(t)) \in \Omega \times \mathbb{R}$ from the probabilistic measure P . Hence from the theorem on change of variables (theorem A.6) we have $Q^2(t+T(t)) \in L^1(\Omega, \mathfrak{M}, P)$. Hence the process $(Q(t+T(t)))_{t \in \mathbb{R}}$ is a stochastic process of the second order. On the other hand, the process $(e(t))_{t \in \mathbb{R}}$ is the second order process too. It follows directly from the fact that the processes $(Q(t))_{t \in \mathbb{R}}$ and $(Q(t+T(t)))_{t \in \mathbb{R}}$ are second order processes. We compute now the mean values and correlation functions of the processes $(Q(t+T(t)))_{t \in \mathbb{R}}$ and $(e(t))_{t \in \mathbb{R}}$. The σ -fields $\mathcal{F}$ and $\sigma(T(t), t \in \mathbb{R})$ are independent then using theorem A.6. we obtain for every $t \in \mathbb{R}$.

$$m_1(t) \stackrel{\text{df}}{=} E(Q(t+T(t))) = E(E(Q(t+T(t)) | \sigma(T(t)))) = E(E(Q(t+t_1)) |_{t_1=T(t)}) = E(m(t+T(t)))$$

$$\begin{aligned} R_1(t_1, t_2) &\stackrel{\text{df}}{=} E(Q(t_1+T(t_1))Q(t_2+T(t_2))) = \\ &= E(E(Q(t_1+T(t_1))Q(t_2+T(t_2)) | \sigma(T(t_1), T(t_2)))) = \end{aligned}$$

$$E(E(Q(t_1+\tilde{t}_1)Q(t_2+\tilde{t}_2)) |_{(\tilde{t}_1, \tilde{t}_2)=(T(t_1), T(t_2))}) = E(R(t_1+T(t_1), t_2+T(t_2))),$$

$$m_2(t) \stackrel{\text{df}}{=} E(e(t)) = E(Q(t+T(t)) - Q(t)) = m_1(t) - m(t),$$

$$\begin{aligned} R_2(t_1, t_2) &\stackrel{\text{df}}{=} E(e(t_1)e(t_2)) = E((Q(t_1+T(t_1)) - Q(t_1))(Q(t_2+T(t_2)) - Q(t_2))) = \\ &= R_1(t_1, t_2) + R(t_1, t_2) - E(Q(t_1+T(t_1))Q(t_2)) - E(Q(t_2+T(t_2))Q(t_1)) = \\ &= R_1(t_1, t_2) + R(t_1, t_2) - E(R(t_1+T(t_1), t_2)) - E(R(t_1, t_2+T(t_2))) \quad \text{Q.E.D.} \end{aligned}$$

Remark 1. Assumption (3.5) can be replaced by one of the following "more physical" assumptions.

- 1) the mean value $m(t)$ and the correlation function $R(t,s)$ are continuous and jitter is bounded i.e. there exists such a $\varepsilon > 0$, that for every $t \in \mathbb{R}$, $|T(t)| < \varepsilon$
- 2) there exists such a $M > 0$, that for every $t \in \mathbb{R}$, $|m(t)|, R(t,t) < M$
- 3) the process $(Q(t))_{t \in \mathbb{R}}$ is bounded i.e. there exists such a $M > 0$, that for every $(t, \omega) \in \mathbb{R} \times \Omega$, $|Q(t, \omega)| < M$
- 4) The stochastic process $(Q(t))_{t \in \mathbb{R}}$ is stationary in the wide sense.

Corollary 3.5 If we sample a random signal uniformly i.e. in time points $k\Delta t$; $k \in \mathbb{Z}$ then (under assumptions of the theorem 3.3) random sequences $(Q(k\Delta t + T(k\Delta t)))_{k \in \mathbb{Z}}$ and $(e(k\Delta t))_{k \in \mathbb{Z}}$ are stochastic process of the second order. Their mean value and correlation function are given for every $k, n \in \mathbb{N}$ by formulae

$$m_1(k) \stackrel{\text{df}}{=} E(Q(k\Delta t + T(k\Delta t))) = E(m(k\Delta t + T(k\Delta t))) \quad (3.10)$$

$$R_1(k, n) \stackrel{\text{df}}{=} E(Q(k\Delta t + T(k\Delta t))Q(n\Delta t + T(n\Delta t))) = E(R(k\Delta t + T(k\Delta t), n\Delta t + T(n\Delta t))) \quad (3.11)$$

$$m_2(k) = m_1(k) - m(k\Delta t) \quad (3.12)$$

$$R_2(k, n) = R(k\Delta t, n\Delta t) + R_1(k, n) - E(R(n\Delta t, k\Delta t + T(k\Delta t))) - \\ - E(R(n\Delta t + T(n\Delta t), k\Delta t)) \quad (3.13)$$

Corollary 3.6 Let a real $\mathcal{F} \otimes B(\mathbb{R})$ measurable process $(Q(t))_{t \in \mathbb{R}}$ (where $\mathcal{F} = \sigma(Q(t); t \in \mathbb{R})$) be stationary in the wide sense with the mean value $m(t)$ and the correlation function $R(t)$. If the stochastic process $(T(t))_{t \in \mathbb{R}}$ describing jitter is stationary and processes $(Q(t))_{t \in \mathbb{R}}$ and $(T(t))_{t \in \mathbb{R}}$ are independent then processes $(Q(t - T(t)))_{t \in \mathbb{R}}$ and $(e(t))_{t \in \mathbb{R}}$ are stationary in the wide sense and their mean values and correlation functions are given for $t \in \mathbb{R}$, by the formulae

$$m_1(t) = m(t) = \text{const}; \quad m_2(t) = 0 \quad (3.14)$$

$$R_1(t) = E(R(t + T(t) - T(0))) = \int_{\mathbb{R}} R(t + x) P_G(dx), \quad (3.15)$$

where $G = T(t) - T(0)$

$$R_2(t) = R_1(t) + R(t) - E(R(t + T(0))) - E(R(t - T(0))). \quad (3.16)$$

Proof. From the corollary 3.4 we obtain that the correlation function $R(\cdot)$ is $(B(\mathbb{R}), B(\mathbb{R}))$ measurable. Because $|R(t)| \leq R(0)$ for every $t \in \mathbb{R}$, then for arbitrary random variable X , we have $R(X) \in L^1(\Omega, \mathfrak{M}, P)$. Then the condition (3.5) and remaining assumptions of the theorem 3.3 are fulfilled. Stationary in the wide sense of processes $(Q(t + T(t)))_{t \in \mathbb{R}}$ and $(e(t))_{t \in \mathbb{R}}$ follows formulae (3.6) – (3.9). Equations (3.14) are direct consequences of (3.6) and (3.8).

Formulae (3.15) and (3.16) follow from equations (3.7), (3.9) and stationarity of the process $(T(t))_{t \in \mathbb{R}}$. Indeed, denoting by $\tilde{R}_1(t_1, t_2)$ and $\tilde{R}_2(t_1, t_2)$ correlation functions respectively of $(Q(t+T(t)))_{t \in \mathbb{R}}$ and $(e(t))_{t \in \mathbb{R}}$ we have

$$\tilde{R}_1(t_1, t_2) = E(R(t_1 + T(t_1) - t_2 - T(t_2))) = \int_{\mathbb{R}} R(t_1 - t_2 + x) P_G(dx) = \int_{\mathbb{R}} R(t_1 - t_2 + x) P_H(dx)$$

where $G = T(t_1) - T(t_2)$ and $H = T(t_1 - t_2) - T(0)$

$$\tilde{R}_2(t_1, t_2) = R(t_1 - t_2) + R_1(t_1 - t_2) - E(R(t_1 - t_2 - T(t_2))) - E(R(t_1 + T(t_1) - t_2))$$

Denoting $t = t_1 - t_2$ we obtain formulae (3.15), (3.16). Q.E.D.

We obtain similar results for uniform sampling.

Corollary 3.7 If a process $(Q(t))_{t \in \mathbb{R}}$ is a real $\mathcal{F} \otimes B(\mathbb{R})$ measurable (where $\mathcal{F} = \sigma(Q(t); t \in \mathbb{R})$), stationary in the wide sense process with the mean value $m(t) = m_0$ and a correlation function $R(t)$ and stochastic process $(T(k\Delta t))_{k \in \mathbb{Z}}$ is stationary and independent from $(Q(t))_{t \in \mathbb{R}}$ the random sequences $(Q(k\Delta t + T(k\Delta t)))_{k \in \mathbb{Z}}$ and $(e(k\Delta t))_{k \in \mathbb{Z}}$ are stationary in then wide sense and their mean value and correlation function are given by the formulae:

$$m_1(k) = m(k\Delta t) = m_0, \quad m_2(k) = 0 \quad \text{for every } k \in \mathbb{Z} \quad \text{and}$$

$$R_1(k) = \int_{\mathbb{R}} R(k\Delta t + x) P_G(dx) \quad (3.17)$$

where $G = T(k\Delta t) - T(0)$

$$R_2(k) = R_1(k) + R(k\Delta t) - E(R(k\Delta t + T(0))) - E(R(k\Delta t - T(0))). \quad (3.18)$$

Remark 1. For $t = 0$ from the formula (3.16) and for $k = 0$ from the formula (3.18) we obtain a formula on a error process power $E(e^2(t)) = R_2(0) = 2(R(0) - E(R(T(0))))$.

Remark 2. If (under assumptions of the corollary 3.7) $(T(k\Delta t))_{k \in \mathbb{Z}}$ is a sequence of independent random variables then

$$R_1(k) = \int_{\mathbb{R}} R(k\Delta t + t) P_{T(0)} * P_{-T(0)}(dt). \quad (3.19)$$

where $*$ denotes convolution of probabilistic measures. If additionally a random variable $T(0)$ has a probability density function f then

$$R_1(n) = \int_{\mathbb{R}} (R(k\Delta t + t) \int_{\mathbb{R}} f(x) f(t - x) l_1(dx)) l_1(dt) = \int_{\mathbb{R}^2} R(k\Delta t + x) f(x) f(t - x) l_2(d(x, t)).$$

Remark 3. If we do not assume in the theorem 3.2 and corollary 3.6 stationarity of the stochastic process describing the STJ then as it follows from the formulae (3.6)–(3.9) processes $(Q(t+T(t)))_{t \in R}$, $(e(t))_{t \in R}$, $(Q(k\Delta t + T(k\Delta t)))_{k \in Z}$ and $(e(k\Delta t))_{k \in Z}$ do not have to be stationary in the wide sense.

Remark 4. If $(Q(t))_{t \in R}$ is a $\mathcal{F} \otimes B(R)$ measurable and stationary (in the wide sense) stochastic process independent from the process $(T(t))_{t \in R}$ describing jitter then from the formula (3.18) we have $E(e^2(t)) = 2(R(0) - E(R(T(t))))$ for every $t \in R$.

Example. Assume a random sequence $(T(k\Delta t))_{k \in Z}$ describing the STJ is a stationary Gaussian stochastic process and random variables $T(k\Delta t)$, $k \in Z$ have the Gaussian probability distribution $N(0, \sigma)$. If the process $(Q(t))_{t \in R}$ is a real stationary in the wide sense and $\mathcal{F} \otimes B(R)$ measurable process then from the theorem 3.7 we obtain that the error process $(e(kt))_{k \in Z}$ is stationary in the wide sense with the mean value equal to 0 and

$$E(e^2(k\Delta t)) = 2(R(0) - E(R(T(0)))) = 2(R(0) - \frac{1}{\sigma \sqrt{2\pi}} \int_R R(t) \exp\left(-\frac{t^2}{2\sigma^2}\right) l_1(dt)).$$

If the stochastic process $(Q(t))_{t \in R}$ has a special density $S(\omega)$ then

$$E(e^2(k\Delta t)) = 2\left(\int_R S(\omega) l_1(d\omega) - \frac{1}{2\pi\sigma\sqrt{2\pi}} \int \left(\int_R \exp(j\omega t - t^2/(2\sigma^2)) S(\omega) l_1(d\omega)\right) l_1(dt)\right) =$$

Using the Fubini theorem we obtain after simple calculations:

$$= \frac{1}{\pi} \int_R S(\omega) (1 - \exp(-(\omega\sigma)^2/2)) l_1(d\omega).$$

If the band of the stochastic process $(Q(t))_{t \in R}$ is limited i.e. $S(\omega) = 0$ for $|\omega| \geq \omega_g$ then

$$E(e^2(k\Delta t)) = (1/\pi) \int_{[-\omega_g, \omega_g]} S(\omega) (1 - \exp(-(\omega\sigma)^2/2)) l_1(d\omega)$$

If additionally $\sigma < \sqrt{2\pi/\omega_g}$ then because of symmetry $S(\omega) = S(-\omega)$ we have (after using series expansion) the following assessment

$$E(e^2(k\Delta t)) \leq (\sigma^4/(8\pi)) \int_{[-\omega_g, \omega_g]} S(\omega) \omega^4 l_1(d\omega).$$

4. SAMPLING TIME JITTER IN SAMPLING FREQUENCY CONVERTERS

A sampling frequency converter scheme (see [1]) is shown in the Fig. 4.1. It works in the following way. A real random analog input signal $(Q(t))_{t \in \mathbb{R}}$ (a stochastic process) is sampled with the sampling circuit 1 with a frequency $1/T_1$, where T_1 is a sampling period. Assume, $(Q(t))_{t \in \mathbb{R}}$ is a real, $\mathcal{F} \otimes B(\mathbb{R})$ measurable¹, stationary in the wide sense process (where $\mathcal{F} = \sigma(Q(t); t \in \mathbb{R})$).

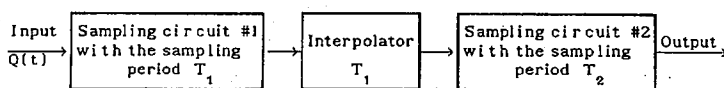


Fig. 4.1 Sampling frequency converter.

Values of the signal between samples of the sampling circuit 1 are computed with the period T_1 with interpolating circuit (T_i is a distance between interpolation points and course $\bar{T}_i < T_1$). The method of sampling and interpolation is shown in the Fig. 4.2. We assume that $T_1/T_i \in \mathbb{N}$. If $T_1/T_i = 1$ the interpolation is not applied. The sampling circuit 2 resamples the input signal with a frequency $1/T_2$ ($T_2 > T_1$). We assume that clock generators which control sampling circuits are independent i.e. unsynchronized. As a sampling value taken by the sampling circuit 2 in the moment t we admit $Q(kT_1)$, where k is an integer fulfilling the equation $|kT_1 - t| = \min_{n \in \mathbb{Z}} |nT_1 - t|$.

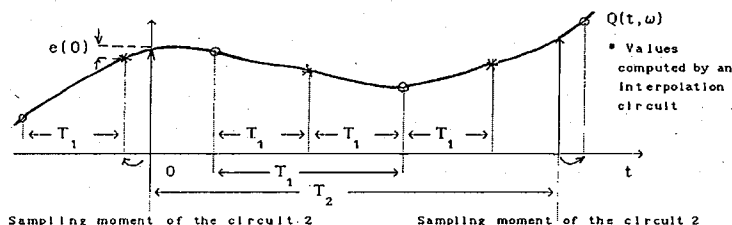


Fig. 4.2 Way of sampling in the sampling frequency converter.

As a result (we neglect interpolation error) caused by deterministic shifting (in time) of the sample can be considered as an error caused by the STJ with the uniform probability distribution in the interval $A = [-T_i/2, T_i/2]$. More precisely, assume the sampling circuit 2 resamples the input signal in the moments kT_2 . Because clock generators controlling the sampling circuits are independent, we can admit that $T(0)^0 \stackrel{\text{an}}{=} \xi$ (the random variable describing the STJ in the moment 0) has the uniform probability distribution on the interval A and the STJ is described by the random sequence $(T(kT_2))_{k \in \mathbb{Z}}$, where $T(kT_2) = [\xi + kT_2 + T_i/2]_{T_i} - T_i/2$, $k \in \mathbb{Z}$, and $[x]_a$ for $x, a \in \mathbb{R}$, $a > 0$ denotes $\min\{x - ka \geq 0; k \in \mathbb{Z}\}$. The random variable

¹ If a real stochastic process $(Q(t))_{t \in \mathbb{R}}$ is continuous then it is $\mathcal{F} \otimes B(\mathbb{R})$ measurable.

ξ has the uniform probability distribution on the interval A then for every $k \in \mathbb{Z}$ a random variable $T(kT_2)$ has the uniform distribution on A . $T(kT_2)$ denotes the random shifting of the k -th real sampling moment $kT_2 + T(kT_2)$, in comparison with the deterministic one kT_2 .

From the theorem A.1 we obtain that $(T(kT_2))_{k \in \mathbb{Z}}$ is a stationary (in the strict and wide sense) sequence of random variables. Because ξ and $(Q(t))_{t \in \mathbb{R}}$ are independent then $(T(kT_2))_{k \in \mathbb{Z}}$ and $(Q(t))_{t \in \mathbb{R}}$ are independent too. Therefore from the theorem 3.7 we have that $(Q(kT_2 + T(kT_2)))_{k \in \mathbb{Z}}$ and $(e(kT_2))_{k \in \mathbb{Z}}$ (where $e(kT_2) = Q(kT_2 + T(kT_2)) - Q(kT_2)$ is an additive noise introduced by the STJ) are stationary (in the wide sense) sequences of random variables. If the ratio $T_1/T_2 \notin \mathbb{Q}$ the probability distribution of the random variable ξ can be motivated as follows. Denote $V(k \cdot T_2) =$

$= [V(0) + T_1/2 + kT_2]_{T_1} - T_1/2$, $k \in \mathbb{Z}$. For $k=0$ we have a fixed value of phase shift $V(0)$ of $T(0)$ (i.e. $V(0) = T(0)(\omega)$). The bijection

$S: [-T_1/2, T_1/2] \ni x \rightarrow [x + T_1/2 + T_2]_{T_1} - T_1/2 \in [-T_1/2, T_1/2]$ is an ergodic function (see [9]) and $S^k = S \circ S \circ \dots \circ S = V(k \cdot T_2)$ then from the ergodic theorem (Theorem A.5 from the appendix) we obtain for an arbitrary Borel subset $B \in \mathcal{B}([-T_1/2, T_1/2])$

$$\lim_{N_0 \rightarrow \infty} \left(\frac{1}{N_0} \right) \left(\sum_{k=0}^{N_0-1} \chi_B(V(k \cdot T_2)) \right) = \frac{l_1(B)}{T_1} \quad (4.1)$$

where l_1 is the Lebesgue measure on $\mathbb{R}$. Assume the sampling circuit 2 can choose (as a start point) an arbitrary sampling moment (among N_0 consecutive sampling moments $k \cdot T_2$), with the same probability $1/N_0$. In this situation, it follows from the equality (4.1) that probability distribution of the introductory phase shift (i.e. ξ) can be considered as uniform.

Principles of two possible solutions of the sampling circuit 2 are shown in the Fig. 4.3. In the first circuit (Fig. 4.3 a) the random variable $T(kT_2)$ has for every $k \in \mathbb{Z}$ the uniform probability distribution on the interval $[-T_1, 0]$. In the second circuit (Fig. 4.3 b) the signal which controls the multiplexer MUX choose the sample

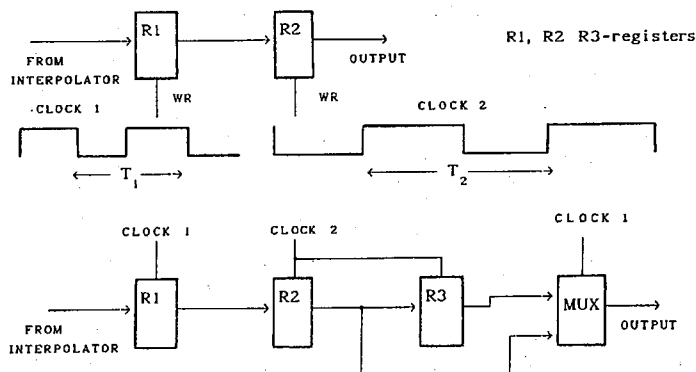


Fig. 4.3 Two possible solutions of the sampling circuit 2.

nearest (in time) to the time point kT_2 . But in this solution a sample on the output is delayed (the delay is even T_i). As a consequence we obtain that the random variable $T(kT_2)$ has for every $k \in \mathbb{Z}$ the uniform probability distribution on the interval $[-T_i/2, T_i/2]$. We admit in the sequel that random variables $T(kT_2)$; $k \in \mathbb{Z}$ have just this probability distribution.

Very similar effects like described above occur in digital oscilloscopes with fixed grid of sampling moments and stroboscopic way of sampling. In such systems a measured signal is repeated (in most cases it is simply a periodic signal) but the moment of his start is randomly situated compared with grid of sampling moments. Detailed analysis is almost the same as in the case of sampling frequency converters.

Our aim in the next section is to assess parameters of the additive error noise $(e(kT_2))_{k \in \mathbb{Z}}$.

5. ADDITIVE NOISE PARAMETERS IN SAMPLING CONVERTERS

In the section 4 we verified that assumptions of the theorem 3.7 are fulfilled. Then according to the formulae (3.17) and (3.18) the mean value $m_1(k) = E(Q(kT_2 + T(kT_2)))$ and autocorrelation function $R_1(k) = E(Q(kT_2 + T(kT_2))Q(T(0)))$ of $(Q(kT_2 + T(kT_2)))_{k \in \mathbb{Z}}$ are given for every $k \in \mathbb{Z}$ by the following formulae: $m_1(k) = m(kT_2) = E(Q(kT_2))$ and

$$R_1(k) = \int_{\mathbb{R}} R(kT_2 + x) P_{Y(k)}(dx), \quad (5.1)$$

where $P_{Y(k)}$ is the probability distribution of the random variable $Y(k) = T(kT_2) - T(0)$. Similarly, if $m_2(k)$ and $R_2(k)$ denote the mean value and autocorrelation function of the error noise $(e(kT_2))_{k \in \mathbb{Z}}$ then $m_2(k) = 0$ and

$$R_2(k) = R_1(k) + R(kT_2) - E(R(kT_2 + T(0))) - E(R(kT_2 - T(0))). \quad (5.2)$$

From the above formula we obtain that the error noise power is even

$$E(e^2(kT_2)) = 2(R(0) - E(R(T(kT_2)))) = 2(R(0) - \frac{1}{T_i} \int_A R(t) l_1(dt)). \quad (5.3)$$

If a stationary in wide sense stochastic process $(Q(t))_{t \in \mathbb{R}}$ has a spectral density $S(\omega)$ then using the Fubini theorem and the property of symmetry $S(\omega) = S(-\omega)$ we have

$$E(e^2(kT_2)) = 2\left(\frac{1}{2\pi} \int_{\mathbb{R}} S(\omega) l_1(d\omega) - \frac{1}{T_i} \int_A \left(\frac{1}{2\pi} \int_{\mathbb{R}} S(\omega) e^{j\omega t} l_1(d\omega)\right) l_1(dt)\right) = \quad (5.4)$$

$$\begin{aligned}
&= \frac{1}{\pi} \int_R S(\omega) l_1(d\omega) - \frac{1}{\pi T_i} \int_R S(\omega) \int_A e^{j\omega t} l_1(dt) l_1(d\omega) = \\
&= \frac{2}{\pi} \left(\int_{R^+} S(\omega) \left(1 - \frac{\sin(\omega T_i/2)}{(\omega T_i/2)} \right) l_1(d\omega) \right) \quad (5.4 \text{ c.d.})
\end{aligned}$$

If the power density function $S(\omega) = 0$ for $|\omega| > \omega_g$ i.e. the band of the random signal is bounded then

$$E(e^2(kT_2)) = \frac{2}{\pi} \left(\int_{[0, \omega_g]} S(\omega) \left(1 - \frac{\sin((\omega T_i)/2)}{(\omega T_i/2)} \right) l_1(d\omega) \right). \quad (5.5)$$

Using the power series expansion of the function $(\sin x)/x$ we have for every $x \in R, x \neq 0$

$$\frac{\sin x}{x} = 1 - \frac{x^2}{3!} + \frac{x^4}{5!} - \dots$$

and for $0 < |x| \leq 1, 0 \leq 1 - \frac{\sin x}{x} \leq \frac{x^2}{6}$. (5.6)

In practice we have $\frac{2\pi}{T_i} > 2\omega_g$ (because usually assumptions of the sampling theorem are fulfilled) and $T_i \ll T_1$, then $\frac{\omega T_i}{2} \ll 1$ for $0 \leq \omega \leq \omega_g$ and using the formula (5.6) we have

$$E(e^2(kT_2)) \leq \frac{T_i^2}{12\pi} \int_{[0, \omega_g]} \omega^2 S(\omega) l_1(d\omega) \quad (5.7)$$

It follows from formulae (5.7) and (5.5) that the power of the error signal is bounded by $\text{const} \cdot T_i^2$ (then roughly speaking less T_i less power.) but small T_i complicates (real time) computations in interpolating circuit. Using the formula (3.18) we can also compute the correlation function $R_2(k)$ of the error process $(e(kT_2))_{k \in Z}$. Because the random variable $T(0)$ is symmetric the formula (3.18) can be written in the following way

$$R_2(k) = \int_R R(kT_2 + x) P_G(dx) + R(kT_2) - 2E(R(kT_2 + T(0))), \quad (5.8)$$

where $G = T(kT_2) - T(0)$. The difference $R(kT_2) - 2E(R(kT_2 + T(0)))$ can be computed in the similar way as in the case $E(e^2(kT_2))$. To compute the first term of the

formula (3.18) we have to find first the probability distribution of the random variable $Y(k)$

$$Y(k) \stackrel{\text{df}}{=} T(kT_2) - T(0) = [\xi + kT_2 + T_i/2]_{T_i} - [\xi + T_i/2]_{T_i} = \\ = [[\xi + T_i/2]_{T_i} + [kT_2]_{T_i} - [\xi + T_i/2]_{T_i}]. \quad (5.9)$$

The random variable $[\xi + T_i/2]_{T_i}$ has the uniform probability distribution on $[0, T_i]$ then the random variable $Y(k)$ has the following two points distribution

$$P(Y(k) = [kT_2]_{T_i}) = \frac{T_i - [kT_2]_{T_i}}{T_i}, \quad P(Y(k) = [kT_2]_{T_i} - T_i) = \frac{[kT_2]_{T_i}}{T_i} \text{ and}$$

$$R_1(k) = \int_R R(kT_2 + x) P_G(dx) = R(kT_2 + [kT_2]_{T_i}) \frac{T_i - [kT_2]_{T_i}}{T_i} + R(kT_2 + [kT_2]_{T_i} - T_i) \frac{[kT_2]_{T_i}}{T_i}$$

where $G = T(kT_2) - T(0)$. If $T_2/T_i = m_1/m_2 \in Q$, where $m_1, m_2 \in Z$ and $GPD(m_1, m_2) = 1$ and we denote $T_2/m_1 = T_i/m_2 = w$ then $[kT_2]_{T_i} = [km_1 w]_{m_2 w} = w[km_1]_{m_2}$. An element $[m_1]_{m_2} \in Z_{m_2} = \{0, 1, \dots, m_2 - 1\}$ is a generator of the additive group of the ring Z_{m_2} . Because $GPD(m_1, m_2) = 1$, then $\{[km_1]_{m_2}; k \in Z\} = Z_{m_2}$ and therefore for every $k \in Z$ there exists such a $k_0 \in Z_{m_2}$, that $Y(k) = Y(k_0)$. Hence there exist only finite number m_2 of different probability distributions of the random variables $Y(k)$ for $k \in Z$. If $T_2/T_i \notin Q$ then the set $\{[kT_2]_{T_i}; k \in Z\}$ is dense on the interval $[0, T_i]$ and $[k_1 T_2]_{T_i} \neq [k_2 T_2]_{T_i}$ for every $k_1, k_2 \in Z$ and $k_1 \neq k_2$. Therefore for every $k_1, k_2 \in Z$ and $k_1 \neq k_2$ probability distributions of the random variables $Y(k_1)$ and $Y(k_2)$ are different.

Remark. It is easily seen that formulae (5.5), (5.7) are true also in a case an arbitrary random sequence $(T(k\Delta t))_{k \in Z}$ describing the STJ if only random variables $T(k\Delta t)$ have the uniform probability distribution on the interval $[-T_i/2, T_i/2]$ for every $k \in Z$. For example the sequence $(T(k\Delta t))_{k \in Z}$ can be a sequence of independent random variables.

CONCLUSIONS

It was shown in the paper that the STJ present in sampling frequency converters has (with probability 1) the uniform probability distribution.

Statistical parameters of jittered input signal and error noise introduced by the STJ can be computed with proved formulae.

² It can be easily proved that if T_2 and T_i are treated as random variables and a two dimensional random variable (T_2, T_i) has the probability density function $f: R^2 \rightarrow R$ then the probability of choosing T_2 and T_i so that $T_2/T_i \in Q$ is equal to zero. In particular it holds the random variables T_i and T_2 are independent with uniform probability distributions.

The STJ occurs in digital oscilloscopes with fixed (clock controlled) grid of sampling moments is very similar to the above described model of the STJ in sampling converters then obtained results can be also useful in accuracy assesment of digital oscilloscopes and waveform recorders.

APPENDIX

Theorem A1. If $f: [0, 2\pi] \rightarrow \mathbb{R}$ is a $B((0, 2\pi), B(\mathbb{R}))$ measurable function, $(A, \psi): \Omega \rightarrow \mathbb{R}^2$ a two dimensional random variable on the probabilistic space $(\Omega, \mathfrak{M}, P)$ and a real random variable φ has the uniform probability distribution on the interval $[0, 2\pi]$ and φ and (A, ψ) are independent then the stochastic process $(Af([\psi t + \varphi]_{2\pi}))_{t \in \mathbb{R}}$, where $[x]_{2\pi} = \min\{x - k \cdot 2\pi \geq 0; k \in \mathbb{Z}\}$ for $x \in \mathbb{R}$ is stationary in the strict sense. If additionally the function f is bounded and the random variable $A \in L^2(\Omega, \mathfrak{M}, P)$ then the stochastic process $(A \cdot f([\psi t + \varphi]_{2\pi}))_{t \in \mathbb{R}}$ is stationary in wide sense.

Proof. 1) For every $t_1, t_2, \dots, t_n \in \mathbb{R}$, $t_1 < t_2 < \dots < t_{2n}$, a function

$G: \mathbb{R}^3 \ni (x_1, x_2, x_3) \rightarrow (x_1 f([x_2 t_1 + x_3]_{2\pi}), x_1 f([x_2 t_2 + x_3]_{2\pi}), \dots, x_1 f([x_2 t_n + x_3]_{2\pi})) \in \mathbb{R}^n$ is $(B(\mathbb{R}^3), B(\mathbb{R}^n))$ measurable because for every $i = 1, 2, \dots, n$, function $\pi_i \circ G$ is $(B(\mathbb{R}^3), B(\mathbb{R}))$ measurable, where $\pi_i: \mathbb{R}^n \ni (x_1, x_2, \dots, x_n) \rightarrow x_i \in \mathbb{R}$ i.e. π_i is a projection.

2) For every $x \in \mathbb{R}$ a random variable $[x + \varphi]_{2\pi}$ has the uniform probability distribution on $[0, 2\pi]$. Similarly, a random variable $\tilde{\varphi} \stackrel{\text{def}}{=} [\psi t + \varphi]_{2\pi}$ for arbitrary fixed $t \in \mathbb{R}$ has the uniform probability distribution on $[0, 2\pi]$. This follows from the Fubini theorem and the assumption that random variables ψ and φ are independent. More precisely, let $B \in B(\mathbb{R})$, and denote $C = \{(x_1, x_2) \in \mathbb{R}^2; [x_1 t + x_2]_{2\pi} \in B\}$ then we obtain

$$\begin{aligned} P([\psi t + \varphi]_{2\pi} \in B) &= E(\chi_{\{[\psi t + \varphi]_{2\pi} \in B\}}) = \int_{\mathbb{R}^2} \chi_C(x_1, x_2) P_{(\psi, \varphi)}(dx_1, dx_2) = \\ &= \int_{\mathbb{R}^2} \chi_C((x_1, x_2)) P_\psi \otimes P_\varphi(dx_1, dx_2) = \int_{\mathbb{R}} \left(\int_{\mathbb{R}} \chi_C(x_1, x_2) P_\varphi(dx_2) \right) P_\psi(dx_1) = \\ &= \int_{\mathbb{R}} P_\varphi(C_{x_1}) P_\psi(dx_1) = \int_{\mathbb{R}} P([x_1 t + \varphi]_{2\pi} \in B) P_\psi(dx_1) = \frac{1_B(B \cap [0, 2\pi])}{2\pi}, \end{aligned}$$

where C_{x_1} denotes the section of the set C for fixed $x_1 \in \mathbb{R}$, i.e. $C_{x_1} = \{x_2 \in \mathbb{R}; [x_1 t + x_2]_{2\pi} \in B\}$.

3) For arbitrary $B_1 \in B(\mathbb{R})$, denote $A \stackrel{\text{def}}{=} \{[\psi t + \varphi]_{2\pi} \in B_1\} \in B_1 \in \sigma([\psi t + \varphi]_{2\pi}) \subset \mathfrak{M}$ and $B_2 = \{(x, y); [xt + y]_{2\pi} \in B_1\}$. The random variables (A, ψ) and φ are independent then we have from the theorem A.6

$$P([\psi t + \varphi]_{2\pi} \in B_1 | \sigma(A, \psi)) = P(A_1 | \sigma(A, \psi)) = E(\chi_{A_1} | \sigma(A, \psi)) = E(\chi_{B_2}(\psi, \varphi) | \sigma(A, \psi)) = \\ = E(\chi_{B_2}(x, \varphi))|_{x=\psi} = P([xt + \varphi]_{2\pi} \in B_1)|_{x=\psi} = \frac{I_1(B_1 \cap (0, 2\pi))}{2\pi} = P(A_1).$$

Hence, for every $A_1 \in \sigma([\psi t + \varphi]_{2\pi})$ and $A_2 \in \sigma(A, \psi)$, we obtain

$$P(A_1 \cap A_2) = \int_{A_2} P(A_1 | \sigma(A, \psi)) dP = P(A_1)P(A_2),$$

then σ -fields $\sigma(\tilde{\varphi})$ i $\sigma(A, \psi)$ are independent.

4) To prove stationary of the stochastic process $(Af[\psi t + \varphi]_{2\pi})_{t \in \mathbb{R}}$ we have to prove, that for each $n \in \mathbb{N}$, each $B \in \mathcal{B}(\mathbb{R}^n)$ and each sequence $t, t_1, t_2, \dots, t_n \in \mathbb{R}$, $t_1 < t_2 < \dots < t_n$ the following formula is true

$$P((Af[\psi t_1 + \varphi]_{2\pi}), Af([\psi t_2 + \varphi]_{2\pi}), \dots, Af([\psi t + \varphi]_{2\pi})) \in B) = \\ = P((Af[\psi(t_1 + t) + \varphi]_{2\pi}), Af[\psi(t_2 + t) + \varphi]_{2\pi}), \dots, Af[\psi(t_n + t) + \varphi]_{2\pi})) \in B)$$

or equivalently we have $P_{(A, \psi, \varphi)}(G^{-1}(B)) = P_{(A, \psi, \tilde{\varphi})}(G^{-1}(B))$ (*).

Because the random variables (A, ψ) and $\tilde{\varphi}$ are independent then $P_{(A, \psi, \tilde{\varphi})} = P_{(A, \psi)} \otimes P_{\tilde{\varphi}}$ and similarly, because (A, ψ) and φ are independent then we have $P_{(A, \psi, \varphi)} = P_{(A, \psi)} \otimes P_{\varphi}$. The random variables $\tilde{\varphi}$ i φ have (as we proved it in the point 2) the same probability distribution, then (*) is true.

5) Stationarity in the wide sense when the function f is bounded and the random variable $A \in L^2(\Omega, \mathcal{M}, P)$ is obvious. Q.E.D.

the interval $[0, 2\pi)$ can be replaced in the above theorem by an arbitrary another one $[0, a)$.

Corollary A.2 Under assumptions from the above theorem the stochastic processes $(A \sin(\psi t + \varphi))_{t \in \mathbb{R}}$ i $(A[\psi t + \varphi]_{2\pi})_{t \in \mathbb{R}}$ are stationary in the strict sense. If additionally $A \in L^2(\Omega, \mathcal{M}, P)$ then these stochastic processes are also stationary in the wide sense.

Theorem A.3 (Birkhoff – Tchinchin ergodic theorem)

Assume $(M, \mathcal{M}, \mu)$ is a probabilistic space with the complete measure μ , $f \in L^1(M, \mathcal{M}, \mu)$ and $T: M \rightarrow M$ is an endomorphism of the space with measure $(M, \mathcal{M}, \mu)$. If a dynamical system $(M, \mathcal{M}, \mu, T)$ is ergodic then for μ almost each $x \in M$, we have

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=1}^{n-1} f(T^k x) = \int_M f(x) \mu(dx),$$

where $T^k = \underbrace{T \circ T \circ \dots \circ T}_k$.

The proof can be found in [8].

Theorem A.4 If $(Q(t))_{t \in R}$ is a real $\mathcal{F} \otimes B(R)$ measurable, stochastic process, where $\mathcal{F} = \sigma(Q(t); t \in R)$, $T_1, T_2, \dots, T_n$ arbitrary real random variables and σ -field $\mathcal{F} = \sigma(Q(t), t \in R)$ and $\sigma(T_1, T_2, \dots, T_n)$ are independent then for every $B \in B(R^n)$

$$P((Q(T_1), Q(T_2), \dots, Q(T_n)) \in B) = \int_{R^n} P((Q(t_1), Q(t_2), \dots, Q(t_n)) \in B) P_{(T_1, T_2, \dots, T_n)}(d(t_1, \dots, t_n))$$

If additionally for every $t_1, t_2, \dots, t_n \in R$, where $t_i \neq t_j$ for $i \neq j$ random variable $(Q(t_1), Q(t_2), \dots, Q(t_n))$ has a density probability function $f_{t_1, t_2, \dots, t_n}: R^n \rightarrow R^+$ (related to Lebesgue measure in R^n and $f_{t_1, t_2, \dots, t_n}(x_1, x_2, \dots, x_n)$ as a function of the random variables $(t_1, t_2, \dots, t_n, x_1, x_2, \dots, x_n) \in R^{2n}$ is $(B(R^{2n}), B(R))$ measurable and for every $i \neq j$, $P(T_i = T_j) = 0$ then a density probability function f of random variable $(Q(T_1), Q(T_2), \dots, Q(T_n))$ is given by the formula

$$f(x_1, x_2, \dots, x_n) = \int_R f_{t_1, t_2, \dots, t_n}(x_1, x_2, \dots, x_n) P_{(T_1, \dots, T_n)}(d(t_1, t_2, \dots, t_n)).$$

Proof. see [10].

Theorem A.5 (theorem on variables change)

If $(X, \mathcal{M}, \mu)$ is a space with a measure μ , $(Y, \mathcal{F})$ a measurable space, functions $T: (X, \mathcal{M}, \mu) \rightarrow (Y, \mathcal{F})$, $f: (Y, \mathcal{F}) \rightarrow (R, B(R))$ are measurable and a ν a measure induced on the measurable space $(Y, \mathcal{F})$ by the function T , then

$$1) f \in L^1(Y, \mathcal{F}, \nu) \equiv f \circ T \in L^1(X, \mathcal{M}, \mu)$$

$$2) \int_X f \circ T d\mu = \int_Y f d\nu$$

Theorem A.6 Let $(\Omega, \mathcal{M}, P)$ be a probabilistic space, $(X, \mathcal{B}_1)$, $(Y, \mathcal{B}_2)$ measurable spaces, $h: (X \times Y, \mathcal{B}_1 \otimes \mathcal{B}_2) \rightarrow (R, B(R))$ a measurable function, ξ, η random variables $\xi: (\Omega, \mathcal{M}, P) \rightarrow (X, \mathcal{B}_1)$, $\eta: (\Omega, \mathcal{M}, P) \rightarrow (Y, \mathcal{B}_2)$, $\mathcal{F}$ σ -field, $\mathcal{F} \subset \mathcal{M}$, $\sigma(\eta) \subset \mathcal{F}$, and $h(\xi, \eta) \in L^1(\Omega, \mathcal{M}, P)$ then

$$1) E(h(\xi, \eta) | \mathcal{F}) = E(h(\xi, y) | \mathcal{F})_{y=\eta}$$

2) if additionally random variable ξ is independent from $\mathcal{F}$ then

$$E(h(\xi, \eta) | \mathcal{F}) = E(h(\xi, y) | \mathcal{F})_{y=\eta}$$

LIST OF SYMBOLS

- $(\Omega, \mathcal{M}, P)$ — probabilistic space
 $\text{dist}(A, x)$ — distance between a set A and a point x in a metric space

N	— set of natural numbers
Z	— set of integers
Q	— set of rational numbers
R	— set of real numbers
R^+	— set of nonnegative real numbers, $R^+ = \{x \in R; x > 0\}$
C	— set of complex numbers
$B(X)$	— σ -field of Borel sets of metric space X
$(X, \mathcal{B})$	— measurable space
l_1	— Lebesgue measure on R
$L^p(\Omega, \mathcal{M}, P)$	— $p \geq 1$, set of all random variable having the p -th moment
$L^p(R, \mathcal{L}, l_1)$	— $p \geq 1$, set of p -th power integrable real functions over R
$D^2(X)$	— variance of the random variable X
$[x]_a$	— real number x modulo $a > 0$ i.e. $[x]_a \stackrel{\text{df}}{=} \min \{x - ka > 0; k \in Z\}$
$GPD(n, m)$	— greatest possible divider of integers n and m
$\mu_1 \otimes \mu_2$	— product of measure μ_1, μ_2
$\mathcal{X}(A, \omega_g)$	— a set of real band limited (by ω_g) functions over R with values in the interval $[-A, A]$

REFERENCES

1. G. Ver kroost: *Noise Caused by Sampling Time Jitter with Applications to Sampling Frequency Conversion*. Signal Processing, Theory and Applications; EURASIP 1986
2. B. Li u: *Effect of Timing Jitter on Performance of MTI Filters*. IEEE Transactions on Aerospace and Electronic Systems, May 1989
3. M. Sh i n a g a w a, Y. A k a z a w a, T. W a k i m o t o: *Jitter Analysis of High Speed Sampling Systems*. IEEE Journal of Solid-State Circuits, Febr. 1990
4. T. M. S o u l d e r s: *The Effects of Timing Jitter in Sampling Systems*. IEEE Transaction on Instrumentation and Measurement, February 1990
5. M. F. W a g d y, S. S. A w a d: *Effect of Sampling Jitter on Some Wave Measurement*. IEEE Transactions on Instrumentation and Measurement, February 1990
6. W. G a n s: *The Measurement and Deconvolution of Time Jitter in Equivalent Time Waveform Samplers*. IEEE Proceedings; March 1983
7. A. W. B a l a k r i s h n a n: *On the Problem of Time Jitter in Sampling*. IRE Trans. IT-8, 1962
8. S. W. F o m i n, I. P. S i n a j: *Ergodic Theory* (in Russian); Nauka; Moskwa 1980
9. T. A d a m s k i: *Jitter Noise in Sampling Frequency Converters*. Prace CVUT w Praze, 1989, vol.3, No 12, Prague
10. T. A d a m s k i: *Jitter in Electronic Circuits and Systems*. Scientific Report; Institute of Electronics Fundamentals, Warsaw University of Technology, Warsaw 1992
11. T. A d a m s k i: *Some Remarks on Phase Detector Circuit*. Electronics and Telecommunications Quarterly; 1992, Vol.38, 4
12. T. A d a m s k i: *Sampling Time Jitter Correction Based on Quantile Filtering; Measurement* (in printing)

T. ADAMSKI

ANALIZA SZUMU ADDYTYWNEGO WPROWADZONEGO PRZEZ NIESTAŁOŚĆ
CHWILOWĄ MOMENTU PRÓBKOWANIA
W KONWERTERACH CZĘSTOTLIWOŚCI PRÓBKOWANIA

Streszczenie

Mimo, że konwertery próbkowania pracują w sposób deterministyczny, w systemach tego typu występują zjawiska, które w swej istocie są równoważne losowej niestałości momentu próbkowania. Niestalość ta wynika z braku synchronizacji pomiędzy sygnałem zegara wpisującym dane do konwertera i sygnałem zegara, w takt którego dane są odczytywane. Pojawienie się niestałości momentu próbkowania prowadzi do wystąpienia błędów wartości próbek odczytywanych, co można zinterpretować jako szum addytywny nakładający się na dyskretny wyjściowy sygnał użyteczny. W pracy zostały przeanalizowane statystyczne własności tego szumu oraz obliczone jego parametry, wartości średnie oraz funkcje korelacyjne. Uzyskane wyniki pozwalają na ocenę dokładności działania konwerterów częstotliwości próbkowania i systemów zbliżonych do nich sposobem działania — takich jak np. oscyloskopy cyfrowe.

Analiza fotoemisji elektronów

JÓZEF WOJAS

Profesor emerytowany, Politechnika Warszawska

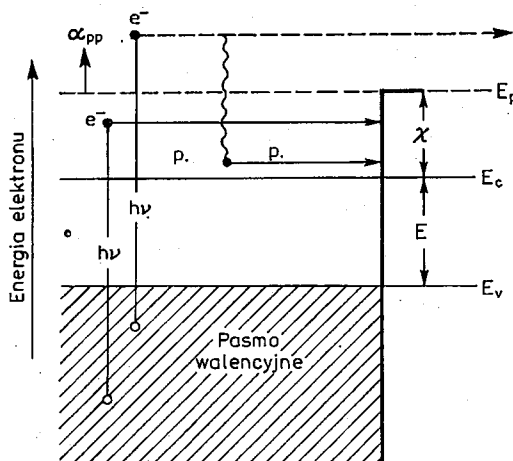
Otrzymano 1993.01.10

Autoryzowano do druku 1993.06.28

W artykule przedstawiono zwięzłą analizę zjawiska fotoemisji i jej korelacje z przejściami optycznymi; bowiem pobudzenie i wyjście elektronów poza kryształ realizuje się prostymi lub nieprostymi przejściami optycznymi. Wyjaśniono, że od wymienionych przejść zależy wydajność kwantowa fotoemisji, która jest funkcją energii fotonów.

1. WSTĘP

Według współczesnych teorii proces fotoemisji składa się z trzech etapów: a) absorpcji optycznej, b) transportu pobudzonego elektronu do powierzchni (elektron może tracić energię podczas tego procesu) oraz c) wyrzucenie elektronu poprzez



Rys. 1. Trzy stopniowy model fotoemisji z półprzewodnika (χ — powinowactwo elektronowe, E_g — przerwa energetyczna, E_p — poziom próżni, α_p — współczynnik absorpcji dla wzbudzenia do stanów powyżej poziomu próżni). Jak pokazano, elektron może tracić energię w poszczególnych etapach transportu zanim osiągnie powierzchnię

powierzchnię do próżni. Te etapy emisji są zilustrowane na rys. 1. Na każdym z tych etapów decydujący jest pewien mechanizm efektów fotoelektrycznych, który wywiera wpływ (ten mechanizm) na przebieg krzywych $N(E)$ [1].

2. EMISJA FOTOELEKTRYCZNA I JEJ ENERGETYCZNE ROZKŁADY

W przypadku powierzchni realnych próbka jest trawiona, lub inaczej obrabiana przed jej zamocowaniem w próżniowym urządzeniu pomiarowym. Czyste, lub odkrywane w próżni powierzchnie kryształu, są wytwarzane w tej samej wysoko-próżniowej komorze. Kilka cykli oczyszczania, lub operacja łupania mogą być przeprowadzone wewnątrz komory połączonej z mieszkciem wysokociśnieniowym.

W pomiarach fotoemisji dobrze odgraniczona monochromatyczna plamka światła jest ogniskowana poprzez okienko kwarcowe na pożądanym obszarze powierzchni próbki. W celu eliminowania fotoemisji z innych części komory wiązka światła jest orientowana prostopadle do powierzchni. W ten sposób światło odbite, jak od zwierciadła, może opuścić komorę tym samym okienkiem, przez które weszło. Pewna część wiązki, (około 10%), jest czasem wydzielona i zogniskowana na fotoczulym detektorze dla zmierzenia natężenia światła.

Dla określenia strumienia fotonów, w tym pomiarze absorbowanego przez próbkę, jest wymagana znajomość współczynnika odbicia R powierzchni w funkcji długości fali. Jest to możliwe dla wielu półprzewodników; w innych przypadkach można mieć zastrzeżenia mierząc odbicie na zewnątrz, podczas gdy próbka znajduje się w komorze próżniowej i jest mało dostępna. Prąd fotoemisji jest często tak mały (10^{-12} A), że potrzebne jest zastosowanie czułych technik pomiarowych, ekranowanie połączeń z elektrometrem, lub modulowane światło ze wzmacniaczem (eliminacja wpływu prądów ciemnych). Dla eliminacji fałszywych prądów emisji z powierzchni innych niż próbka, montuje się ją na izolacyjnym podłożu, a wszystkie inne części układu są na potencjale ziemi i służą jako kolektor. Nie jest konieczne, jeśli próg emisji jest w zakresie widzialnym, bo wówczas emisja z innych powierzchni jest do zaniechania.

Rozkład energetyczny emitowanych elektronów $N(E)$ może być określony kilkoma sposobami, między innymi przez magnetyczne odchylenia torów elektronów [2]. Na osiągnięcie przebiegu $N(E)$ nakłada się kilka efektów i kilka mechanizmów, przy czym jedne mechanizmy odgrywają dużą rolę na prawym zboczku wykreślanego rozkładu, a inne w obszarze niskich energii — czyli na lewym zboczku, gdzie najsilniejszym „tłumikiem” elektronów jest ich rozpraszanie, szczególnie z głębszej geometrycznie warstwy. Niskoenergetyczne elektrony są o wiele silniej rozpraszane, niż te z pobliża szczytu pasma walencyjnego.

Przy nałożeniu się na siebie kilku funkcji, z których jedna decyduje o wzroście liczby elektronów z coraz głębszych poziomów, inne zaś wpływają na zmniejszanie się $N(E)$, musi istnieć przedział, w którym wystąpi maksimum tego przebiegu $N(E)$ pozostającego pod wpływem kilku czynników. I jest to przyczyną maksimum

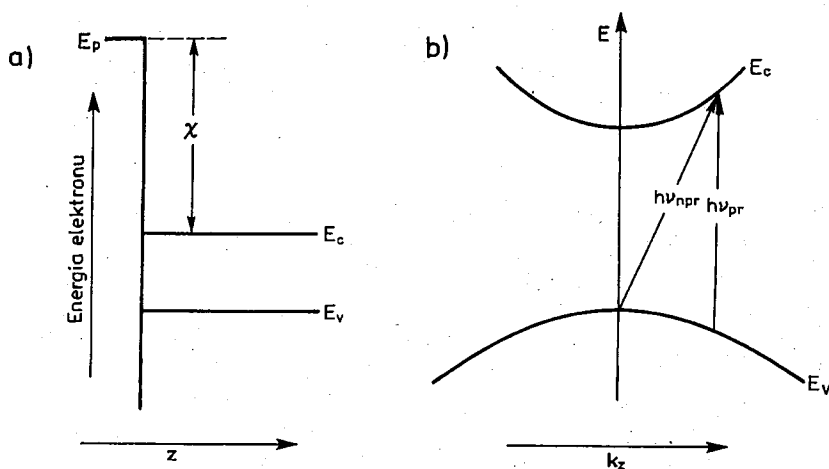
krzywej. Jeśli lewe zbocze jest w przybliżeniu prostoliniowe, to głównym mechanizmem rozpraszania jest objętościowe (na drodze elektronu do powierzchni) rozproszenie elektronów, których średnia droga swobodna jest proporcjonalna do E . Jeśli zbocza są wklęsłe — dołącza się przejście przez barierę.

W literaturze naogół szacuje się odległość energii maksimum krzywej ($E_{N_{\max}}$) od zera energii kinetycznej. Według autora, bardziej istotnych informacji udzielają położenia $E_{N_{\max}}$ względem wysokoenergetycznych końców rozkładów $E_{\max}$ [3–6].

Widma energetyczne elektronów w funkcji ich energii kinetycznych odzwierciedlają więcej cech fotoemisji niż rozkłady wydajności $Y(h\nu)$; są one bardziej czułe na subtelności struktury pasmowej oraz na strukturę elektronową powierzchni [7–12].

3. PRZEJŚCIA OPTYCZNE I WYDAJNOŚĆ FOTOELEKTRYCZNA

Wydajność z funkcji długości fali padającego światła można zmierzyć eksperymentalnie. Pobudzenie i wyjście elektronów z pasma walencyjnego przez fotony $h\nu$ realizuje się prostym lub skośnym przejściem optycznym [13–22]. Wykazano [23], że w nieobecności bariery potencjału przy powierzchni pasma danego półprzewodnika są niezagięte (rys. 2a). Rys. 2b przedstawia wykres energii elektronów w funkcji składowej z wektora falowego k dla pasma walencyjnego i pasma przewodnictwa.



Rys. 2a. Model płaskich pasm energetycznych półprzewodnika

Rys. 2b. Wykres energii elektronów w funkcji składowej z wektora falowego k . Schematyczne przedstawienie przejść prostych i nieprostych biorących udział w fotoemisji w warunkach pasm płaskich

Z przejściami prostymi i nieprostymi wiążą się progi proste i nieproste $h\nu_{t,pr}$ i $h\nu_{t,np}$. Dla przejścia prostego próg emisji nazywany jest progiem prostym, lub progiem przejść prostych, a dla przejścia nieprostego wprowadzono nazwę progu niepros-

tego. Dla przejść nieprostych próg $E_{t_{npr}}$ fotoemisji odpowiada energii fotonów $h\nu_{npr}$ równej energetycznej odległości pomiędzy szczytem pasma walencyjnego i poziomem próżni. Przy czym energia fotonu wymagana w koniecznym przeniesieniu momentu pędu jest do zaniedbania. Stąd $E_{t_{npr}} = \kappa + E_G$. Dla przejść prostych wektor falowy k jest zachowany i w konsekwencji (rys. 2b) ogólnie wymaga się do wzudzenia fotoemisji wyższej energii fotonu $h\nu_d$. Na wielkość energii progowej przejść prostych $E_{t_{pr}}$ może wpływać fakt, czy pobudzone elektrony podlegają możliwym do oszacowania rozproszeniom na fotonach porzedzającym ich wyjście z kryształu [24]. Jeśli nie ma rozprożeń, $E_{t_{pr}}$ będzie większe niż $E_{t_{npr}}$, o energię dziury wytworzonej poniżej szczytu pasma walencyjnego, z dodaniem energii kinetycznej emitowanego elektronu w kierunku równoległym do powierzchni, który to kierunek nie jest modulowany przez emisję. Rozproszenia mogą zmniejszyć próg przejść prostych, ponieważ nie spełniają one wymagania stałości prostopadłego momentu. Rozproszenia mogą czasem przeorientować kierunki elektronów o niższych energiach umożliwiając im wyjście z kryształu. Można oczekiwać, że fotoelektryczna wydajność wolno rośnie jeśli energia $h\nu$ przewyższa próg przejść nieprostych, a gdy zaś próg prosty jest przekroczony wtedy pokazuje gwałtowny wzrost wydajności i bardziej wydajny proces pobudzenia zaczyna być dominujący [25; 26; 11].

Kane [27] wyprowadził teoretycznie funkcję wydajności w pobliżu punktów progowych dla ogólnej postaci struktury pasmowej i dla różnych procesów pobudzenia i rozproszenia. Teoria ta jest oparta na rozpatrzeniu przebiegu w funkcji gęstości stanów, przy czym przyjmuje się, że absorpcja optyczna i prawdopodobieństwa przejść nie zmieniają się ostro blisko progu. Analiza wykazuje, że każdy z wprowadzonych procesów daje wydajność powyżej progu E_t w postaci prawa potęgowego

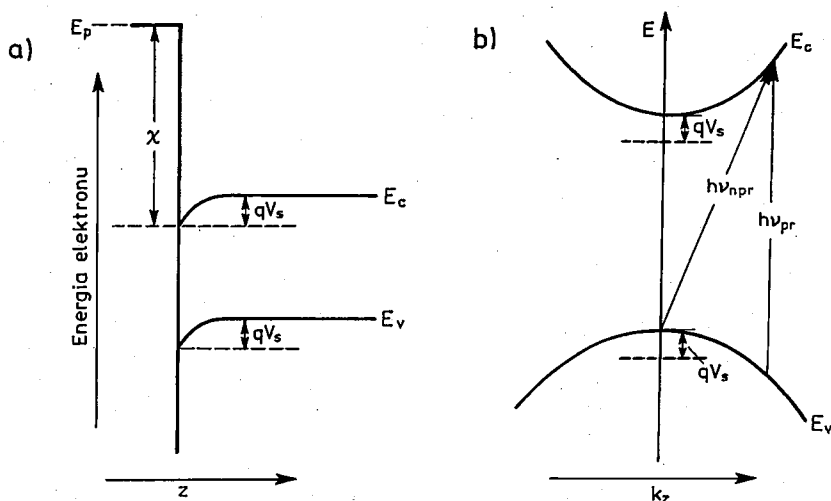
$$Y \sim (h\nu - E_t)^m \quad (1)$$

gdzie m przyjmuje wartości całkowite lub połówkowe pomiędzy 1 i 5/2 w zależności od typu pobudzenia i rozproszenia.

Odpowiednie energie progowe będą wyraźnie różnić się w przypadku istnienia bariery potencjału przy powierzchni w porównaniu z pasmami płaskimi; będą niższe dla zagięcia pasm w dół $V_s > 0$, a wyższe dla zagięcia pasm w górę $V_s < 0$, co jest przedstawione na rys. 3a i 3b. Przyjmuje się, że przejścia pochodzą z punktu tuż poniżej warstwy ładunku przestrzennego przy założeniu, że penetracja światła i głębokości wyrzutu elektronów są dostatecznie duże dla takich przejść.

Dla przejść nieprostych próg $E_{t_{npr}}$ będzie zredukowany o wartość qV_s . Z drugiej strony dla przejść prostych należy wziąć pod uwagę, że pobudzenie energią progową wprowadza teraz mniejszą wartość $|k_z|$ odpowiadającą niższej energii dziury pozostawionej w pasmie walencyjnym. Zredukowanie wartości $E_{t_{pr}}$ będzie większe niż qV_s o wielkość związaną ze szczegółową strukturą obu pasm. Próg prosty może więc być zmniejszony w przybliżeniu o bqV_s , gdzie współczynnik liczbowy $b > 1$ i zazwyczaj jest określany eksperymentalnie [12].

W celu obliczenia całkowitej wydajności należy całkować w głąb od powierzchni kryształu rozważając zarówno efekt zagięcia pasm, a także osłabienie światła



Rys. 3a. Model zagiętych pasm energetycznych półprzewodnika

Rys. 3b. Schematyczne przedstawienie przejść prostych i nieprostych biorących udział w fotoemisji w warunkach warstwy akumulacyjnej

(zmniejszenie natężenia światła) i liczby emitowanych elektronów w funkcji głębokości. Pierwszy efekt jest określony przez ustalenie prostych i nieprostych progów w pewnym punkcie „z” o potencjale V jako: $E_{t_{npr}} + |e|(V - V_s)$ dla przejścia nieprostego i $E_{t_{pr}} + b|e|(V - V_s)$ dla prostego [28–31; 23; 24]. Współczynnik b przyjmuje się jako niezależny od „z”.

Oslabienie emisji wraz z głębokością jest wprowadzane przez parametr l określony jako $\left(\frac{1}{l_v} + \frac{1}{l_e}\right)^{-1}$, gdzie l_v oraz l_e są odpowiednio głębokością absorpcji optycznej i głębokością wyrzutu elektronów; obie wartości przyjmuje się jako stałe w zakresie stosowanych długości fal. Stąd, jeśli rozpatruje się nierozproszone przejścia proste oraz nieproste, wydajność może być wyrażona w postaci:

$$Y(h\nu) = a_{npr} \int_0^{\infty} [h\nu - E_{npr} - q(V - V_s)]^{5/2} e^{-z/l} dz + \\ + a_{pr} \int_0^{\infty} [h\nu - E_{pr} - bq(V - V_s)] e^{-z/l} dz$$

Całkowanie jest przeprowadzone tylko w tym zakresie wartości „z”, w którym dla danego $h\nu$ funkcja podcałkowa daje dodatnie wartości.

Współczynniki a_{npr} i a_{pr} są odpowiednio proporcjonalnymi mnożnikami dla przejść prostych i nieprostych. Te współczynniki podobnie jak mnożnik b mogą być określone przez liczbowe obliczenia całek i przystosowanie otrzymanej krzywej do danych doświadczalnych. Oczywiście tylko dla $h\nu$ mniejszych niż próg prosty, słabszy mechanizm przejść nieprostych będzie rozwiązywalny doświadczalnie; w do-

Tablica 1.

Energie prądu i energie zależne od fotoelektrycznej wydajności dla różnych wzbudzeń i procesów rozpraszania (wg Kane) [27]

<u>Wzbudzenie z.</u>	<u>Przejście</u>	<u>Procesy rozpraszania</u>	<u>Próg $h\nu_i$</u>	<u>m</u>	
PROCESY OBJĘTOŚCIOWE					
Pasmo walencyjne	skośne	nierozproszone	$h\nu_i = \chi + E_c - E_v$	5/2	
		rozproszone	$h\nu_i \geq \chi + E_c - E_v$	1	
	proste	nierozproszone	$h\nu_i \geq \chi + E_c - E_v$	2	
		rozproszone			
PROCESY POWIERZCHNIOWE					
Pasmo walencyjne		Rozprzeszczone rozproszenie powierzchniowe	$h\nu_i = \chi + E_c - E_v$	5/2	
				3/2	
		Powierzchniowe rozproszenie odbiciowe			
Dyskretne położenie stanów powierzchniowych poniżej E_F			$h\nu_i = \chi + E_c - E_{s_i}$	1	
	skośne		$h\nu_i = \chi + E_c - E_F$	2	
	proste		$h\nu_i > \chi + E_c - E_F$	2	
			$h\nu_i > \chi + E_c - E_F$	1	
	skośne		$h\nu_i = \chi + E_c - E_F$	5/2	
					3/2

datku do pobudzenia objętościowego i procesów rozproszenia należy rozważyć podobne procesy wprowadzające udział powierzchni (niższa część tablicy 1). W tej sytuacji dwa główne rodzaje mogą być rozróżnione. Pierwszy zawiera pobudzenie objętościowe jak poprzednio, lecz teraz powierzchnia bardziej niż objętość uczestniczy w wymiarze momentu. Próg dla tego typu przejścia będzie taki sam jak próg przejść nieprostych, objętościowych, czyli $\chi + E_G$; jednakże wartości wykładnika „ m ” będą zależeć od tego czy powierzchnia jest dyfuzyjnym czy zwierciadlanym „rozpraszaczem” elektronów, jak pokazano w tablicy 1.

Drugi rodzaj procesów powierzchniowych zawiera fotoemisję z zajętych stanów powierzchniowych. W związku ze względnie małą gęstością tych stanów wydajność fotoelektryczna w tym przypadku będzie znacznie mniejsza niż dla pobudzenia objętościowego. Jednakże, skoro próg emisji ze stanów powierzchniowych leży poniżej $\chi + E_G$, taka emisja może być wykrywalna jako mały „ogon” na wykresie $Y(h\nu)$.

Dla fotoemisji z pojedynczego dyskretnego rozłożenia stanów powierzchniowych zlokalizowanych przy energii E_S , próg będzie dokładnie równy: $\chi + E_c - E_S$. W przypadku ciągłego rozkładu stanów zlokalizowanych próg będzie odpowiadał emisji z najwyższych wierzchołkowych poziomów (jeśli wszystkie stany leżą poniżej E_F (i stąd są w pełni zajęte), lub będzie odpowiadał emisji z poziomu Fermiego, jeśli on przechodzi poprzez stany [32–35]. Zależność funkcjonalna wydajności od $h\nu$ będzie różna w tych dwóch przypadkach (Tablica 1).

Kane rozważał także przejścia z pasma powierzchniowego, które może być obecne na idealnej powierzchni. Charakterystyki fotoemisji zależą w tym przypadku od tego czy wektor momentu K równoległy do powierzchni jest, czy nie jest zachowany oraz czy próg pojawia się przy poziomie Fermiego czy daleko. Te cztery procesy są przedstawione w ostatnich wierszach tablicy 1.

4. STANY ELEKTRONOWE NA REALNEJ POWIERZCHNI PÓŁPRZEWODNIKA

Na powierzchni realnej, na przykład krzemu, różnorodność wiązań między atomami krzemu, tlenu i pary wodnej oraz obecności defektów i innych zanieczyszczeń, prowadzą do powstania złożonej struktury stanów elektronowych, które dla swobodnych nośników ładunku mogą stanowić centra pułapkowania lub rekombinacji [36; 37].

Wykazano, że na powierzchni realnej krzemu (o przeciętnej koncentracji $\sim 10^{15} \text{ m}^{-2}$) zlokalizowane są (na granicy półprzewodnik — warstwa tlenu) stany „szybkie”, [38–41]. Natomiast stany „powolne” na powierzchniach realnych krzemu ułożone we wnętrzu tlenu charakteryzują się znacznie dłuższym czasem wymiany nośników sięgającym nawet do kilku godzin. Opracowano już kilka modeli struktury energetycznej tych stanów elektronowych [43–45].

UWAGI KOŃCOWE

W praktyce, szczególnie w procesach technologicznych stosowanych w mikroelektronice, występuje powierzchnia realna półprzewodnika otrzymana w wyniku różnorodnej obróbki. Ponieważ pozostaje ona zazwyczaj w kontakcie z atmosferą powietrza, obojętego gazu technicznego, lub niewysokiej próżni, pokrywa się cienką warstwą naturalnego tlenku o złożonej strukturze fizyko-chemicznej [42].

Granica fazowa półprzewodnik – tlenek naturalny wywiera znaczny wpływ na przebieg zjawisk elektronowych w obszarze przypowierzchniowym, wykorzystywanym w konstrukcji przyrządów półprzewodnikowych [43; 44].

Wieloletnie prowadzone przez autora badania fotoemisji z powierzchni realnych pozwalają stwierdzić, że uzyskiwane wartości wyznaczanych parametrów nie zawsze są powtarzalne, zwłaszcza przy braku właściwej kontroli i stabilności fizyko-chemicznej struktury takiej powierzchni, a częstokroć zaniżone. Takie pomiary przeprowadza się raczej dla celów technicznych niż naukowych. Dlatego, dużo dokładniejsze wyniki uzyskuje się stosując badania tzw. powierzchni atomowo czystej uzyskiwanej poprzez łupanie kryształów w ultrawysokiej próżni rzędu 10^{-10} Tr.

BIBLIOGRAFIA

1. W.E. Spicer, Appl. Phys., 12, 115 (1977)
2. J. Wojas, Postępy Fizyki, Tom XX, Zeszyt 4, 441 (1969)
3. J. Wojas, Acta Phys. Pol., A60, 767 (1981)
4. J. Wojas, Elektronika, Z. 9, 17 (1982)
5. J. Wojas, Archiwum Elektrotechniki, Tom XXXII, Z. 123/124 – 3/4/1983, 677 (1987)
6. J. Wojas, Archiwum Elektroniki, Tom XXX, Zeszyt 125/126 – 3/4/1983, 697 (1985)
7. R.E. Palmer, The Review of Scientific Instruments, 42, Nr 10 (1971)
8. D.E. Eastman, W.D. Grobman, Phys. Rev. B, 9, Nr 8, 3473 (1974)
9. W.E. Spicer, C.N. Berglund, The Review of Scientific Instruments, 35, Nr 12, 1665 (1964)
10. M.L. Cohen, J.C. Phillips, Phys. Rev., 139, Nr 3A, A912 (1965)
11. T. Fischer, Phys. Rev., 139, Nr 4A, A1228 (1965)
12. G.W. Gobeli, F.G. Allen, Phys. Rev., 127, nr 1, 141 (1962)
13. M. Leleyter, P. Joes, Surf. Sci., 156, 800 (1965)
14. J.E. Rowe, H. Ibach, Phys. Rev. Lett., 31, Nr 2, 102 (1973)
15. T. Radoń, Ch. Kleint, Surf. Sci., 144, 638 (1984)
16. B.Z. Olshansky i inni, Pisma Zh. Eks. Teor. Fiz., 25, No 4, 195 (1977)
17. B.Z. Olshansky, A.A. Shklyayev, Surf. Sci., 67, No 1, 581 (1977)
18. A.Yu. Mityagin, V.V. Panteleev, Izv. Akad. nauk SSSR. Ser. Fiz., 40, No 8, 1566 (1976)
19. H. Ibach, J.E. Rowe, Phys. Rev. B, 9, No 4, 1951 (1974)
20. M. Cardona, L. Ley, Topics in Applied Physics, Photoemission in Solids, Vol. 26, 16 (1978) – Springer Verlag – Berlin
21. L. Ley, M. Cardona, Topics in Applied Physics, Photoemission in Solids II, Vol. 27, 1 (1979); L. Ley, M. Cardona, R.A. Pollak, Vol. 27, 11 (1979)
22. V.D. Shadrin, Fiz. Tverdogo Tela (USSR), 21, 2495 (1979) [received 1980]
23. F.G. Allen, G.W. Gebeli, Proc. Int. Conf. Semiconductor Physics, Exter, 1962, s. 818, (the Institute of Physics, London, 1962); W. Shockley, Elektrony i dziury w półprzewodnikach, PWN, Warszawa 1956

24. T.E. Fischer, F.G. Allen, G.W. Gobeli, *Phys. Rev.*, 163, No 3, 703 (1967)
25. T. Lopez-Rios, G. Hincelin, *Phys. Rev. B.*, 38, No 5, 3561 (1988)
26. B. Seroczyńska-Wojas, *Prace Instytutu Fizyki PW*, Z. 22, 5 (1980)
27. E.O. Kane, *Phys. Rev.*, 127, 131 (1962)
28. P.E. Gregory, W.E. Spicer, *Phys. Rev. B.*, 13, Nr 2, 725 (1976)
29. R. Garron, *Surf. Sci.*, 146, 527 (1984)
30. J.P. Buissoniinni, *Surf. Sci.*, 120, L477 (1982)
31. R.E. Alleniinni, *Surf. Sci.*, 110, L625 (1981)
32. Xiong Shi-Jie, Tang Ming-Qiu, *Acta Phys. Sin. (China)*, 36, No 9, 1230 (1987)
33. A. Gold, *J. Phys. Colloq. (France)*, 48, no C-5, 255 (1987)
34. H.J. Zandvliet, A. Van Silfhout, *Surf. Sci.*, 195, no 1-2, 138 (1988)
35. Y. Bor-Yam, J.D. Joannopoulos, *Phys. Rev. B. (USA)*, 30, 1844 (1984)
36. A. Ismailiinni, *Surf. Sci.*, 157, 319 (1985)
37. Ming-fu Li, Shang-yuan Ren, De-qiang Mao, *Chin. Phys. (USA)*, 4, No 2, 294 (1984)
38. V.A. Zueviinni, *Ukr. Fiz. Zh.*, 20, 1147 (1975)
39. H. Flietner, *Surf. Sci.*, 46, 251 (1974)
40. H. Flietner, N. Duong Singh, *Phys. Stat. Sol. (a)*, 37, 533 (1976)
41. V.N. Ovsyuk, A.V. Rzhanov, *Fiz. Tekh. Poluprovodn.* 3, 294 (1969)
42. B.A. Nesterenko, O.V. Smirko, *Fizicheskiye svoystva atomarnoczystoy poverkhnosti poluprovodnikov*, Naukova Dumka, Kiev. 1983
43. Y. Bor-Yam, J.D. Joannopoulos, *Phys. Rev.* 13 (USA), 30, 1844 (1984)
44. G. Allan, *Proceeding of the NATO Advanced Study Institute, Acquafredda di Maratea, Italy*, 14-25 Juli 1986, s. 9
45. R.S. Bauer [Editor], "Surfaces and Interfaces": Physics and Electronics, North-Holland Second Trieste, ICTP-IUPAP, Semiconductor Symposium (1982). North-Holland Publishing Company — Amsterdam
46. V.A. Zujeviinni, *Nieravnovesnye pripoverkhnostnye processy v poluprovodnikakh i poluprovodnikovyykh priborakh*, Sovetskoye Radio, Moskva 1977
47. B. Adamowicz, *Praca doktorska. Badania procesów elektronowych na realnej powierzchni krzemu metodą fotonapięcia powierzchniowego*. Politechnika Śląska Gliwice 1988

J. WOJAS

ANALYSIS OF PHOTOEMISSION OF ELECTRONS

Summary

In the paper there is presented concise analysis of external photoemission effect and its correlation with optical transitions, for electron excitation and escaping out of crystal is realized by direct or indirect optical transitions. There is explained that photoemission quantum yield which is function of photon energy depends on above mentioned transitions in this paper.

FFT jako algorytm postprocessingu zadań syntezy pola

EUGENIUSZ KORNATOWSKI

Instytut Elektroniki i Informatyki, Politechnika Szczecińska

Otrzymano 1993.04.28

Autoryzowano do druku 1993. 06.02

Praca jest uzupełnieniem i rozwinięciem metody oceny jakości pól przedstawionej w [5]. Istotą proponowanego obecnie algorytmu jest zastosowanie szybkiej transformaty Fouriera do oceny jakości pola będącego przedmiotem optymalizacji.

1. WSTĘP

Celem zadań syntezy pola elektromagnetycznego jest uzyskanie rozkładu potencjału wektorowego (składowych wektora indukcji) o ściśle określonym kształcie. Odpowiednie postawienie problemu, przy wykorzystaniu metod teorii pola, prowadzi do rozwiązania bliskiego założeniom. Uzyskane rozwiązanie można zoptymalizować w sensie odpowiedniego kryterium.

W pracy zaproponowano zastosowanie algorytmu FFT do iteracyjnej optymalizacji rozkładu pola, uzyskanego jedną z klasycznych metod teorii pola. Przedstawiona metoda może być także z powodzeniem zastosowana (po niewielkich modyfikacjach) do oceny jakości innych przebiegów $N-D$.

2. ISTOTA METODY

Proponowana metoda, w swoich podstawowych założeniach, jest zbliżona do algorytmu przedstawionego w [5]. W [5] wskaźnik jakości pola konstruuje się w oparciu o kryterium zawartości pasożytniczych częstotliwości harmoniczných, uzyskanych w wyniku filtracji cyfrowej 2--D. Przebiegiem wejściowym algorytmu był rozkład błędu będącego różnicą między założonym, a uzyskanym rozkładem potencjału wektorowego. W niniejszej pracy natomiast proponuje się zastosowanie

FFT do zdefiniowania wskaźnika jakości pola będącego przedmiotem syntezy, a następnie optymalizacji.

Obliczenie wskaźnika jakości przebiega wg następującego schematu:

1. Określenia rozkładu $I'(x, y)$ będącego przedmiotem syntezy

$$I'(x, y) = I(x, y) + \Phi(x, y) \quad (1)$$

gdzie: $I'(x, y)$ — rozkład pola uzyskany w wyniku syntezy

$I(x, y)$ — założony rozkład pola

$\Phi(x, y)$ — przebieg odkształcający $I(x, y)$ — zakłócenie.

2. Analiza Fouriera szybką transformatą przebiegu $I'(x, y)$

$$\mathcal{F}[I'(m, n)] = \mathcal{F}[I(m, n)] + \mathcal{F}[\Phi(m, n)] = \mathcal{F}_I(m, n) + \mathcal{F}_\Phi(m, n) \quad (2)$$

gdzie: $m = \frac{x}{\Delta x}, \quad m = 0, 1, 2, \dots, M$

$n = \frac{y}{\Delta y}, \quad n = 0, 1, 2, \dots, N.$

3. Określenie macierzy modułów harmonicznych przestrzennych

$$F_I(m, n) = |\mathcal{F}_I(m, n)|, \quad F_\Phi(m, n) = |\mathcal{F}_\Phi(m, n)|. \quad (3)$$

4. Określenie wskaźnika jakości

$$h_I = \frac{\left(\sum_{m=0}^M \sum_{n=0}^N f_\Phi^2(m, n) \right)^{\frac{1}{2}}}{\left(\sum_{m=0}^M \sum_{n=0}^N f_I^2(m, n) \right)^{\frac{1}{2}}}. \quad (4)$$

gdzie $f_\Phi(m, n), f_I(m, n)$ — elementy macierzy F_Φ i F_I odpowiednio.

Należy zauważyć, że obecnie nie jest niezbędny dokładny opis rozkładu $I(x, y)$. Wystarczy tylko informacja, które harmoniczne przestrzenne w rozkładzie $I'(x, y)$ są szkodliwe, a które użyteczne.

3. PRZYKŁAD OBLICZENIOWY

Proponowaną metodę zastosowano do oceny jednorodności pola magnetycznego w przykładzie optymalizacji układu przewodów pokazanego na rys. 2 w [5]. Układ ten był szczegółowo analizowany w [7].

Zadanie optymalizacji, przykładów dotyczy przypadku $\mu = 10 \mu_0$ i takiego doboru kąta α , przy którym w obszarze:

$$\delta = \{x, y : 0 \leq x \leq 0.6 R_2, \quad 0 \leq y \leq 0.6 R_2\} \quad (5)$$

pole będzie jednorodne (tj. składowe wektora indukcji $B_x(x, y) = \text{const}$ i $B_y(x, y) = 0$). Przykładowy przebieg linii $A = \text{const}$ przy $\mu = 10 \mu\text{o}$ i $\alpha = 60^\circ$ pokazano na rys. 3 w [5].

Równania opisujące rozkład potencjału wektorowego w obszarze określonym przez (5) zostały przedstawione w [7]. Stanowią one podstawę do obliczenia składowych wektora indukcji B_x i B_y . Składowe te obliczono metodą różniczkowania numerycznego funkcji potencjału wektorowego $A(x, y)$, wykorzystując jednowymiarową interpolację czteropunktową:

$$y'(\varphi) = \frac{1}{h_\varphi} \left[-\frac{3P_\varphi^2 - 6P_\varphi + 2}{2} y(\varphi_0 - h_\varphi) + \frac{3P_\varphi^2 - 2P_\varphi - 1}{2} y(\varphi_0) - \right. \\ \left. - \frac{3P_\varphi^2 - 2P_\varphi - 2}{2} y(\varphi_0 + h_\varphi) + \frac{3P_\varphi^2 - 1}{6} y(\varphi_0 + 2h_\varphi) \right], \quad (6)$$

gdzie: y — funkcja różniczkowana

φ — zmienna x lub y

h_φ — $\Delta\varphi$

P_φ — $(\varphi - \varphi_0)/h_\varphi$.

Uzyskane macierze $B_x(m, n)$ i $B_y(m, n)$ stanowią dane wejściowe algorytmu opisanego w punkcie 2.

Algorytmu tego użyto dwukrotnie:

określając h_{B_x} i h_{B_y} , przy czym przyjęto następujące oznaczenia w(4):

$$a) I'(x, y) \triangleq B_x(x, y) \quad (7)$$

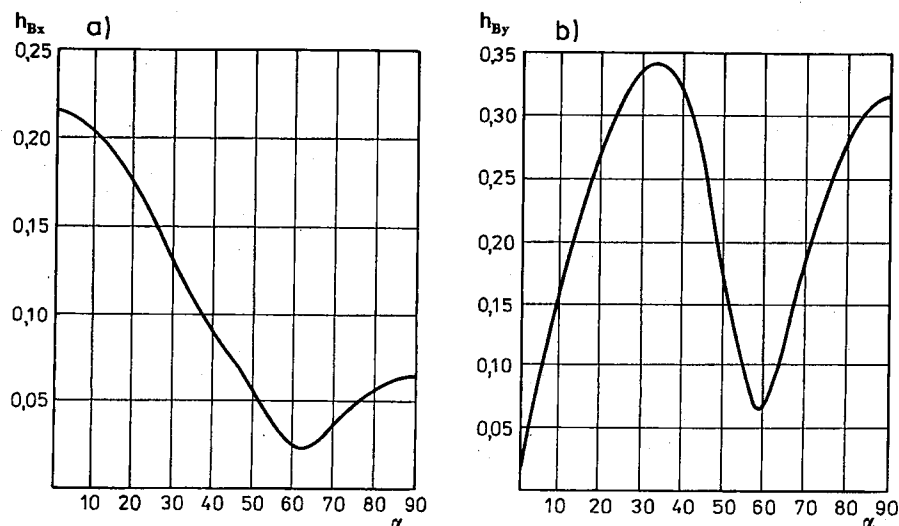
$$\left(\sum_{m=0}^M \sum_{n=0}^N f_I^2(m, n) \right)^{\frac{1}{2}} = f_I(0, 0) \triangleq |\mathcal{F}[B_x(m, n)]|, \quad m=n=0 \quad (8)$$

$$\left(\sum_{m=0}^M \sum_{n=1}^N f_\varphi^2(m, n) \right)^{\frac{1}{2}} = \left(\sum_{m=0}^M \sum_{n=0}^N f_\varphi^2(m, n) \right)^{\frac{1}{2}} \triangleq \left(\sum_{m=1}^M \sum_{n=0}^N |\mathcal{F}[B_x(m, n)]|^2 \right)^{\frac{1}{2}} \quad (9)$$

$$b) I'(x, y) \triangleq B_y(x, y) \quad (10)$$

$$H_{B_y} = \left(\sum_{m=0}^M \sum_{n=0}^N f_{\varphi'}^2(m, n) \right)^{\frac{1}{2}} \triangleq \left(\sum_{m=0}^M \sum_{n=0}^N |\mathcal{F}[B_y(m, n)]|^2 \right)^{\frac{1}{2}} \quad (11)$$

W rozpatrywanym przykładzie przyjęto $M=N=31$ poszukując optymalnego α , przy którym h_{B_x} i h_{B_y} przyjmują wartości minimalne. Uzyskane wyniki iteracyjnej optymalizacji α pokazano w postaci wykresów $h_{B_x}(\alpha)$ i $h_{B_y}(\alpha)$ na rys. 1.



Rys. 1. Przebieg zawartości pulsacji zakłócających wektor $\vec{B}$: a) w kierunku OX, b) w kierunku OY

4. WNIOSKI I UWAGI

W wyniku przeprowadzonych obliczeń uzyskano minimalne wartości h_B dla $\alpha = 59^\circ$ i h_{B_x} dla $\alpha = 62^\circ$. Biorąc pod uwagę charakter obu krzywych: wyraźne minimum h_{B_y} dla $\alpha = 59^\circ$ i przebieg h_{B_x} w zakresie α od 60° do 65° z bardzo małą pochodną, przyjęto α optymalne, równe $60,5^\circ$. Należy podkreślić, że w czasie iteracyjnego poszukiwania optymalnego α w pierwszym kroku przyjęto $M=N=7$ w celu skrócenia czasu obliczeń algorytmu oceny jednorodności pola. Następnie stopniowo siatkę zagęszczono, aż do $M=N=31$.

Można zauważyć, że przedstawiona metoda jest jedną z wielu. Przykładem podobnego algorytmu jest metoda opisana w [5] wykorzystująca do analizy jednorodności pola pasmowo-przepustowy filtr cyfrowy. Algorytm ten jest jednak dużo bardziej czasochłonny od przedstawionego w niniejszej pracy.

PODSUMOWANIE

W pracy przedstawiono metodę oceny jakości pól magnetycznych z wykorzystaniem szybkiej transformaty Fouriera. Metoda ta została zastosowana w przykładowym zadaniu optymalizacji układu przewodów generujących pole.

Jako kryterium optymalizacji przyjęto jednorodność pola w określonym obszarze, którego wskaźnikiem była zawartość przestrzennych częstotliwości zakłócających składowe B_x i B_y . Przedstawione rozważania i uzyskane wyniki obliczeń numerycznych wskazują na celowość stosowania nowoczesnych metod cyfrowego przetwarzania sygnałów jako algorytmów postprocessingu zadań syntezy pola.

BIBLIOGRAFIA

1. K. A d a m i a k: *Method of the magnetic field synthesis of the axis of cylinder solenoid*. J. Appl. Phys. 1978, No 16
2. J.M. Costa: A.N. Venetsanopoulos, M. Trefler: *Digital tomographic filtering of radiographs*. IEEE Trans. on Medical Imaging, vol. Mi—2, No 2, June 1983
3. J.M. Costa: A.N. Venetsanopoulos: *Design of circularly symmetric two-dimensional recursive filters*. IEEE Trans. Acoust., Speech Signal Processing, vol. ASSSP—22 No 6, Dec 1974
4. E. K o r n a t o w s k i, J. R a d e c k i: *Dwuwymiarowe pasmowe filtry o kołowej symetrii charakterystyki amplitudowej*. Archiwum Elektrotechniki, 1986, z. 3/4
5. E. K o r n a t o w s k i: *Zastosowanie filtracji cyfrowej 2-D do oceny jakości pól elektromagnetycznych*. Kwartalnik Elektroniki i Telekomunikacji, 1992, z. 4, s. 133—140
6. R. P a ł k a: *Synthesis of magnetic field due to the direct current*. Etz — Archiv, 1985, Bd. 7, H. 9
7. R. S i k o r a, J. P u r c z y Ń s k i, W. L i p i Ń s k i: *Pole magnetyczne prądów płynących wzdłuż powierzchni walcowych*. Archiwum Elektrotechniki, 1973, t. 22, z. 3

E. KORATOWSKI

FFT AS THE ALGORITHM OF POSTPROCESSING OF AN ASSIGNMENT
OF THE FIELD SYNTHESIS

S u m m a r y

In this work the supplement and display of the method [5] of valuation quality the fields is presented. The essence of the proposed algorithm is the application of the Fourier transformation for an assignment of the field that is object of optymalization.

On linear adaptive filtering with constant step of adaptation

MARIUSZ ŻÓŁTOWSKI

Instytut Telekomunikacji, Politechnika Gdańska

Received 1992.01.03

Authorized 1993.09.03

Square loss function while minimized in recursive way results in adaptive algorithms which are relatively simple. Among the different kinds of linear adaptive digital filtering algorithms there are optimal ones which are called fast [3], due to the sophisticated way of performing calculations. However, in this paper the linear adaptive filter with constant step of adaptation is of concern alone in review form. Though being suboptimal only the constant step filter is of importance. The reasons in favour of it are; the simplest implementation among all linear adaptive filtering algorithms and the lack of divergence. Though many brilliant authors have paid attention to this kind of problems [1–31], the topic of adaptation is just the one of these areas of interest which deserve continued attention in view of impressive developments in hardware and architecture of digital systems. In contrast to the recent approach of Bellanger [3], the comparison with Kalman's filter has been added, though similar one, one can find in Haykin [25]. The conditions under which the variance of adaptive filter coefficients is computed have been risen to the conditions of adaptation well posed. The effects of quantization on filter's performance have been considered showing the existence of optimal step of adaptation well established. The concept of normed space can be used to show the convergence of adaptive algorithm in pleasant way. The approach of the paper is aimed towards the refinement of recursive schemes in view of their importance in communication, control, system identification etc.

1. INTRODUCTION

The square loss function, being related to the notion of energy is the very reasonable one. Moreover, it results in adaptive algorithms rather simple, especially in linear case. New trends in the architecture of digital systems towards high signal processing ability, including parallel processing via systolic arrays, wave front processors etc., while offered at low prices make the improvements of performance of existing systems by even complex algorithms possible. Being considered in this

paper the adaptive digital filter with constant step of adaptation is well suited for meeting the real time of performance reasons. It is relatively simple and not divergent. The least squares method in digital adaptive filtering is characterized in chapter 2. The comparison with Kalman's filter is done in chapter 3. The linear adaptive filter with constant step of adaptation is considered in chapter 4 in detail. Quantization and blocking effects are described in chapter 5, while end comments are provided in last chapter.

2. METHOD OF LEAST SQUARES

The task of adaptation is [1, 3, 33] to find the minimum in respect to $D(n)$ of:

$$\hat{C}_0\{D(n)\} = \sum_{k=1}^n L[x(k), y(k), D(n)] = \sum_{k=1}^n w^{n-k} [y(k) - D^T(n)X(k)]^2, \quad (1a)$$

while

1^o. L is a loss function and $0 < w \leq 1$ ($w \cong 1$) positive scalar

2^o. $x(k)$, $k=1, \dots$ is the input state waveform of system under consideration, $x(k)=0$ for $k < 1$.

3^o. $D(n) = \text{col}[D_0(n), \dots, D_{N-1}(n)]$ is the vector of FIR coefficients [4].

4^o. $y(k)$, $k=1, \dots$ is the observed waveform of one dimension or reference one.

5^o. $X(k) = \text{col}[x(k), \dots, x(k-N+1)]$.

If $\hat{y}(k) = D^T(n)X(k)$ is given an interpretation of FIR filter response then $x(k)$ is an excitation obviously. In other words the task of (1) is to match the output $\hat{y}(n)$ of FIR filter while excited by $x(k)$ to the observation or reference $y(k)$, $k=1, \dots$

$$\hat{y}(k) = D^T(n)X(k) = \sum_{i=0}^{N-1} D_i(n)x(k-i). \quad (1b)$$

By equalizing the gradient of $\hat{C}_0(n)$ in respect to $D(n)$ to zero one obtains

$$\nabla_{D(n)} \hat{C}_0(n) = - \sum_{k=1}^n w^{n-k} [y(k) - D^T(n)X(k)]X(k) = 0 \quad (2a)$$

or

$$\sum_{k=1}^n w^{n-k} [y(k) - D^T(n)X(k)]X(k)^T = 0. \quad (2b)$$

Let

$$R_{xx}(n) \stackrel{\text{df}}{=} \sum_{k=1}^n w^{n-k} X(k)X(k)^T \quad (2c)$$

and

$$r_{yx}(n) \stackrel{\text{df}}{=} \sum_{k=1}^n w^{n-k} y(k) X(k). \quad (2d)$$

Then from (2b)

$$r_{yx}(n)^T = D^T(n) R_{xx}(n) \quad (2e)$$

or

$$\begin{aligned} (R_{xx}(n))^T &= R_{xx}(n) \\ R_{xx}(n) D(n) &= r_{yx}(n) \end{aligned} \quad (2f)$$

and

$$D(n) = R_{xx}^{-1}(n) r_{yx}(n), \quad (3a)$$

$$R_{xx}(n) = [r_{xx}^{ij}(n)]_{N \times N}, \quad i, j = 0, \dots, N-1 \quad (3b)$$

$$r_{xx}^{ij}(n) = \sum_{k=1}^n x(k+i)x(k+j)w^{n-k}, \quad i = 0, \dots, N-1. \quad (3c)$$

On the other hand

$$r_{yx}(n) = [r_{yx}^i(n)]_{N \times 1}, \quad i = 0, \dots, N-1 \quad (3d)$$

$$r_{yx}^i(n) = \sum_{k=1}^n y(k)x(k+i)w^{n-k}, \quad i = 0, \dots, N-1. \quad (3e)$$

Now, let y and x be stationary in wide sense and

$$R_{xx} \stackrel{\text{df}}{=} E\{X(k)X(k)^T\}, \quad r_{yx} = E\{y(k)X(k)\}. \quad (4a)$$

Then

$$E\{R_{xx}(n)\} = \frac{1 - w^n}{1 - w} R_{xx} \quad (4b)$$

$$E\{r_{yx}(n)\} = \frac{1 - w^n}{1 - w} r_{yx}. \quad (4c)$$

That is why $R_{xx}(n)$ has the meaning of the estimate of autocorrelation matrix, while $r_{yx}(n)$ is the estimate of crosscorrelation vector.

If $n \rightarrow \infty$ (ergodic case¹) the vector of the optimal coefficients of the filter is:

$$D_{\text{opt}} = R_{xx}^{-1} r_{yx} \quad (5)$$

One can notice that the weighting factor w has no effect on optimal solution in this case. The (5) and (3) are called Yule–Walker equations for evolving and

¹ in this case the ergodicity means the existence of the estimates of the expressions under consideration while n tends to the infinity equal to their space counterparts of (4a) form in mean squared sense.

stationary filter respectively. However recursive realations are required in adaptive filtering. First

$$R_{xx}(n+1) = w R_{xx}(n) + X(n+1)X^T(n+1) \quad (6a)$$

$$r_{yx}(n+1) = w r_{yx}(n) + y(n+1)X(n+1). \quad (6b)$$

Secondly

$$D(n+1) = R_{xx}^{-1}(n)wr_{yx}(n) + X(n+1)y(n+1) \quad (7a)$$

$$wr_{yx}(n) = wR_{xx}(n)D(n) = [R_{xx}(n+1) - X(n+1)X^T(n+1)]D(n) \quad (7b)$$

Eventually

$$D(n+1) = D(n) + R_{xx}^{-1}(n+1)X(n+1)[y(n+1) - X^T(n+1)D(n)]. \quad (8)$$

This is the recursion for filter's coefficients. However (8) is the special case of a recursion considered in [1, 33], while the sequence of adaptation gain is optimal, even globally optimal in this linear case. The following expressions are known respectively

$$\epsilon_1(n+1) \stackrel{\text{df}}{=} y(n+1) - X^T(n+1)D(n) \quad (9a)$$

as "a priori" error

$$\epsilon_2(n+1) \stackrel{\text{df}}{=} y(n+1) - X^T(n+1)D(n+1) \quad (9b)$$

as "a posteriori" error and

$$G(n+1) = R_{xx}^{-1}(n+1)X(n+1) \quad (9c)$$

as adaptation gain.

3. COMPARISON WITH KALMAN'S FILTER

Linear filtering while optimal in probabilistic meaning [2] is given in terms of;

1°. State equation

$$x_{k+1} = A(k+1, k)x_k + w_k$$

$$x_{k \times 1}, A_{n \times n}, w_{k \times 1}$$

$w_k \sim N(0, Q_k)$, $N(0, K)$ stands for the normal distribution of zero mean and K matrix of correlation.

$$E\{x_0 w_k\} = 0, \quad k = 0, 1, \dots$$

2°. Observation equation

$$y_k = C(k)x_k + v_k$$

$$y_{k \times 1}, C(k)_{m \times n}, v_{k \times 1}$$

$$E(v_k w_l) = 0, \quad k, l = 0, 1, \dots$$

$$v_k \sim N(0, R_k).$$

3°. State prediction equation

$$\hat{x}_k^{k-1} = A \hat{x}_{k-1}^{k-1}, \quad A = A(k, k-1)$$

$$\hat{x}_0^0 = x_0$$

4°. Variance equation "a priori"

$$P_k^{k-1} = A P_{k-1}^{k-1} A^T + Q_k.$$

5°. Variance equation "a posteriori"

$$P_k^{k-1} = P_k^{k-1-1} + C^T R_k^{-1} C, \quad P_0^0 = P_0,$$

$$C = C(k).$$

6°. Filter gain equation

$$K_k = P_k^{k-1} C^T (C P_k^{k-1} C^T + R_k)^{-1} = P_k^k C^T R_k^{-1}.$$

7°. State estimate equation

$$\hat{x}_k^k = \hat{x}_k^{k-1} + K_k (y_k - C \hat{x}_k^{k-1}), \quad k = 0, 1, \dots$$

Now, let the state equation be

$$D(n+1) = D(n) \tag{10a}$$

and the observation or reference one

$$y(n) = X^T(n) D(n) + v(n). \tag{10b}$$

Then from 1° – 7°

$$\hat{D}(n) = \hat{D}(n-1) + K(n) [y(n) - X^T(n) \hat{D}(n-1)] \tag{10c}$$

$$K(n) = P_n^n X(n) R_n^{-1} \tag{10d}$$

$$P_n^{n-1} = P_{n-1}^{n-1-1} + X(n) R_n^{-1} X(n)^T. \tag{10e}$$

One can compare the least squares approach and the one based on the optimal linear filtering theory. It turns out that (8) and (10c) are the same if

1°. (10a) is the state equation

2°. (10b) is the observation equation

3°. $P_n^n = R_{xx}^{-1}(n)$

4°. effective $R_n = 1$

5°. weighting factor $w = 1$.

In view of comparison the disturbing noise can be weighted by effective $R_n \neq 1$ on one hand and a kind of forgetting can be introduced on the other by $w < 1$.

The effective R_n equals 1 in view of shortening (see 10e and 10d) while the noise $v(n)$ is stationary and no knowledge about the initial state is assumed (Fisher's model is valid [5]). Besides, the identity of 3^o may arise interest. However, if the normal distributions were defined with reciprocal to P_{n-1}^n matrix ($n=0, 1, \dots$) the complete accord would be gained. Moreover P_n^n can be given the interpretation of the information matrix in Fisher's sense [5]. Also any initial condition is introduced in natural way within probabilistic approach by an initial normal distribution. In view of the results of information theory the normal, gaussian model is of the worst case type approach in any situation. Adaptive filter with constant step of adaptation is considered next.

4. ADAPTIVE FILTER WITH CONSTANT STEP

The performance of adaptive filter with constant step of adaptation is given in terms of recursion

$$D(n+1) = D(n) - \frac{1}{2} \delta \nabla_{D(n)} \epsilon_1^2(n+1) \quad (11a)$$

$0 < \delta$ is scalar,
or

$$D(n+1) = D(n) + \delta X(n+1) \epsilon_1(n+1). \quad (11b)$$

While comparing with (8) one can notice that the simplification in respect to optimal case relies on substitution

$$R_{xx}^{-1}(n+1) = \delta I. \quad (I_{N \times N} \text{ is an identity matrix}) \quad (11c)$$

From algorithmic point of view the performance of adaptive constant step filter is at cost of

- 1^o. $2N+1$ multiplications
- 2^o. $2N$ additions
- 3^o. $2N$ active memory cells.

First, the problem of filter's stability can be discussed in terms of "a priori" ϵ_1 and "a posteriori" ϵ_2 errors.

$$\begin{aligned} \epsilon_2(n+1) &= y(n+1) - D^T(n+1)X(n+1) = \\ &= y(n+1) - D(n) + [X(n+1)\epsilon_1(n+1)]^T X(n+1) = \\ &= \epsilon_1(n+1) [1 - \delta X(n+1)^T X(n+1)] \end{aligned} \quad (12a)$$

The criterion of less expected "a posteriori" than "a priori" error value [3] can be used as stability condition.

Under the hypothesis of no correlation between $\epsilon_1(n+1)$ and the second factor of (12a);

$$E\{|\epsilon_2|\} = E\{|\epsilon_1|\} E\{|1 - \delta X(n+1)^T X(n+1)|\}. \quad (12b)$$

As a matter of fact the filter is stable if

$$0 < \delta/2 < 1/(N\sigma_x^2) \quad (12c)$$

while $\sigma_x^2 = E\{x(n)^2\}$ is the power of waveform $x(n)$ ($n=1, \dots$).

There is a safety margin used in practice [3]:

$$0 < \delta/2 < 1/(cN\sigma_x^2) \quad (12d)$$

and c is about 3.

The hypothesis about the statistical correlation of factors of (12) becomes valid while near optimal performance is being gained by the filter. It is known that steady

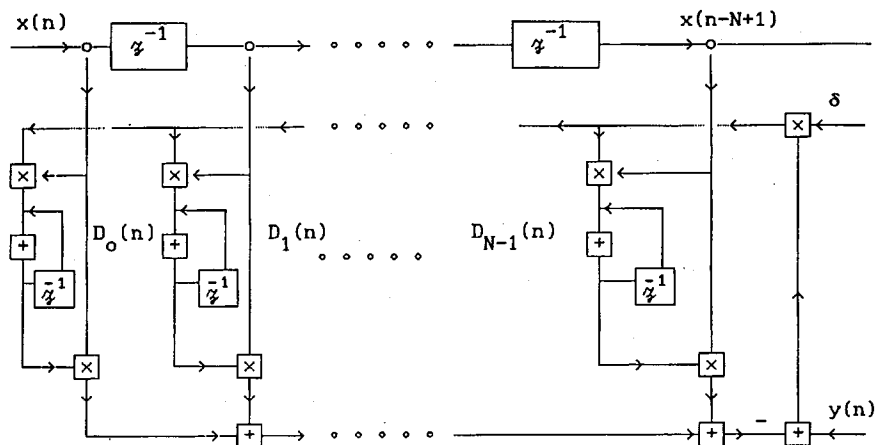


Fig. 1. Adaptive transversal filter with constant step of adaptation.

state error is introduced by constant step algorithm [1, 3] which results also in residual error. If n tends to the infinity and the data $x(n)$ (input signal) and filter's coefficient are not correlated then from (11b)

$$r_{yx} - R_{xx} E D(n) = 0 \quad (13)$$

$$\text{or } [E\{\cdot\}]_{n \rightarrow \infty} = E \lim_{n \rightarrow \infty} \{\cdot\}$$

$$D_{\text{opt}} = E\{D(n)\}_{n \rightarrow \infty} = R_{xx}^{-1} r_{yx}. \quad (14)$$

That is why the mean values of filter's coefficients are equal to their optimal values in steady state. As a matter of fact the steady state mean squared value of filter's coefficients is of interest. Suppose n tends to the infinity. Then

$$\begin{aligned}
 \sigma_{\epsilon_1}^2(n|D(n)) &\stackrel{\text{df}}{=} E\{\epsilon_1^2(n)|D(n-1)\} = E\{[y(n) - D^T(n-1)X(n)]^2|D(n-1)\} = \\
 &= E\{y(n)^2\} - 2D^T(n-1)r_{yx}(n) + D^T(n-1)R_{xx}(n)D(n-1) = \\
 &= E\{y(n)^2\} - D^T(n-1)R_{xx}(n)D(n-1)
 \end{aligned} \quad (15a)$$

in view of

$$R_{xx}(n)D(n) = r_{yx}(n) \quad (15b)$$

and

$$D(n-1) \cong D(n). \quad (15c)$$

If the optimal coefficients are used instead of actual ones then according to (14)

$$\begin{aligned}
 \sigma_{\epsilon_{1\min}}^2 &\stackrel{\text{df}}{=} E\{\epsilon_1^2(n)|D_{\text{opt}}\} = E\{y(n)^2\} - D_{\text{opt}}^T R_{xx}(n) D_{\text{opt}} = \\
 E\{y(n)^2\} - D_{\text{opt}}^T r_{yx}(n) &= E\{y(n)^2\} - r_{yx}^T(n) R_{xx}^{-1}(n) r_{yx}(n).
 \end{aligned} \quad (15d)$$

Moreover

$$\begin{aligned}
 [D_{\text{opt}} - D(n)]^T R_{xx}(n) [D_{\text{opt}} - D(n)] &= \\
 = D_{\text{opt}}^T R_{xx}(n) D_{\text{opt}} + D(n)^T R_{xx}(n) D(n) - 2D_{\text{opt}}^T R_{xx}(n) D(n) &= \\
 = D_{\text{opt}}^T R_{xx}(n) D_{\text{opt}} + D(n)^T R_{xx}(n) D(n) - 2D_{\text{opt}}^T r_{yx}(n) &= \\
 = D(n)^T R_{xx}(n) D(n) - D_{\text{opt}}^T R_{xx}(n) D_{\text{opt}}.
 \end{aligned} \quad (16a)$$

That is why

$$\sigma_{\epsilon_1}^2(n|D(n)) = \sigma_{\epsilon_{1\min}}^2 + (D_{\text{opt}} - D(n))^T R_{xx}(n) (D_{\text{opt}} - D(n)) \quad (16b)$$

Now, let

$$\begin{aligned}
 1^\circ. R_{xx}(n) &= R_{xx}; \quad n \rightarrow \infty \\
 2^\circ. R_{xx} &= U^T \text{diag}(\lambda_i) U, \quad UU^T = I,
 \end{aligned}$$

$$\text{diag}(\lambda_i) = \begin{bmatrix} \lambda_0 & & & 0 \\ & \ddots & & \\ & & \ddots & \\ 0 & & & \lambda_{N-1} \end{bmatrix}_{N \times N}$$

λ_i is the i^{th} eigen value of R_{xx} , $i = 0, \dots, N-1$.

$$3^\circ. D'(n) = U(D_{\text{opt}} - D(n)).$$

Then

$$\sigma_{\epsilon_1}^2(n|D(n)) = \sigma_{\epsilon_{1\min}}^2 + D'(n)^T \text{diag}(\lambda_i) D'(n) \quad (16c)$$

or

$$\sigma_{\epsilon_1}^2(n|D(n)) = \sigma_{\epsilon_{1\min}}^2 + \sum_{i=0}^{N-1} \lambda_i D'(n)^2 \quad (16d)$$

Now, the equation (11b) can be transformed into

$$-UD(n+1) + UD_{\text{opt}} = -UD(n) + UD_{\text{opt}} - \delta UX(n+1)\epsilon_1(n+1) \quad (17a)$$

or

$$D'(n+1) = D'(n) - \delta UX(n+1)\epsilon_1(n+1). \quad (17b)$$

In addition

$$\begin{aligned} D'(n+1)D'(n+1)^T &= [D'(n) - \delta UX(n+1)\epsilon_1(n+1)][D'(n) - \delta UX(n+1)\epsilon_1(n+1)]^T = \\ &= D'(n)D'(n)^T - \delta \epsilon_1(n+1)D'(n)X(n+1)^T U^T + \\ &\quad - \delta UX(n+1)\epsilon_1(n+1)D'(n)^T + \delta^2 UX(n+1)\epsilon_1^2(n+1)X(n+1)^T U^T = \\ &= D'(n)D'(n)^T - \delta UX(n+1)\epsilon_1(n+1)D'(n)^T + \\ &\quad - \delta D'(n)\epsilon_1(n+1)X(n+1)^T U^T + \delta^2 UX(n+1)\epsilon_1^2(n+1)X(n+1)^T U^T. \end{aligned} \quad (18)$$

Also

$$\begin{aligned} \epsilon_1(n+1) &= y(n+1) - D^T(n)X(n+1) = \\ &= y(n+1) - D_{\text{opt}}^T X(n+1) + (D_{\text{opt}} - D(n))^T X(n+1) = \\ &= y(n+1) - D_{\text{opt}}^T X(n+1) + D'(n)^T UX(n+1) = \\ &= y(n+1) - D_{\text{opt}}^T X(n+1) + X(n+1)^T U^T D'(n) = \\ &= v(n+1) + X(n+1)^T U^T D'(n) = \\ &= v(n+1) + D'(n)^T UX(n+1) \end{aligned} \quad (19a)$$

and

$$v(n+1) \stackrel{\text{df}}{=} y(n+1) - D_{\text{opt}}^T X(n+1). \quad (19b)$$

With aid of (19) the (18) is

$$\begin{aligned} D'(n+1)D'(n+1)^T &= D'D'^T - \delta UX(v + X^T U^T D')D'^T + \\ &\quad - \delta D'(v + D'^T UX)X^T U^T + \delta^2 UX\epsilon_1^2 X^T U^T. \end{aligned} \quad (20)$$

Suppose [3]:

{The conditions of adaptation well posed}

- 1°. $E\{vX^T U^T\} = E\{v\}E\{X^T U^T\} = 0$
- 2°. $E\{D'D'^T UXX^T U^T\} = E\{D'D'^T\}E\{UXX^T U^T\}$
- 3°. $E\{UXX^T U^T D'D'^T\} = E\{UXX^T U^T\}E\{D'D'^T\}$
- 4°. $E\{\epsilon_1^2 UXX^T U^T\} = E\{\epsilon_1^2\}E\{UXX^T U^T\}$

Then

$$\begin{aligned} E\{D'(n+1)D'(n+1)^T\} &= E\{D'(n)D'(n)^T\} - \delta \text{diag}(\lambda_i) E\{D'(n)D'(n)^T\} + \\ &\quad - \delta E\{D'(n)D'(n)^T\} \text{diag}(\lambda_i) + \delta^2 E\{\epsilon_1^2\} \text{diag}(\lambda_i) = \\ &= [I - 2\delta \text{diag}(\lambda_i)] E\{D'(n)D'(n)^T\} + \delta^2 E\{\epsilon_1^2\} \text{diag}(\lambda_i), \end{aligned} \quad (21)$$

In Bellanger [3], 1^0-4^0 have been formulated as hypotheses. One can strengthen these assumptions and refer to 1^0-4^0 as to the conditions of adaptation well posed. As a matter of fact the lack of the correlations of 1^0-4^0 reflects the fact of optimal performance of adaptive procedure.

Similar reasoning due to Kailath in filtering theory case gave rise to the innovation formulation of Kalman's filter.

Next, the expected value operation can be taken of both sides of (16c) in respect to $D(n)$:

$$\sigma_{\epsilon_i}^2(n) = \sigma_{\epsilon_{\min}}^2 + \Lambda^T E\{D'^2(n)\} \quad \Lambda = \text{col}(\lambda_0, \dots, \lambda_{N-1}) \quad (22)$$

In view of (22) the equation (21) is

$$\begin{aligned} E\{D'(n+1)D'^T(n+1)\} &= [I - 2\delta \text{diag}(\lambda_i)] E\{D'(n)D'(n)^T\} + \\ &\quad + \delta^2 \sigma_{\epsilon_{\min}}^2 \text{diag}(\lambda_i) + \delta^2 \Lambda^T E\{D'(n)^2\} \text{diag}(\lambda_i) \end{aligned} \quad (23)$$

While the main diagonal of both sides of (23) is of interest one gets

$$\begin{aligned} E\{D'^2(n+1)\} &= [I - 2\delta \text{diag}(\lambda_i)] E\{D'^2(n)\} + \\ &\quad + \delta^2 \sigma_{\epsilon_{\min}}^2 \Lambda + \delta^2 \Lambda \Lambda^T E\{D'(n)^2\} \end{aligned} \quad (24)$$

Or equivalently

$$E\{D'^2(n+1)\} = A E\{D'^2(n)\} + \delta^2 \sigma_{\epsilon_{\min}}^2 \Lambda, \quad (25a)$$

where

$$A = I - 2\delta \text{diag}(\lambda_i) + \delta^2 \sigma_{\epsilon_{\min}}^2 \Lambda, \quad (25b)$$

The concept of norm can be used to establish convergence:

$$\|A\|_{l^2} \leq \|A\|_{l^1} = \max_j \sum_i |a_{ij}| \quad (25c)$$

$\{l^1, l^2$ are the spaces of the sequences summable in absolute and in square sense, of finite dimension in this case}.

The condition of convergence which follows from the requirement

$$\|A\|_{l^1} < 1 \text{ is;} \quad (26a)$$

$$0 < 1 - 2\delta \lambda_i + \delta^2 \lambda_i \left(\sum_{j=0}^{N-1} \lambda_j \right) < 1 \quad (26b)$$

$$0 \leq i \leq N-1$$

or

$$0 < \delta/2 < \frac{1}{\sum_{j=0}^{N-1} \lambda_j} = \frac{1}{N \sigma_x^2}. \quad (26c)$$

Moreover from (21)

$$E\{D'(\infty)D'^T(\infty)\} = \frac{\delta}{2} \sigma_{\epsilon_1}^2(\infty) I \quad (27)$$

Also from (22)

$$\sigma_{\epsilon_1}^2(\infty) = \sigma_{\epsilon_{1\min}}^2 + \frac{\delta}{2} N \sigma_x^2 \sigma_{\epsilon_1}^2 \quad (28)$$

Finally

$$\sigma_{\epsilon_1}^2(\infty) = \frac{\sigma_{\epsilon_{1\min}}^2}{1 - \frac{\delta}{2} N \sigma_x^2}. \quad (29)$$

If

$R_{xx} = U^T \text{diag}(\lambda_i) U$ and $y = Ux$ then

$y^T y = x^T U^T U x = x^T x$ and $E\{yy^T\} = \text{diag}(\lambda_i)$.

That is why

$$E\{x^T x\} = N \sigma_x^2 = E\{yy^T\} = \sum_i \lambda_i.$$

Moreover (see (24))

$$E\{D'(n+1)^2\} \cong \text{diag}(1 - \delta\lambda_i)^2 E\{D'(n)^2\} + \delta^2 \sigma_{\epsilon_{1\min}}^2 \Lambda \quad (30a)$$

and

$$E\{D'(n)^2\} \cong \text{diag}(1 - \delta\lambda_i)^{2n} E\{D'(0)\} + \delta^2 \sigma_{\epsilon_{1\min}}^2 \Lambda. \quad (30b)$$

Also ((22), (30b))

$$\sigma_{\epsilon_1}^2(n) - \sigma_{\epsilon_{1\min}}^2 = \Lambda^T \text{diag}(1 - \delta\lambda_i)^{2n} E\{D'^2(0)\} \quad (30c)$$

and

$$\sigma_{\epsilon_1}^2(0) \cong \Lambda^T E\{D'(0)^2\}. \quad (30d)$$

That is why if

$$\sigma_{\epsilon_1}^2(n) = \sigma_{\epsilon_1}^2(0) e^{-n/(\tau_\epsilon/2)} \quad (30e)$$

then

$$\Lambda^T E\{D'^2(0)\} [1 - 1/(\tau_\epsilon/2)] \cong \Lambda^T \text{diag}(1 - 2\delta\lambda_i) E\{D'^2(0)\}. \quad (30f)$$

From (30f)

$$\Lambda^T E\{D'^2(0)\} \frac{1}{\tau_\epsilon} \cong \delta \Lambda^T \text{diag}(\lambda_i) E\{D'^2(0)\} \quad (30g)$$

and

$$\tau_\epsilon \cong \frac{\sum_{i=0}^{N-1} \lambda_i E\{D'^2(0)\}}{\sum_{i=0}^{N-1} \lambda_i^2 E\{D'^2(0)\}} \frac{1}{\delta}. \quad (30h)$$

Thus

$$\frac{\lambda_{\min}}{\lambda_{\max}^2} \frac{1}{\delta} \leq \tau_\epsilon \leq \frac{\lambda_{\max}}{\lambda_{\min}^2} \frac{1}{\delta} \quad (30i)$$

If the eigenvalues are not widespread then

$$\tau_\epsilon \cong \frac{1}{\sigma_x^2 \delta}. \quad (30j)$$

In view of (25b) the greatest speed of convergence (minimal norm of A) is for

$$\delta = \delta_{\min} = \frac{1}{\sum_{j=0}^{N-1} \lambda_j} = \frac{1}{N\sigma_x^2}. \quad (30k)$$

The time constant for this step of adaptation is

$$\tau_{\epsilon_{\min}} \cong N. \quad (30l)$$

5. QUANTIZATION AND BLOCKING EFFECTS IN THE ADAPTIVE FILTER WITH CONSTANT STEP OF ADAPTATION

When the additional quantization noise is taken into account the equation (11b) should be changed into

$$D(n+1) = D(n) + \delta X(n+1)[y(n+1) - X^T(n+1)D(n) + \sum_{i=0}^{N-1} w_{q_i}^i] + W_{q_1} \quad (31)$$

By vector $W_{q_1} = \text{col}(w_{q_1}^0, \dots, w_{q_1}^{N-1})$ the effective quantization noise of the multipliers marked as Q_1 in fig. 2 is being modelled. Furthermore

$W_{q_2} = \text{col}(w_{q_2}^0, \dots, w_{q_2}^{N-1})$ is the effective quantization noise of the multipliers marked as Q_2 in fig. 2.

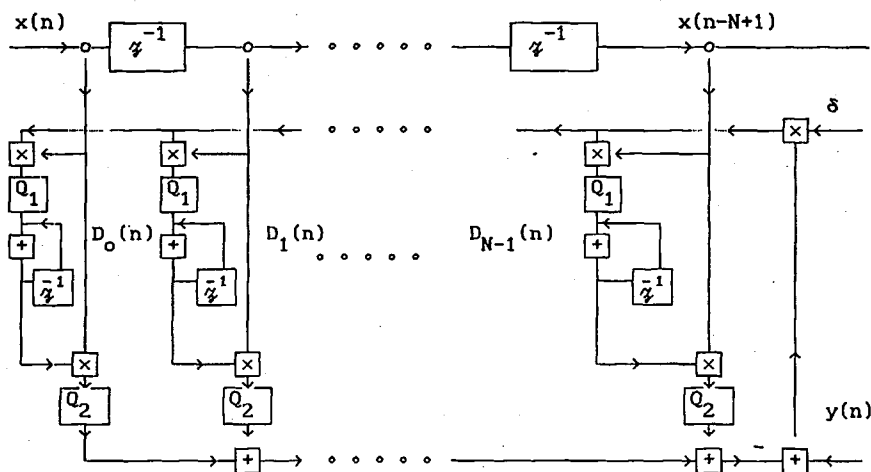


Fig. 2. Adaptive transversal filter with constant step of adaptation and quantization effects.

The equation (31) can be changed into

$$\begin{aligned}
 -UD(n+1) + UD_{\text{opt}} &= -UD(n) + \delta UD_{\text{opt}} + \\
 &\quad - \delta UX(n+1)[\epsilon_1(n+1) + \sum_{i=0}^{N-1} w_{q_2}^i] + W_{q_1}
 \end{aligned} \quad (32a)$$

or

$$D'(n+1) = D'(n) - \delta UX(n+1)[\epsilon_1(n+1) + \sum_{i=0}^{N-1} w_{q_2}^i] + w_{q_1}. \quad (32b)$$

In comparison to (21) the additional term is introduced.

Suppose

1°. The quantization noise sequence of step q is white noise like and uniformly distributed [3,4].

2°. There is no correlation between noise, data and coefficients of the filter.

Under these hypotheses the additional term is

$$\begin{aligned}
 E\{UX(n+1) \sum_{i=0}^{N-1} w_{q_2}^i \delta X(n+1)^T U^T \sum_{i=0}^{N-1} w_{q_2}^i\} + E\{W_{q_1} W_{q_1}^T\} &= \\
 &= \delta^2 \text{diag}(\lambda_i) N \frac{q_2^2}{12} + \frac{q_1^2}{12} I.
 \end{aligned} \quad (33)$$

Moreover from (21) (see also (27))

$$E\{D'^{(\alpha)}(\infty)^T\} = \frac{\delta}{2} \sigma_{\epsilon_{\min}}^2(\infty) I + \frac{\delta}{2} N \frac{q_2^2}{12} I + \frac{1}{2\delta} \frac{q_1^2}{12} \text{diag}(\lambda_i^{-1}). \quad (34)$$

Furthermore the modification of (22) is at addition of quantization power $N \frac{q_2^2}{12}$.

That is why

$$\sigma_{\epsilon_i}^2(\infty) \underset{(22)}{=} \sigma_{\epsilon_{\min}}^2 + N \frac{q_2^2}{12} + N \sigma_x^2 \sigma_{\epsilon_i}^2(\infty) \frac{\delta}{2} + N^2 \sigma_x^2 \frac{\delta}{2} \frac{q_2^2}{12} + \frac{1}{2\delta} \frac{q_1^2}{12} N. \quad (35a)$$

or

$$\begin{aligned} \sigma_{\epsilon_i}^2(\infty) &= \frac{1}{1 - \frac{\delta}{2} N \sigma_x^2} \left[\sigma_{\epsilon_{\min}}^2 + \frac{N q_2^2}{12} + \frac{N}{2\delta} \frac{q_1^2}{12} + N^2 \sigma_x^2 \frac{\delta}{2} \frac{q_2^2}{12} \right] = \\ &= \frac{1}{1 - \frac{\delta}{2} N \sigma_x^2} \left[\sigma_{\epsilon_{\min}}^2 + \frac{N q_2^2}{12} \left(1 + \frac{N \sigma_x^2 \delta}{2} + \frac{N q_1^2}{2\delta} \frac{1}{12} \right) \right], \end{aligned} \quad (35b)$$

While

$$\frac{1}{1 - \frac{\delta}{2} N \sigma_x^2} \cong 1 + \frac{\delta}{2} N \sigma_x^2 \quad (36a)$$

one gets

$$\begin{aligned} \sigma_{\epsilon_i}^2(\infty) &\cong \sigma_{\epsilon_{\min}}^2 \left(1 + \frac{\delta}{2} N \sigma_x^2 \right) + \frac{N q_2^2}{12} \left(1 + \frac{\sigma_x^2 N}{2} \delta \right) + \frac{N q_1^2}{2\delta} \frac{1}{12} + \\ &+ \frac{\delta}{2} \frac{N^2 \sigma_x^2 q_2^2}{12} \left(1 + \frac{N \sigma_x^2 \delta}{2} \right) + \frac{\delta}{2} N \sigma_x^2 \frac{N}{2\delta} \frac{q_1^2}{12} \cong \\ &\cong \sigma_{\epsilon_{\min}}^2 \left(1 + \frac{\delta}{2} N \sigma_x^2 \right) + \frac{N q_2^2}{12} (1 + N \sigma_x^2 \delta) + \\ &+ \frac{N}{2\delta} \frac{q_1^2}{12} + \frac{N^2 \sigma_x^2}{4} + \frac{q_1^2}{12}. \end{aligned} \quad (36b)$$

By the calculation of the derivative of (36) in respect to δ and equating to zero one gets

$$\sigma_{\epsilon_{1\min}}^2 \sigma_x^2 + \frac{2N\sigma_x^2 q_2^2}{12} = \frac{1}{\delta^2} \frac{q_1^2}{12}. \quad (36c)$$

That is why

$$\delta_{\text{opt}} \cong \frac{q_1}{2\sqrt{3}} \frac{1}{\sigma_{\epsilon_{1\min}} \sigma_x} \frac{1}{\sqrt{1 + \frac{Nq_2^2}{6\sigma_{\epsilon_{1\min}}^2}}}. \quad (36d)$$

or

$$\delta_{\text{opt}} \cong \frac{q_1}{2\sqrt{3}} \frac{1}{\sigma_{\epsilon_{1\min}} \sigma_x} \text{ while } \frac{Nq_2^2}{6\delta\sigma_{\epsilon_{1\min}}^2} \ll 1. \quad (36e)$$

The existence of optimal step of adaptation σ_{opt} can be explained as follows. The greater the step of adaptation the greater residual error. On the other hand the less the step of adaptation the greater impact of quantization noise. Namely, the results of multipliers with Q_1 quantizing effects are embedded in greater quantization noise relatively, giving rise to blocking phenomenon.

The blocking effects, due to the nonlinear mechanism of quantization can take place in the digital circuitry of the filter if [3]:

$$|\delta x(n-i)\epsilon_1(n)| < q_1/2. \quad (37a)$$

Suppose that;

- 1°. The elements of $\delta X(n)\epsilon_1(n)$ vector are not correlated each with the other,
- 2°. Uniformly distributed in $[-q_1/2, q_1/2]$ closed interval and
- 3°. The input signal is the white noise like.

Then the blocking condition is [3]:

$$E\{\delta^2 \epsilon_1^2(n) X(n)X(n)^T\} = \frac{q_1^2}{12} I = \delta^2 E\{E_{n \rightarrow \infty} [\epsilon_1(n)X(n)] E_{|n \rightarrow \infty} [\epsilon_1(n)X(n)^T]\} \quad (37b)$$

$$E_{|n \rightarrow \infty} \{\cdot(n)\} = E\{\cdot(n) | n \text{ large}\}.$$

While input signal and filter's coefficients are not correlated and the discrete time n tends to the infinity:

$$\begin{aligned} R_{xx}^{-1} E_{|n \rightarrow \infty} \{\epsilon_1(n)X(n)\} &= R_{xx}^{-1} E_{|n \rightarrow \infty} \{X(n)[y(n) - X^T(n)D(n-1)]\} = \\ &= R_{xx}^{-1} r_{yx}(n) - R_{xx}^{-1} R_{xx} E_{|n \rightarrow \infty} \{D(n-1) \cong D(n) - D_{\text{opt}}\}. \end{aligned} \quad (37c)$$

That is why

$$\begin{aligned} E \{ R_{xx}(D(n) - D_{\text{opt}}) [R_{xx}(D(n) - D_{\text{opt}})]^T \} &= E \{ \varepsilon_1^2(n) X(n) X(n)^T \} = \\ &= \frac{q_1^2}{12 \delta^2} I. \end{aligned} \quad (38a)$$

Now, from (38a)

$$E \{ (D(n) - D_{\text{opt}})(D(n) - D_{\text{opt}})^T R_{xx}^{-1} \frac{q_1^2}{\delta^2 12} I R_{xx}^{-1} \} = \frac{q_1^2}{\delta^2 12} U \text{diag}(\lambda_i^{-2}) U^T, \quad (38b)$$

That is why

$$E \{ [D(n) - D_{\text{opt}}][D(n) - D_{\text{opt}}]^T \} = \left(\frac{q_1}{12 \delta} \right)^2 \sum_{i=0}^{N-1} \lambda_i^{-2}. \quad (38c)$$

The last quantity in hand can be called the blocking radius [3]

$$\rho^2 \stackrel{\text{df}}{=} (q_1/12\delta)^2 \sum_{i=0}^{N-1} \lambda_i^{-2}, \quad (38d)$$

If $N\sigma_x^2\delta \ll 1$ (as usually valid in practice) then from (36b)

$$\sigma_{\varepsilon_1}^2(\infty) \cong \sigma_{\varepsilon_{1\min}}^2 \left(1 + \frac{\delta}{2} N\sigma_x^2 \right) + \frac{Nq_2^2}{12} + \frac{N}{2\delta} \frac{q_1^2}{12}. \quad (39)$$

To express the error while $\sigma_{\varepsilon_{1\min}}^2$ is dominant it is reasonable to assume [3]:

$$\frac{1}{4} \sigma_{\varepsilon_{1\min}}^2 \delta N\sigma_x^2 = N \frac{q_2^2}{12} = \frac{N}{2\delta} \frac{q_1^2}{12}. \quad (40)$$

The kind of balanced expression can be gained this way.

Now, let the binary representation of q_1 be

$$q_1 = D_{\max} 2^{-(b_1-1)} \quad (41a)$$

while b_1 is the number of bits and

$$D_{\max} = \|D\|_c = \max_{0 \leq i \leq N-1} |D_i| \quad (41b)$$

stands for the greatest filter's coefficients. According to (40)

$$2^{2b_1} = \frac{2}{3} \frac{D_{\max}^2}{\delta^2 \sigma_{\varepsilon_{1\min}}^2 \sigma_x^2} = \frac{2}{3} D_{\max}^2 \tau_{\varepsilon_1}^2 \frac{\sigma_x^2}{\sigma_{\varepsilon_{1\min}}^2} \cong \tau_{\varepsilon_1}^2 A_f^2 D_{\max}^2 \frac{\sigma_x^2}{\sigma_y^2}. \quad (42a)$$

The meaning of $A_f \stackrel{\text{df}}{=} \frac{\sigma_y^2}{\sigma_{\epsilon_{\text{Imin}}}^2}$ is the amplification factor.

Next, from (42b)

$$b_1 \cong \log_2 \tau_{\epsilon_1} + \log_2 \left(D_{\max} \frac{\sigma_x}{\sigma_y} \right) + \log_2 A_f. \quad (43)$$

Moreover $\{y(n) \cong D(n)X^T(n)\}$.

$$\sigma_y^2 = E\{y^2(n)\} \cong D^T(n)R_{xx}D(n) \geq \lambda_{\min} D^T(n)D(n) \geq \lambda_{\min} D_{\max}^2. \quad (44a)$$

That is why

$$D_{\max} \frac{\sigma_x}{\sigma_y} \leq \sigma_x / \sqrt{\lambda_{\min}}. \quad (44b)$$

On the other hand the number of quantization levels with q_2 step can be estimated as follows

$$q_2 = \max\{x(n), y(n)\} 2^{-(b_2-1)}. \quad (45a)$$

If

$$1^0) \sigma_x^2 \geq \sigma_y^2 \quad \text{and}$$

$$2^0) |x(n)|_{\max} = 4\sigma_x$$

then

$$q_2 = 4\sigma_x 2^{-(b_2-1)}. \quad (45b)$$

From (40)

$$2^{2b_2} = 2^{2*2} \frac{4}{3} \frac{1}{\sigma_{\epsilon_{\text{Imin}}}^2 \delta} \cong 2^{2*2} \tau_{\epsilon} \frac{\sigma_x^2}{\sigma_{\epsilon_{\text{Imin}}}^2} \cong 2^{2*2} \tau_{\epsilon} A_f^2 \frac{\sigma_x^2}{\sigma_y^2}, \quad (45c)$$

That is why

$$b_2 \cong 2 + \frac{1}{2} \log_2 \tau_{\epsilon} + \log_2 A_f + \log_2 \frac{\sigma_x}{\sigma_y}. \quad (45d)$$

FINAL REMARKS

The linear adaptive filter with constant step of adaptation has been described in detail. It is rather simple and not sensitive to divergence. The effects of quantization have been explained while showing the existence of the optimal step of adaptation. Also the blocking phenomenon has been considered. The detailed expressions being derived can be very useful on designing stage. Moreover the case of the filter with constant step of adaptation has been embedded in the problem of optimal linear filtering. Just the knowledge of relations between the different approaches to the similar kind of problems can give meaningful insight into understanding the

algorithms, making their modifications and improvements. The linear filter of this paper can be used while designing adaptive digital filters, fast and optimal on each stage of adaptation and possessing the properties of short term and long term convergence [32].

REFERENCES

1. J.Z. Cypkin: *Podstawy teorii układów uczących się*. WNT, Warszawa 1973
2. A.H. Jazwiński: *Stochastic processes and filtering theory*. Academic Press, 1970
3. M. Bellanger: *Adaptive digital filtering*. ECCTD 1987, Paris
4. R.W. Schaffer, A.V. Oppenheim: *Cyfrowe przetwarzanie sygnałów*. WNT, Warszawa 1979
5. F.C. Schweppe: *Układy dynamiczne w warunkach losowych*. WNT, Warszawa, 1978
6. J. Koronacki: *Aproksymacja stochastyczna: metody optymalizacji w warunkach losowych*. Warszawa, WNT, 1989
7. D. Falconer, L. Ljung: *Application of Fast Kalman Estimation to Adaptive Equalization*. IEEE Trans., COM-26, no 10, October 1978, pp. 1439-1446
8. G. Carayannis, D. Manolakis, N. Kalauptsidis: *A Fast Sequential Algorithm for L.S. Filtering and Prediction*. IEEE Trans., vol. ASSP-31, no 6, Dec. 1983, pp. 1394-1402
9. S.M. Kay, J.R. Marple: *Spectrum analysis - a modern perspective*. Proceedings of the IEEE, vol. 69, no 11, November 1981
10. S.W. Fomin, I.P. Kornfeld, J.G. Sinaj: *Teoria ergodyczna*. Warszawa 1987
11. P. De Larminat, Y. Thomas: *Automatyka - układy liniowe*, PWN, Warszawa 1983, (Polish translation from French)
12. W. Szlenk: *Wstęp do teorii gładkich układów dynamicznych*. PWN, Warszawa 1982
13. A.V. Oppenheim, R.W. Schaffer: *Cyfrowe przetwarzanie sygnałów*. WKiŁ, Warszawa 1975 (Polish translation from English)
14. F.C. Schweppe: *Układy dynamiczne w warunkach losowych*. WNT, Warszawa 1978 (Polish translation from English)
15. Gauss: *Theoria Motus Corporum Coelestium*, 1809
16. N. Wiener: *Extrapolation and smoothing of stationary time series with engineering applications*. New York, Technology Press and Wiley, 1949
17. R.E. Kalman: *A new approach to linear filtering and prediction problems*. Journal of Basic Engineering, vol. 82, pp. 34-35, 1960
18. N. Levinson: *The Wiener RMS (Root Mean Square) error criterion in filter prediction*. Journal of Mathematics and Physics, 1947, vol. 25, pp. 261-278
19. H. Robbins, S.A. Monro: *A stochastic approximation method*. Ann. Math. Stat., 1954, vol. 25, pp. 382-386
20. T. Kailath: *A view on three decades of linear filtering theory*. IEEE Transaction on Information Theory, 1974, vol. IT-20, pp. 145-181
21. P. Eykhoff: *Identyfikacja w układach dynamicznych*. PWN, Warszawa 1980
22. J.G. Proakis: *Digital Communications*. New York, Mc Graw Hill, Inc., 1983
23. B. Widrow, S.D. Stearns: *Adaptive Digital processing*. Prentice Hall, Englewood Cliffs, New York, USA, 1985
24. M. Bellanger: *Adaptive Digital Filtering and Signal Analysis*. Marcel Dekker Inc., New York, 1987
25. S. Haykin: *Adaptive Filtering Theory*. Prentice-Hall, Englewood Cliffs, New York, 1986
26. L.Ljung, T. Soderstrom: *Theory and Practice of recursive identification*. MIT Press, Cambridge 1983
27. L. Ljung, M. Morf, D.D. Falconer: *Fast calculation of gain matrices for recursive estimation schemes*. Int. J. of Control, 1978, vol. 27, no 1, pp. 1-19

28. G. Carayannis, D. Manolakis, N. Kaloupsidis: *Fast Kalman-type algorithms for sequential signal processing*. ICASSP 83, Boston, USA
29. J. Cioffi and T. Kailath: *Fast Recursive Least Squares Filters for Adaptive Filtering*. IEEE Trans., Vol. ASSP-32, no 2, April 1984, pp. 304-337
30. F. Ling, D. Manolakis and J. Proakis: *Numerically Robust L.S Lattice Ladder Algorithms with Direct Updating of the Coefficients*. IEEE Trans., vol. ASSP-34, no 4, August 1986, pp. 837-845
31. M.L. Honig, D.G. Messerschmitt: *Adaptive Filters - Structures, Algorithms and Applications*. Kluwer Academic Publishers, Boston, USA, 1984.
32. M. Żółtowski: *On the problem of divergence in adaptive filtering* (in progress)
33. M. Żółtowski: *Circuit theory oriented approach to the short term convergence of digital filtering algorithms*. Kwartalnik Elektroniki i Telekomunikacji, Electronics and Telecommunications - Quarterly, 1993, vol. 39, z. 3

M. ŻÓŁTOWSKI

O LINIOWEJ ADAPTACYJNEJ FILTRACJI ZE STAŁYM KROKIEM ADAPTACJI

Streszczenie

Rezultatem rekursyjnej minimalizacji kwadratowej funkcji strat są względnie proste adaptacyjne algorytmy. Wśród nich można wyróżnić szybkie algorytmy optymalne [3], szybkie dzięki pomysłowej metodzie obliczania. Tym niemniej w tym artykule przedmiotem rozważań w formie przeglądowej jest liniowy filtr adaptacyjny ze stałym krokiem adaptacji. Chociaż jest to filtr suboptymalny to jednak jest ważnym filtrem. Za jego znaczeniem przemawiają: najprostsza realizacja spośród wszystkich liniowych algorytmów filtracji adaptacyjnej i brak rozbieżności. Chociaż wielu świetnych autorów poświęciło uwagę tej klasie problemów [1-31], problem adaptacji leży w tym obszarze wiedzy, który zasługuje na stałą uwagę wobec znacznego rozwoju możliwości i architektury systemów cyfrowego przetwarzania sygnałów. W porównaniu z niedawnym ujęciem Bellangera [3], zostało dodane porównanie z filtrem Kalmana, chociaż podobne można znaleźć u Haykina [25]. Warunki, przy których oblicza się wariancję błędu współczynników filtru zostały podniesione do warunków adekwatnie sformułowanego problemu adaptacji. Rozważając efekty kwantowania pokazano istnienie dobrze umotywowanego optymalnego kroku adaptacji. Można użyć pojęcia przestrzeni unormowanej w celu wykazania zbieżności adaptacyjnego algorytmu w sposób elegancki. Artykuł ten jest skierowany w stronę ulepszania rekursyjnych algorytmów wobec ich istotnego znaczenia w telekomunikacji, sterowaniu, identyfikacji itp.

Circuits theory oriented approach to the short term convergence of adaptive digital filtering algorithms

MARIUSZ ŻÓŁTOWSKI

Instytut Telekomunikacji, Politechnika Gdańska

Received 1992.05.10

Authorized 1993.09.03

While representing the task of adaptation the solution to the cost function optimization by the numerical algorithm which performs the adaptation procedure is considered. Adaptive filtering, running the adaptive procedure within digital circuitry, can be referred to the related problems being formulated as four different approaches; 1) The method of Recursive Least Squares, 2) The Filtering Theory, 3) Stochastic approximation, 4) Model Reference Adaptive Systems [1–39]. Though many brilliant authors have paid attention to this kind of problems they are still of interest for while being at the advent of numerous practical implementations should be weighted carefully and intensively. Just the knowledge of relation between the different approaches to the system interference, well understood, in order to obtain required information in terms of the set of parameters can give considerable insight into understanding the algorithms, making their improvements and modifications. This idea is being partly compromised in this paper within unified approach in review form. The proof of short term (transient) convergence of stochastic identification type algorithm has been originated by the author within circuits theory oriented approach using the notion of pulse response and norm space concept. Though similar can be found in Haykin [29] and Bellanger [11] in linear and constant step of adaptation case. The extension to include nonsmooth case within one approach is proposed by using sigmoidal function, which proved useful for neural networks [37]. Linear adaptive filtering is characterized at the end.

1. INTRODUCTION

New trends in the architecture of microprocessors towards high signal processing ability including parallel processing via systolic arrays, wave front processors etc., while offered at low prices make chance for practical implementation of even

complex algorithms in order to improve the performance of existing systems or realization of new ones not feasible before. The problem of adaptation is just the one of these areas of interest which deserve continued attention in view of recent developments in hardware and architecture of digital systems. In this paper the problem of adaptation is considered in terms of optimization of cost functional according to new data available for processing. The problem is stated in chapter II and the recursive, stochastic approximation type scheme is presented to provide the solution. The constant step of adaptation case of this scheme is given in chapter III, which requires less hardware complexity and exhibits very good convergence properties but at the expense however of convergence speed. The algorithms with step diminishing, including the optimal ones, are treated in chapter 4, where the original proof of the convergence in white and coloured disturbing error sequences cases is provided. The considerations of chapters 2 and 3 are performed using familiar to electric and electronic engineer related notions of the discrete time system transfer function and pulse response function. In chapter V the problem of algorithm optimization is addressed from the two points of view: the minimization of the coefficients of the filter error variance sequence and the one of the recursive scheme of form (3) gain sequence. The equivalence of these approaches in linear case is pointed out twice; in this chapter and chapter 7 in relation to Kalman's filter. The effects of nonsmoothness are held in chapter 6, where the way of treatment within the unified approach of previous chapters using the neural network example is shown. Eventually, the linear adaptive filtering is characterized in chapter 7. The last chapter provides end comments.

2. THE TASK OF ADAPTATION

One can consider the following problem. Find the minimum of functional

$$\min_{D(n)} C\{D(n)\} = \min_{D(n)} E [L\{x, y(n), n, D(n)\}] = \min_{D(n)} \int_X L\{x, y(n), n, D(n)\} p(x, n|y) dx \quad (1)$$

$$= \min_{D(n)} \int_X L\{x, y(n), n, D(n)\} p(x, n|y) dx$$

$n = 1, \dots$

while

1°. $D(n)_{N \times 1}$ is the unknown vector being looked for (of N dimension)

2°. L is the convex loss function

3°. $p(x, n|y)$ is the density of conditional probability of vector $x \in X$ (X is the state space) while $y \in Y$ (Y is the space of observations or reference one).

4°. The parameter n stands for a discrete time sequence.

5°. $E \{ \}$ stands for the conditioned on Y mean value of the expression in $X|Y$ brackets in respect to X .

Moreover the case of unknown $p(x, n|y)$ is of interest. If x is a stationary process then the estimator $\hat{p}(x, n|y)$ of this density can be formulated:

$$\hat{p}(x, n|y) = \frac{1}{n} \sum_{k=1}^n \delta(x - x(k)|y(k)) \quad (2)$$

The $\delta(u)$ of (2) stands for Dirac function.

With p of (2) in hand the criterion of optimization is

$$\hat{C}\{D(n)\} = \frac{1}{n} \sum_{k=1}^n L[x(k), y(k), D(n)] \quad (3)$$

or

$$\hat{C}_0\{D(n)\} = \sum_{k=1}^n L[x(k), y(k), D(n)]. \quad (4)$$

The optimization of (3) in respect to $D(n)$ should be performed in recursive way while new observations are taken into account subsequently to modify $D(n)$ from previous stage at each step. Moreover, in view of ergodic theorem it is important to assure that

$$\min_{D(n)} \hat{C}\{D(n)\} \rightarrow \min_{n \rightarrow \infty} C\{D(n)\} \quad (5)$$

i.e. the minimum of the original cost function is gained in stationary, ergodic regime while n tends to infinity.

The conditions, under which $\hat{p}(x, n|y) \rightarrow p(x|y)$, while n tends to the infinity are described in ergodic theorems [39, 38]. For the stationary distribution $p(x|y)$ exists by adequate statement of identification task, the convergence takes place in $\mathcal{L}^p$ space in view of statistical ergodic theorem [39]. In view of this theorem also, the mean in respect to time of (3) is convergent to the one in respect to space. The minimization of both sides results in (5).

Unfortunately, the convergence just considered is generally in mean and may be very slow one from the practical point of view. The additional conditions, which are known as mixing properties are required [38]. For the mixing properties are closely related to the spectral ones of the dynamical system the convergence in mean squared sense of the error sequence of filter's coefficients is of concern in this paper. The spectral properties of the system are taken into account in terms of the transfer function concept in z or $j\omega$ variable which is familiar to electric and electronic engineer. The notion of the functional space is used in the proof of convergence.

The unimodal loss function has also been assumed, while the practical treatment of many external points can be performed by the decomposition into unimodal parts.

The latter criterion $\hat{C}\{D(n)\}$ of (3) can be optimized while the sequence of $D(n)$, $n=1, \dots$ evolves according to

$$D(n) = D(n-1) - \text{diag } \gamma(n) \nabla_D L[x(n), y(n), D(n-1)] \quad (6a)$$

$$\text{diag } \gamma(n) = \begin{bmatrix} \gamma_1(n) & & 0 \\ & \ddots & \\ 0 & & \gamma_N(n) \end{bmatrix} \quad (6b)$$

$$D(n)_{N \times 1} = \text{col } [D_1(n), \dots, D_N(n)] \quad (6c)$$

and the gradient of L is

$$\nabla L_{N \times 1} = \text{col } \left[\frac{\partial L}{\partial D_1}, \dots, \frac{\partial L}{\partial D_N} \right]. \quad (6d)$$

New correction of $D(n)$ is performed to make the minimum of $\hat{C}$ be gained at every step of this recursive algorithm. Though this form is the most suitable one for numerical processing it is worth while to consider another form of (6) in order to address the question of the short term stability:

$$\begin{aligned} D(n) = D(n-1) - \text{diag } \gamma(n) \mathop{\mathbb{E}}_{X|Y} \{ \nabla_D L[x(n), y(n), D(n-1)] \} + \\ + \text{diag } \gamma(n) \{ \mathop{\mathbb{E}}_{X|Y} \nabla_D L[x(n), y(n), D(n-1)] - \nabla_D L[x(n), y(n), D(n-1)] \} \end{aligned} \quad (7a)$$

In view of

$$\mathop{\mathbb{E}}_{X|Y} \nabla_D L[x(n), y(n), D(n-1)] = \nabla_D C\{D(n-1)\} \quad (7b)$$

(7) is

$$D(n) = D(n-1) - \text{diag } \gamma(n) \nabla_D C\{D(n-1)\} + \text{diag } \gamma(n) Z(n-1) \quad (7c)$$

and

$$\begin{aligned} Z(n-1) &\stackrel{\text{df}}{=} \nabla_D C\{D(n-1)\} - \nabla_D L[x(n), y(n), D(n-1)] = \\ &= \mathop{\mathbb{E}}_{X|Y} \nabla_D L[x(n), y(n), D(n-1)] - \nabla_D L[x(n), y(n), D(n-1)]. \end{aligned} \quad (7d)$$

¹⁾ While the loss function is convex and smooth (with bounded derivatives) this interchange of limit and integration is justified by Lebesgue theorem. On the other hand the any difference such that the equality of (7b) does not hold true exactly can be accounted for in $Z(n)$ of (7d).

If the term $Z(n)$ of (7c) is neglected the optimization following (7c) is in accord with minimization of $C\{D\}$ by gradient algorithm. That is why $Z(n)$ of (7c) can be given the meaning of disturbance while n tends to the infinity.

Suppose D^* is the optimal, stationary limit of $D(n)$ while $n \rightarrow \infty$. Then the next form of (7c) is

$$\begin{aligned} D(n) - D^* &= D(n-1) - D^* - \text{diag } \gamma(n) \nabla_D^2 C\{D^*\} [D(n-1) - D^*] + \\ &+ \text{diag } \gamma(n) \{ \nabla_D^2 C\{D^*\} [D(n-1) - D^*] - \nabla_D C[D^* + D(n-1) - D^*] \} + \\ &+ \text{diag } \gamma(n) Z(n-1) \end{aligned} \quad (8b)$$

$$\nabla_D^2 C\{D^*\} = \frac{\partial^2 C\{D^*\}}{\partial D_i \partial D_j}, \quad i, j = 1, \dots, N \quad (8b)$$

Let

$$\Delta D(n) \stackrel{\text{def}}{=} D(n) - D^* \quad (8c)$$

$$\varepsilon(n-1) \stackrel{\text{def}}{=} \nabla_D^2 C[D^*] \Delta D(n-1) - \nabla_D C[D^* + \Delta D(n-1)] \quad (8d)$$

while

$$\lim_{D(n) \rightarrow D^*} \|\varepsilon(n)\| = 0; \quad \varepsilon(n) = O^2(\|\Delta D(n-1)\|^2). \quad (8e)$$

Then around D^*

$$\Delta D(n) = \Delta D(n-1) - \text{diag } \gamma(n) \nabla_D^2 C[D^*] \Delta D(n-1) + \text{diag } \gamma(n) \{ \varepsilon(n-1) + Z(n-1) \} \quad (9)$$

Furthermore the spectral representation of matrix $\nabla_D^2 C[D^*]$ of (9) while symmetric and positively definite is;

$$\nabla_D^2 C[D^*] = U \text{diag } \lambda_i U^T; \quad U : U U^T = I \text{ (identity matrix) is unitary}$$

$$\text{diag } \lambda_i = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_N \end{bmatrix} \quad (10)$$

That is why (9) can be transformed to

$$\begin{aligned} \Delta D(n) &= \Delta D(n-1) - \text{diag } \gamma(n) U \text{diag } \lambda_i U^T \Delta D(n-1) + \\ &+ \text{diag } \gamma(n) [\varepsilon(n-1) + Z(n-1)] \end{aligned} \quad (11a)$$

²⁾ 0 stands for the quantity of the same order of value as given in parenthesis.

and

$$\begin{aligned} \mathbf{U}^T \Delta \mathbf{D}(n) &= \mathbf{U}^T \Delta \mathbf{D}(n-1) - \mathbf{U}^T \text{diag } \gamma(n) \mathbf{U} \text{diag } \lambda_i \mathbf{U}^T \Delta \mathbf{D}(n-1) + \\ &+ \mathbf{U}^T \text{diag } \gamma(n) [\varepsilon(n-1) + \mathbf{Z}(n-1)]. \end{aligned} \quad (11b)$$

Moreover, while

$$\Delta \mathbf{D}'(n) \stackrel{\text{df}}{=} \mathbf{U}^T \Delta \mathbf{D}(n) \quad (11c)$$

$$\Delta \mathbf{D}'(n) = [\mathbf{I} - \text{diag } \gamma_i(n) \lambda_i] \Delta \mathbf{D}'(n-1) + \text{diag } \gamma_i(n) \mathbf{U}^T [\varepsilon(n-1) + \mathbf{Z}(n-1)]. \quad (11d)$$

In view of (11d) each of N scalar cases instead of the vector one can be considered:

$$\Delta D'_i(n) = [1 - \gamma_i(n) \lambda_i] \Delta D'_i(n-1) + \gamma_i(n) \varepsilon'_i(n-1) \quad (12a)$$

$$i = 1, \dots, N$$

$$\varepsilon'_i(n-1) \stackrel{\text{df}}{=} [\mathbf{U}^T \{ \varepsilon(n-1) + \mathbf{Z}(n-1) \}]_i \quad (12b)$$

and

$[x]_i$ of (12b) stands for the i -th coordinate of x .

3. ALGORITHMS WITH CONSTANT STEP

While adaptation is performed at constant rate the algorithms performing the task of adaptation are called the ones with constant step. In this case

$$\gamma_i(n) = \gamma_i, \quad i = 1, \dots, N \quad (13a)$$

$$\Delta D'_i(n+1) = (1 - \gamma_i \lambda_i) \Delta D'_i(n) + \gamma_i \varepsilon'_i(n) \quad (13b)$$

$$n = 0, 1, \dots$$

It is convenient to introduce $\mathcal{Z}$ -Laurent transformation next. It is given for the time sequence $\{a(n)\}$ by

$$\mathcal{A}(\mathfrak{z}) = \mathcal{Z}\{a(n)\} \stackrel{\text{df}}{=} \sum_{n=0}^{\infty} a(n) \mathfrak{z}^{-n}. \quad (14)$$

Now, let

$$\mathcal{Z}\{\Delta D'_i(n)\} = \Delta \mathcal{D}'_i(\mathfrak{z}). \quad (15)$$

Then one gets by transforming the both sides of equation (13b)

$$\mathfrak{z} \Delta \mathcal{D}'_i(\mathfrak{z}) - \mathfrak{z} \Delta D'_i(0) = (1 - \gamma_i \lambda_i) \Delta \mathcal{D}'_i(\mathfrak{z}) + \gamma_i \mathcal{Z}\{\varepsilon'_i(n)\} \quad (16a)$$

or

$$\Delta \mathcal{D}'_i(\mathfrak{z}) = \frac{\mathfrak{z} \Delta D'_i(0)}{\mathfrak{z} - (1 - \gamma_i \lambda_i)} + \frac{\gamma_i}{\mathfrak{z} - (1 - \gamma_i \lambda_i)} \mathcal{Z}\{\varepsilon'_i(n)\}. \quad (16b)$$

On the other hand, the time domain form of (16b) after the inverse $\mathcal{Z}^{-1}$ -Laurent transformation is

$$\Delta D'_i(n) = \Delta D'_i(0) k_1(n) + \gamma_i \sum_{k=0}^{n-1} k_1(n-1-k) \varepsilon'_i(k) \quad (17a)$$

with

$$k_1(n) = (1 - \gamma_i \lambda_i)^n \mathbf{1}(n) = \mathcal{Z}^{-1}(\mathbf{K}_1(\mathfrak{z})) = \mathcal{Z}^{-1} \left\{ \frac{\mathfrak{z}}{1 - (1 - \gamma_i \lambda_i)} \right\} \quad (17b)$$

and

$$\mathbf{1}(n) = \begin{bmatrix} 0 & n < 0 \\ 1 & n \geq 0 \end{bmatrix}.$$

Now one can easily notice from (17) that if

$$0 < \gamma_i \cong 0 \quad \text{then} \quad 1 \cong 1 - \gamma_i \lambda_i < 1 \quad (\lambda_i > 0). \quad (18a)$$

Also:

1⁰. The initial deviation from the optimal vector of coefficients diminishes at rate given in (18a):

$$\Delta D'_i(0) k_1(n) \rightarrow 0. \quad (18b)$$

$n \rightarrow \infty$

2⁰. The steady state fluctuations while n tends to the infinity are given by

$$\Delta D'_i(n) = \gamma_i \sum_{n=-\infty}^{n-1} k_1(n-1-k) \varepsilon'_i(k) \quad (18c)$$

or

$$\Delta \mathcal{D}'_i(\mathfrak{z}) = \gamma_i \mathfrak{z}^{-1} \mathbf{K}_1(\mathfrak{z}) \mathcal{Z}\{\varepsilon'_i(n)\} \quad (18d)$$

in $\mathfrak{z}$ variable.

3⁰. If $P_{\varepsilon'_i}(\omega)$ is the spectral density function of $\varepsilon'_i(k)$ in stationary case then the one of $\Delta D'_i(n)$ in steady state is:

$$P_{\Delta D'_i}(\omega) = \gamma_i^2 |\mathbf{K}_1(e^{j\omega})| P_{\varepsilon'_i}(\omega). \quad (18e)$$

Thus the steady state variance $\sigma_{\Delta D'_i}^2$ of oscillations (the source of residual error [12]) is

$$\sigma_{\Delta D'_i}^2 = \frac{1}{2\pi} \int_{-\pi}^{\pi} P_{\Delta D'_i}(\omega) d\omega. \quad (18f)$$

Also from (13b) directly:

$$\sigma_{\Delta D_i}^2 \cong \frac{\gamma_i^2 \sigma_{\varepsilon_i}^2}{1 - (1 - \gamma_i \lambda_i)^2} \quad (18g)$$

and $\sigma_{\varepsilon_i}^2$ is the variance of $\varepsilon'_i(n)$ when white sequence of random variables in stationary case.

4°. The following estimation is in hand as well

$$\begin{aligned} \sigma_{\Delta \bar{D}_i}^2 &= \overline{\lim}_{n \rightarrow \infty} \sum_{k=0}^n \frac{\Delta D'_i(k)^2}{n} = \\ &= \gamma_i^2 \frac{1}{2\pi} \int_{-\pi}^{\pi} |K_1(e^{j\omega})| \cdot \overline{\lim}_{n \rightarrow \infty} \frac{|\mathcal{L}\{\varepsilon'_{i(n)}(k)\}|^2}{n} \bigg|_{z=e^{j\omega}} d\omega \end{aligned} \quad (18h)$$

$\overline{\lim}$ stands for limes superior and

$$\varepsilon'_{i(n)} = \begin{cases} \varepsilon'_i(k) & k \leq n \\ 0 & k > n \end{cases} \quad (18i)$$

In stationary case

$$\sigma_{\Delta \bar{D}_i}^2 = \sigma_{\Delta D_i}^2 \quad (18j)$$

and

$$P_{\varepsilon_i}(\omega) = \overline{\lim}_{n \rightarrow \infty} \frac{|\mathcal{L}\{\varepsilon'_{i(n)}\}|^2}{n} \bigg|_{z=e^{j\omega}} \quad (18k)$$

It follows from the presented approach that the attenuation of residual error is always at the cost of the prolonged transients of adaptation. It is more advantageous but more difficult for hardware requirements reason to implement the variable sequence of coefficients instead of constant one. However better trade off between the steady and the transient behaviour can be gained this way.

4. ALGORITHMS WITH STEP DIMINISHING

While the sequence of adaptation gain $\{\gamma(n)\}$ is diminishing the algorithm under consideration is

$$\begin{aligned} \Delta D'_i(n+1) &= [1 - \gamma_i(n+1)\lambda_i] \Delta D'_i(n) + \gamma_i(n+1)\varepsilon'_i(n) \\ i &= 1, \dots, N, \quad n = 0, 1, \dots \end{aligned} \quad (19a)$$

The homogeneous equation can be studied first;

$$\Delta D'_i(k+1) = [1 - \gamma_i(k+1)\lambda_i] \Delta D'_i(k). \quad (19b)$$

Let

$$\delta(k-m) \stackrel{\text{df}}{=} \begin{cases} 1 & k = m \\ 0 & k \neq m. \end{cases} \quad (19c)$$

Then the unit pulse response at time instant n due to the unit pulse $\delta(k-m)$ excitation view of (19b) is

$$H(n, m) = \begin{cases} 0 & n \leq 0 \\ 1 & n = m + 1 \\ \prod_{k=m+1}^{n-1} 1 - \gamma_i(k+1)\lambda_i, & n > m + 1. \end{cases} \quad (20a)$$

Now, according to superposition principle

$$\Delta D'_i(n) = H(n, -1) \Delta D'_i(0) + \sum_{m=0}^{n-1} H(n, m) \gamma_i(m+1) \varepsilon'_i(m) \quad (21)$$

$$n = 0, 1, \dots \quad \left(\sum_0^{-1} () = 0 \right).$$

Suppose

$$E\varepsilon'_i(k) = 0, \quad \sigma_{\varepsilon'_i}^2 = E\{\varepsilon'_i{}^2\} < \infty \quad (22a)$$

and $\{\varepsilon'_i(k)\}$ sequence is white, $k=0, 1, \dots$.

Then

$$\sigma_{\Delta D'_i}^2(n) = H(n, -1)^2 \Delta D'_i(0)^2 + \sum_{m=0}^{n-1} H(n, m)^2 \gamma_i^2(m+1) \sigma_{\varepsilon'_i}^2(m). \quad (22b)$$

As a matter of fact the following estimations are valid:

Lemma 1

If $\gamma_1(m) = 0$ ($1/m$) then (23a)

$$H(n, m)^2 = \frac{1}{n^2} 0(m^2) \quad (23b)$$

and

$$\sum_{m=0}^{n-1} H(n, m)^2 \gamma_i(m+1) = 0(1). \quad (23c)$$

Lemma 1 can be readily verified by the examination of (23b) and (23c) in view of (23a).

Lemma 2

Suppose

$$1^\circ. \|b(n)\|_0 = \overline{\lim}_{n \rightarrow \infty} |b(n)|, \overline{\lim} \equiv \lim \text{ es superior} \quad (24a)$$

$$2^\circ. b(k) = \sum_{m=0}^{k-1} H'(k, m) a(m) \quad (24b)$$

$$k=1, \dots$$

or shortly

$$b = H'a$$

3°.

$$\forall m \sum_{k=0}^{\infty} |H'(k, m)| < \infty, k > m \quad (24c)$$

Then the norm $\|H'\|_0$ of the operator H' in the normed space of sequences with $\|\cdot\|_0$ norm is

$$\|H'\|_0 = \overline{\lim}_{k \rightarrow \infty} \sum_{m=0}^{k-1} |H'(k, m)|. \quad (24d)$$

Proof

$$\begin{aligned} \|b\|_0 &= \left\| \sum_{m=0}^{Ent(k/2)} H'(k, m) a(m) + \sum_{m=Ent(k/2)+1}^{k-1} H'(k, m) a(m) \right\|_0 \leq \\ &\leq \overline{\lim}_{k \rightarrow \infty} \sum_{m=0}^{k-1} |H'(k, m)| \|a(m)\|_0, \quad Ent(n) \equiv Entier(n). \end{aligned}$$

On the other hand if

$$a^*(m) = \text{signum } H'(k, m) \text{ and } \|a^*\|_0 = 1 \text{ then}$$

$$b^*(k) = \sum_{m=0}^{k-1} |H'(k, m)| \text{ and } \|b^*\| = \overline{\lim}_{k \rightarrow \infty} b^*(k) = \overline{\lim}_{k \rightarrow \infty} \sum_{m=0}^{k-1} |H'(k, m)|.$$

That is why the preposition of lemma 2 holds true. The following theorem can be stated in view of lemma2;

Theorem 1

Suppose

1°. The sequence of variances evolves according to (22) while (21) is valid.

$$2^\circ. \gamma_i(m) = 0(1/m).$$

Then

$$\lim_{n \rightarrow \infty} \sigma_{\Delta D'_i}^2(n) = 0, \quad (25)$$

and

$$\lim_{n \rightarrow \infty} P^{1)} (|\Delta D'_i(n)| > \epsilon) = 0. \quad (26)$$

Proof²⁾

Let $H'(n, n-m) = H(n, m)^2 \gamma_i(m+1)$. Then a view of (21, 22) and lemma 2

$$\|\sigma_{\Delta D'_i}^2\|_0 \leq \|H'\|_0 \|\gamma_i(n) \sigma_{\varepsilon_i}^2(n-1)\|_0 = 0. \quad \square$$

The norm of H' is finite in view of (23c). Moreover the convergence in the mean squared sense implies the one in probability; (Chebychev inequality)

$$0 \leq \epsilon^2 P(|x - x_0| > \epsilon) = \int_{|x-x_0|} \epsilon^2 p(x) dx \leq \int_{-\infty}^{\infty} (x - x_0)^2 p(x) dx = \sigma_x^2. \quad ^3)$$

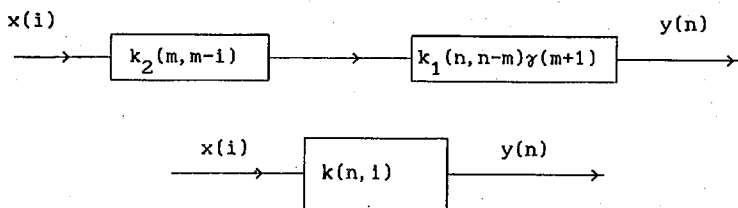


Fig. 1. Two systems connected in series being modelled by their individual unit pulse functions and the equivalent system.

Now, let find the equivalent unit pulse function of two systems of fig.1 connected in series while modelled by their unit pulse functions.

Lemma 3

The equivalent unit pulse function is:

$$k(n, n-m) = \gamma(m+1) \sum_{i=0}^m k_2(m, m-i) k_1(n, i). \quad (27)$$

¹⁾ P stands for the probability that...

²⁾ The expression of (24a) is the example of the norm in quotient space. Zero of this space with $\|\cdot\|_0$ norm is the equivalence class of all sequences convergent to zero in normal meaning. To complete the proof within the normed space concept it is enough to show that the sequence of variances is of zero norm.

³⁾ $p(x)$ stands for the probability density function.

Proof

In view of fig. 1:

$$\begin{aligned} y(n) &= \sum_{m=0}^n k_1(n, n-m) \gamma(m+1) \sum_{i=0}^m k_2(m, m-i) x(i) \underset{n-m=i}{=} \\ &= \sum_{k=0}^n \gamma(m+1) \sum_{i=0}^m k_2(m, m-i) k_1(n, i) x(n-m) = \sum_{m=0}^n k(n, n-m) x(n-m) \end{aligned}$$

and

$$k(n, m-m) = \gamma(m+1) \sum_{i=0}^m k_2(m, m-i) k_1(n, i).$$

Consider also following;

Lemma 4

Suppose

$$1^0. k_1(n, i) = 0(i/n) \quad (28a)$$

$$2^0. k_2(m, m-i) = 0(e^{-a(m-i)}) \quad (28b)$$

Then

$$\sum_{i=0}^m k_2(m, m-i) k_1(n, i) = 0(m/n). \quad (28c)$$

Proof

$$\begin{aligned} \sum_{i=0}^m k_2(m, m-i) k_1(n, i) &= 0 \left(\sum_{i=0}^m e^{-a(m-i)} i/n \right) \cong \\ &\cong 0 \left(\int_{i=1}^m e^{-a(m-i)} i/n \, di \right) \cong 0(m/n). \end{aligned}$$

Conclusion

The proposition of theorem 1 can be extended for coloured noise sequence $\{\varepsilon'_i(m)\}$:

$$\varepsilon'_i(m) = \sum_{k=0}^m k_2(m, m-k) \varepsilon''_i(k) \quad (29)$$

with $\{\varepsilon''_i(k)\}$ white and of finite variance and $k_2(m, m-k) \cong 0(e^{-a(m-k)})$. In view of lemma 4 (23c) still holds true with

$$\sum_{i=0}^m k_2(m, m-i) H(n, i) \cong 0(m/n)$$

instead of $H(n, m)$ of (20a).

5. ON THE OPTIMAL GAIN SEQUENCE

Suppose

1°. $\{\varepsilon'_i(n)\}$ of (19a) is white and stationary:

$$E\{\varepsilon'_i(n)^2\} = \sigma^2$$

2°. $E\{\Delta D'_i(n)^2\} = \sigma_i^2(n)$.

Then from (19a)

$$\sigma_i^2(n) = [1 - \lambda_i \gamma_i(n)]^2 \sigma_i^2(n-1) + \gamma_i^2(n) \sigma^2 \quad (30a)$$

$$n = 1, \dots, \quad i = 1, \dots, N.$$

The variance of the vector of the unknown coefficients $\Delta D'(n)$ evolves according to this equation. Now the optimal gain sequence can be obtained by the minimization of variance sequence. From (30a) by equating the first derivative to zero in respect to $\gamma_i(n)$ one gets

$$1/\gamma_i^{\text{opt}}(n) = \frac{\sigma^2}{\lambda_i \sigma_i^2(n-1)} + \lambda_i \quad (30b)$$

and

$$1/\sigma_i^2(n) = 1/\sigma_i^2(n-1) + a^2; \quad 0 < a \cong (\lambda_i/\sigma)^2. \quad (30c)$$

That is why

$$1/\sigma_i^2(n) \cong n a^2. \quad (30d)$$

As a matter of fact the optimal gain sequence changes like

$$\gamma^{\text{opt}}(n) \cong 1/n. \quad (30e)$$

This is just the case being considered.

Second method for the definition of optimal gain sequence can be based on the cost function $\hat{C}[D(n)]$. While the first method is of theoretical meaning rather in this case the second is the practical one. The knowledge of the state and the observation or reference vectors is required only instead of the density being estimated so far as the algorithm proceeds. The known conditions under which these two approaches are equivalent will be reminded next. The cost function of (3) is

$$\hat{C}[D(n)] = \frac{1}{n} \sum_{k=1}^n L[x(k), y(k), D(n-1) -$$

$$- \text{diag } \gamma(n) \nabla_D L[x(n), y(n), D(n-1)] \quad (31a)$$

The condition for the optimality is readily obtained from (31a) by gradient calculations:

$$\nabla_{\gamma} \hat{C}[D(n)] = -\frac{1}{n} \sum_{k=1}^n \nabla_D L[x(k), y(k), D(n)] \cdot \text{diag} \nabla_D L[x(n), y(n), D(n-1)] = 0. \quad (31c)$$

Significantly the minimum of the cost function is gained in respect to adaptation gain and coefficients' vector together. However, in view of (31) the optimal gain sequence can be calculated generally in approximate way only by the linearization of the cost function while

$$\sum_{k=1}^n \nabla_D L[x(k), y(k), D(n-1)] - \text{diag} \gamma(n) \nabla_D L[x(n), y(n), D(n-1)] = 0 \quad (32a)$$

$$n = 1, \dots$$

Namely, suppose

$$\|D(n) - D(n-1)\| \ll 1 \text{ and}$$

$$\sum_{k=1}^n \nabla_D L[x(k), y(k), D(n-1)] - \text{diag} \gamma(n) \nabla_D L[x(n), y(n), D(n-1)] = 0 \quad (32b)$$

Then by Taylor series expansion

$$\sum_{k=1}^n \nabla_D L[x(k), y(k), D(n-1)] - \sum_{k=1}^n \nabla_D^2 L[x(k), y(k), D(n-1)] \cdot$$

$$\cdot \text{diag} \gamma^{\text{opt}}(n) \nabla_D L[x(n), y(n), D(n-1)] \cong 0. \quad (32c)$$

(32a)

and

$$\nabla_D^2 L \stackrel{\text{df}}{=} \left[\frac{\partial^2 L(x, y, D)}{\partial D_i \partial D_j} \right]_{N \times N} \quad i, j = 1, \dots, N. \quad (32d)$$

While

$$K(n) \stackrel{\text{df}}{=} \sum_{k=1}^n \nabla_D^2 L[x(k), y(k), D(n-1)] \quad (32e)$$

(32c) is

$$K(n) \text{diag} \gamma^{\text{opt}}(n) \nabla_D L[x(n), y(n), D(n-1)] = \nabla_D L[x(n), y(n), D(n-1)] \quad (32f)$$

Eventually

$$\gamma^{\text{opt}}(n) \cong [\text{diag} \nabla_D L[x(n), y(n), D(n-1)]^{-1} \cdot K^{-1}(n) \nabla_D L[x(n), y(n), D(n-1)]] \quad (32g)$$

It is well known [5] that the minimization of variance within probabilistic approach to filtering problem (method I) and the minimization of square cost function (method II) are equivalent in linear case. As far as the optimal gain sequence is of concern the methods under consideration are equivalent also when n tends to infinity and linear model resulting from linearization is valid. Within simplifications in hand the optimal gain sequence

$$\gamma^{\text{opt}}(n) = 0(1/n) \quad (33)$$

asymptotically in both cases. However the problem of filter divergence arise in the filters with diminishing gain. The origin of divergence and improvements can be pointed out while the comparison with Kalman's filter is being done within probabilistic approach [5]. Roughly, to avoid divergence filter's gain should be kept constant after its small value has been achieved. However there are some differences between the tasks of filtration and identification and that is why this problem deserves further treatment [17].

6. THE EFFECTS OF NONSMOOTH CASE OF OPTIMIZATION PROBLEM

The effects of quantization can be taken into account by additional disturbance called quantization noise in any algorithm. It is a rule of thumb to make the effects of quantization negligible.

One can also try to approximate any nonsmooth problem by smooth one while the problem of convergence is to be answered. To be more specific let consider the information processing by neuron of fig. 2 in brief way.

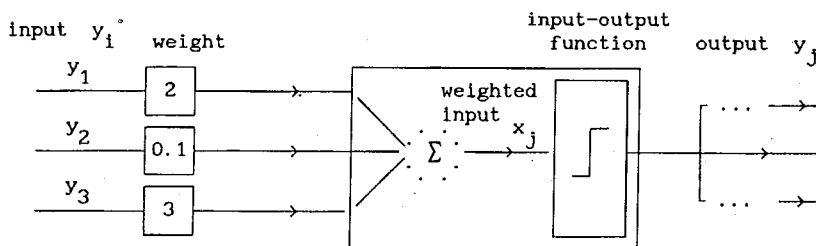


Fig. 2. Information processing in neuron.

The cost criterion is usually expressed as

$$C = \frac{1}{2} \sum_j (y_j - d_j)^2, \quad (34a)$$

where y_j of (34) is the excitation level of j element in upper layer of $j-s$ and d_j is the desired output from the layered structure of neural network. The weighted input x_j by w_{ij} of fig. 2 is

$$x_j = \sum_i y_i w_{ij}. \quad (34b)$$

The threshold calculation by the input-output function of fig. 2 is modelled by smooth sigmoidal function yielding the output y_j

$$y_j = \frac{1}{1 + e_j^{-x}}. \quad (34c)$$

According to backpropagation algorithm the gradient calculation in respect to the weights is calculated backward from the last output layer of neurons to the first, input one, according to

1. $\mathcal{E}A_j \stackrel{\text{def}}{=} \frac{\partial C}{\partial y_j} = y_j - d_j$
 2. $\mathcal{E}I_j \stackrel{\text{def}}{=} \frac{\partial C}{\partial x_j} = \frac{\partial C}{\partial y_j} \frac{dy_j}{dx_j} = \mathcal{E}A_j y_j (1 - y_j)$
 3. $\mathcal{E}W_{ij} \stackrel{\text{def}}{=} \frac{\partial C}{\partial w_{ij}} = \frac{\partial C}{\partial x_j} \frac{\partial x_j}{\partial w_{ij}} = \mathcal{E}I_j y_i$
 4. $\mathcal{E}A_i \stackrel{\text{def}}{=} \frac{\partial C}{\partial y_i} = \sum_j \frac{\partial C}{\partial x_j} \frac{\partial x_j}{\partial y_i} = \sum_j \mathcal{E}I_j w_{ij}.$
- (35)

The gradient is calculated at step 3 for j -layer and the calculations propagate backward while using the result of step 4 for upper layer. With gradient in hand one can update the neural network with an algorithm of (6) form. In view of (1) the probabilistic meaning of neural adaptation can be imposed either.

7. LINEAR ADAPTIVE FILTERING

Linear adaptive filtering usually invokes the programmable filter structure of fig. 3. and is remarkably well developed for Communication and Control applications [11, 27, 28, 29, 35].

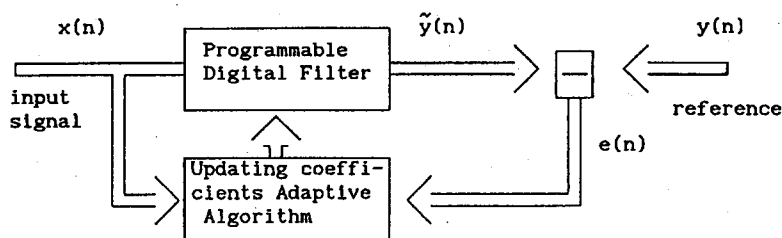


Fig. 3. The principle of operation of adaptive filter.

The error sequence $e(n)$ of fig. 3, produced by subtracting the output $\tilde{y}(n)$ of the digital filter from a reference waveform $y(n)$, is used together with the input sequence $x(n)$ to update the filter coefficients while following the optimization procedure. The adaptive filters are classified according to [11]

- 1) the optimization criterion
- 2) the algorithm for coefficients updating
- 3) the programmable filter structure
- 4) the type and dimensionality of signals processed.

In the case of vector $D(n)$ of FIR filter coefficients and vector $X(n)$ of input data:

$$D(n) \stackrel{\text{df}}{=} \text{col}\{D_0(n), \dots, D_{N-1}(n)\}$$

$$X(n) \stackrel{\text{df}}{=} \text{col}\{x(n), \dots, x(n-N+1)\}$$

the error $e(n)$ sequence is; $e(n) = y(n) - D^T(n)X(n)$.

The optimization criterion may be of mean square form of (1) with

$$L\{x, y(n), n, D(n)\} = [y(n) - D^T(n)X(n)]^2. \quad (36a)$$

The criterion

$$\hat{C}_0\{D(n)\} = \sum_{k=1}^n w^{n-k} \{y(k) - D^T(n)X(k)\}^2 \quad (36b)$$

can be viewed as the one of form (4) with overweighting the last data by the scalar $w \cong 1 : 0 \ll w \leq 1$. The calculation of the optimal gain sequence of paragraph V or direct optimization of (36b) in respect to $D(n)$ yields the normal, Yule–Walker equation

$$D(n) = R_N^{-1}(n)r_{yx}(n) \quad (37a)$$

with

$$R_N(n) = \sum_{k=1}^n w^{n-k} X(k)X(k)^T \quad (37b)$$

$$r_{yx}(n) = \sum_{k=1}^n w^{n-k} X(k)y(k). \quad (37c)$$

The criterion (36b) is often used as prior one without any statistical and probabilistic interpretation attached. It turns out that this form of adaptive filter is also optimal in probabilistic sense as the linear Kalman's filter is [5], resulting thus in procedure optimal at the each step of adaptation in deterministic and probabilistic meaning [34].

There are two main classes of the application of adaptive filtering; system correction and system identification (see fig. 4 and fig. 5).

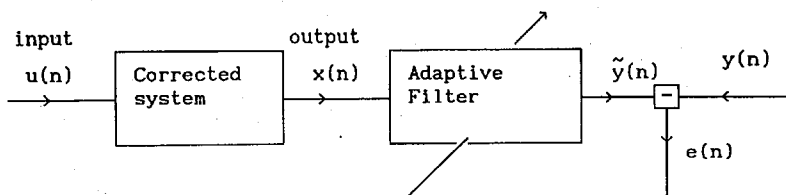


Fig. 4. The application of adaptive filter to system correction.

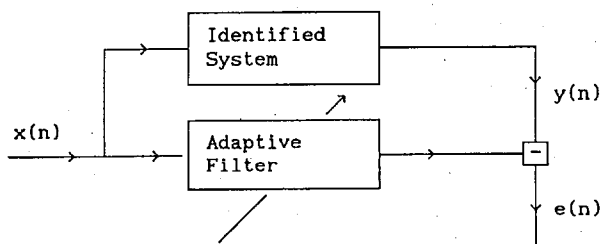


Fig. 5. The application of adaptive filter to system identification.

FINAL REMARKS

The problem of adaptation has been stated as the cost function optimization task. The question of convergence has been answered in circuit theory oriented way being familiar to electric and electronic engineer. The problem of the short term convergence, on the transient stage of adaptive device, while seeking for solution, has been of concern. The algorithms with constant and diminishing steps have been considered. The approach of this paper is aimed towards the refinement of recursive algorithms in view of their increasing importance for practical applications due to the recent developments in hardware. With the knowledge of different approaches and results for this kind of problems one can try to answer further ones of practical interest [17].

REFERENCES

1. H. Robbins, S. Monro: *Ann. Math. Statist.* 22, 400, 1951
2. J. Kiefer, J. Wolfowitz: *Ann. Math. Statist.* 23, 462, 1952
3. D. Sokrison: *Efficient recursive estimation; Application to estimating the parameters of a covariance function.* Int., J. Engng. Sci. vol. 3, pp. 461—483, Pergamon Press 1965
4. P. Eykhoff: *Identyfikacja w układach dynamicznych.* W—wa 1980
5. A.H. Jazwiński: *Stochastic processes and filtering theory.* Academic Press 1970
6. J.Z. Cypkin: *Podstawy teorii układów uczących się.* WNT, W—wa 1973
7. H.J. Kushner: *Convergence of Recursive Adaptive and Identification procedures via weak convergence theory.* IEEE Trans. Automatic Control, vol. AC—22, no 6, December 1977
8. H.J. Kushner, Huang Hai: *Rates of convergence for stochastic approximation type algorithms.* SIAM J. Control and Optimization, vol. 17, no 5, September 1979
9. L. Ljung: *Analysis of Recursive Stochastic Algorithms.* IEEE Trans. Automatic Control, vol. AC-22, no 4, August 1977
10. A. Benveniste, M. Goursat, G. Ruget: *Analysis of Stochastic Approximation Schemes with Discontinuous and Dependent Forcing Terms with Application to Data Communication Algorithms.* IEEE Trans. on Automatic Control, vol. AC—25, no 6, December 1980
11. M. Bellanger: *Adaptive digital filtering.* ECCTD 1987, Paris
12. E.I. Jury: *Przekształcenie Z i jego zastosowania.* WNT, Warszawa 1970
13. D. Falconer, L. Ljung: *Application of Fast kalman Estimation to Adaptive Equalization.* IEEE Trans., COM—26, no 10, October 1978, pp. 1439—1446
14. G. Carayannis, D. Manolakis, N. Kalauptsidis: *A Fast Sequential Algorithm for L.S. Filtering and Prediction.* IEEE Trans., vol. ASSP—31, no 6, Dec. 1983, pp. 1394—1402
15. S.M. Kay, J.R. Marple: *Spectrum analysis — A modern perspective.* Proceedings of the IEEE, vol.69, no 11, November 1981, pp. 1395—1396
16. P. De Larminat, Y. Thomas: *Automatyka — układy liniowe.* PWN, Warszawa 1983, (Polish translation from french)
17. M. Żółtowski: *On the problem of divergence in adaptive filtering (in progress)*
18. A.V. Oppenheim, R.W. Schaffer: *Cyfrowe przetwarzanie sygnałów.* WKiŁ, Warszawa 1975, (Polish translation from english)
19. F.C. Schweppe: *Układy dynamiczne w warunkach losowych.* WNT, Warszawa 1978, (Polish translation from english)
20. Gauss: *Theoria Motus Corporum 'Caelestium.* 1809
21. N. Wiener: *Extrapolation and smoothing of stationary time series with engineering applications.* New York, Technology Press and Wiley, 1949
22. R.E. Kalman: *A new approach to linear filtering and prediction problems.* Journal of Basic Engineering, 1960, vol. 82, pp. 34—35
23. N. Levinson: *The Wiener RMS (Root Mean Square) error criterion in filter design and prediction.* Journal of Mathematics and Physics, 1947, vol. 25, pp. 261—278
24. T. Kailath: *A view on three decades of linear filtering theory.* IEEE Translation on Information Theory, 1974, vol. IT—20, pp. 145—181
25. P. Eykhoff: *Identyfikacja w układach dynamicznych.* PWN Warszawa 1980
26. J.G. Proakis: *Digital Communications.* New York, Mc Graw Hill, Inc., 1983
27. B. Widrow, S.D. Stearns: *Adaptive Digital Processing.* Prentice Hall, Englewood Cliffs, New York, USA, 1985
28. M. Bellanger: *Adaptive Digital Filtering and Signal Analysis.* Marcel Dekker Inc., New York, 1987
29. S. Haykin: *Adaptive Filtering Theory.* Prentice—Hall, Englewood Cliffs, New York, 1986
30. L. Ljung, T. Soderstrom: *Theory and practice of recursive identification.* MIT Press, Cambridge 1983.
31. L. Ljung, M. Morf, D. Falconer: *Fast calculation of gain matrices for recursive estimation schemes.* Int. J. of Control, 1978, vol.27, no 1, pp. 1—19

32. G. Carayannis, D. Manolakis, N. Kaloupsidis: *Fast Kalman-type algorithms for sequential signal processing*. ICASSP 83, Boston, USA
33. J. Cioffi, T. Kailath: *Fast recursive least Squares Filters for Adaptive Filtering*. IEEE Trans., vol. ASSP—32, no 2, April 1984, pp. 304—337
34. M. Żółtowski: On linear adaptive filtering with constant step of adaptation. Kwart. Elektr. i Telekom. 1993, 38, z. 3
35. F. Ling, D. Manolakis and J. Proakis: *Numerically Robust L.S. Lattice Ladder Algorithms with Direct Updating of the Coefficients*. IEEE Trans., vol. ASSP—34, no 4, August 1986, pp. 837—845
36. M.L. Honig, D.G. Messerschmitt: *Adaptive Filters — Structures, Algorithms and Applications*. Kluwer Academic Publishers, Boston, USA, 1984
37. J. Kornacki: *Aproksymacja stochastyczna: metody optymalizacji w warunkach losowych*. Warszawa, WNT, 1989
38. G.E. Hinton: *Jak sieci neuropodobne uczą się na podstawie doświadczeń*. Scientific American, wydanie polskie, listopad 1992, nr 11
39. S.W. Fomin, I.P. Kornfeld, J. Sinaj: *Teoria ergodyczna*. Warszawa 1987
40. W. Szlenk: *Wstęp do teorii gładkich układów dynamicznych*. PWN, Warszawa 1982

M. ŻÓŁTOWSKI

KRÓTKOCZASOWA ZBIEŻNOŚĆ ALGORYTMÓW ADAPTACYJNEJ FILTRACJI W TEORIOOBWODOWOPODOBNYM UJĘCIU

Streszczenie

Przedstawiono rozwiązanie problemu adaptacji jako zadanie optymalizacji, które jest wykonywane przy pomocy numerycznego algorytmu. Adaptacyjna filtracja realizowana przez procesor cyfrowy posiada związki z różnymi powiązаныmi ze sobą dziedzinami takimi jak: 1) Rekursyjna metoda najmniejszych kwadratów, 2) Teoria filtracji, 3) Aproksymacja stochastyczna, 4) Identyfikacja adaptacyjna z modelem—odniesieniem [1—39]. Chociaż wielu świetnych autorów poświęciło uwagę tej problematyce to jest ona nadal interesująca, gdyż wkracza w obszar licznych zastosowań praktycznych, wymagających dalszych przemyśleń koncepcyjnych i szczegółowych. Dlatego znajomość związków pomiędzy różnymi metodami wydobywania informacji w postaci zbioru parametrów odnośnie badanego układu może pogłębić rozumienie algorytmów, dać wgląd w ich ulepszanie i modyfikację. Taka idea przyświeca tej pracy, która ma charakter przeglądowy. Dowód twierdzenia dotyczącego krótkoczasowej zbieżności (zbieżności prowadzącej do stanu ustalonego) algorytmu identyfikacji typu aproksymacji stochastycznej jest pomysłem autora i wykorzystywane są w nim pojęcia zaczerpnięte z teorii obwodów takie jak odpowiedź impulsowa, czy również pojęcie przestrzeni unormowanej, chociaż podobne ujęcia można zauważyć u Haykina [29], lub Bellangera [11] w przypadku liniowym i stałego kroku adaptacji. Zaproponowano rozszerzenie rozważań w ramach jednolitego, spójnego rozwiązania, gdy pojawiają się nieciągłości, w oparciu o funkcję sigmoidalną, którą z powodzeniem stosuje się do opisu sieci neuro-nowych [37]. Na końcu scharakteryzowano liniową adaptacyjną filtrację.

681.372.4:621.3.012.1

D-równoważność struktur układów realizujących funkcje wymierne

ZYGMUNT WRÓBEL

*Zakład Elektroniki i Automatyki, Instytut Problemów Techniki
Uniwersytet Śląski*

Otrzymano 1993.03.15

Autoryzowano do druku 1993.09.01

D-równoważność struktur układów to ich równoważność w sensie przekształcenia izomorficznego i 2-izomorficznego. W pracy przedstawiono algebraiczne kryteria oraz algorytm D-równoważności struktur układów realizujących funkcje wymierne.

1. WPROWADZENIE

Dla klasy układów zawierających skupione, liniowe i stacjonarne elementy (SLS) [4, 10] dowolna funkcja układowa taka jak admitancja, transmitancja, impedancja jest funkcją wymierną parametru sterującego. Jak wiadomo, uzyskiwane w procesie syntezy układów rozwiązanie jest w większości przypadków niejednoznaczne [1, 2, 7, 8]. Może bowiem istnieć wiele układów, o różnych strukturach topologicznych i różnych wartościach tworzących go elementów, które spełniają postawione zadanie syntezy, czyli są to układy równoważne względem przyjętych kryteriów.

W procesie syntezy topologicznej [14, 15] interesuje nas tylko rodzaj i sposób połączenia poszczególnych rodzajów elementów zakładając, że kiedy każdy element układu ma wartość równą jeden. W procesie syntezy topologicznej generowane są struktury układów, które mogą być między sobą równoważne. Istnieje więc konieczność podziału wygenerowanych struktur układów na klasy struktur równoważnych względem przyjętych kryteriów równoważności. W ogólnym przypadku równoważność struktur układów można sformułować następująco: „Dwie struktury układów

G_I^v i G_{II}^v są równoważne (względem zadanych kryteriów), jeśli istnieje funkcja równowartościowa, która odwzorowuje strukturę akładu G_I^v w strukturę układu G_{II}^v , czyli $f: G_I^v \Rightarrow G_{II}^v$ i odwrotnie $f^{-1}: G_{II}^v \Rightarrow G_I^v$.

Jak wiadomo [10] układ SLS można odwzorować w postaci grafu, którego struktura jest konsekwencją określonej konfiguracji połączeń elementów układu. W ogólnym przypadku każdemu układowi SLS można przyporządkować graf skierowany odwzorujący jego strukturę. Dwie struktury układów SLS odwzorowane w postaci grafów skierowanych są w tej samej klasie równoważności, wtedy i tylko wtedy, gdy pozostają ze sobą w relacji izomorfizmu. W niniejszej pracy tym typem równoważności układów nie będziemy się zajmować.

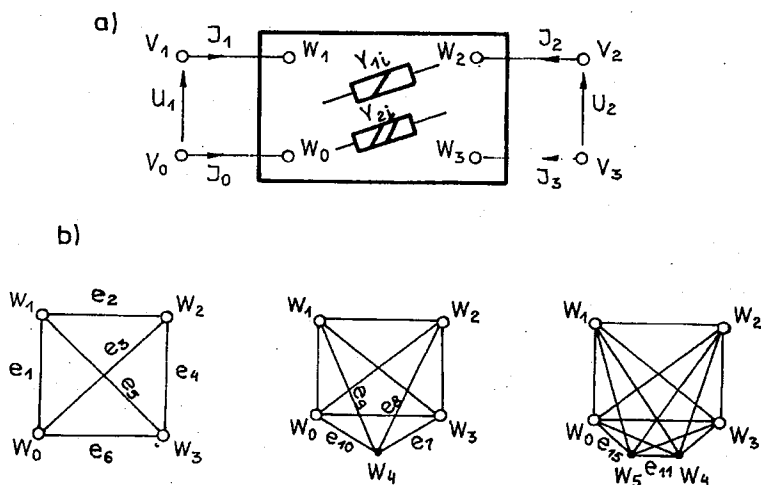
Ważną podgrupę układów zawierających elementy SLS stanowią układy, które zawierają tylko elementy odwracalne i pasywne SLSOP. Układowi SLSOP można przyporządkować graf nieskierowany odwzorujący ich strukturę. Jak pokazano w pracach [12–14] „Dwie struktury układów SLSOP odwzorowane w postaci grafów skierowanych są w tej samej klasie równoważności, wtedy i tylko wtedy gdy pozostają ze sobą w relacji izomorfizmu lub 2-izomorfizmu”. Ten typ równoważności będziemy nazywać D-równoważnością rozpatrywanych struktur układów. Zagadnienie izomorfizmu grafów (odwzorujących struktury układów SLSOP) znalazło swoje rozwiązanie w postaci szerokiej palety algorytmów pozwalających stwierdzić równoważność grafów w sensie przekształcenia izomorficznego [4, 10]. Dotychczas nie opracowano algorytmu pozwalającego stwierdzić 2-izomorfizm struktur układów [12–14].

Niniejsza praca jest kontynuacją prac [12, 13]. Przedstawiono w niej nowe spojrzenie na zagadnienie D-równoważności struktur układów SLSOP. Rozważania dotyczące D-równoważności struktur układów w niniejszej pracy zilustrowano na przykładzie układów czwórników składających się z dwóch rodzajów SLSOP elementów.

2. ELEMENTY ANALIZY TOPOLOGICZNEJ CZWÓRNIKÓW

Można wyróżnić w zasadzie trzy grupy układów elektrycznych składających się z dwóch rodzajów elementów [4]. Są to: układy RC , RL ; układy składające się z elementów typu R , G i elementów sterowanych typu $R \cdot q$, $G \cdot q$, układy zawierające pojemność C i przełączaną pojemność SC . W pracy [4] wykazano, że istnieje wzajemna transformacja wyżej wymienionych typów układów. Jeden rodzaj elementów w tych układach będziemy nazywali elementami typu $\{t_1\}$ i charakteryzowali admitancją $\{Y_{1i} \cdot t_1\}$, natomiast drugi rodzaj elementów będziemy nazywali elementami typu $\{t_2\}$ i charakteryzowali admitancją $\{Y_{2j} \cdot t_2\}$. Należy zwrócić uwagę, że parametry $\{t_1\}$ i $\{t_2\}$ są wielkościami bezwymiarowymi i służą tylko do wyodrębnienia danego typu elementów w rozpatrywanych układach.

Schemat blokowy układu czwórnika (czterokońcówkowego) pokazano na rysunku 1. Na rysunku tym zaproponowano również sposób tworzenia pełnego grafu



Rys. 1. Czwórnik; a – schemat blokowy, b – numeracja gałęzi w jego unigrafie

(tzw: unigrafu [4]) odwzorującego układ czwórnik w funkcji liczby jego wierzchołków. Taki „rozszerzający się” w zależności od liczby, wierzchołków unigraf czwórnik ma tę zaletę, że gałąź o danym numerze jest zawsze incydentna do tej samej pary wierzchołków układu, co ma istotne znaczenie w procesie ich syntezy topologicznej. Dla czwórnik (rys. 1.) przyjmując za punkt odniesienia węzeł w_0 , czyli $V_0 = 0$ równania określające zależności między prądami a napięciami [1, 3] mają postać:

$$V_{1,0} = \frac{1}{\Delta} \cdot (\Delta_{11} \cdot I_1 + \Delta_{21} \cdot I_2 - \Delta_{31} \cdot I_3) \quad (1)$$

$$V_{2,0} = \frac{1}{\Delta} \cdot (\Delta_{12} \cdot I_1 + \Delta_{22} \cdot I_2 - \Delta_{32} \cdot I_3) \quad (2)$$

$$V_{3,0} = \frac{1}{\Delta} \cdot (\Delta_{13} \cdot I_1 + \Delta_{23} \cdot I_2 - \Delta_{33} \cdot I_3) \quad (3)$$

Ponieważ

$$U_1 = V_{1,0} \quad (4)$$

$$U_2 = V_{2,0} - V_{3,0} \quad (5)$$

to równanie impedancyjne czwórnik można zapisać w postaci:

$$\begin{vmatrix} U_1 \\ U_2 \end{vmatrix} = \frac{1}{\Delta} \cdot \begin{vmatrix} \Delta_{11} & \Delta_{21} - \Delta_{31} \\ \Delta_{12} - \Delta_{13} & \Delta_{22} + \Delta_{33} - \Delta_{32} - \Delta_{23} \end{vmatrix} \begin{vmatrix} I_1 \\ I_2 \end{vmatrix}. \quad (6)$$

Ponieważ macierz impedencji jest symetryczna ($\Delta_{ij} = \Delta_{ji}$), dlatego można ją zapisać w postaci:

$$Z(t_1, t_2) = \frac{1}{\Delta} \cdot \begin{vmatrix} \Delta_{11} & \Delta_{12} - \Delta_{13} \\ \Delta_{12} - \Delta_{13} & \Delta_{22} + \Delta_{33} - 2\Delta_{23} \end{vmatrix}. \quad (7)$$

Macierz admitancji zwarciovych czwórnik otrzymuje się przez inwersję i ma ona postać:

$$Y(t_1, t_2) = \frac{1}{\Delta_{11,12} + \Delta_{11,33} - 2\Delta_{11,23}} \cdot \begin{vmatrix} \Delta_{22} + \Delta_{33} - 2\Delta_{23} & \Delta_{13} - \Delta_{12} \\ \Delta_{13} - \Delta_{12} & \Delta_{11} \end{vmatrix}. \quad (8)$$

Aby wyrazić elementy macierzy $Z(t_1, t_2)$ oraz $Y(t_1, t_2)$ układu czwórnik poprzez odpowiednie iloczyny admitancji $\{Y_{ij} \cdot t_i\}$ elementów typu $\{t_1\}$ i admitancji $\{Y_{ij} \cdot t_j\}$ elementów typu $\{t_2\}$ skorzystamy z rezultatów otrzymanych w pracy [4, 5].

Celem skrócenia zapisu wprowadzimy następujące oznaczenia:

$$H_1^v(t_1, t_2) = \sum_{\substack{\text{wszystkie} \\ \text{dendryty} \\ \text{1-drzewowe}}} \left\{ \begin{array}{l} \text{iloczyn admitancji elementów} \\ \text{tworzących dendryt 1-drzewowy} \\ T_1^v(t_1, t_2) \end{array} \right\} \quad (9)$$

$$H_{ij,kl}^v(t_1, t_2) = \sum_{\substack{\text{wszystkie} \\ \text{dendryty} \\ \text{2-drzewowe}}} \left\{ \begin{array}{l} \text{iloczyn admitancji elementów} \\ \text{tworzących dendryt 2-drzewowy} \\ T_{ij,kl}^v(t_1, t_2) \end{array} \right\} \quad (10)$$

$$H_{ijk,0}^v(t_1, t_2) = \sum_{\substack{\text{wszystkie} \\ \text{dendryty} \\ \text{3-drzewowe}}} \left\{ \begin{array}{l} \text{iloczyn admitancji elementów} \\ \text{tworzących dendryt 3-drzewowy} \\ T_{ijk,0}^v(t_1, t_2) \end{array} \right\} \quad (11)$$

Poszczególne elementy macierzy $Z(t_1, t_2)$ można zapisać w postaci:

$$\Delta_{11} = H_{1,0}^v(t_1, t_2) \quad (12)$$

$$\Delta_{12} - \Delta_{13} = H_{12,0}^v(t_1, t_2) - H_{13,0}^v(t_1, t_2) \quad (13)$$

$$\Delta_{22} + \Delta_{33} - 2 \cdot \Delta_{23} = H_{2,0}^v(t_1, t_2) + H_{3,0}^v(t_1, t_2) + 2 \cdot H_{23,0}^v(t_1, t_2). \quad (14)$$

Aby wyrazić otrzymane zależności poprzez dendryty 2-drzewowe związane z wierzchołkami wejściowymi $\{w_0, w_1\}$ i wyjściowymi $\{w_2, w_3\}$ układu czwórnik, należy zwrócić uwagę, że każdy dendryt 2-drzewowy $T_{ij,0}^v(t_1, t_2)$ układu posiadającego wierzchołek w_n jest albo dendrytem 2-drzewowym $T_{ijn,0}^v(t_1, t_2)$ lub dendrytem 2-drzewowym $T_{ij,n0}^v(t_1, t_2)$. Stąd dla układu czwórnik parametry $H_{12,0}^v(t_1, t_2)$, $H_{2,0}^v(t_1, t_2)$, $H_{3,0}^v(t_1, t_2)$ można zapisać w postaci:

$$H_{12,0}^v(t_1, t_2) = H_{12,30}^v(t_1, t_2) + H_{123,0}^v(t_1, t_2) \quad (15)$$

$$H_{2,0}^v(t_1, t_2) = H_{23,0}^v(t_1, t_2) + H_{3,20}^v(t_1, t_2) \quad (16)$$

$$H_{3,0}^v(t_1, t_2) = H_{32,0}^v(t_1, t_2) + H_{3,20}^v(t_1, t_2) \quad (17)$$

Korzystając z powyższych zależności równania (13) i (14) można zapisać w postaci:

$$\Delta_{12} - \Delta_{13} = H_{12,30}^v(t_1, t_2) + H_{123,0}^v(t_1, t_2) - H_{13,20}^v(t_1, t_1) - H_{123,0}^v(t_1, t_2) = H_{12,30}^v(t_1, t_1) - H_{13,20}^v(t_1, t_2) \quad (18)$$

$$\begin{aligned} \Delta_{22} - \Delta_{33} - 2 \cdot \Delta_{23} &= H_{2,0}^v(t_1, t_2) + H_{3,0}^v(t_1, t_2) - 2 \cdot H_{23,0}^v(t_1, t_2) = \\ &= H_{23,0}^v(t_1, t_2) + H_{2,30}^v(t_1, t_2) + H_{32,0}^v(t_1, t_2) + H_{3,20}^v(t_1, t_2) - \\ &- 2 \cdot H_{23,0}^v(t_1, t_2) = H_{2,30}^v(t_1, t_2) - H_{3,20}^v(t_1, t_2) = H_{2,3}^v(t_1, t_2) \end{aligned} \quad (19)$$

Tak więc macierz impedencji rozwarciowych układu czwórnika można zapisać w postaci:

$$Z(t_1, t_2) = \frac{1}{H(t_1, t_2)} \cdot \begin{vmatrix} H_{1,0}^v(t_1, t_2) & H_{12,3}^v(t_1, t_2) - H_{13,20}^v(t_1, t_1) \\ H_{12,30}^v(t_1, t_2) - H_{13,20}^v(t_1, t_2) & H_{2,3}^v(t_1, t_2) \end{vmatrix} \quad (20)$$

Aby uzyskać analogiczną formę zapisu dla macierzy admitancji zwarciovych $Y(t_1, t_2)$ (wzór 8) należy dopełnienia algebnaiczne 2-ego rzędu $\Delta_{i,j,k}$ wyrazić przez sumę iloczynów admitancji elementów wszystkich dendrytów 3-drzewowych $H_{i,j,k,0}^v(t_1, t_2)$ co można zapisać w postaci:

$$\Delta_{i,j,k} = H_{i,j,k,0}^v(t_1, t_2) \quad (21)$$

Z definicji dendrytu 3-drzewowego [5, 4] wynika, że dendryt 2-drzewowy $T_{jk,0}^v(t_1, t_2)$ po zwarcu jego w_i wierzchołka z wierzchołkiem w_0 (i wykreśleniu pętli powstałych po tym zwarcu) tworzy dendryt 3-drzewowy $T_{i,j,k,0}^v(t_1, t_2)$. Korzystając z powyższych rozważań można zapisać:

$$\Delta_{11,22} = H_{1,2,0}^v(t_1, t_2) = H_{13,2,0}^v(t_1, t_2) + H_{1,23,0}^v(t_1, t_2) + H_{1,2,30}^v(t_1, t_2) \quad (22)$$

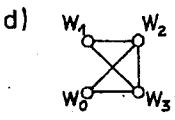
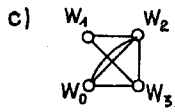
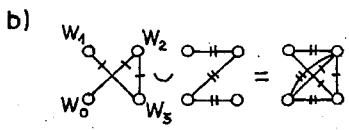
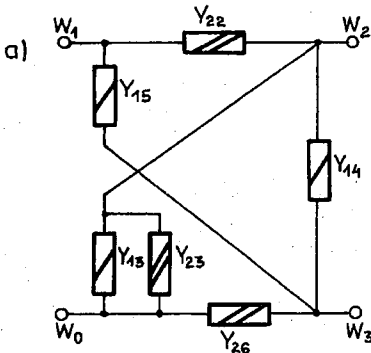
$$\Delta_{11,33} = H_{1,3,0}^v(t_1, t_2) = H_{12,3,0}^v(t_1, t_2) + H_{1,23,0}^v(t_1, t_2) + H_{1,3,20}^v(t_1, t_2) \quad (23)$$

Stąd sumę dopełnień algebraicznych 2-ego rzędu występującą w zależności (1.8) można zapisać:

$$\begin{aligned} \Delta_{11,22} + \Delta_{11,33} + 2 \cdot \Delta_{11,23} &= H_{13,2,0}^v(t_1, t_2) + H_{1,23,0}^v(t_1, t_2) + \\ &+ H_{1,2,30}^v(t_1, t_2) + H_{12,3,0}^v(t_1, t_2) + H_{1,23,0}^v(t_1, t_2) + \\ &+ H_{1,3,30}^v(t_1, t_2) - 2 \cdot H_{1,23,0}^v(t_1, t_2) = H_{13,2,0}^v(t_1, t_2) + \\ &+ H_{1,2,30}^v(t_1, t_2) + H_{12,3,0}^v(t_1, t_2) + H_{1,3,20}^v(t_1, t_2) \end{aligned} \quad (24)$$

Tak więc macierz admitancji zwarciovych układu czwórnika, można zapisać w postaci:

$$Y(t_1, t_2) = \frac{1}{H_{13,2,0}^v(t_1, t_2) + H_{1,2,30}^v(t_1, t_2) + H_{12,3,0}^v(t_1, t_2) + H_{1,3,20}^v(t_1, t_2)} \cdot \begin{vmatrix} H_{2,3}^v(t_1, t_2) & H_{13,20}^v(t_1, t_2) - H_{12,30}^v(t_1, t_2) \\ H_{13,20}^v(t_1, t_2) - H_{12,30}^v(t_1, t_2) & H_{1,0}^v(t_1, t_2) \end{vmatrix} \quad (25)$$



struktura	$e_6 e_5 e_4 e_3 e_2 e_1$	LT_K^4
$T_K^4(t_1, t_2)$	$0 t_1 t_1 t_1 0 0$ $t_2 0 0 t_2 t_2 0$	
A	B	C
$T_1^4(t_1^3, t_2^0)$	$0 t_1 t_1 t_1 0 0$	1
$T_1^4(t_1^2, t_2^1)$	$0 0 t_1 t_1 t_2 0$ $0 t_1 0 t_1 t_2 0$ $0 t_1 t_1 t_2 0 0$ $t_2 t_1 0 t_1 0 0$ $t_2 t_1 t_1 0 0 0$	5

A	B	C
$T_1^4(t_1^1, t_2^2)$	$0 t_1 0 t_2 t_2 0$ $t_2 0 0 t_1 t_2 0$ $t_2 0 t_1 0 t_2 0$ $t_2 t_1 0 0 t_2 0$ $t_2 0 t_1 t_2 0 0$ $t_2 t_1 0 t_2 0 0$	6
$T_1^4(t_1^0, t_2^3)$	$t_2 0 0 t_2 t_2 0$	1
$T_{1,0}^4(t_1^2, t_2^0)$	$0 0 t_1 t_1 0 0$ $0 t_1 0 t_1 0 0$ $0 t_1 t_1 0 0 0$	3
$T_{1,0}^4(t_1^1, t_2^1)$	$0 0 t_1 0 t_2 0$ $0 t_1 0 0 t_2 0$ $0 0 t_1 t_2 0 0$ $0 t_1 0 t_2 0 0$ $t_2 0 0 t_1 0 0$ $t_2 0 t_1 0 0 0$	6
$T_{1,0}^4(t_1^0, t_2^2)$	$t_2 0 0 0 t_2 0$ $t_2 0 0 t_2 0 0$	2
$T_{2,3}^4(t_1^2, t_2^0)$	$0 t_1 0 t_1 0 0$	1
$T_{2,3}^4(t_1^0, t_2^1)$	$0 0 0 t_1 t_2 0$ $0 t_1 0 t_2 0 0$ $t_2 t_1 0 0 0 0$	3
$T_{2,3}^4(t_1^0, t_2^2)$	$0 0 0 t_2 t_2 0$ $t_2 0 0 0 t_2 0$	2
$T_{12,30}^4(t_1^2, t_2^0)$		0
$T_{12,30}^4(t_1^0, t_2^1)$		0
$T_{12,30}^4(t_1^0, t_2^2)$	$t_2 0 0 0 t_2 0$	1
$T_{13,20}^4(t_1^2, t_2^0)$	$0 t_1 0 t_1 0 0$	1
$T_{13,20}^4(t_1^1, t_2^1)$	$0 t_1 0 t_2 0 0$	1
$T_{13,20}^4(t_1^0, t_2^2)$		0
$T_{14,2,30}^4(t_1^1, t_2^0)$		0
$T_{14,2,30}^4(t_1^0, t_2^1)$	$t_2 0 0 0 0 0$	1
$T_{14,2,30}^4(t_1^1, t_2^2)$		0
$T_{12,3,0}^4(t_1^0, t_2^1)$	$0 0 0 0 t_2 0$	1
$T_{14,3,20}^4(t_1^1, t_2^0)$	$0 0 0 t_1 0 0$	1
$T_{14,3,20}^4(t_1^0, t_2^1)$	$0 0 0 t_2 0 0$	1
$T_{13,2,0}^4(t_1^1, t_2^0)$	$0 t_1 0 0 0 0$	1
$T_{13,2,0}^4(t_1^0, t_2^1)$		0

Rys. 2. Czworńmiki składające się z dwóch rodzajów elementów;
a – schemat; b – bi-graf; c – multigraf; d – unigraf;
tabela – dendryty k-drzewowe konieczne do określenia jego parametrów

Korzystając z powyższych rozważań w tablicy 1 zestawiono pozostałe parametry $a_{ij}(t_1, t_2)$, $b_{ij}(t_1, t_2)$, $h_{ij}(t_1, t_2)$, $g_{ij}(t_1, t_2)$ charakteryzujące własności czwórników.

Z otrzymanych zależności wynika, że poszczególne parametry $x_{ij}(t_1, t_2)$ [$x_{ij}(t_1, t_2) \in z_{ij}(t_1, t_2) \vee y_{ij}(t_1, t_2) \vee a_{ij}(t_1, t_2) \vee b_{ij}(t_1, t_2) \vee h_{ij}(t_1, t_2) \vee g_{ij}(t_1, t_2)$] charakteryzujące własności czwórników są wyrażone tylko poprzez odpowiednie sumy iloczynów admitancji elementów tworzących dendryty k-drzewowe ($k=1 \vee 2 \vee 3$) $H_1^v(t_1, t_2)$, $H_{ij,kl}^v(t_1, t_2)$, $H_{i,jk,0}^v(t_1, t_2)$.

Na rysunku 2 pokazano przykład czwórnika składającego się z dwóch rodzajów elementów. W tabeli na tym rysunku zestawiono wszystkie dendryty 1- 2- i 3-drzewowe, składające się z elementów typu $\{t_1\}$ i $\{t_2\}$ konieczne do wyznaczenia parametrów rozpatrywanego układu. Istnienie w dendrycie k-drzewowym elementu typu $\{t_1\}$ zaznaczono literą t_1 , natomiast występowanie elementu typu $\{t_2\}$ – literą t_2 .

Jak wynika z powyższych rozważań oraz prac [12] dany parametr charakteryzujący układ czwórnika składającego się z dwóch rodzajów elementów można w ogólnym przypadku zapisać w postaci funkcji wymiernej typu:

$$X_{ij}(t_1, t_2) = \frac{H_{ka}^v(t_1, t_2)}{H_{kb}^v(t_1, t_2)} \quad (26)$$

gdzie:

$H_{ka}^v(t_1, t_2)$; $H_{kb}^v(t_1, t_2)$ – odpowiednio sumy iloczynów admitancji elementów typu $\{t_1\}$ i $\{t_2\}$ tworzących dendryty ka -drzewowe w liczniku i dendryty kb -drzewowe w mianowniku funkcji wymiernej opisującej parametry czwórnika.

W ogólnym przypadku $H_k^v(t_1, t_2)$ można zapisać w postaci:

$$H_k^v(t_1, t_2) = \sum_{n_{1k}=v-k}^0 \left\{ \sum_{T_k^v(t_1^{n_{1k}}, t_2^{n_{2k}})} \left[\prod_{i=1}^{n_{1k}} (Y_{1i} t_1) \cdot \prod_{j=1}^{n_{2k}} (Y_{2j} t_2) \right] \right\} \quad (27)$$

gdzie:

n_{1k} – liczba elementów typu $\{t_1\}$ występujących w dendrycie k-drzewowym.

$n_{2k} = v-k-n_{1k}$ – liczba elementów typu $\{t_2\}$ występujących w dendrycie k-drzewowym.

$T_k^v(t_1^{n_{1k}}, t_2^{n_{2k}})$ – dendryt k-drzewowy składający się z n_{1k} ($n_{1k}=v-k, v-k-1, \dots, 1, 0$) elementów typu $\{t_1\}$ oraz n_{2k} ($n_{2k}=v-k-n_{1k}$) elementów typu $\{t_2\}$.

Parametr $H_k^v(t_1, t_2)$ można również przedstawić w postaci:

$$H_k^v(t_1, t_2) = \sum_{n_{1k}=v-k}^0 A_{n_{1k}, n_{2k}} \cdot t_1^{n_{1k}} \cdot t_2^{n_{2k}} \quad (28)$$

gdzie:

$$A_{n_{1k}, n_{2k}} = \sum_{T_k^v(t_1^{n_{1k}}, t_2^{n_{2k}})} \left[\prod_{i=1}^{n_{1k}} (Y_{1i}) \cdot \prod_{j=1}^{n_{2k}} (Y_{2j}) \right] \quad (29)$$

Należy zwrócić uwagę, że występujące we wzorze 1.28 wyrazy $t_i^{n_i}$ są wielkościami bezwymiarowymi i sygnalizują nam że w danym iloczynie admitancji występuje n_{ik} elementów rodzaju $\{t_i\}$. Wynika stąd, że licznik i mianownik funkcji wymiernej

opisującej parametry $X_{ij}(t_1, t_2)$, to wielomiany zmiennej t_1 (lub t_2). dla układów RC, RL, to wielomian zmiennej s , dla układów składających się elementów R , G i elementów sterowanych typu $R \cdot q$, $G \cdot q$, to wielomian zmiennej q , dla układów zawierających pojemność C i przełączaną pojemność SC , to wielomian zmiennej z .

W procesie syntezy topologicznej interesuje nas tylko rodzaj i sposób połączenia poszczególnych rodzajów elementów (zakłada się, że każdy element układu ma wartość równą jeden). Sposób połączenia poszczególnych elementów określa strukturę topologiczną układu.

Można wyróżnić trzy podstawowe sposoby prezentacji danego obwodu w postaci grafu [4]. Są to: unigraf, multigraf i bi-graf. Na rysunku 2 przedstawiono modele czwórnik składającego się z dwóch rodzajów elementów $\{t_1\}$ i $\{t_2\}$ w postaci bi-grafu (rysunek 2.c), multigrafu (rysunek 2.d) i unigrafu (rysunek 2.e).

Struktury układów będziemy zapisywać w postaci ich macierzy.

Macierz unigrafu odwzorującego układ ma postać:

$$G^v(t) = [g_{ij}]_{1 \times e} \quad (30)$$

gdzie:

$$g_{ij} = \begin{cases} 1; & \text{jeśli unigraf zawiera } j\text{-tą gałąź.} \\ 0; & \text{w przeciwnym przypadku.} \end{cases}$$

e — liczba kolumn

Liczba kolumn jest równa liczbie gałęzi w grafie zupełnym o v wierzchołkach $[e = (v-1) \cdot v/2]$.

Macierz multigrafu odwzorującego układ ma postać:

$$G^v(t_2) = [g_{ij}]_{1 \times e} \quad (31)$$

gdzie:

$$g_{ij} = \begin{cases} p; & \text{jeśli multigraf zawiera } p \text{ równoległych } j\text{-tych gałęzi.} \\ 0; & \text{w przeciwnym przypadku.} \end{cases}$$

Macierz bi-grafu odwzorującego układ elektryczny ma postać:

$$G^v(t_1, t_2) = [g_{ij}]_{2 \times e} \quad (32)$$

gdzie:

$$g_{ij} = \begin{cases} 1; & \text{jeśli } i\text{-ty rodzaj elementu } \{t_i\} \text{ jest } j\text{-tą gałąź bi-grafu.} \\ 0; & \text{w przeciwnym przypadku.} \end{cases}$$

Dla układu zaprezentowanego na rysunku 2. macierze $G^v(t)$, $G^v(t_2)$ i $G^v(t_1, t_2)$ mają postać:

$$\begin{aligned} G^4(t) &= \begin{vmatrix} e_6 & e_5 & e_4 & e_3 & e_2 & e_1 \\ 1 & 1 & 1 & 1 & 1 & 0 \end{vmatrix} \\ G^4(t_2) &= \begin{vmatrix} 1 & 1 & 1 & 2 & 1 & 0 \end{vmatrix} \\ G^4(t_1, t_2) &= \begin{vmatrix} 0 & 1 & 1 & 1 & 0 & 0 \\ 1 & 0 & 0 & 1 & 1 & 0 \end{vmatrix} \end{aligned}$$

W procesie syntezy topologicznej struktur kanonicznych zakłada się, że impedencja każdego elementu układu ma wartość równą jeden. Wówczas wzór 1.27 przyjmuje postać:

$$H_k^v(t_1, t_2) = \sum_{n_{1k}=v-k}^0 \{LT_k^v(t_1^{n_{1k}}, t_2^{n_{2k}})\} \quad (33)$$

gdzie:

$$LT_k^v(t_1^{n_{1k}}, t_2^{n_{2k}}) = \sum_{T_k^1(t_1^{n_{1k}}, t_2^{n_{2k}})} [\prod^{n_{1k}}(t_1) \cdot \prod^{n_{2k}}(t_2)] \quad (34)$$

$LT_k^v(t_1^{n_{1k}}, t_2^{n_{2k}})$ – liczba dendrytów k -drzewowych składających się z n_{1k} elementów typu $\{t_1\}$ i $n_{2k}=v-k-n_{1k}$ elementów typu $\{t_2\}$, istniejących w bi-grafie układu zawierającego v wierzchołków.

Dla układu z rysunku 1.2 parametr $H_1^4(t_1, t_2)$ ma postać:

$$H_1^4(t_1, t_2) = 1 \cdot (t_1^3 \cdot t_2^0) + 5 \cdot (t_1^2 \cdot t_2^1) + 6 \cdot (t_1^1 \cdot t_2^2) + 1 \cdot (t_1^0 \cdot t_2^3)$$

Struktury kanoniczne układów będziemy charakteryzować wyrażeniem:

$$\langle LT_k^v(t_1^{n_{1k}}, t_2^{n_{2k}}) \rangle = \sum_{n_{1k}=v-k}^0 \{LT_k^v(t_1^{n_{1k}}, t_2^{n_{2k}})\} \quad (35)$$

$\langle LT_k^v(t_1^{n_{1k}}, t_2^{n_{2k}}) \rangle$ – uporządkowany zbiór (od $n_{1k}=v-k$ do $n_{1k}=0$) liczb dendrytów k -drzewowych złożonych z n_{1k} elementów typu $\{t_1\}$ i $n_{2k}=v-k-n_{1k}$ elementów typu $\{t_2\}$, (w skrócie $\langle LT_k^v(t_1, t_2) \rangle$).

Dla układu z rysunku 2 parametr $\langle LT_k^v(t_1, t_2) \rangle$ ma postać:

$$\langle LT_1^4(t_1, t_2) \rangle = \langle 1, 5, 6, 1 \rangle$$

Z przytoczonych rozważań wynika, że parametry czwórnika można charakteryzować w postaci funkcji wymiernej (1.26) zawierającej w liczniku i mianowniku odpowiednie wielomiany $H_k^v(t_1, t_2)$, lub w postaci uporządkowanego zbioru $\langle LT_k^v(t_1, t_2) \rangle$ (1.35).

Należy zwrócić uwagę, że zapis parametrów czwórnika w postaci uporządkowanego zbioru $\langle LT_k^v(t_1, t_2) \rangle$, daje nam dodatkową informację o liczbie wierzchołków w układzie. Uporządkowany zbiór $\langle LT_k^v(t_1, t_2) \rangle$ zawiera zawsze $v-k+1$ elementów, przy czym jeśli któryś ze współczynników wielomianu $H_k^v(t_1, t_2)$ jest równy zeru, to w $\langle LT_k^v(t_1, t_2) \rangle$ wpisujemy w to miejsce zero [12, 14].

3. MACIERZ DENDRYTÓW 1-DRZEWOWYCH

Z definicji dendrytów 2- i 3-drzewowych [12, 14, 15] wynika, że są one tylko „fragmentami” dendrytów 1-drzewowych istniejących w danej strukturze układu. Mając więc wygenerowane dla danego czwórnika dendryty 1-drzewowe można łatwo zapisać wszystkie jego parametry (tabela 1). Zanim przejdziemy do sprecyzo-

wania warunków D-równoważności struktur układów zdefiniujemy macierze dendrytów 1-drzewowych unigrafów, multigrafów i bi-grafów.

Macierz dendrytów 1-drzewowych $D\{T_1^v(t)\}$ unigrafu $G^v(t)$ odwzorującego strukturę układu nazywamy:

$$D\{T_1^v(t)\} = [d_{ij}]_{q \times e} \quad (36)$$

gdzie:

$$d_{ij} = \begin{cases} 1; & \text{jeśli } i\text{-ty dendryt 1-drzewowy zawiera } j\text{-tą gałąź unigrafu.} \\ 0; & \text{w przeciwnym przypadku.} \end{cases}$$

q — liczba wierszy

e — liczba kolumn.

Liczba wierszy macierzy $D\{T_1^v(t)\}$ jest równa liczbie dendrytów 1-drzewowych $T_1^v(t)$ istniejących w rozpatrywanej strukturze układu.

Liczba kolumn macierzy $D\{T_1^v(t)\}$ jest równa liczbie gałęzi w grafie zupełnym o v wierzchołkach.

Macierz dendrytów 1-drzewowych $D\{T_1^v(t_2)\}$ multigrafu $G^v(t_2)$ odwzorującego strukturę układu nazywamy:

$$D\{T_1^v(t_2)\} = [d_{ij}]_{q \times e} \quad (37)$$

gdzie:

$$d_{ij} = \begin{cases} p; & \text{jeśli } i\text{-ty dendryt 1-drzewowy zawiera } p \\ & \text{równoległych } j\text{-tych gałęzi multigrafu} \\ 0; & \text{w przeciwnym przypadku.} \end{cases}$$

Macierz dendrytów 1-drzewowych $D\{T_1^v(t)\}$ bi-grafu $G^v(t_1, t_2)$ odwzorującego strukturę układ ma postać:

$$D\{T_1^v(t_1, t_2)\} = \begin{vmatrix} G(t_1, t_2) \\ D(T_1^v(t_2)) \end{vmatrix} \quad (38)$$

Dla układu zaprezentowanego na rysunku 2. czwórnika macierze dendrytów 1-drzewowych $D\{T_1^v(t)\}$ $D\{T_1^v(t_2)\}$ i $D\{T_1^v(t_1, t_2)\}$ mają postać:

$$\begin{array}{lll} D\{T_1^4(t)\} = & D\{T_1^4(t_2)\} = & D\{T_1^4(t_1, t_2)\} = \\ e_6, e_5, e_4, e_3, e_2, e_1 & e_6, e_5, e_4, e_3, e_2, e_1 & e_6, e_5, e_4, e_3, e_2, e_1 \\ & & \begin{vmatrix} 0, & 1, & 1, & 1, & 0, & 0 \\ 1, & 0, & 0, & 1, & 1, & 0 \end{vmatrix} \end{array}$$

0, 0, 1, 1, 1, 0
0, 1, 0, 1, 1, 0
0, 1, 1, 1, 0, 0
1, 0, 0, 1, 1, 0
1, 0, 1, 0, 1, 0
1, 1, 0, 0, 1, 0
1, 1, 0, 1, 0, 0
1, 1, 1, 0, 0, 0

0, 0, 1, 2, 1, 0
0, 1, 0, 2, 1, 0
0, 1, 1, 2, 0, 0
1, 0, 0, 2, 1, 0
1, 0, 1, 0, 1, 0
1, 1, 0, 0, 1, 0
1, 1, 0, 2, 0, 0
1, 1, 1, 0, 0, 0

0, 0, 1, 2, 1, 0
0, 1, 0, 2, 1, 0
0, 1, 1, 2, 0, 0
1, 0, 0, 2, 1, 0
1, 0, 1, 0, 1, 0
1, 1, 0, 0, 1, 0
1, 1, 0, 2, 0, 0
1, 1, 1, 0, 0, 0

Należy zwrócić uwagę, że z macierzy $D\{T_1^v(t_1, t_2)\}$ łatwo uzyskać poszczególne dendryty 1-drzewowe $T_1^v(t_1, t_2)$ struktury czwórnika (tabela na rysunku 1.2) poprzez „przypisanie” poszczególnym gałęziom podmacierzy $D\{T_1^v(t_2)\}$ elementów $\{t_1\}$ i $\{t_2\}$ zgodnie z podmacierzą $G^v(t_1, t_2)$. W przypadku gdy w podmacierzy $d\{T_1^v(t_2)\}$ występuje cyfra 2 (w układzie występują równoległe elementy), to z danego dendrytu 1-drzewowego podmacierzy $D\{T_1^v(t_2)\}$ powstają dwa dendryty 1-drzewowe $T_1^v(t_1, t_2)$ bi-grafu.

4. D-RÓWNOWAŻNOŚĆ STRUKTUR UKŁADÓW

Równoważność struktur dwóch układów SLSOP będziemy definiować następująco:

„Dwie struktury układów $G_I^v(t_1, t_2)$ i $G_{II}^v(t_1, t_2)$, są równoważne, jeśli każdemu elementowi danego rodzaju $\{t_i\}$ układu $G_I^v(t_1, t_2)$, można przyporządkować element takiego samego rodzaju układu $G_{II}^v(t_1, t_2)$, oraz jeśli strukturalne funkcje wymierne $X_{ij}(t_1, t_2)$ opisujące poszczególne parametry układu $G_I^v(t_1, t_2)$ mają analogiczną postać jak strukturalne funkcje wymierne $X_{ij}(t_1, t_2)$ opisujące własności układu $G_{II}^v(t_1, t_2)$ ”.

Z powyższej definicji równoważności struktur układów oraz zależności 26 i 27 wynika, że dwie struktury są równoważne jeśli licznik $H_{ka}^v(t_1, t_2)$ i mianownik $H_{kb}^v(t_1, t_2)$ strukturalnej funkcji wymiernej $X_{ij}(t_1, t_2)$ opisującej dany parametr $X_{ij}(t_1, t_2)$ (wyrażony poprzez sumy iloczynów admitancji elementów tworzących odpowiednie dendryty k-drzewowe) układu $G_I^v(t_1, t_2)$ jest równy licznikowi $H_{ka}^v(t_1, t_2)$ i mianownikowi $H_{kb}^v(t_1, t_2)$ strukturalnej funkcji wymiernej $X_{ij}(t_1, t_2)$ opisującej dany parametr $X_{ij}(t_1, t_2)$ układu $G_{II}^v(t_1, t_2)$. Innymi słowy, każdy dendryt k-drzewowy $T_k^v(t_1, t_2)$ występujący w liczniku lub mianowniku funkcji wymiernej opisującej dany parametr układu $G_I^v(t_1, t_2)$ ma swojego odpowiednika w funkcji wymiernej opisującej analogiczny parametr układu $G_{II}^v(t_1, t_2)$. Ten typ równoważności struktur układów będziemy nazywać D-równoważnością (równoważnością dendrytów k-drzewowych).

W procesie syntezy struktur kanonicznych układów SLSOP metodą przeglądu hierarchicznego [14, 15], należy stwierdzić równoważność poszczególnych struktur odwzorowanych w postaci unigrafu, multigrafu i m-graflu. Zagadnienie D-równoważności dwóch struktur układów dotyczy więc:

1. D-równoważności unigrafów układów.
2. D-równoważności multigrafów układów.
3. D-równoważności m-grafów układów.

Wykażemy, że pojęcie równoważność struktur układów SLSOP realizujących funkcję wymierną obejmuje zarówno równoważność w sensie przekształcenia izomorficznego jak i 2-izomorficznego grafów nieskierowanych odwzorujących ich struktury.

W teorii układów SLSOP [1, 8] równoważność definiuje się następująco: „Dwa układy są równoważne, jeśli ich strukturalne macierze obwodów różnią się tylko permutacją wierszy i kolumn”.

W teorii grafów [3, 11] wykazano że: „Dwa grafy mają tę samą strukturalną macierz obwodu, wtedy i tylko wtedy, gdy są one 2-izomorficzne lub izomorficzne”.

Z powyższych twierdzeń oraz z faktu, że macierz podstawowych obwodów powstaje z dowolnie wybranego dendrytu 1-drzewowego i cięciw układu wynika, że D-równoważność struktur układów SLSOP dotyczy zarówno przekształcenia izomorficznego jak i 2-izomorficznego.

Twierdzenie 1.

Dwie struktury układów $G_I^v(t_1, t_2)$ i $G_{II}^v(t_1, t_2)$, są D-równoważne, jeśli ich macierze dendrytów 1-drzewowych różnią się tylko permutacją wierszy i (lub) kolumn, (kiedy mają odpowiedniość macierzy dendrytów 1-drzewowych).

Dowód:

1. Konieczność:

Z definicji D-równoważności wynika, że dwie struktury są równoważne jeśli opisujące ich parametry funkcje wymierne mają identyczną postać. Jak pokazano tabeli 1 każdy parametr opisujący fizyczne własności czwórników, wyraża się poprzez funkcję wymierną której licznik i mianownik to sumy iloczynów admitancji poszczególnych elementów tworzących dendryty k-drzewowe. Jak pokazano wyżej macierz $D\{T_1^v(t_1, t_2)\}$ zawiera wszystkie występujące w danej strukturze układu dendryty k-drzewowe, a tym samym można na jej podstawie wyznaczyć wszystkie parametry charakteryzujące dany układ.

2. Dostateczność:

Macierz $D_I\{t_1^v(t_1, t_2)\}$ zawiera pełny zestaw dendrytów k-drzewowych koniecznych do określenia parametrów struktury układu $g_I^v(t_1, t_2)$. Macierz $D_{II}\{t_1^v(t_1, t_2)\}$ zawiera pełny zestaw dendrytów k-drzewowych koniecznych do określenia parametrów struktury układu $G_{II}^v(t_1, t_2)$. Jeśli macierz $D_I\{T_1^v(t_1, t_2)\}$ struktury $G_I^v(t_1, t_2)$, można permutacją kolumn i (lub) wierszy przekształcić w macierz $D_{II}\{T_1^v(t_1, t_2)\}$ struktury $G_{II}^v(t_1, t_2)$, to parametry układów $G_I^v(t_1, t_2)$ i $G_{II}^v(t_1, t_2)$ są opisywane za pomocą identycznych funkcji topologicznych, czyli układy te są równoważne.

c.n.d.

Parametry $x_{ij} = \{x_{ij} \in z_{ij} \vee y_{ij} \vee a_{ij} \vee b_{ij} \vee h_{ij} \vee g_{ij}\}$
 charakteryzujące własności czwórników (czterokońcówkowych)

x	ij			
	11	12	21	22
z	$\frac{H_{1,0}}{H_1}$	$\frac{H_{12,30} - H_{12,30}}{H_1}$	$\frac{H_{13,20} - H_{12,30}}{H_1}$	$\frac{H_{2,3}}{H_1}$
y	$\frac{H_{2,3}}{\Sigma H_{i,jk,0}}$	$\frac{H_{13,20} - H_{12,30}}{\Sigma H_{i,jk,0}}$	$\frac{H_{12,30} - H_{13,20}}{\Sigma H_{i,jk,0}}$	$\frac{H_{1,0}}{\Sigma H_{i,jk,0}}$
a	$\frac{H_{1,0}}{H_{13,20} - H_{12,30}}$	$\frac{\Sigma H_{i,jk,0}}{H_{13,20} - H_{12,30}}$	$\frac{H_1}{H_{13,20} - H_{12,30}}$	$\frac{H_{2,3}}{H_{13,20} - H_{12,30}}$
b	$\frac{H_{2,3}}{H_{12,30} - H_{13,20}}$	$\frac{\Sigma H_{i,jk,0}}{H_{12,30} - H_{13,20}}$	$\frac{H_1}{H_{12,30} - H_{13,20}}$	$\frac{H_{1,0}}{H_{12,30} - H_{13,20}}$
h	$\frac{\Sigma H_{i,jk,0}}{H_{2,3}}$	$\frac{H_{12,30} - H_{13,20}}{H_{2,3}}$	$\frac{H_{12,30} - H_{13,20}}{H_{2,3}}$	$\frac{H_1}{H_{2,3}}$
g	$\frac{H}{H_{1,0}}$	$\frac{H_{13,20} - H_{12,30}}{H_{1,0}}$	$\frac{H_{13,20} - H_{12,30}}{H_{1,0}}$	$\frac{\Sigma H_{i,jk,0}}{H_{1,0}}$

gdzie:

$$\begin{aligned}
 H_1 &= H_1(t_1, t_2) \\
 H_{1,0} &= H_{1,0}(t_1, t_2) \\
 H_{2,3} &= H_{2,3}(t_1, t_2) \\
 H_{12,30} &= H_{12,30}(t_1, t_2) \\
 H_{13,20} &= H_{13,20}(t_1, t_2) \\
 \Sigma H_{i,jk,0} &= H_{12,3,0}(t_1, t_2) + H_{1,2,3,0}(t_1, t_2) + H_{13,2,0}(t_1, t_2) + H_{1,3,2,0}(t_1, t_2)
 \end{aligned}$$

Podziału wygenerowanych w procesie syntezy topologicznej metodą przeglądu hierarchicznego unigrafów i multigrafów na klasy ich struktur topologicznie równoważnych będziemy dokonywać korzystając z następujących twierdzeń.

Twierdzenie 2.

Dwa struktury układów w postaci multigrafów $G_I^v(t_2)$ i $G_{II}^v(t_2)$, są topologicznie równoważne jeśli ich macierze dendrytów 1-drzewowych $D\{T_1^v(t_2)\}$ różnią się tylko permutacją wierszy i (lub) kolumn, (kiedy mają odpowiedniość macierzy dendrytów 1-drzewowych).

Dowód:

Dowód powyższego twierdzenia jest analogiczny do dowodu twierdzenia 1.

Twierdzenie 3.

Dwa struktury układów w postaci unigrafów $G_I^v(t)$ i $G_{II}^v(t)$, są topologicznie równoważne jeśli ich macierze dendrytów 1-drzewowych $D\{T_1^v(t)\}$ różnią się tylko permutacją wierszy i (lub) kolumn, (kiedy mają odpowiedniość macierzy dendrytów 1-drzewowych).

Dowód:

Dowód powyższego twierdzenia jest analogiczny do dowodu twierdzenia 1.1.

5. ALGORYTM — „D-RÓWNOWAŻNOŚĆ STRUKTUR UKŁADÓW”

Z twierdzeń 1, 2 i 3 wynika, że aby stwierdzić D-równoważność dwóch struktur (unigrafów, multigrafów i bigrafów) układów należy wygenerować ich macierze dendrytów 1-drzewowych, a następnie wykonać w najgorszym przypadku $e! \cdot q!$ permutacji wierszy i kolumn (e — liczba wierszy, q — liczba kolumn), co może się okazać trudne do wykonania nawet przy użyciu szybkich komputerów.

W procesie syntezy topologicznej metodą przeglądu hierarchicznego, w pierwszym etapie generuje się struktury kanoniczne unigrafów a dopiero w drugim etapie tworzy się w nich struktury bi-grafów układów. Wynika stąd, że w pierwszym etapie syntezy istnieje tylko konieczność stwierdzenia D-równoważności unigrafów. W praktyce dla unigrafów reprezentujących struktury układów, liczbę kroków koniecznych do stwierdzenia ich D-równoważności można wyraźnie zmniejszyć, biorąc pod uwagę następujące fakty. Z definicji D-równoważności dwóch struktur unigrafów układów wynika, że dwie struktury są D-równoważne jeśli:

1. Moce zbiorów węzłów odpowiednio unigrafów $G_I^v\{t\}$ i $G_{II}^v\{t\}$, są równoliczne.

2. Moce zbiorów gałęzi odpowiednio unigrafów $G_I^v\{t\}$ i $G_{II}^v\{t\}$, są równoliczne.

3. Moce zbiorów gałęzi incydentnych do odpowiadających sobie wierzchołków wejściowych $\{w_1, w_0\}$ i wyjściowych $\{w_3, w_2\}$ odpowiednio unigrafów $G_I^v\{t\}$ i $G_{II}^v\{t\}$, są równoliczne. Warunek ten wynika z faktu, że w dowolnych przekształceniach struktur układów tylko gałęzie incydentne do wierzchołków brzegowych przekształcają się same w siebie.

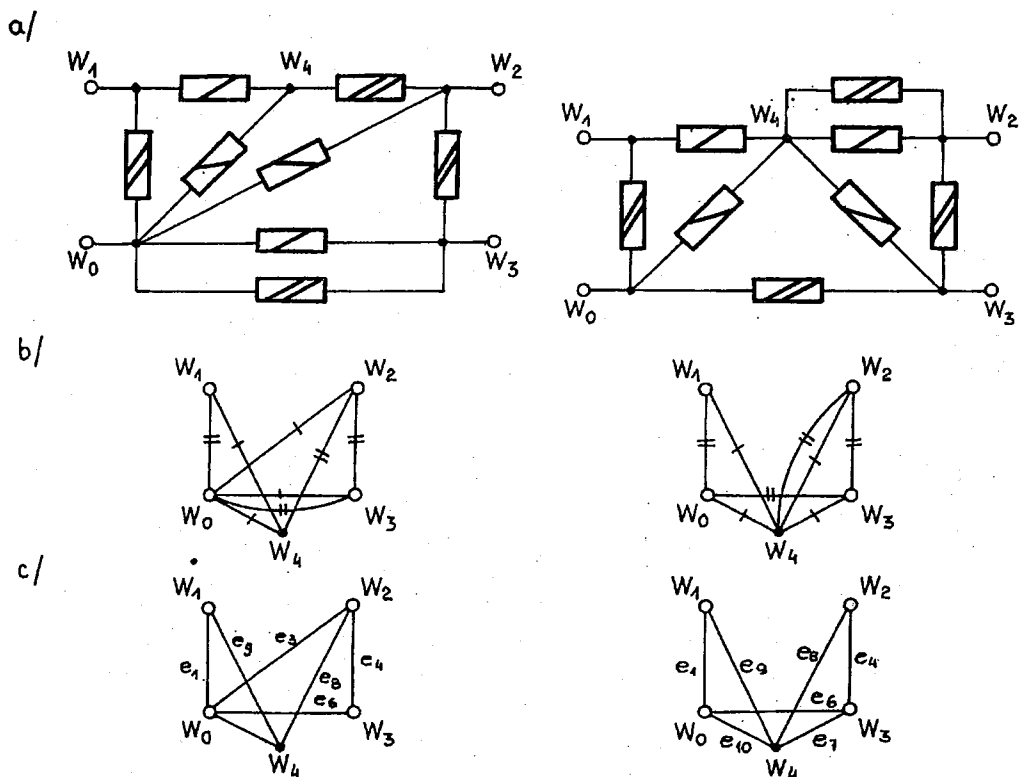
Również liczbę koniecznych permutacji zarówno wierszy jak i kolumn macierzy $D\{T_1^v(t)\}$ można w istotny sposób zredukować. Zdefiniujemy pomocniczą macierz jednowierszową $A = [a_j]$

$$A = [a_j]_{1..e} \quad (39)$$

gdzie:

$$a_j = \sum_{i=0}^q d_{ij} \quad (40)$$

Sposób redukcji liczby koniecznych permutacji kolumn macierzy dendrytów 1-drzewowych, zilustrujemy na przykładzie dwóch D-równoważnych struktur unigrafów układów pokazanych na rysunku 3. Dla rozpatrywanych struktur układów wygenerowano również ich macierze $D_I\{t_1^v(t)\}$ i $D_{II}\{t_1^v(t)\}$ oraz macierze A_I i A_{II} . Poszczególne wyrazy macierzy A to sumy „jedynek” w każdej kolumnie macierzy $D\{T_1^v(t)\}$.

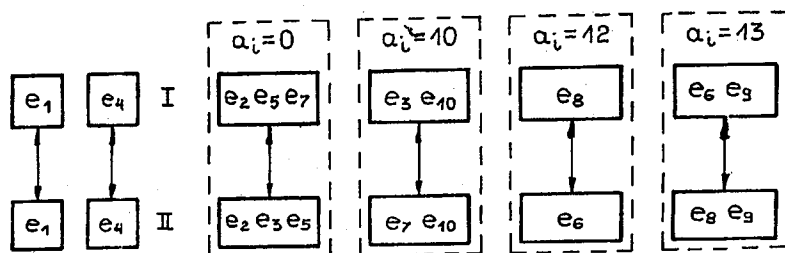


$$G^S(t) = \begin{matrix} & e_{10} & e_9 & e_8 & e_7 & e_6 & e_5 & e_4 & e_3 & e_2 & e_1 \\ \begin{matrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \end{matrix} & \begin{bmatrix} 1 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 0 & 0 & 1 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 1 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 1 \\ 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 1 \\ 0 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 & 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 1 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 1 \\ 1 & 0 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 1 \\ 1 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 1 \\ 1 & 0 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \\ 1 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \end{bmatrix} \end{matrix}$$

$$G^S(t) = \begin{matrix} & e_{10} & e_9 & e_8 & e_7 & e_6 & e_5 & e_4 & e_3 & e_2 & e_1 \\ \begin{matrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \end{matrix} & \begin{bmatrix} 1 & 1 & 1 & 1 & 1 & 0 & 1 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & 1 & 0 & 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \\ 0 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 1 \\ 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 \\ 1 & 0 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 1 \\ 1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix} \end{matrix}$$

$$A^S(t) = |10, 13, 12, 0, 13, 0, 13, 10, 0, 13| \quad A^S(t) = |10, 13, 13, 10, 12, 0, 13, 0, 0, 13|$$

Rys. 3. Topologicznie równoważne czwórnik; a, b — czwórnik, c, d — bi-grafy struktur czwórników, e, f — unigrafy struktur czwórników, $G^S(t)$, $D^S(t) = D\{T_1^S(t_1, t_2)\}$ $A^S(t)$ — macierze charakteryzujące unigrafy



Rys. 4. Ilustracja permutacji elementów a_j w macierzy A^5 charakteryzujące unigrafy z rysunku 1.3

Rozważmy permutacje kolumn macierzy $D_I\{T_1^v(t)\}$ w $D_{II}\{T_1^v(t)\}$. Permutacje kolumn macierzy $D_I\{T_1^v(t)\}$ w $D_{II}\{T_1^v(t)\}$, mogą dotyczyć tylko tych kolumn które mają identyczne wartości elementów a_j w macierzach A_I i A_{II} . Ilustruje to rysunek 4. Pętle własne dotyczą zawsze kolumn d_{ij} i d_{ji} które reprezentują gałąź wejściową $e_1 \in \{w_0, w_1\}$ i gałąź wyjściową $e_4 \in \{w_2, w_3\}$ w macierzach $D_I\{T_1^v(t)\}$ i $D_{II}\{T_1^v(t)\}$. Tym sposobem całkowitą permutację $f: D_I\{T_1^v(t)\} \Rightarrow D_{II}\{T_1^v(t)\}$, można podzielić na cykliczne permutacje dotyczące tylko tych kolumn które mają taką samą wartość elementu a_j w macierzach A_I i A_{II} .

Należy zwrócić uwagę na fakt, że zaprezentowana metoda redukcji liczby permutacji kolumn macierzy $D\{T_1^v(t)\}$ jest szczególnie efektywna dla struktur o większej liczbie wierzchołków {liczba kolumn macierzy dendrytów 1-drzewowych zależy od liczby wierzchołków $e = v \cdot (v-1)/2$ } kiedy prawdopodobieństwo, że wiele kolumn ma identyczną wartość elementu a_j macierzy A wyraźnie maleje.

Liczbę koniecznych permutacji wierszy macierzy $D\{T_1^v(t)\}$ można również w istotny sposób zredukować, poprzez tzw: uporządkowanie leksykograficzne [2] poszczególnych dendrytów 1-drzewowych.

Uporządkowanie leksykograficzne występuje wtedy kiedy każdy następny wiersz w macierzy, traktowany jako reprezentacja liczby naturalnej w systemie dwójkowym, ma większą wartość.

Uporządkowanie leksykograficzne szczególnie łatwo daje się zastosować do sortowania wierszy w macierzy $D\{T_1^v(t)\}$. Analogiczny algorytm można zastosować do stwierdzenia topologicznej równoważności multigrafów i bi-grafów struktur czwórników, korzystając z macierzy odpowiednio $D\{t_1^v(t_2)\}$ i $D\{t_1^v(t_1, t_2)\}$. Uporządkowanie leksykograficzne w macierzy $D\{T_1^v(t_1, t_2)\}$ dotyczy tylko wierszy jej podmacierzy $D\{T_1^v(t_2)\}$.

BIBLIOGRAFIA

1. S. B o k o w s k i: *Elektrotechnika teoretyczna — teoria obwodów elektrycznych*. Warszawa; WNT, 1982
2. W.K. C h e n: *Applied graph theory — graphe and electric networks*. Amsterdam North-Holland P.C. 1976
3. N. D e o: *Teoria grafów i jej zastosowanie w technice i informatyce*. Warszawa, PWN 1980
4. A. K o Ń c z y k o w s k a, J. W o j c i e c h o w s k i: *Podstawy topologicznych metod analizy układów elektrycznych*. Warszawa, PWN 1983

5. N.R. Malik, G.L. Jackson, Y.S. Kim: *Theory and application of resistor, linear controlled resistor, linear controlled networks*. IEEE. Trans. Circuits and Systems, 1976, vo. CAS—23, pp.222—228
6. N.G. Maksimowicz: *Teoria grafów i elektryczkie sieci*. Lwów Wisza Szkoła, 1987
7. K. Mikołajuk Z. Trzaska: *Elektrotechnika teoretyczna — analiza i synteza elektrycznych obwodów liniowych*. Warszawa: PWN, 1984
8. M.E. Reingold, J. Nievergelt, N. Deo: *Algorytmy kombinatoryczne*. Warszawa: PWN, 1985
9. M.N.S. Swamy, K. Thulasiraman: *Graphs, networks and algorithms*. New York: Wiley — Interscience, 1981
10. Z. Wróbel: *Funkcje strukturalne grafów układów dwójników i czwórników składających się z dwóch rodzajów elementów*. Kwartalnik Elektroniki i Telekomunikacji, 1991, 37, z.1—2, ss.35—52
11. Z. Wróbel: *Homeomorfizm grafów układów dwójników i czwórników składających się z dwóch rodzajów elementów*. Kwartalnik Elektroniki i Telekomunikacji, 1991, 37, z.1—2, ss.53—65
12. Z. Wróbel: *Synteza topologiczna struktur kanonicznych bigrafów układów dwójników, Część I. Metoda przeglądu bezpośredniego*. Kwartalnik Elektroniki i Telekomunikacji, 1991, 37, z.1—2, ss.67—83
13. Z. Wróbel: *Synteza topologiczna struktur kanonicznych bi-grafów układów dwójników, Część II. Metoda przeglądu hierarchicznego*. Kwartalnik Elektroniki i Telekomunikacji, 1991, 37, z.1—2, ss.85—97

Z. WRÓBEL

D-EQUIVALENCE STRUCTURE CIRCUITS REALIZING RATIONAL FUNCTION

Summary:

D-equivalence structure circuits is equivalence in the sense of isomorphic and 2-isomorphic transformation. Algebraic criteria and algorithm D-equivalence structure circuits realizing rational function were presented in this paper.

Synteza topologiczna struktur kanonicznych czwórników składających się z dwóch rodzajów elementów

ZYGMUNT WRÓBEL

*Zakład Elektrotechniki i Automatyki, Instytut Problemów Techniki
Uniwersytet Śląski*

Otrzymano 1993.03.15

Autoryzowano do druku 1993.09.01

W pracy przedstawiono procedurę syntezy topologicznej struktur kanonicznych metodą przeglądu hierarchicznego, pozwalającą wygenerować wszystkie możliwe struktury topologiczne czwórników. Przedstawiono warunki realizacji minimalno-elementowych struktur czwórników. Procedurę syntezy topologicznej zilustrowano na przykładzie minimalno-elementowych struktur czwórników.

1. WPROWADZENIE

Synteza układów jest zwykle rozumiana jako zbiór metod umożliwiających wyznaczenie rodzaju i wartości elementów oraz sposobu takiego ich połączenia, aby otrzymany w efekcie układ fizyczny realizował żadaną zależność funkcyjną [3].

Procedurę syntezy układów można w zasadzie podzielić na dwa odrębne grupy zagadnień:

- synteze topologiczną układów.
- synteze parametryczną układów.

Pod pojęciem syntezy topologicznej układów będziemy rozumieć generację wszystkich możliwych struktur topologicznych tych układów, oraz ich podział według zadanych kryteriów na klasy struktur równoważnych. W procesie syntezy topologicznej interesuje nas tylko rodzaj i sposób połączenia poszczególnych rodzajów elementów (przyjmując, że każdy element układu ma wartość równą jeden). Sposób połączenia poszczególnych elementów określa strukturę topologiczną układu.

Pod pojęciem syntezy parametrycznej układów, będziemy rozumieć opracowanie algorytmów umożliwiających takie wyznaczenie rodzaju wartości poszczególnych elementów tworzących wygenerowane struktury topologiczne, aby skonstruowany w efekcie układ posiadał pożądane własności zewnętrzne. Należy zwrócić uwagę, że w prachach dotyczących syntezy układów zawierających skupione, liniowe, stacjonarne odwracalne i pasywne (SLSOP) elementy. pierwszemu etapowi syntezy czyli generacji struktur kanonicznych poświęca się niewiele uwagi. W większości prac strukturę układu zadaje się a priori (np: struktury drabinkowe, mostkowe itp.) koncentrując uwagę przede wszystkim na drugim etapie syntezy poszukując zwykle rozwiązań optymalnych ze względu na zadany parametr funkcji celu procesu syntezy.

W procesie syntezy topologicznej generowane są struktury układów które mogą być między sobą równoważne. Istnieje więc konieczność podziału wygenerowanych struktur układów na klasy struktur równoważnych względem przyjętych kryteriów równoważności. Pojęcie równoważności układów można rozpatrywać w różnych aspektach. Dla układów SLSOP równoważność definiuje się następująco: „Dwa układy są równoważne, jeśli ich strukturalne macierze obwodów różnią się tylko permutacją wierszy i kolumn”. Ten typ równoważności w pracy [9] zdefiniowano jako D-równoważność (równoważność struktur układów w sensie przekształcenia izomorficznego lub 2-izomorficzne). Problem generacji pełnego zbioru struktur kanonicznych układów SLSOP, to ważny problem teorii obwodów leżący na pograniczu teorii grafów [1, 5] i kombinatoryki [2]. Dotychczas nie opracowano kompleksowo zagadnienia generacji struktur kanonicznych układów SLSOP. Związane jest to przede wszystkim z faktem, że dotychczas nie opracowano kryteriów równoważności pozwalających podzielić generowane w procesie syntezy struktury na klasy struktur D-równoważnych (w sensie przekształcenia izomorficznego i 2-izomorficznego). Opracowany w pracy [9] algorytm pozwalający stwierdzić D-równoważność dwóch struktur układów pozwolił na nowo spojrzeć na zagadnienie syntezy struktur kanonicznych układów.

Niniejsza praca jest kontynuacją prac [7, 8]. Przedstawiono w niej nowe nie standartowe kombinacyjno-strukturalne podejście do zagadnienia syntezy topologicznej układów. Wyniki syntezy struktur kanonicznych zilustrowano na przykładzie czwórników (czterokońcówkowych) składających się z dwóch rodzajów elementów, ze szczególnym uwzględnieniem generacji struktur minimalno-elementowych.

2. GENERACJA STRUKTUR KANONICZNYCH CZWÓRNIKÓW

W procesie syntezy topologicznej interesuje nas tylko rodzaj i sposób połączenia poszczególnych rodzajów elementów (przyjmując, że każdy element układu ma wartość równą jeden). Sposób połączenia poszczególnych elementów określa strukturę topologiczną układu.

W procesie syntezy topologicznej generowane struktury układów, mogą być między sobą topologicznie równoważne. Istnieje więc konieczność podziału wygenerowanych struktur na klasy struktur topologicznie równoważnych, co można dokonać korzystając z twierdzeń 1, 2 3 pracy [9].

Strukturę znajdującą się na pierwszym miejscu liniowo uporządkowanego ciągu struktur tworzących daną klasę równoważności (podobieństwa) będziemy nazywał kanonicznym reprezentantem klasy równoważności struktur. Kanoniczny reprezentant klasy równoważności struktur jest oczywiście zależny od przyjętej zasady liniowego uporządkowania struktur, a więc własność „bycia kanonicznym reprezentantem klasy równoważności struktur” nie jest immanentną cechą struktury. Struktura kanoniczna jest niezmiennikiem przekształceń topologicznych struktur układów.

Strukturę kanoniczną unigrafu będziemy oznaczali $G^v(t)$, multigrafu $G^v(t_2)$ a bi-grafu $G^v(t_1, t_2)$.

Strukturę kanoniczną unigrafu będziemy charakteryzować liczbami dendrytów k -drzewowych $LT_k^v(t)$, multigrafu $LT(t_2)$ natomiast strukturę kanoniczną bigrafu czwórnika będziemy charakteryzowali uporządkowanymi zbiorami $\langle LT_k^v(t_1, t_2) \rangle$ [6].

Typem struktury kanonicznej układu będziemy nazywamy przedstawiciela tych struktur kanonicznych, które charakteryzują się analogicznym typem uporządkowanego zbioru $\langle LT_k^v(t_1, t_2) \rangle$.

Typ uporządkowanego zbioru $\langle LT_k^v(t_1, t_2) \rangle$ zawiera tylko trzy wartości liczb dendrytów k -drzewowych $\langle 1, 1, 0 \rangle$. Jeśli: $LT_k^v(t_1, t_2) > 1 \Rightarrow 1$, $LT_k^v(t_1, t_2) = 1 \Rightarrow 1$, $LT_k^v(t_1, t_2) = 0 \Rightarrow 0$.

Syntezę topologiczną struktur kanonicznych czwórników można prowadzić metodą przeglądu bezpośredniego lub przeglądu hierarchicznego [7, 8]. Jak wynika z pracy [8], bardziej efektywna jest metoda przeglądu hierarchicznego. Niżej przedstawiono procedurę syntezy topologicznej struktur kanonicznych czwórników metodą przeglądu hierarchicznego.

Algorytm syntezy topologicznej struktur kanonicznych czwórników metodą przeglądu hierarchicznego zawiera dwa etapy:

1. Generację struktur unigrafów układów i podział ich na klasy struktur kanonicznych $G(t)$.
2. Tworzenie bigrafów w poszczególnych klasach struktur kanonicznych unigrafów $G(t)$, oraz podział wygenerowanych struktur czwórników na klasy ich struktur kanonicznych $G(t_1, t_2)$.

2.a. GENERACJA STRUKTUR KANONICZNYCH UNIGRAFÓW CZWÓRNIKÓW

W procesie syntezy topologicznej czwórników metodą przeglądu hierarchicznego, zgodnie z założeniami przedstawionymi w pracy [8], będziemy rozpatrywać tylko te unigrafy, których liczba gałęzi 1_c spełnia zależność:

$$v \leq 1_c \leq 2 \cdot (v - 1) \quad (1)$$

Ograniczenie od dołu liczby gałęzi w unigrafie czwórnika wynika z faktu, że unigraf jest cyklicznie spójny względem obu par wierzchołków brzegowych $\{w_0, w_1\}$ i $\{w_2, w_3\}$ gdy ma co najmniej v gałęzi.

Ograniczenie od góry liczby gałęzi w unigrafie czwórnika dotyczy jak pokazano niżej, struktur minimalno- elementowych tworzonych jako sumy mnogościowe

dwóch dendrytów i-drzewowych o $v-1$ gałęziach. Z wygenerowanych struktur unigrafów spełniających warunek 1, należy więc wydzielić te struktury które są cyklicznie spójne względem obu par wierzchołków brzegowych. Procedurę pozwalającą stwierdzić spójność unigrafu czwórnika względem obu par wierzchołków brzegowych wejściowych $\{w_0, w_1\}$ i wyjściowych $\{w_2, w_3\}$ można zapisać w postaci:

$$P_i^v(w_0, w_1; t) \cap G^v(t) = P_i^v(w_0, w_1; t) \quad (2)$$

$$P_i^v(w_2, w_3; t) \cap G^v(t) = P_i^v(w_2, w_3; t) \quad (3)$$

$$\bigcup_{i=1}^{n1} P_i^v(w_0, w_1; t) = G^v(t) \quad (4)$$

$$\bigcup_{j=1}^{n2} P_j^v(w_2, w_3; t) = G^v(t) \quad (5)$$

gdzie:

$P_i^v(w_0, w_1; t)$, $P_i^v(w_2, w_3; t)$ – macierze ścieżek odpowiednio względem wierzchołków $\{w_0, w_1\}$ i $\{w_2, w_3\}$ unigrafu układu.

Równania (2) i (3) dają odpowiedź na pytanie, czy i -ta ścieżka wyznaczona względem wierzchołków $\{w_0, w_1\}$ oraz j -ta ścieżka wyznaczona względem wierzchołków $\{w_2, w_3\}$, są ścieżkami unigrafu czwórnika. Równania (4) i (5) dają natomiast odpowiedź na pytanie, czy otrzymane z dwóch pierwszych równań ścieżki, zawierają wszystkie gałęzie występujące w danym unigrafie.

Tabela 1

Struktury unigrafów $G^4(t)$ czwórników

Nr	$G^4(t)$	Model topol.	$LT_k^4(t)$					l_e
			$LT_{12,30}^4$	$LT_{13,20}^4$	$LT_{1,0}^4$	$LT_{2,3}^4$	LT_1^4	
1	1, 1, 1, 1, 1, 1		1	1	8	8	16	6
2	1, 1, 1, 1, 0, 1		0	1	5	5	8	5
2a	0, 1, 1, 1, 1, 1							
2b	1, 0, 1, 1, 1, 1		1	0				
2c	1, 1, 1, 0, 1, 1							
3	1, 1, 0, 1, 1, 1		1	1	4	8	8	
4	1, 1, 1, 1, 1, 0				8	4		
5	1, 1, 0, 1, 1, 0		1	1	4	4	4	4
6	1, 0, 1, 0, 1, 1		1	0	3 3			
6a	0, 1, 1, 1, 0, 1		0	1				

Struktury kanoniczne unigrafów $G^4(t)$ zaznaczono pogrubioną czcionką.

W tabeli 1 zestawiono wszystkie struktury unigrafów $G^4(t)$ spełniających równania 1–5, dzieląc je na klasy ich struktur kanonicznych (zaznaczonych pogrubioną czcionką). Należy zwrócić uwagę, że w procesie przekształcania wierzchołków $\{w_0, w_1\} \Rightarrow \{w_1, w_0\}$ lub $\{w_2, w_3\} \Rightarrow \{w_3, w_2\}$ unigrafu $G^v(t)$, dendryty $T_{12,30}(t) \Rightarrow T_{13,20}(t)$ i $T_{13,20}(t) \Rightarrow T_{12,30}(t)$. dlatego też do danej klasy struktur kanonicznych będziemy zaliczać struktury $G^v(t)|_I$, i $G^v(t)|_{II}$ których $T_{12,30}(t)|_I = T_{13,20}(t)|_I$ i $T_{13,20}(t)|_{II} = T_{12,30}(t)|_{II}$.

2.b. GENERACJA STRUKTUR KANONICZNYCH CZWÓRNIKÓW $G^v(t_1, t_2)$

Mając wygenerowane struktury kanoniczne unigrafów czwórników $G(t)$, można przejść do tworzenia w nich struktur zawierających poszczególne rodzaje elementów $\{t_1\}$ i $\{t_2\}$.

Możliwe są różne warianty „wypełnienia” zadanej struktury kanonicznej unigrafu $G^v(t)$ elementami typu $\{t_1\}$ i $\{t_2\}$. Zbiór elementów danego rodzaju $\{t_i\}$ wchodzących w skład struktury czwórnika będziemy nazywali mono-strukturą i oznaczali $E(t_i)$. Procedurę „wypełnienia” zadanej struktury kanonicznej unigrafu $G^v(t)$ elementami typu $\{t_1\}$ i $\{t_2\}$, można zapisać w postaci:

$$E^v(t_i) \cap G^v(t) = G^v(t_i) \quad (6)$$

$$\bigcup_{i=1}^2 E^v(t_i) = G^v(t) \quad (7)$$

Pierwsze równanie daje odpowiedź na pytanie – czy mono-struktura $E^v(t_1)$ zawiera się w strukturze $G^v(t)$. Drugie równanie daje odpowiedź na pytanie – czy utworzona jako suma mnogościowa mono-struktur $E^v(t_1)$, struktura czwórnika $G^v(t_1, t_2)$ zawiera wszystkie gałęzie struktury kanonicznej $G^v(t)$.

W procesie syntezy topologicznej funkcja celu jest generacją wszystkich możliwych struktur układów i ich podział na klasy struktur topologicznie równoważnych (struktur kanonicznych). Taki cel ma bardziej charakter poznawczy, pozwala stwierdzić że przy danym naborze elementów istnieje skończony zbiór struktur kanonicznych układów. Korzystając z wyników prac [8, 9], łatwo można wygenerować pełny zbiór struktur kanonicznych czwórników. W praktyce istnieje często konieczność generacji struktur które realizowałyby funkcje wymierne o maksymalnym stopniu licznika i mianownika przy minimalnej liczbie elementów tzw. struktury minimalno-elementowe.

3. MINIMALNO-ELEMENTOWE TYPY STRUKTUR KANONICZNYCH CZWÓRNIKÓW

Strukturę czwórnika nazywamy minimalno-elementową, jeśli realizuje parametry $x_{ij}(t_1, t_2)$ [$x_{ij}(t_1, t_2) \in z_{ij}(t_1, t_2) \vee y_{ij}(t_1, t_2) \vee a_{ij}(t_1, t_2) \vee b_{ij}(t_1, t_2) \vee h_{ij}(t_1, t_2) \vee g_{ij}(t_1, t_2)$] (tabela 1 w pracy [9]), w postaci funkcji wymiernej o maksymalnym

stopniu wielomianu licznika i mianownika względem $\{t_1\}$ lub $\{t_2\}$ przy minimalnej liczbie poszczególnych rodzajów elementów $\{t_1\}$.

Z tabeli 1 pracy [9] wynika, że parametry opisujące parametry czwórnika są funkcjami wymiernymi, których licznik i mianownik to wielomiany iloczynów admitancji poszczególnych elementów układu tworzących:

1. $H_1^v(t_1, t_2)$.
2. $H_{1,0}^v(t_1, t_2)$.
3. $H_{2,3}^v(t_1, t_2)$.
4. $H_{12,30}^v(t_1, t_2)$, $H_{13,20}^v(t_1, t_2)$.
5. $H_{1,2,30}^v(t_1, t_2)$, $H_{12,3,0}^v(t_1, t_2)$, $H_{1,3,20}^v(t_1, t_2)$, $H_{13,2,0}^v(t_1, t_2)$.

Niżej przedstawiono warunki jakie muszą spełniać dwie monostruktury $E^v(t_1)$ i $E^v(t_2)$ aby utworzona z nich (zgodnie z równaniami 6 i 7) struktura czwórnika, realizowała funkcję wymierną o maksymalnym stopniu licznika lub mianownika w postaci wielomianów $H_k^v(t_1, t_2)$, przy minimalnych liczbach $n(t_1)$ i $n(t_2)$ poszczególnych rodzajów elementów.

3.a. STRUKTURY MINIMALNO-ELEMENTOWE REALIZUJĄCE $H_1^v(t_1, t_2)$

Twierdzenie 1

Czwórnik o v wierzchołkach zawierający dwa rodzaje elementów o admitancjach $\{Y_{1i}, t_1\}$ i $\{Y_{2j}, t_2\}$ realizuje wielomian $H_1^v(t_1, t_2)$ licznika lub mianownika funkcji wymiernej $X_{ij}(t_1, t_2)$ o maksymalnym stopniu względem $\{t_1\}$ typu:

$$H_1^v(t_1, t_2) = A_{v-1,0} \cdot t_1^{v-1} \cdot t_2^0 + A_{v-2,1} \cdot t_1^{v-2} \cdot t_2^1 + \dots + A_{1,v-1} \cdot t_1^1 \cdot t_2^{v-2} + A_{0,v-1} \cdot t_1^0 \cdot t_2^{v-1} \quad (8)$$

co w postaci jego uporządkowanego zbioru $\langle LT_1^v(t_1, t_2) \rangle$ można zapisać:

$$\langle 1, LT_1^v(t_1^{v-2}, t_2^1), \dots, LT_1^v(t_1^1, t_2^{v-2}), 1 \rangle \quad (9)$$

wtedy i tylko wtedy gdy:

$$G^v(t_1, t_2) = T_1^v(t_1) \cup T_1^v(t_2) \quad (10)$$

gdzie:

$T_1^v(t_1)$ – dendryt 1-drzewowy o v wierzchołkach składający się tylko z elementów typu $\{t_1\}$.

$T_1^v(t_2)$ – dendryt 1-drzewowy o v wierzchołkach składający się tylko z elementów typu $\{t_2\}$.

Dowód:

Konieczność:

Z zależności 29 pracy [9] wynika, że $A_{v-1,0}$ jest sumą iloczynów admitancji tylko elementów typu $\{t_1\}$, tworzących dendryt 1-drzewowy, a zatem maksymalny stopień wielomianu w stosunku do elementów $\{t_1\}$ będzie miał wartość t_1^{v-1} . Analogicznie $A_{0,v-1}$ to suma iloczynów admitancji tylko elementów typu $\{t_2\}$ tworzących dendryt

1-drzewowy, w związku z tym maksymalny stopień wielomianu w stosunku do elementu $\{t_2\}$ będzie miał wartość $t_1^{v-1} \cdot t_2^0$ czwórnik musi zawierać co najmniej $v-1$ elementów typu $\{t_1\}$ tworzących minimum jeden dendryt 1-drzewowy składający się tylko z elementów typu $\{t_1\}$. Aby zaś wielomian wyrażenia (8) miał stopień $(t_1^0 \cdot t_2^{v-1})$ czwórnik musi zawierać co najmniej $v-1$ elementów typu $\{t_2\}$ tworzących minimum jeden dendryt 1-drzewowy, składający się tylko z elementów typu $\{t_2\}$.

Tabela 2

Dendryty k -drzewowe w pełnym unigrafie $G^4(t)$ konieczne do wyznaczania parametrów charakteryzujących czworniki

		$T_1^4(t)$	Model topol.	$T_{1,0}^4$	$T_{2,3}^4$	$T_{12,30}^4$	$T_{13,20}^4$	$\frac{T_{12,3,0}^4}{T_{1,2,30}^4}$	$\frac{T_{13,2,0}^4}{T_{1,3,20}^4}$
A	A ₁	100011		100010	100010	100010		000010 100000	
		101010							
	A ₂	010101		010100	010100		010100		000100 010000
		011100							
B	B ₁	010011		010010	000011 010001	000010	010000	000010	010000
		101010							
	B ₂	001110		001010 001100	000110	000010	000100	000010	000100
		101010							
C	C ⁺	001011		001010	000011	000010		000010	
		101001							
	C ⁻	001101		001100	000101		000100		000100
		011001							
D	D ⁺	100110		100010 100100	100010 000110	100010	000100	000010 100000	000100
		110010							
	D ⁻	010110		010010 010100	010110 010110	000010	010100	000010 010000	000100 010000
		110100							

Kursywą oznaczono „niepełne” dendryty k -drzewowe

Dostateczność:

Warunkiem dostatecznym aby czwórnik realizował wielomian 8 przy minimalnej liczbie elementów obu typów jest minimalna liczba dendrytów 1-drzewowych czyli $LT_1^v(t_1^{v-1}, t_2^0) = 1$ oraz $LT_1^v(t_1^0, t_2^{v-1}) = 1$. Aby były spełnione powyższe warunki, czwórnik musi zawierać tylko $v-1$ elementów typu $\{t_1\}$ tworzących dendryt 1-drzewowy $T_1^v(t_1)$, (tym samym $LT_1^v(t_1^{v-1}, t_2^0) = 1$) oraz zawierać tylko $v-1$ elementów typu $\{t_2\}$ tworzących dendryt 1-drzewowy $T_1^v(t_1)$, (tym samym $LT_1^v(t_1^0, t_2^{v-1}) = 1$).
c.n.d.

Rozpatrzmy współzależności konfiguracji dendrytów 1-, 2- i 3-drzewowych w pełnym unigrafie czwornika, ilustrując rozważania dendrytami k -drzewowymi czwornika (czterokońcówkowego) o $v=4$ wierzchołkach (Tab. 2).

W pełnym unigrafie wszystkie dendryty 2-drzewowe typu $T_{1,0}^v(t)$ powstają tylko z tych dendrytów 1-drzewowych $T_1^v(t)$ które zawierają gałąź między wierzchołkami w_1 i w_0 . Po „zwarciu” (opuszczeniu) tej gałęzi (w naszym przypadku $e_1 \in \{w_1, w_0\}$) z dendrytów 1-drzewowych $T_1^v(t)$ powstają dendryty 2-drzewowe typu $T_{1,0}^v(t)$. Analogicznie wszystkie dendryty 2-drzewowe typu $T_{2,3}^v(t)$ powstają tylko z tych dendrytów 1-drzewowych $T_1^v(t)$ które zawierają gałąź między wierzchołkami w_2 i w_3 . Po „zwarciu” (opuszczeniu) tej gałęzi (w naszym przypadku $e_4 \in \{w_2, w_3\}$) z dendrytów 1-drzewowych $T_1^v(t)$ powstają dendryty 2-drzewowe typu $T_{2,3}^v(t)$. Dendryty 2-drzewowe typu $T_{12,30}^v(t)$, $T^{v13,20}(t)$ są dendrytami „tranzytowymi” przenoszącymi informację z wejścia (wierzchołki $\{w_1, w_0\}$) do wyjścia (wierzchołki $\{w_2, w_3\}$).

Jak łatwo zauważyć, część dendrytów 2-drzewowych typu $T_{1,0}^v(t)$, ma analogiczną strukturę jak dendryty 2-drzewowe typu $T_{2,3}^v(t)$ (Tab. 2).

Najprostszy sposób generacji dendrytów 2-drzewowych typu $T_{12,30}^v(t)$ i $T_{13,20}^v(t)$ w pełnym unigrafie można zapisać w postaci:

$$T_{1,0}^v(t) = T_{2,3}^v(t) \Rightarrow T_{12,30}^v(t) \vee T_{13,20}^v(t) \quad (11)$$

Dendryty 2-drzewowe typu $T_{12,30}^v(t)$ przenoszącymi informację z wejścia do wyjścia z znakiem „+”. Dendryty 2-drzewowe typu $T_{13,20}^v(t)$ przenoszącymi informację z wejścia do wyjścia z znakiem „-”.

Wszystkie dendryty 3-drzewowe typu $T_{1,2,30}^v(t)$, $T_{12,3,0}^v(t)$, $T_{1,3,20}^v(t)$, $T_{13,2,0}^v(t)$ powstają z dendrytów 2-drzewowych typu $T_{12,30}^v(t)$, $T_{13,20}^v(t)$ po „opuszczeniu”, pewnych gałęzi.

Dendryty 3-drzewowe $T_{1,2,30}^v(t)$, $T_{12,3,0}^v(t)$ powstają z dendrytów 2-drzewowych $T_{12,30}^v(t)$ po „opuszczeniu” gałęzi $e_2 \in \{w_1, w_2\}$ lub $e_6 \in \{w_3, w_0\}$, co można zapisać w postaci:

$$T_{12,30}^v(t) - e_2 \Rightarrow T_{1,2,30}^v(t) \quad (12)$$

$$T_{12,30}^v(t) - e_6 \Rightarrow T_{12,3,0}^v(t) \quad (13)$$

Dendryty 3-drzewowe $T_{1,3,20}^v(t)$, $T_{13,2,0}^v(t)$ powstają z dendrytów 2-drzewowych $T_{13,20}^v(t)$ po „opuszczeniu” gałęzi $e_5 \in \{w_1, w_3\}$ lub $e_3 \in \{w_2, w_0\}$, co można zapisać w postaci:

$$T_{13,20}^v(t) - e_5 \Rightarrow T_{1,3,20}^v(t) \quad (14)$$

$$T_{13,20}^v(t) - e_3 \Rightarrow T_{13,2,0}^v(t) \quad (15)$$

Z powyższych rozważań wynika, że obecność lub brak w dendrycie 1-drzewowym gałęzi brzegowych $e_1 \in \{w_1, w_0\}$ i lub $e_4 \in \{w_2, w_3\}$, nie ma wpływu na strukturę dendrytów 2-drzewowych i dendrytów 3-drzewowych. Dendryty 1-drzewowe tworzące zgodnie z równaniem (10) struktury czwórników, można podzielić na cztery grupy, ze względu na „uczestnictwo” elementu wejściowego $\{e_1\}$ i elementu wyjściowego $\{e_4\}$ w dendrycie $T_1^v(t)$, oraz sposobu tworzenia dendrytów 2-drzewowych typu $T_{1,0}^v(t)$, $T_{2,3}^v(t)$, $T_{12,30}^v(t)$, $T_{13,20}^v(t)$, $T_{1,2,30}^v(t)$, $T_{12,3,0}^v(t)$, $T_{1,3,20}^v(t)$, $T_{13,2,0}^v(t)$, (Tab. 3).

Grupę A tworzą te dendryty 1-drzewowe które zawierają element wejściowy $\{e_1\}$ lub element wyjściowy $\{e_4\}$, oraz utworzone z nich dendryty 2-drzewowe $T_{1,0}^v(t)$ i $T_{2,3}^v(t)$ są sobie równe. Grupę tę dzielimy jeszcze na dwie podgrupy A_1 i A_2 (Tab. 3 i 2). Dendryty 1-drzewowe $T_1^v(t)$ podgrupy A_1 zawierają $v-2$ elementów, które uczestniczą w tworzeniu dendrytów 2-drzewowych typu $T_{1,0}^v(t)$, $T_{2,3}^v(t)$, $T_{12,30}^v(t)$. Dendryty $T_1^v(t)$ podgrupy A_2 zawierają $v-2$ elementów, które uczestniczą w tworzeniu $T_{1,0}^v(t)$, $T_{2,3}^v(t)$, $T_{13,20}^v(t)$.

Grupę B tworzą te dendryty 1-drzewowe które zawierają element wejściowy $\{e_1\}$ lub element wyjściowy $\{e_4\}$, ale utworzone z nich dendryty 2-drzewowe $T_{1,0}^v(t)$ i $T_{2,3}^v(t)$ nie są sobie równe. Grupę tę dzielimy na dwie podgrupy B_1 i B_2 . Podgrupę B_1 tworzą te dendryty 1-drzewowe, które zawierają element wejściowy $\{e_1\}$. Zawierają one $v-2$ elementy które uczestniczą w tworzeniu tylko $T_{1,0}^v(t)$, oraz $v-3$ elementy które uczestniczą w tworzeniu $T_{2,3}^v(t)$. Podgrupę B_2 tworzą te dendryty 1-drzewowe, które zawierają element wyjściowy $\{e_4\}$. Zawierają one $v-2$ elementy które uczestniczą w tworzeniu tylko $T_{2,3}^v(t)$, oraz $v-3$ elementy które uczestniczą w tworzeniu $T_{1,0}^v(t)$ (Tab. 3 i 2.2). Dendryty $T_1^v(t)$ grupy B nie tworzą samodzielnie dendrytów $T_{12,30}^v(t)$ i $T_{13,20}^v(t)$. Zawierają bowiem tylko $l=v-3 \div 1$ gałęzi mogących uczestniczyć w tworzeniu dendrytu $T_{12,30}^v(t)$ i $l=1 \div v-3$ gałęzi mogących uczestniczyć w tworzeniu dendrytu $T_{13,20}^v(t)$ (Tab. 3 i 2).

Grupę C tworzą te dendryty 1-drzewowe które zawierają element wejściowy $\{e_1\}$ i element wyjściowy $\{e_4\}$. Dendryty $T_1^v(t)$ grupy C zawierają $v-3$ elementy, które uczestniczą w tworzeniu „niepełnych” dendrytów $T_{1,0}^v(t)$ i $T_{2,3}^v(t)$ (Tab. 3 i 2). Dendryty $T_1^v(t)$ grupy C nie tworzą samodzielnie dendrytów $T_{12,30}^v(t)$ i $T_{13,20}^v(t)$. Zawierają bowiem tylko $l=v-3 \div 0$ gałęzi mogących uczestniczyć w tworzeniu dendrytu $T_{12,30}^v(t)$ i $l=0 \div v-3$ gałęzi mogących uczestniczyć w tworzeniu dendrytu $T_{13,20}^v(t)$ (Tab. 3 i 2). Dendryty $T_1^v(t)$ grupy C zawierają $v-3$ elementy tworzące „niepełne” dendryty $T_{12,30}^v(t)$ będziemy oznaczali C^+ , a dendryty $T_1^v(t)$ zawierające $v-3$ elementy tworzące „niepełne” dendryty $T_{13,20}^v(t)$ będziemy oznaczali C^- .

Grupę D tworzą te dendryty 1-drzewowe które nie zawierają elementu wejściowego $\{e_1\}$ i elementu wyjściowego $\{e_4\}$. W tworzeniu dendrytów $T_{1,0}^v(t)$, $T_{2,3}^v(t)$, $T_{12,30}^v(t)$ i $T_{13,20}^v(t)$ uczestniczą wszystkie $v-1$ elementy dendrytów $T_1^v(k t)$. Dendryty $T_1^v(t)$ grupy D zawsze tworzą $t_{1,0}^v(t)$, $T_{2,3}^v(t)$. Dendryty $T_1^v(t)$

Tabela 3

Liczby gałęzi uczestniczących w tworzeniu dendrytów k -drzewowych $T_k^v(t_i)$ których podstawą jest:

a) Dendryt $T_1^v(t_i)$

		T_1^4	$T_{1,0}^4$	$T_{2,3}^4$	$T_{12,30}^4$	$T_{13,20}^4$	$T_{12,3,0}^4$ $T_{1,2,30}^4$	$T_{13,2,0}^4$ $T_{1,3,20}^4$
$LT^{4k}(t_i)$		$v-1$	$v-2$	$v-2$	$v-2$	$v-2$	$v-3$	$v-3$
A	A_1	$v-1$	$v-2$	$v-2$	$v-2$		$v-3$	
	A_2					$v-2$		$v-3$
B	B_1	$v-1$	$v-2$	$v-3$	$v-3$	1	$v-3$	1
	B_2		$v-3$	$v-2$	$\vdots$ 1	$\vdots$ $v-3$	$\vdots$ 1	$\vdots$ $v-3$
C	C^+	$v-1$	$v-2$	$v-2$	$v-3$	0	$v-3$	0
	C^-				$\vdots$ 0	$\vdots$ $v-3$	$\vdots$ 0	$\vdots$ $v-3$
D	D^+	$v-1$	$v-2$	$v-2$	$v-2$	1	$v-3$	1
	D^-				$\vdots$ 1	$\vdots$ $v-2$	$\vdots$ 1	$\vdots$ $v-3$

b) Dendryt $T_{1,0}^v(t_i)$

		T_1^4	$T_{1,0}^4$	$T_{2,3}^4$	$T_{12,30}^4$	$T_{13,20}^4$	$T_{12,3,0}^4$ $T_{1,2,30}^4$	$T_{13,2,0}^4$ $T_{1,3,20}^4$
$LT^{4k}(t_i)$		$v-1$	$v-2$	$v-2$	$v-2$	$v-2$	$v-3$	$v-3$
A	A_1	$v-2$	$v-2$	$v-2$	$v-2$		$v-3$	
	A_2					$v-2$		$v-3$
B	B_1	$v-2$	$v-2$	$v-3$	$v-3$	1	$v-3$	1
	B_2		$v-3$	$v-3$	$\vdots$ 1	$\vdots$ $v-3$	$\vdots$ 1	$\vdots$ $v-3$
C	C^+	$v-2$	$v-2$	$v-3$	$v-3$	0	$v-3$	0
	C^-				$\vdots$ 0	$\vdots$ $v-3$	$\vdots$ 0	$\vdots$ $v-3$

c) Dendryt $T_{2,3}^v(t_i)$

		T_1^4	$T_{1,0}^4$	$T_{2,3}^4$	$T_{12,30}^4$	$T_{13,20}^4$	$T_{12,3,0}^4$ $T_{1,2,30}^4$	$T_{13,2,0}^4$ $T_{1,3,20}^4$
$LT^{4k}(t_i)$		$v-1$	$v-2$	$v-2$	$v-2$	$v-2$	$v-3$	$v-3$
A	A_1	$v-2$	$v-2$	$v-2$	$v-2$		$v-3$	
	A_2					$v-2$		$v-3$
B	B_1	$v-2$	$v-2$	$v-3$	$v-3$	1	$v-3$	1
	B_2		$v-3$	$v-3$	$\vdots$ 1	$\vdots$ $v-3$	$\vdots$ 1	$\vdots$ $v-3$
C	C^+	$v-2$	$v-3$	$v-2$	$v-3$	0	$v-3$	0
	C^-				$\vdots$ 0	$\vdots$ $v-3$	$\vdots$ 0	$\vdots$ $v-3$

grupy **D** mogą samodzielnie tworzyć $T_{12,30}^v(t_i)$ i $T_{13,20}^v(t_i)$. Dendryty $T_1^v(t_i)$ grupy **D** zawierające $v-2$ elementy tworzące dendryty $T_{12,30}^v(t_i)$ będziemy oznaczali $\mathbf{D}^+$, a dendryty $T_1^v(t_i)$ zawierające $v-2$ elementy tworzące dendryty $T_{13,20}^v(t_i)$ będziemy oznaczali $\mathbf{D}^-$ (Tab. 3 i 2).

Ponieważ obecność lub brak w dendrycie 1-drzewowym gałęzi brzegowych e_1 i e_4 , nie ma wpływu na strukturę dendrytów 2-drzewowych $T_{1,0}^v(t_i)$, $T_{2,3}^v(t_i)$, $T_{12,30}^v(t_i)$ i $T_{13,20}^v(t_i)$, czwórnik w ogólnym przypadku można utworzyć jako:

$$G^v(t_1, t_2) = T_{k1}^v(t_1) \cup T_{k2}^v(t_2) \quad (16)$$

gdzie:

$T_{ki}^v(t_i)$ – dendryt k -drzewowy ($T_1^v(t_i) \vee T_{1,0}^v(t_i) \vee T_{2,3}^v(t_i)$) składający się tylko z elementów typu $\{t_i\}$.

Można sformułować jeszcze trzy twierdzenia analogiczne do twierdzenia 1 i udowodnić, że istnieją jeszcze trzy typy minimalno-elementowych struktur kanonicznych czwórników względem parametru $H_1^v(t_1, t_2)$. Typy minimalno-elementowych struktur kanonicznych czwórników realizujących wielomian $H_1^v(t_1, t_2)$ o maksymalnym stopniu zmiennej t_1 , przy minimalnej liczbie elementów $\{t_1\}$ i $\{t_2\}$ zestawiono w tabeli 4a.

3.b. STRUKTURY MINIMALNO-ELEMENTOWE REALIZUJĄCE $H_{1,0}^v(t_1, t_2)$

Zgodnie z zależnością 27 pracy [9], poszczególne składniki wielomianu $kH_{1,0}^v(t_1, t_2)$, powstają jako iloczyny admitancji elementów tworzących dendryty $kT_{1,0}^v(t_1, t_2)$ w strukturze $G^v(t_1, t_2)$ czwórnika. Dendryty k -drzewowe $T_{k1}^v(t_1)$ tworzące zgodnie z zależnością 16 struktury czwórników, można ze względu na sposób uczestnictwa ich w $H_{1,0}^v(t_1, t_2)$ podzielić na trzy klasy.

Do pierwszej klasy zaliczamy dendryty 1-drzewowe $T_1^v(t_i)$ grupy **D**. Tworzą one zawsze więcej niż jeden dendryt $T_{1,0}^v(t_i)$, a zatem dla $t_i = t_1$, $LT_{1,0}^v(t_1^{v-2}, t_2^0) > 1$, oraz dla $t_i = t_2$, $LT_{1,0}^v(t_1^0, t_2^{v-2}) > 1$.

Do drugiej klasy zaliczamy dendryty 1-drzewowe $T_1^v(t_i)$ grupy **A** $\vee$ **B**₁ $\vee$ **C**, oraz dendryty 2-drzewowe $T_{1,0}^v(t_i)$ [**A** $\vee$ **B**₁ $\vee$ **C**]. Tworzą one zawsze tylko jeden dendryt $T_{1,0}^v(t_i)$ a zatem dla $t_i = t_1$, $LT_{1,0}^v(t_1^{v-2}, t_2^0) = 1$, oraz dla $t_i = t_2$, $LT_{1,0}^v(t_1^0, t_2^{v-2}) = 1$.

Do trzeciej klasy zaliczamy $T_1^v(t_i)$ [**B**₂], oraz $T_{2,3}^v(t_i)$ [**B**₂]. Nie tworzą one samodzielnie dendrytu $t_{1,0}^v(t_i)$, a zatem dla $t_i = t_1$, $LT_{1,0}^v(t_1^{v-2}, t_2^0) = 0$, oraz dla $t_i = t_2$, $LT_{1,0}^v(t_1^0, t_2^{v-2}) = 0$. Mając po trzy klasy dendrytów $T_{k1}^v(t_1)$ i $T_{k2}^v(t_2)$, można zgodnie z równaniem 16 utworzyć dziewięć (3^2) typów minimalno-elementowych struktur kanonicznych czwórników. Minimalno-elementowe typy struktur kanonicznych czwórników, realizujące wielomian $H_{1,0}^v(t_1, t_2)$, uszeregowane względem postaci typu uporządkowanego zbioru $\langle LT_{1,0}^v(t_1, t_2) \rangle$ zestawiono w tabeli 4b.

3.c. STRUKTURY MINIMALNO-ELEMENTOWE REALIZUJĄCE $H_{2,3}^v(t_1, t_2)$

Analogicznie można wykazać, że istnieje dziewięć typów minimalno-elementowych struktur kanonicznych czwórników realizujących wielomian $H_{2,3}^v(t_1, t_2)$.

W tabeli 4 zestawiono warunki jakie muszą spełniać dendryty k -drzewowe

Typy minimalno-elementowych struktur kanonicznych czwórników realizujących:

a) $H_1^V(t_1, t_2)$

Lp	$G^V(t_1, t_2) =$	$\langle LT_1^V(t_1, t_2) \rangle$
1	$T_1^V(t_1)[AvBvCvD] \cup T_1^V(t_2)[AvBvCvD]$	$\langle 1, I, \dots, I, 1 \rangle$
2	$T_2^V(t_1)[AvBvC] \cup T_1^V(t_2)[AvBvCvD]$	$\langle 0, I, \dots, I, 1 \rangle$
3	$T_1^V(t_1)[AvBvCvD] \cup T_2^V(t_2)[AvBvC]$	$\langle 1, I, \dots, I, 0 \rangle$
4	$T_2^V(t_1)[AvBvC] \cup T_2^V(t_2)[AvBvC]$	$\langle 0, I, \dots, I, 0 \rangle$

gdzie:

$$T_2^V(t_1, t_2) = T_{1,0}^V(t_1, t_2) \vee T_{2,3}^V(t_1, t_2)$$

 $T_k^V(t_i)[AvBvC]$ – dendryt 2-drzewowy utworzony z dendrytów 1-drzewowych grupy $[AvBvC]$.
b) $H_{1,0}^V(t_1, t_2)$ i $H_{2,3}^V(t_1, t_2)$

Lp	$G^V(t_1, t_2) =$	$\langle LT_2^V(t_1, t_2) \rangle$	$G^V(t_1, t_2) =$
1	$T_1^V(t_1)[D] \cup T_1^V(t_2)[D]$	$\langle I, I, \dots, I, I \rangle$	$T_1^V(t_1)[D] \cup T_1^V(t_2)[D]$
2	$T_k^V(t_1)[E_1] \cup T_1^V(t_2)[D]$	$\langle 1, I, \dots, I, I \rangle$	$T_k^V(t_1)[E_2] \cup T_1^V(t_2)[D]$
3	$T_1^V(t_1)[D] \cup T_k^V(t_2)[E_1]$	$\langle I, I, \dots, I, 1 \rangle$	$T_1^V(t_1)[D] \cup T_k^V(t_2)[E_2]$
4	$T_k^V(t_1)[E_1] \cup T_k^V(t_2)[E_1]$	$\langle 1, I, \dots, I, 1 \rangle$	$T_1^V(t_1)[E_2] \cup T_k^V(t_2)[E_2]$
5	$T_{2,3}^V(t_1)[B_2] \cup T_1^V(t_2)[D]$	$\langle 0, I, \dots, I, I \rangle$	$T_{1,0}^V(t_1)[B_1] \cup T_1^V(t_2)[D]$
6	$T_{2,3}^V(t_1)[B_2] \cup T_k^V(t_2)[E_1]$	$\langle 0, I, \dots, I, 1 \rangle$	$T_{1,0}^V(t_1)[B_1] \cup T_k^V(t_2)[E_2]$
7	$T_1^V(t_1)[D] \cup T_{2,3}^V(t_2)[B_2]$	$\langle I, I, \dots, I, 0 \rangle$	$T_1^V(t_1)[D] \cup T_{1,0}^V(t_2)[B_1]$
8	$T_k^V(t_1)[E_1] \cup T_{2,3}^V(t_2)[B_2]$	$\langle 1, I, \dots, I, 0 \rangle$	$T_k^V(t_1)[E_2] \cup T_{1,0}^V(t_2)[B_1]$
9	$T_{2,3}^V(t_1)[B_2] \cup T_{2,3}^V(t_2)[B_2]$	$\langle 0, I, \dots, I, 0 \rangle$	$T_{1,0}^V(t_1)[B_1] \cup T_{1,0}^V(t_2)[B_1]$

gdzie: $\langle LT_2^V(t_1, t_2) \rangle = \langle LT_{1,0}^V(t_1, t_2) \rangle \vee \langle LT_{2,3}^V(t_1, t_2) \rangle$

$$T_k^V(t_i)[E_1] = [T_1^V(t_i) \vee T_{1,0}^V(t_i)][AvB_1vC]$$

$$T_k^V(t_i)[E_2] = [T_1^V(t_i) \vee T_{2,3}^V(t_i)][AvB_2vC]$$

d) $H_{12,30}^V(t_1, t_2)$ i $H_{13,20}^V(t_1, t_2)$

Lp	$\langle LT_{12,30}^V(t_1, t_2) \rangle$	$G^V(t_1, t_2) =$	$\langle LT_{13,20}^V(t_1, t_2) \rangle$
1	$\langle 1, I, \dots, I, 1 \rangle$	$T_k^V(t_1)[A_1] \cup T_k^V(t_2)[A_1]$	$\langle 0, 0, \dots, 0, 0 \rangle$
2	$\langle 1, I, \dots, 0, 0 \rangle$	$T_k^V(t_1)[A_1] \cup T_k^V(t_2)[A_2]$	$\langle 0, 0, \dots, I, 1 \rangle$
3	$\langle 1, I, \dots, I, 0 \rangle$ $\langle 1, I, \dots, I, 0 \rangle$	$T_k^V(t_1)[A_1] \cup T_k^V(t_2)[BvC]$	$\langle 0, 0, \dots, I, 0 \rangle$ $\langle 0, 0, \dots, I, 0 \rangle$
4	$\langle 1, I, \dots, I, I \rangle$ $\langle 1, I, \dots, I, 0 \rangle$	$T_k^V(t_1)[A_1] \cup T_k^V(t_2)[D]$	$\langle 0, I, \dots, I, 0 \rangle$ $\langle 0, I, \dots, I, I \rangle$
5	$\langle 0, 0, \dots, I, 1 \rangle$	$T_k^V(t_1)[A_2] \cup T_k^V(t_2)[A_1]$	$\langle 1, I, \dots, 0, 0 \rangle$
6	$\langle 0, 0, \dots, 0, 0 \rangle$	$T_k^V(t_1)[A_2] \cup T_k^V(t_2)[A_2]$	$\langle 1, I, \dots, I, 1 \rangle$
7	$\langle 0, 0, \dots, I, 0 \rangle$ $\langle 0, 0, \dots, I, 0 \rangle$	$T_k^V(t_1)[A_2] \cup T_k^V(t_2)[BvC]$	$\langle 1, 0, \dots, I, 0 \rangle$ $\langle 1, I, \dots, I, 0 \rangle$
8	$\langle 0, I, \dots, I, I \rangle$ $\langle 0, I, \dots, I, I \rangle$	$T_k^V(t_1)[A_2] \cup T_k^V(t_2)[D]$	$\langle 1, I, \dots, I, I \rangle$ $\langle 1, I, \dots, I, I \rangle$
9	$\langle 0, I, \dots, I, 1 \rangle$ $\langle 0, I, \dots, I, 1 \rangle$	$T_k^V(t_1)[BvC] \cup T_k^V(t_2)[A_1]$	$\langle 0, I, \dots, I, 0 \rangle$ $\langle 0, I, \dots, I, 0 \rangle$

c.d. tabeli 4

10	$\langle 0, I, \dots, I, 0 \rangle$ $\langle 0, I, \dots, I, 0 \rangle$	$T_k^v(t_1)[BvC] \cup T_k^v(t_2)[A_2]$	$\langle 0, I, \dots, I, 1 \rangle$ $\langle 0, I, \dots, I, 1 \rangle$
11	$\langle 0, I, \dots, I, 0 \rangle$ $\langle 0, I, \dots, I, 0 \rangle$	$T_k^v(t_1)[BvC] \cup T_k^v(t_2)[BvC]$	$\langle 0, I, \dots, I, 0 \rangle$ $\langle 0, I, \dots, I, 0 \rangle$
12	$\langle 0, I, \dots, I, 1 \rangle$ $\langle 0, I, \dots, I, 1 \rangle$	$T_k^v(t_1)[BvC] \cup T_k^v(t_2)[D]$	$\langle 0, I, \dots, I, 1 \rangle$ $\langle 0, I, \dots, I, 1 \rangle$
13	$\langle 1, I, \dots, I, 1 \rangle$ $\langle 0, I, \dots, I, 1 \rangle$	$T_k^v(t_1)[D] \cup T_k^v(t_2)[A_1]$	$\langle 0, I, \dots, I, 0 \rangle$ $\langle 1, I, \dots, I, 0 \rangle$
14	$\langle 1, I, \dots, I, 0 \rangle$ $\langle 0, I, \dots, I, 0 \rangle$	$T_k^v(t_1)[D] \cup T_k^v(t_2)[A_2]$	$\langle 0, I, \dots, I, 1 \rangle$ $\langle 1, I, \dots, I, 1 \rangle$
15	$\langle 1, I, \dots, I, 0 \rangle$ $\langle 0, I, \dots, I, 0 \rangle$	$T_k^v(t_1)[D] \cup T_k^v(t_2)[BvC]$	$\langle 0, I, \dots, I, 0 \rangle$ $\langle 1, I, \dots, I, 0 \rangle$
16	$\langle 1, I, \dots, I, 1 \rangle$ $\langle 1, I, \dots, I, 1 \rangle$	$T_k^v(t_1)[D] \cup T_k^v(t_2)[D]$	$\langle 1, I, \dots, I, 1 \rangle$ $\langle 1, I, \dots, I, 1 \rangle$

Kursywą oznaczono możliwość istnienia w uporządkowanym zbiorze $\langle LT_k^v(t_1, t_2) \rangle$ danej liczby $LT_k^v(t_1^{v-k-1}, t_2)$.

$T_{kl}^v(t_i)$, aby utworzone struktury czwórników, były minimalno-elementowe względem parametru $H_{2,3}^v(t_1, t_2)$.

Jak wynika z powyższych rozważań w wielomianach $H_1^v(t_1, t_2)$, $H_{1,0}^v(t_1, t_2)$ i $H_{2,3}^v(t_1, t_2)$, tylko współczynniki $Ap_{v-k,0}$ i $a_{0,v-k}$ mogą przyjmować wartość zero. Pozostałe współczynniki rozpatrywanych wielomianów są zawsze większe od zera.

3.d. STRUKTURY MINIMALNO-ELEMENTOWE REALIZUJĄCE $H_{12,30}^v(t_1, t_2)$ i $H_{13,20}^v(t_1, t_2)$

Należy zwrócić uwagę, że w tworzeniu dendrytów $T_{12,30}^v(t_1, t_2)$ nie biorą udziału zawsze elementy $e_1 \in \{w_1, w_0\}$, $e_4 \in \{w_2, w_3\}$, $e_3 \in \{w_0, w_2\}$ i $e_5 \in \{w_1, w_3\}$ czyli wszystkie kombinacje gałęzi incydentnych do dwóch zbiorów wierzchołków $\{w_1, w_2\}$ i $\{w_3, w_0\}$. Natomiast w tworzeniu dendrytów $T_{13,20}^v(t_1, t_2)$ nie biorą udziału elementy $e_1 \in \{w_1, w_0\}$, $e_4 \in \{w_2, w_3\}$, $e_2 \in \{w_1, w_2\}$ i $e_6 \in \{w_0, w_3\}$ czyli wszystkie kombinacje gałęzi indycentnych do dwóch zbiorów wierzchołków $\{w_1, w_3\}$ i $\{w_2, w_0\}$.

Dendryty k-drzewowe $T_{kl}^v(t_i)$ tworzące zgodnie z zależnością 16 struktury czwórników, można ze względu na sposób uczestnictwa ich w tworzeniu wielomianu $H_{12,30}^v(t_1, t_2)$ i $H_{13,20}^v(t_1, t_2)$ podzielić na cztery rozmyte klasy.

Do pierwszej klasy zaliczamy dendryty $T_k^v(t_i)$ grupy A_1 (gdzie: $T_k^v(t_i) \in T_1^v(t_i) \vee T_{1,0}^v(t_i) \vee T_{2,3}^v(t_i)$). Tworzą one zawsze dendryt $T_{12,30}^v(t_i)$, a więc dla $t_i = t_1$, $LT_{12,30}^v(t_1^{v-2}, t_2^0) = 1$, oraz dla $t_i = t_2$, $LT_{12,30}^v(t_1^0, t_2^{v-2}) = 1$. Dendryty $T_k^v(t_i)$ grupy A_1 nigdy nie tworzą samodzielnie $T_{13,20}^v(t_1, t_2)$, a więc dla $t_i = t_1$, $LT_{13,20}^v(t_1^{v-2}, t_2^0) = 0$ i $LT_{13,20}^v(t_1^{v-3}, t_2^1) = 0$, oraz dla $t_i = t_2$, $LT_{13,20}^v(t_1^0, t_2^{v-2}) = 0$ i $LT_{13,20}^v(t_1^1, t_2^{v-3}) = 0$.

Do drugiej klasy zaliczamy dendryty $T_k^v(t_i)$ grupy A_1 . Tworzą one zawsze dendryt $T_{p_{13,20}}^v(t_i)$, a więc dla $t_i = t_1$, $LT_{13,20}^v(t_1^{v-2}, t_2^0) = 1$, oraz dla $t_i = t_2$,

$LT_{13,20}^v(t_1^0, t_2^{v-2})=1$. Dendryty $T_k^v(t_i)$ grupy A_1 nigdy nie tworzą samodzielnie $T_{12,30}^v(t_1^1, t_2)$, a więc dla $t_i=t_1$, $LT_{12,30}^v(t_1^{v-2}, t_2^0)=0$ i $LT_{12,30}^v(t_1^{v-3}, t_2^1)=0$, oraz dla $t_i=t_2$, $LT_{12,30}^v(t_1^0, t_2^{v-2})=0$ i $LT_{12,30}^v(t_1^1, t_2^{v-3})=0$.

Do trzeciej klasy zaliczamy dendryty $T_k^v(t_i)$ $[B \vee C]$ (gdzie: $T_k^v(t_i) \in T_1^v(t_i) \vee T_{2,3}^v(t_i)$). Dendryty $T_k^v(t_i)$ grupy B i C nie tworzą samodzielnie dendrytu $T_{12,30}^v(t_i)$ ani $T_{13,20}^v(t_i)$. Zawierają bowiem w B_1 i C^+ maksymalnie $v-3$ gałęzie mogące uczestniczyć w tworzeniu dendrytu $T_{12,30}^v(t_i)$ oraz w B_2 i C maksymalnie $v-3$ gałęzie mogące uczestniczyć w tworzeniu dendrytu $T_{13,20}^v(t_i)$ (Tab. 3). Dla $t_i=t_1$ zawsze $LT_{12,30}^v(t_1^{v-2}, t_2^0)=0$ i $LT_{13,20}^v(t_1^{v-2}, t_2^0)=0$, oraz dla $t_i=t_2$, $LT_{12,30}^v(t_1^0, t_2^{v-2})=0$, i $LT_{13,20}^v(t_1^0, t_2^{v-2})=0$.

Do czwartej klasy zaliczamy $T_1^v(t_i)$ $[D]$. Dendryty $T_1^v(t_i)$ $[D]$ mogą tworzyć samodzielnie dendryt $T_{12,30}^v(t_i)$ lub $T_{13,20}^v(t_i)$. Zawierają w $[D^+]$ $v-2$ gałęzie uczestniczące w tworzeniu dendrytu $T_{12,30}^v(t_i)$ [wtedy dla $t_i=t_1$ $LT_{12,30}^v(t_1^{v-2}, t_2^0)=1$ i $LT_{13,20}^v(t_1^{v-2}, t_2^0)=0$], oraz (D^-) $v-2$ gałęzie uczestniczące w tworzeniu dendrytu $T_{13,20}^v(t_i)$ [wtedy dla $t_i=t_1$ $LT_{12,30}^v(t_1^{v-2}, t_2^0)=0$ i $LT_{13,20}^v(t_1^{v-2}, t_2^0)=1$] (Tab. 3).

Mając po cztery klasy dendrytów $T_{k1}^v(t_1)$ i $T_{k2}^v(t_2)$, można zgodnie z równaniem 16 utworzyć szesnaście (4^2) grup typów minimalno-elementowych struktur kanonicznych czwórników realizujących wielomiany $H_{12,30}^v(t_1, t_2)$ i $H_{13,20}^v(t_1, t_2)$. Minimalno-elementowe typy struktur kanonicznych czwórników, realizujące wielomian $H_{12,30}^v(t_1, t_2)$ i $H_{13,20}^v(t_1, t_2)$, zestawiono w tabeli 4d.

3.c. STRUKTURY MINIMALNO-ELEMENTOWE REALIZUJĄCE $H_{12,30}^v(t_1, t_2)$, $H_{12,30}^v(t_1, t_2)$, $H_{13,20}^v(t_1, t_2)$ i $H_{13,20}^v(t_1, t_2)$

W pełnym unigrafie z dendrytów 2-drzewowych $T_{12,30}^v(t)$ po „opuszczeniu” gałęzi $e_2 \in \{w_1, w_2\}$ powstają dendryty 3-drzewowe $T_{1,2,30}^v(t)$, natomiast po „opuszczeniu” gałęzi $e_6 \in \{w_3, w_0\}$ powstają dendryty 3-drzewowe $T_{12,3,0}^v(t)$.

Analogicznie w pełnym unigrafie z dendrytów 2-drzewowych $T_{13,20}^v(t)$ po „opuszczeniu” gałęzi $e_3 \in \{w_1, w_3\}$ powstają dendryty 3-drzewowe $T_{1,3,20}^v(t)$, natomiast po „opuszczeniu” gałęzi $e_3 \in \{w_2, kw_0\}$ powstają dendryty 3-drzewowe $T_{13,2,0}^v(t)$. Zgodnie z tabelą 3 dendryty 3-drzewowe $T_{1,2,30}^v(t)$, $T_{12,3,0}^v(t)$ tworzą analogicznie jak dendryty 2-drzewowe $T_{12,30}^v(t)$ cztery klasy, i podobnie dendryty 3-drzewowe $T_{1,3,20}^v(t)$, $T_{13,2,0}^v(t)$ tworzą analogicznie jak dendryty 2-drzewowe $T_{13,20}^v(t)$ cztery klasy.

4. PRZYKŁAD SYNTEZY TOPOLOGICZNEJ STRUKTUR KANONICZNYCH CZWÓRNIKÓW

Jako ilustrację rozważań niniejszej pracy, niżej przedstawiono wyniki syntezy topologicznej struktur czwórników o $v=4$ wierzchołkach. Strukturę kanoniczną $G^4(t_1, t_2)$, będziemy charakteryzować uporządkowanymi zbiorami $\langle LT_k^v(t_1, t_2) \rangle$. W tabeli 5 zestawiono przykładowo struktury kanoniczne czwórników (cztero-

Tabela 5.

Typy struktur kanonicznych czwórników typu

$$G^4(t_1, t_2) = T_1^4(t_1) \cup T_1^4(t_2)$$

a	$G^4(t_1)$	$LT_k^4(t_1, t_2)$					
b	$G^4(t_2)$	LT_1^4	$LT_{1,0}^4$	$LT_{2,3}^4$	$LT_{13,20}^4$	$LT_{12,30}^4$	LT_3^4
Nr	$G^4(t_1, t_2)$		$LT_{2,3}^4$	$LT_{1,0}^4$	$LT_{12,30}^4$	$LT_{13,20}^4$	
A	B	C	D	E	F	G	H
a	1,1,1,1,1,1	16	8	8	1	1	4
b	1,1,1,1,1,1	16	8	8	1	1	4
1	001011 110100	1.7.7.1	1.5.2	1.5.2	0.1.0	0.0.1	1.3
2	010101 101010	1.7.7.1	1.4.3	3.4.1	0.0.1	1.0.0	2.2
3	010110 101001	1.7.7.1	2.5.1	2.5.1	0.1.0	1.0.0	3.1
a	0,1,1,1,1,1	8	5	5	0	1	3
b	0,1,1,1,1,2	13	5	8	0	1	3
4	001011 010101	1.5.6.1	1.3.1	1.4.3	0.0.0	0.0.1	1.2
5	001101 010011	1.6.5.1	1.3.1	1.5.2	0.0.0	0.1.0	1.2
6	010011 001101	1.5.6.1	1.3.1	2.5.1	0.0.0	0.1.0	2.1
7	010101 001011	1.6.5.1	1.3.1	3.4.1	0.0.0	1.0.0	2.1
b	0,1,1,1,2,1	12	7	7	0	1	4
8	001011 010110	1.5.5.1	1.4.2	1.4.2	0.0.0	0.0.1	1.3
9	001110 010011	1.5.5.1	2.4.1	1.4.2	0.0.0	0.1.0	2.2
10	010110 001011	1.5.5.1	2.4.1	2.4.1	0.0.0	1.0.0	3.1
b	0,1,1,2,1,1	13	7	8	0	2	4
11	001101 010110	1.6.5.1	1.4.2	1.5.2	0.0.0	0.1.1	1.3
12	001110 010101	1.5.6.1	2.4.1	1.4.3	0.0.0	0.1.1	2.2
13	010101 001110	1.6.5.1	1.4.2	3.4.1	0.0.0	1.1.0	2.2
14	010110 001101	1.5.6.1	2.4.1	2.5.1	0.0.0	1.1.0	3.1
a	1,1,0,1,1,1	8	5	8	1	1	4
b	1,1,0,1,1,2	12	4	12	1	1	4

c.d. – tabela 5.

A	B	C	D	E	F	G	H
15	0 1 0 0 1 1 1 0 0 1 0 1	1.5.5.1	1.2.1	2.8.2	0.1.0	0.1.0	2.2
16	0 1 0 1 0 1 1 0 0 0 1 1	1.5.5.1	1.2.1	3.6.3	1.0.0	0.0.1	2.2
b	1,1,0,1,2,1	13	6	11	2	1	5
17	0 1 0 0 1 1 1 0 0 1 1 0	1.5.6.1	1.3.2	2.7.2	0.1.1	0.1.0	2.3
18	1 0 0 0 1 1 0 1 0 1 1 0	1.5.6.1	1.3.2	3.6.2	0.0.1	1.1.0	2.3
19	0 1 0 1 1 0 1 0 0 0 1 1	1.6.5.1	2.3.1	2.6.3	0.1.1	1.0.0	3.2
20	1 0 0 1 1 0 0 1 0 0 1 1	1.6.5.1	2.3.1	2.7.2	0.1.0	1.1.0	3.2
a	0,1,1,1,0,1	4	3	3	0	1	2
b	0,1,1,2,0,2	12	5	8	0	2	3
21	0 0 1 1 0 1 0 1 0 1 0 1	1.5.5.1	1.3.1	1.4.3	0.0.0	0.1.1	1.2
22	0 1 0 1 0 1 0 0 1 1 0 1	1.5.5.1	1.3.1	3.4.1	0.0.0	1.1.0	2.1
b	0,1,2,1,0,2	12	5	5	0	1	2
23	0 0 1 1 0 1 0 1 1 0 0 1	1.5.5.1	1.3.1	1.3.1	0.0.0	0.1.0	1.1
b	0,2,1,2,0,1	12	8	8	0	4	4
24	0 1 0 1 0 1 0 1 1 1 0 0	1.5.1	1.4.3	3.4.1	0.0.0	1.2.1	2.2
a	1,1,0,1,1,0	4	4	4	0	1	4
b	1,1,0,2,2,0	12	8	9	2	2	6
25	0 1 0 1 1 0 1 0 0 1 1 0	1.5.5.1	2.4.2	2.5.2	0.1.1	1.1.0	3.3
b	1,2,0,2,1,0	12	9	9	1	4	6
26	0 1 0 1 1 0 1 1 0 1 0 0	1.5.5.1	2.5.2	2.5.2	0.1.0	1.2.1	3.3

końcówkowych), typu $G^4(t_1, t_2) = T_1^4(t_1) \cup T_1^4(t_2)$. Wygenerowane typy struktur kanonicznych czwórników pogrupowano według trzech kryteriów:

Typu struktury kanonicznej unigrafu $G^4(t)$.

Typu struktury kanonicznej multigrafu $G^4(t_2)$.

Uporządkowanych zbiorów $\langle LT_k^4(t_1, t_2) \rangle$.

W praktyce, często interesuje nas wartość tylko jednego parametru $\langle LT_k^4(t_1, t_2) \rangle$ struktury układu. W tabeli 6 zestawiono przykłady typów struktur

Tabela 6.

Struktury kanoniczne czwórników $G^4(t_1, t_2)$ realizujących wielomian $H_{12,30}^4(t_1, t_2)$ w postaci $\langle LT_{12,30}^4(t_1, t_2) \rangle = \langle 1, 1, 0 \rangle$

Nr	$G^4(t_1, t_2)$	$\langle LT_{12,30}^4(t_1, t_2) \rangle$				
		LT_1^4	$LT_{1,0}^4$	$LT_{2,3}^4$	$LT_{12,30}^4$	$LT_{13,20}^4$
1	1 0 0 1 1 0 0 1 0 1 1 0	1.5.5.1	2.4.2	2.5.2	1.1.0	0.1.1
2	1 1 0 0 1 0 0 1 0 1 1 0	1.5.5.1	2.5.2	2.4.2	1.1.0	0.1.1
3	1 0 0 1 1 0 0 1 0 0 1 1	1.5.6.1	2.3.1	2.7.2	1.1.0	0.1.0
4	1 1 0 0 1 0 0 0 1 1 1 0	1.5.6.1	2.7.2	2.3.1	1.1.0	0.1.0
5	1 1 0 0 1 0 0 0 0 1 1 0	1.4.2.0	2.4.0	2.3.1	1.1.0	0.1.0
6	1 0 0 1 1 0 0 1 0 0 1 0	1.4.2.0	2.3.1	2.4.0	1.1.0	0.1.0
7	1 0 0 0 1 1 0 1 0 1 1 0	1.6.5.1	1.3.2	3.6.2	1.1.0	0.0.1
8	1 0 1 0 1 0 0 1 0 1 1 0	1.6.5.1	3.6.2	1.3.2	1.1.0	0.0.1
9	1 0 0 0 1 0 0 1 0 1 1 0	0.2.4.1	1.3.2	1.3.2	1.1.0	0.0.1
10	1 0 0 1 1 0 0 0 1 0 1 1	1.5.6.1	2.3.2	2.4.2	1.1.0	0.0.0
11	1 1 0 0 1 0 0 0 1 0 1 1	1.6.5.1	2.4.2	2.3.2	1.1.0	0.0.0
12	1 0 0 0 1 1 0 0 1 1 1 0	1.6.5.1	1.4.2	3.4.1	1.1.0	0.0.0
13	1 0 1 0 1 0 0 1 0 0 1 1	1.5.6.1	3.4.1	1.4.2	1.1.0	0.0.0
14	1 0 0 0 1 0 0 0 1 0 1 1	1.5.5.1	1.3.1	3.4.1	1.1.0	0.0.0
15	1 0 1 0 1 0 0 0 1 0 1 1	1.5.5.1	3.4.1	1.3.1	1.1.0	0.0.0
16	1 0 0 0 1 1 0 0 1 0 1 0	1.2.2.0	1.3.1	3.2.0	1.1.0	0.0.0
17	1 0 1 0 1 0 0 0 0 0 1 1	1.2.2.0	3.2.0	1.3.1	1.1.0	0.0.0
18	1 0 0 0 1 0 0 0 1 0 1 1	0.2.2.1	1.3.1	1.3.1	1.1.0	0.0.0

kanonicznych czwórników, które charakteryzują się jednakową postacią wielomianu $H_{12,30}^v(t_1, t_2)$, którego uporządkowany zbiór ma wartość $\langle LT_{12,30}^v(t_1, t_2) \rangle = \langle 1.1.0 \rangle$.

5. ALGORYTM SYNTEZY TOPOLOGICZNEJ

Algorytm generacji struktur kanonicznych czwórników zawiera dwa podstawowe etapy:

- I. Generację struktur unigrafów czwórników i ich podział na klasy struktur kanonicznych.
- II. „Wypełnienie” struktury kanonicznej unigrafu dendrytami $T_{k1}^v(t_1)$ i $T_{k2}^v(t_2)$ zgodnie z równaniem 16, a następnie podział otrzymanych struktur na klasy ich struktur kanonicznych.

Przystępując do generacji struktur kanonicznych unigrafów, należy zwrócić uwagę na następujące fakty:

1. D – równoważne struktury unigrafów mają jednakową liczbę wierzchołków.
2. D – równoważne struktury unigrafów mają jednakową liczbę gałęzi.
3. D – równoważne struktury unigrafów mają odpowiadające sobie gałęzie brzegowe (dla czwórników $\{e_1, e_4\}$).

4. Liczba gałęzi w unigrafie spełnia warunek 1.

Biorąc pod uwagę ww fakty, ogólnie proces generacji struktur kanonicznych unigrafów zawiera następujące etapy:

- Dla zadanej liczby wierzchołków, generujemy wszystkie możliwe macierze jednowierszowe $G^v(t)$ o stałej liczbie kolumn ($e = v \cdot (v-1)/2$) i stałej liczbie jedynek – odpowiadających poszczególnym gałęziom.
- Sprawdzamy obecność lub brak gałęzi brzegowych $\{e_1\}$, $\{e_4\}$.
- Sprawdzamy spójność unigrafu zgodnie z zależnościami 2 i 3 (korzystając z wygenerowanych macierzy ścieżek w pełnym unigrafie).
- Generujemy macierz $D\{T_1^v(t)\}$ [] i sprawdzamy równoważność topologiczną unigrafu w stosunku do już istniejących klas unigrafów, korzystając z algorytmu opisanego w pierwszej części niniejszej pracy.

Algorytm generacji struktur bi-grafów $G^v(t_1, t_2)$ w danej strukturze kanonicznej unigrafu $G^v(t)$ i ich podziału na klasy struktur kanonicznych $G^v(t_1, t_2)$ zawiera następujące etapy.

Pierwszy etap to znalezienie w danej macierzy $D\{T_1^v(t)\}$ wszystkich par dendrytów 1-drzewowych które spełniałyby równanie 16.

Drugi etap to generacja macierzy $D\{T_1^v(t_1, t_2)\}$, poprzez wstawienie w macierzy $D\{T_1^v(t)\}$ cyfry 2 w miejscu gdzie dendryty 1-drzewowe $T_1^v(t_1)$ i $T_1^v(t_2)$ mają gałęzie równoległe, oraz dopisując macierz $G_1^v(t_1, t_2)$.

Trzeci etap to podział wygenerowanych bi-grafów $G_1^v(t_1, t_2)$, na klasy ich struktur kanonicznych, metodą opisaną w pracy [].

W praktyce często należy zachodzić konieczność generacji struktur czwórnika charakteryzujących się zadaną postacią parametru $H_k^v(t_1, t_2)$. Dla parametrów $H_1^v(t_1, t_2)$, $H_{1,0}^v(t_1, t_2)$ i $H_{2,3}^v(t_1, t_2)$ ich typy uporządkowanych zbiorów $\langle LT_1^v(t_1, t_2) \rangle$, $\langle LT_{1,0}^v(t_1, t_2) \rangle$ i $\langle LT_{2,3}^v(t_1, t_2) \rangle$, mają ściśle określoną postać (tabela 4a i 4b). Struktury kanoniczne $G^v(t_1, t_2) = T_{k1}^v(t_1, t_2) \cup T_{k2}^v(t_1, t_2)$ charakteryzujące się danym typem uporządkowanego zbioru $\langle LT_k^v(t_1, t_2) \rangle$ generuje się zgodnie z wyrażeniami zestawionymi w tabeli 4a i 4b.

Sytuacja jest bardziej złożona w przypadku parametrów $H_{12,30}^v(t_1, t_2)$ i $H_{13,20}^v(t_1, t_2)$ struktury $G^v(t_1, t_2)$ tworzonych z $T_1^v(t_i)$ klas B, C, D dla których uporządkowane zbiory $\langle LT_{12,30}^v(t_1, t_2) \rangle$, $\langle LT_{13,20}^v(t_1, t_2) \rangle$ nie są jednoznacznie zdefiniowane, ponieważ zgodnie z tabelą 3 w tworzeniu dendrytów $T_{12,30}^v(t_1, t_2)$ i $T_{13,20}^v(t_1, t_2)$ może uczestniczyć różna liczba elementów dendrytów $T_1^v(t_1, t_2)$.

BIBLIOGRAFIA

1. N. Deo: *Teoria grafów i jej zastosowania w technice i informatyce*. Warszawa: PWN, 1980
2. W. Lipski: *Kombinatoryka dla programistów*. Warszawa, WNT, 1982
3. K. Mikołajuk, Z. Trzaska: *Elektrotechnika teoretyczna – analiza i synteza elektrycznych obwodów liniowych*. Warszawa PWN 1984
4. M.E. Reingold, J. Nievergeli, N. Deo: *Algorytmy kombinatoryczne*. Warszawa: PWN, 1985
5. M.N.S. Swamy, K. Thulasiraman: *Graphs, networks and algorithms*. New York: Wiley – Interscience, 1981
6. Z. Wróbel: *Funkcje strukturalne grafów układów dwójników i czwórników składających się z dwóch rodzajów elementów*. Kwartalnik Elektroniki i Telekomunikacji, 1991, 37, z.1–2, ss.35–52
7. Z. Wróbel: *Synteza topologiczna struktur kanonicznych bigrafów układów dwójników, Część I. Metoda przeglądu bezpośredniego*. Kwartalnik Elektroniki i Telekomunikacji, 1991, 37, z.1–2, ss.67–83
8. Z. Wróbel: *Synteza topologiczna struktur kanonicznych bigrafów układów dwójników, Część II. Metoda przeglądu hierarchicznego*. Kwartalnik Elektroniki i Telekomunikacji, 1991, 37, z.1–2, ss.85–97
9. Z. Wróbel: *D-równoważność struktur układów realizujących funkcje wymierne*. Kwartalnik Elektroniki i Telekomunikacji, 1993, t.XXXIX, z.3

Z. WRÓBEL

TOPOLOGICAL SYNTHESIS OF CANONICAL STRUCTURES OF BIGRAPHS TWO-PORT CIRCUITS CONSISTING OF TWO KINDS OF ELEMENTS

Summary.

Procedure of topological synthesis of canonical structures by means of hierarchical reviewing method, allowing to generate all possible topological structures of two-port circuits was shown in this paper. Condition of realizing minimal-elements of two-port circuits was shown in this paper. This synthesis procedure was illustrated by the example of minimal-elements of two-port circuits.

MICAP – The computer program for the analysis and design of the linear microwave integrated circuits

ZYGMUNT MUSIAŁOWSKI

Instytut Telekomunikacji i Akustyki. Politechnika Wroclawska

Received 1993.02.25

Authorized 1993.09.25

The paper presents the computer program MICAP assigned to the analysis and design of linear microwave integrated circuits. The frequency-domain analysis is based on the nodal method together with sparse-matrix method. MICAP analyzes the circuits with distributed, lumped, and active elements of any topology as well as handles resistors, conductors, inductors, capacitors, voltage-controlled current sources, and multiport elements, such as transistors, described by their scattering or admittance parameters. Furthermore it handles the built-in models of bipolar and unipolar microwave transistors as well as microstrip transmission lines and built-in models of some frequently used microstrip discontinuities. The effects of dispersion and the thickness of the strip film are also included. MICAP is therefore particularly useful to analysis and design of linear hybrid microwave integrated circuits where the technology, circuit lay-out and circuit functions are simultaneously considered.

1. INTRODUCTION

Computer-aided analysis and design of microwave circuits, particularly when very fast and personal computers with large memory are easily accessible, is nowadays a widely spread habit in microwave engineering and scientific practice. The existing computer programs deal with the analysis of the non-linear microwave circuits in the time-domain [6, 25] and in the frequency-domain [3, 4, 6, 15, 22] as well as the analysis, optimization and synthesis of linear ones [3, 5, 6, 13, 18, 19, 23, 24].

Furthermore noise and statistical analyses and multiparameter optimization together with lay-out problems, modeling and real-time measurements of the element parameters as well as full workstations are possible [3, 6, 13]. Usually programs for the analysis and optimization of microwave circuits of arbitrary

topology are based on the S scattering matrix or Y nodal-admittance matrix [3, 5, 6, 10, 13, 16, 18, 20, 24]. The transfer $[T]$ and chain $[A]$ matrices are also widely used in the case of cascade circuit structures [5, 10, 16, 19, 23]. The inspiration to develop the program described below was to create an instrumental tool for an analysis and design of hybrid microwave integrated circuits (MICs). For this reason the features of the microstrip line have been preferentially included. MICAP (Microwave Integrated Circuits Analysis Program), developed for this purpose, analyzes the linear MICs in the frequency domain using the nodal description of the circuit and exploiting the sparsity of the Y nodal-admittance matrix. The sparse-matrix method results in computer memory savings, higher efficiency and accuracy of calculations. The program is capable of handling (among others) things the built-in models of the microwave transistors and the microstrip line discontinuities as well as including effects of dispersion and the thickness of the strip film.

2. NODAL SPARSE-MATRIX APPROACH

MICAP is based on the nodal analysis of the microwave circuit, which results in the matrix description (1):

$$Y V = I. \quad (1)$$

where Y is the nodal-admittance matrix of order $\omega \times \omega$, V and I are ω -element column vectors of nodal voltages and currents respectively and ω is the number of the circuit nodes. The method is efficient and non-specific enough to analyze a wide variety of circuits of arbitrary topology. Furthermore the structural symmetry and sparsity, numerically dominant main diagonal and facility of direct construction of Y result in the wide usage of the nodal description in computer applications [1, 2, 5, 10]. From a mathematical point of view the equation (1) is the sparse system of linear complex equations of form (2)

$$A x = b \quad (2)$$

The solution to (2) is very often obtained by means of a Doolittle algorithm, where the A matrix is divided into a two-factor formula of the form $A = LU$ and L , U , which are the lower and upper triangular matrices respectively [2, 21]. The solution of (2) is based on L and U matrices as below:

$$Lz = b \quad z = L^{-1}b \quad (3)$$

$$Ux = z \quad x = U^{-1}L^{-1}b. \quad (4)$$

Before reordering the Y nodal-admittance matrix its non-zero structure is stored using the static storage system called a row pointer/column index scheme as illustrated in Figure 1a. This scheme uses two integer arrays. The first IUR contains

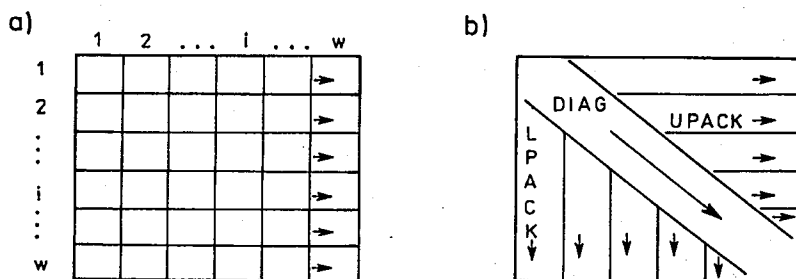


Fig. 1. Static storage schemes of the Y matrix non-zero structure before (a) and after its ordering (b)

w pointers and the pointer $IUR(k)$ represents the starting location in the second pointer array IUC of nonzeros associated with row k of the admittance matrix. The array IUC includes the column indices of all off-diagonal nonzeros in each row of the original Y matrix [1]. Both arrays are determined by means of the localization of the numerical values inserted into Y by the microwave circuit elements. The Y matrix is reordered using any of the near-optimal algorithms proposed in [1, 17]. After reordering the Y it is very convenient to store nonzeros by means of the storage scheme given in [2] and shown in Fig. 1b. **DIAG** array stores all of the elements of the main diagonal, the **UPACK** array stores in packed form, the nonzeros placed in the upper triangular part of the Y and the **LPACK** array stores the nonzeros in the lower triangular part of the Y .

The elements of **UPACK** are stored by rows and the column indices of nonzeros associated with every row are placed in ascending order.

Such an approach demands the renumbering of the element nodes but it considerably increases the effectiveness of the LU factorization and the forward and back substitutions. For the localization of the nonzeros of **UPACK** the new arrays of pointers $IUR1$ and $IUC1$ have to be determined in a similar way to those of IUR and IUC above. The elements of **LPACK** are stored by columns and, because of its structural symmetry to **UPACK**, the same arrays of pointers $IUR1$ and $IUC1$ can be used. The values of the impressed sources b in (2) are stored in array X , which is also used to store the values of vectors z and x during the forward and back substitutions.

3. THE PROGRAM STRUCTURE

The overall structure of **MICAP** is shown in Fig. 2, where its main stages are: input, preparation for analysis, numerical analysis and output.

The input stage consists of algorithms dealing with the processing of input data and its modifications during microwave circuit design. **MICAP** input language has been specifically designed for use by circuit designers and contains a dozen diagnostic messages designed to assist the user in correcting input errors.

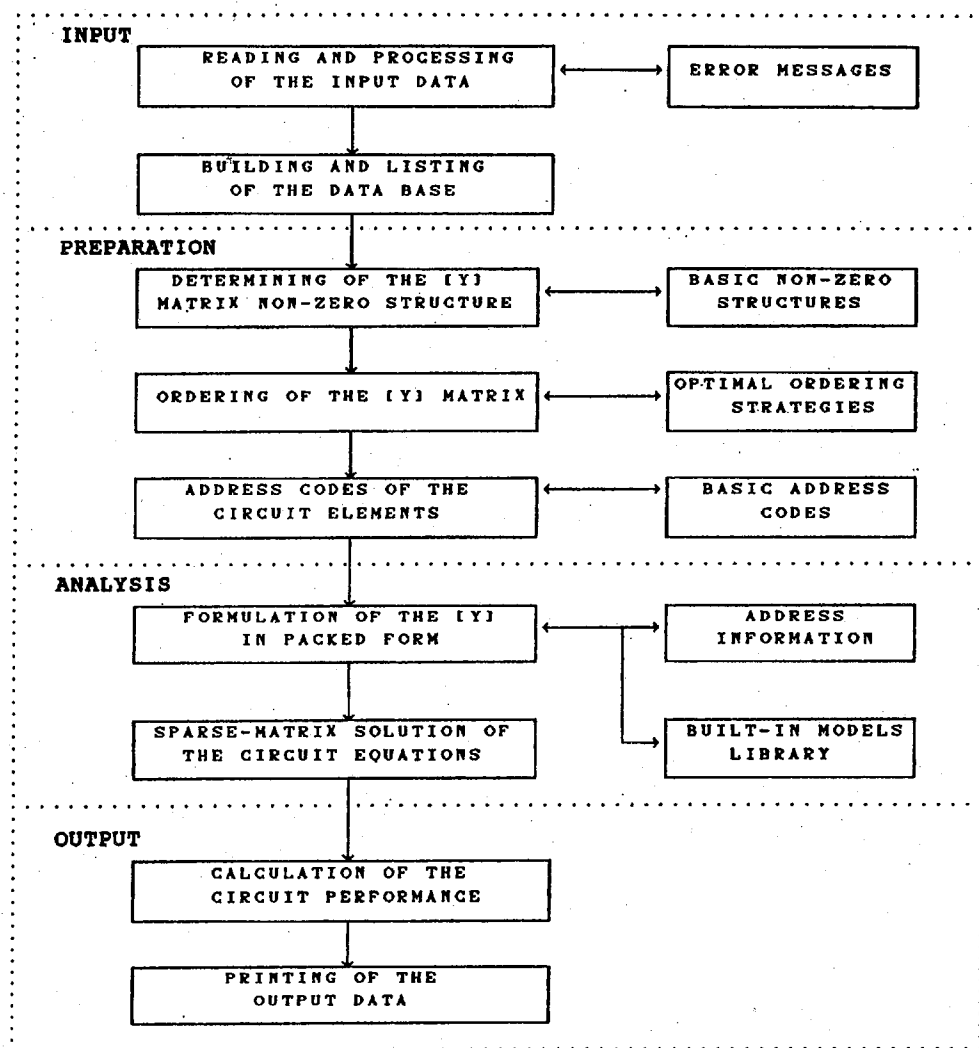


Fig. 2. Overall MICAP structure

The preparation stage refers to the sparse-matrix approach. The nonzero structure of the Y admittance matrix is determined and next it is reordered with respect to the minimum number of new non-zero elements arising during the ordering and the minimum number of long, time-consuming numerical operations. In order to effectively formulate the ordered nodal equations in packed form, address codes have been introduced. All the above is executed once only for a given microwave circuit structure and, despite its complexity, the time required to accomplish this is usually negligible in comparison to the one necessary to repeated circuit analysis.

The analysis stage comprises the formulation and solving of the packed nodal equations in order to determine the circuit functions in the given frequency band. As has been stated above, the rapid formulation of the packed nodal equations is assured owing to the address codes. Furthermore it should be noted, that the rearrangement of non-zero pointers in order to satisfy the rule of ascending order of the column indicies in array IUC1 also makes the procedure of numerical solving of (2) easier and considerably faster.

The output stage includes the procedures dealing with calculating the circuit functions and printing output data.

4. PROGRAM FUNCTIONS

MICAP is assigned to steady-state analysis of the linear hybrid microwave integrated circuits of arbitrary topology in the frequency domain.

For the given frequency range the program:

- leads the analyzed microwave circuit to the equivalent two-port between the given input and output ports as shown in Fig. 3,

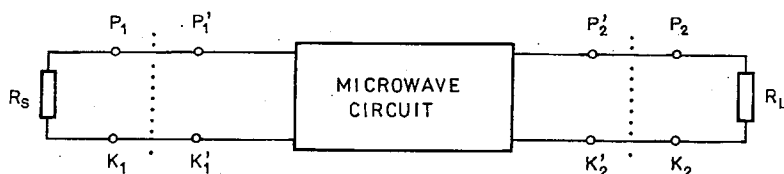


Fig. 3. Equivalent two-port of analyzed microwave circuit

- calculates and prints parameters of the equivalent two-port as well as the circuit functions,
- calculates the dispersion of microwave elements,
- exploits the built-in models of bipolar and unipolar microwave transistors,
- exploits the built-in models of microstrip line discontinuities.

The program enables both calculations for a given network and modifications of its elements values and/or structure.

MICAP handles resistors (**R**) and conductors (**G**), capacitors (**C**) inductors (**L**), voltage controlled current sources (**S**), two- and three-terminal elements (**DF** and **TF** respectively), bipolar and unipolar microwave transistors (**MT**) and transmission lines (**TL**), as well as microstrip line discontinuities (**DS**).

The elements of **DF** and **TF** types enable their measured **[S]**, **[Y]** and **[A]** parameters to be included directly into the calculations. The opportunity to define the parameters of the microstrip lines and their discontinuities by means of geometry, dimensions and material properties as well as automatic calculation of the dispersion are useful in the analysis and design of hybrid MICs.

5. PREPARING OF MICROWAVE CIRCUIT TO ANALYSIS

MICAP allows an analysis of microwave circuits whose distributed parameter elements can be described by means of the substrate parameters (h , ϵ) and the dimensions of the microstrip lines (w , l , t), see Fig. 6. This makes it possible to recognize MICs in conjunction with their hybrid technology by taking into account the relevant circuit lay-out. In this case in order to make the equivalent circuit of the MIC it is necessary to determine the reference planes showing the places of connections the distributed and lumped circuit elements.

A designer denotes the planes of reference arbitrary. However, in the case of microstrip line discontinuities, the reference planes have to be the same as the ones of the built-in program models. It deals particularly with microstrip line bends and T-junctions (see Table 5 and Table 6).

It is assumed that the signal source and load impedances (R_s and R_L respectively) are two-terminal elements of type R, G or DF. The analyzed microwave network is of arbitrary topology and it can incorporate any elements handled by the program. Its nodes and all of its elements have to be numbered separately, beginning from 1. The order of the circuit node numbers is optional because of the applied sparse-matrix algorithms.

Furthermore the desired frequency range, circuit functions and prospect circuit modifications should be also given.

6. THE STRUCTURE OF THE INPUT DATA

The general view of the input data structure, with the main commands that are displayed and the data blocks, is shown in Fig. 4.

6.1. COMMANDS

The main commands are obligatory and occur as separate items of meaning given below.

- | | |
|--------------|---|
| MICAP | — Opens the data for a given circuit and starts the program. |
| DATA | — Gives information about the introductory commands and begins the input of the data blocks. |
| MODY | — Enables modification of the input data. The main rules of the circuit data modification are given later, with a description of the circuit element records. |
| ANAL | — Starts the numerical analysis of the microwave circuit. |
| STOP | — Stops the program. |

The introductory commands, which are all optional, are located before the data blocks and mean the following:

- | | |
|-------------|--------------------------------------|
| TITL | — Precedes the title of the problem. |
|-------------|--------------------------------------|

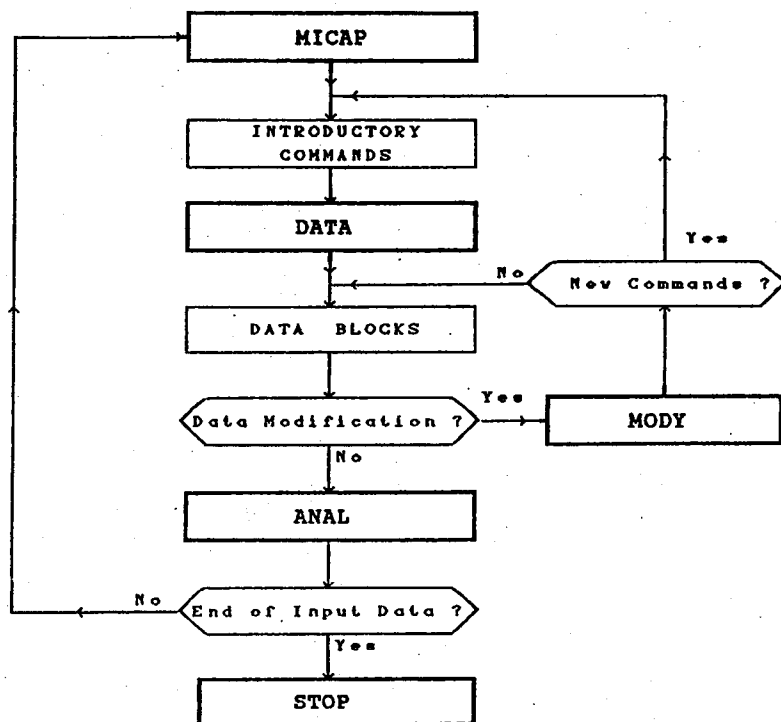


Fig. 4. General view of MICAP input data structure

LIST	— Prints the input data in its original form.
TABL	— Lists the content of the program arrays dealing with the input data.
DISP	— Runs the calculations of microstrip line dispersion.
JENS	— The dispersion of microstrip line is calculated according to the formulas proposed by Hammerstadt and Jansen in [12] or by Getsinger in [8]. When both commands are omitted the first of them is assumed by the program.
GETS	
BERRY	— The Y nodal-admittance matrix reordering is performed using the ordering strategies given by Berry [1] or Markowitz [17].
MARK	
ZN Z_0	— Characteristic impedance.
EPS.R ϵ	— Dielectric constant of the substrate.
HEIGHT h	— Substrate height.
THICK t	— Thickness of the microstrip film.

It is assumed that ϵ_r , h , t are the same for all network elements and if omitted, then the typical values of $\epsilon_r=10$, $h=0,7$ mm, $t=5$ μ m are accepted by the program. Z_0 given here deals with the elements when their characteristic impedance is omitted on the input list. By default $Z_0=50$ Ω .

6.2. DATA BLOCKS

There are five data blocks in the program and each of them starts with the relevant heading which means the following.

SOURCE The data blocks deal with the signal source and load impedances of the microwave circuit, respectively. It is assumed that both of them are two-terminal elements described by means of **R**, **G** or **DF**.

LOAD

MICNET Deals with the data concerning the microwave circuit and describes its elements and topology using any of the circuit elements handled by the program (**R**, **G**, **C**, **L**, **S**, **DF**, **TF**, **MT**, **TL**, **DS**). Furthermore the circuit input and output ports as well as the reference node (numbered as zero) have to be given.

FREQUENCY Identifies the desired frequency range of the analysis. There are three possibilities for defining the frequency points: the first permits a choice of individual points and the remaining two deal with the linear and logarithmic steps in the frequency band. **MICAP** sets all the frequency points in the ascending order.

OUTPUT **MICAP** leads the analyzed microwave circuit to the equivalent two-port between required input and output ports. For the given signal source and load impedances (R_s , R_L) the following circuit functions are calculated: transducer and operating power gains (G_p , G_p), insertion loss (**L**), voltage gain (K_u), current gain (K_i) and stability factor (**K**). Furthermore input and output immittances (Y_{in} , Z_{in} , Y_{out} , Z_{out}), reflection coefficients (Γ_i , Γ_{out}) and voltage standing wave ratios (ρ_{in} , ρ_{out}) are computed. Additionally, microwave circuit equivalent two-port [**S**], [**Y**] and [**A**] matrices are available.

Each of the above data blocks ends with keyword **END**.

7. DESCRIPTION OF THE CIRCUIT ELEMENTS

The data dealing with the elements of any kind (**R**, **G**, **C**, **L**, **S**, **MT**, **DF**, **TF**, **TL**, **DS**) appear on the input list as groups of relevant elements. These groups can be given in any order and the sequence of the elements creating them is also arbitrary.

7.1. LUMPED ELEMENTS (**R**, **G**, **C**, **L**, **S**)

MICAP holds lumped resistors (**R**), conductors (**G**), capacitors (**C**), inductors (**L**) and voltage-controlled current sources (**S**). The records of lumped elements are given in Table 1, and their structure is of the form:

$$In, (P, K, R, T) W = v;$$

where **I** is replaced by the relevant element identifier (**R**, **G**, **C**, **L**, **S**); **W** is the element value identifier and *n*, **P**, **K**, **R**, **T**, *v* are element numbers, nodes and value. The data fields are separated by commas and the record ends with a semicolon. In

Table 1

Lumped Elements of MICAP

Lumped Elements			Record Form
1	VCCS		$S_n, (P-K, R-T) W = g_s c_s l_s$
2	G		$G_n, (P-K) W = g;$
	R		$R_n, (P-K) W = r;$
3	C		$C_n, (P-K) W = c;$
4	L		$L_n, (P-K) W = l;$

the case of VCCS, its transadmittance is equal to $y_m = g_s + j(\omega C_s - 1/\omega L_s)$, where ω is the considered angular frequency and v consists of g_s, C_s, L_s . The conductors and resistors use the same set of numbers. When modification of the element is necessary, apart from its identifier and number, only modified data fields have to be given.

7.2. TWO- AND THREE-TERMINAL ELEMENTS (DF, TF)

MICAP holds two- and three-terminal elements which can be described by their admittance [Y], scattering [S] or transfer [A] parameters. The parameters $x(f_i)$ obtained at the i th frequency f_i are complex and can be given using polar or exponential form. The description of two-terminals (DF) consists of one main record:

$$\boxed{\text{DFn, XXXX, (P - K) RN} = r_n;}$$

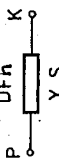
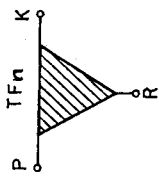
and some additional records of the form:

$$\boxed{\text{F} = f_i \quad \text{Re}[x(f_i)] \quad \text{Im}[x(f_i)]} \quad \text{or} \quad \boxed{\text{F} = f_i \quad |x(f_i)| \quad \arg[x(f_i)]}$$

where DF, F, RN are the identifiers of the main and additional records as well as characteristic impedance Z_0 ; XXXX is the element code given in Table 2 and refers to the admittance or scattering element parameters; n, P, K, r_n, f_i are the values of

Table 2

Two- and Three-terminal Elements of MICAP

Element		Description type		Description form	Main record form	Additional records form
DF		1	Y	$\text{Re}\{x(f)\}, \text{Im}\{x(f)\}$	DFn, Y2RI, (P-K)	$F=f_i \text{Re}\{x(f)\} \text{Im}\{x(f)\}$
		2	S		DFn, S2RI, (P-K) RN=r _n	
		3	Y	$ x(f) , \arg\{x(f)\}$	DFn, Y2MF, (P-K)	$F=f_i x(f) \arg\{x(f)\}$
		4	S		DFn, S2MF, (P-K)	
TF	 [Y], [S], [A]	5	[Y]	$\text{Re}\{x_{PK}(f)\}$ $\text{Im}\{x_{PK}(f)\}$	TFn, Y3RI, (P-K-R);	According to the formula (*) in chapter 7.2.
		6	[S]		TFn, S3RI, (P-K-R) RN=r _n	
		7	[A]		TFn, A3RI, (P-K-R);	
		8	[Y]	$ x_{PK}(f) $ $\arg\{x_{PK}(f)\}$	TFn, Y3MF, (P-K-R);	According to the formula (**) in chapter 7.2.
		9	[S]		TFn, S3MF, (P-K-R) RN=r _n	
		10	[A]		TFn, A3MF, (P-K-R);	
		11	[S]	Real / Imag · / arg	TFn, S3MX, (P-K-R) RN=r _n	According to the formula (***) in chapter 7.2.

the element number, nodes, Z_o and the i th frequency. Similarly the description of three-terminals (TF) consists of the main record

$$\boxed{\text{TF}n, \text{XXXX}, (\text{P}-\text{K}-\text{R}) \text{RN} = r_n;}$$

and some additional records of the form

$$\begin{array}{l} \mathbf{F} = f_i \\ \text{Re}[x_{11}(f_i)] \quad \text{Im}[x_{11}(f_i)] \quad \text{Re}[x_{12}(f_i)] \quad \text{Im}[x_{12}(f_i)] \\ \text{Re}[x_{21}(f_i)] \quad \text{Im}[x_{21}(f_i)] \quad \text{Re}[x_{22}(f_i)] \quad \text{Im}[x_{22}(f_i)] \end{array} \quad (*)$$

or

$$\begin{array}{l} \mathbf{F} = f_i \\ |x_{11}(f_i)| \quad \arg[x_{11}(f_i)] \quad |x_{12}(f_i)| \quad \arg[x_{12}(f_i)] \\ |x_{21}(f_i)| \quad \arg[x_{21}(f_i)] \quad |x_{22}(f_i)| \quad \arg[x_{22}(f_i)] \end{array} \quad (**)$$

or

$$\begin{array}{l} \mathbf{F} = f_i \\ \text{Re}[x_{11}(f_i)] \quad \text{Im}[x_{11}(f_i)] \quad |x_{12}(f_i)| \quad \arg[x_{12}(f_i)] \\ |x_{21}(f_i)| \quad \arg[x_{21}(f_i)] \quad \text{Re}[x_{22}(f_i)] \quad \text{Im}[x_{22}(f_i)] \end{array} \quad (***)$$

where TF is the identifier of the main record and the meaning of the other symbols is similar to the DF element case. The TF element code (see Table 2) refers to the form of [Y], [S], and [A] two-port parameters. The TF element can be recognized to be a two-port or three-port as shown in Fig. 5. It should be noted that the order of the P, K, R element nodes refers to the two-port parameters as follows:

$$\begin{bmatrix} x_{11} & x_{12} \\ x_{21} & x_{22} \end{bmatrix} \div \begin{bmatrix} x_{PP} & x_{PK} \\ x_{KP} & x_{KK} \end{bmatrix}$$

In the case of the DF and TF element modifications, it is only necessary to repeat the element identifier, its number and code. Next only the altered additional records, either nodes or Z_o , have to be given.

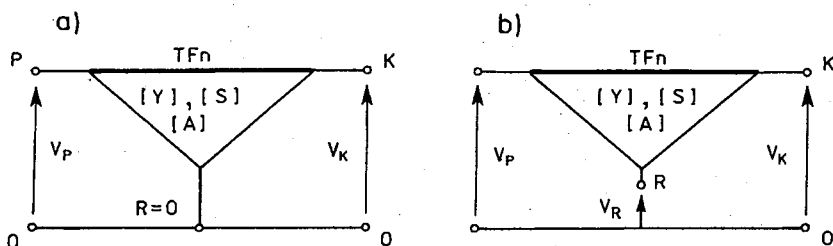


Fig. 5. MICAP TF-type element as two-port (a) or three-port (b) element

Table 3

Built-in Models of Microwave Transistors of MICAP

Transistor type	Built-in model	Model parameters	The record form
1 Bipolar Junction Transistor (BJT)	$Y_m = g_m e^{-j\omega\tau_0}$	$g_m, \tau_0, r_{bb'}, r_{bb'}, r_{ce}, r_e, C_{bc}, C_{ce}, C_{tc}, L_{bb}, L_e, L_c$	MTn, MBJT, (B-C-E) $g_m, \tau_0, r_{bb'}, r_{bb'}, r_{ce}, r_e, C_{bc}, C_{ce}, C_{tc}, L_{bb}, L_e, L_c$;
2 Schottky-Barrier Field Effect Transistor (MESFET)	$Y_m = g_m e^{-j\omega\tau_0}$	$g_m, \tau_0, R_i, r_{ds}, R_g, R_d, R_s, R_d, C_{gs}, C_{dc}, C_{ds}, L_g, L_d, L_s$	MTn, MFET, (G-D-S) $g_m, \tau_0, R_i, r_{ds}, R_g, R_d, R_s, C_{gs}, C_{dc}, C_{ds}, L_g, L_d, L_s$;

7.3. BUILT-IN MODELS OF MICROWAVE TRANSISTORS (MT)

Although previously described lumped elements (**R**, **G**, **C**, **L**, **S**) enable modeling of microwave transistors, their built-in models are very useful. **MICAP** holds two built-in models of bipolar junction transistors **BJTs** and unipolar **GaAs MESFETs**, as shown in Table 3. Their record is of the form:

$$\text{MT}n, \text{XXXX}, (\text{P}-\text{K}-\text{R}), \text{par}(1), \dots \text{par}(i), \dots \text{par}(k);$$

where **MT** is the identifier of the microwave transistor; n is its number; **XXXX** is the transistor code and **P**, **K**, **R**, $\text{par}(i)$ are the numbers of transistor external terminals and its i th parameter value respectively. The detailed description of the transistor code as well as the order and meaning of its i th parameter are given in Table 3.

7.4. TRANSMISSION LINES

From among many elements of distributed parameter **MICAP** holds transmission lines, open- and short-circuited stubs as well as coupled microstrip lines. The parameters of this group can be given in various ways. Very useful is the installation of microstrip line described by means of its dimensions w , t , l as well as height h and permittivity ϵ_r of the substrate, see Fig. 6a. The record dealing with this element is as follows.

$$\text{TL}n, \text{XXXX}, (\text{P}-\text{O}, \text{K}-\text{O}) \quad \text{W} = w, \text{L} = l;$$

where **TL**, **W**, **L** are the identifiers of transmission line, microstrip width and length respectively; **XXXX** is the element code shown in Table 4; n , **P**, **K**, w , l are the values of the element number, ports, microstrip width and length. The remaining parameters ϵ_r , t , h are assumed to be the same for all microstrips and given as part of the introductory commands. The admittance parameters of these elements are calculated using the formulas given in [12], where the effects of the microstrip thickness are included. Dispersion is calculated with the formulas given in [8, 12]. The microstrip thickness may be also neglected i.e. $t=0$.

When the designer defines the **TLs** using their characteristic impedance Z_o , length l and propagation constant $\gamma = \alpha + j\beta$, the record is given as follows:

$$\text{TL}n, \text{XXXX}, (\text{P}-\text{O}, \text{K}-\text{O}) \quad \text{Z} = Z_o, \quad \text{L} = l, \quad \text{ALFA} = \alpha, \quad \text{BETA} = \beta;$$

where **TL**, **Z**, **L**, **ALFA**, **BETA** are the identifiers of relevant quantities and the element code **XXXX** is given in Table 4.

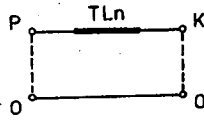
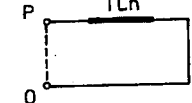
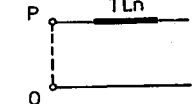
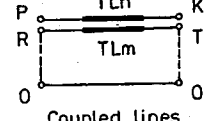
Furthermore another record is possible:

$$\text{TL}n, \text{XXXX}, (\text{P}-\text{O}, \text{K}-\text{O}) \quad \text{Z} = Z_o, \quad \text{L} = l, \quad \text{ALFA} = \alpha, \quad \text{VFW} = v_{fw}$$

where the **VFW**, v_{fw} are the identifier and value of the relative phase velocity defined as $v_{fw} = v_f/c$, whereas v_f and c are phase velocity and the speed of light respectively.

Table 4

Transmission Lines of MICAP

Transmission line type		TL Data	The record form			
1	 Transmission line	ϵ_r, h, w, l, t	TLn,	NLPT	, (P-O,K-O)	$W=w, L=l;$
2		Z_0, l, α, β		DCCT		$Z=Z_0, L=l, ALFA=\alpha, BETA=\beta$
3		Z_0, l, α, v_{fw}		DCVT		$Z=Z_0, L=l, ALFA=\alpha, VFW=v_{fw};$
4	 Short stub	ϵ_r, h, w, l, t	TLn,	NLPZ	, (P-O)	$W=w, L=l;$
5		Z_0, l, α, β		DCCZ		$Z=Z_0, L=l, ALFA=\alpha, BETA=\beta$
6		Z_0, l, α, v_{fw}		DCVZ		$Z=Z_0, L=l, ALFA=\alpha, VFW=v_{fw};$
7	 Open stub	ϵ_r, h, w, l, t	TLn,	NLPR	, (P-O)	$W=w, L=l;$
8		Z_0, l, α, v_{fw}		DCCR		$Z=Z_0, L=l, ALFA=\alpha, BETA=\beta$
9		Z_0, l, α, v_{fw}		DCVR		$Z=Z_0, L=l, ALFA=\alpha, VFW=v_{fw};$
10	 Coupled lines	ϵ_r, h, w, l, t	TLn,	NLPC	,	$TC=n m, D=d ;$

It should be pointed out, that Z_0 and γ depend on frequency and if dispersion is to be taken into account both these quantities should also be updated by the user for every frequency point. It is, of course not necessary in the case when the geometry and dimensions of microstrip line are given (the first of the records above). In the case of stubs, possessing only one input port (P-O), the other port (K-O) should be removed from the above records. MICAP also holds coupled microstrip lines (Fig. 6b). For these lines the record is given as follows:

$$TLk, NLPC, TC = n m, D = d;$$

where NLPC, k are the coupled line code and number; TC, D are the identifiers of the separate microstrips TLn and TLm, creating the coupled line and the distance d between them.

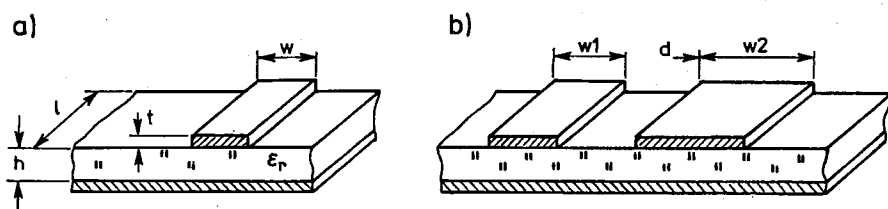


Fig. 6. Pictorial representation of microstrip (a) and coupled microstrip (b) cross-sectional geometry

When TL group elements are modified only their identifier, number and code are to be repeated and only the altered values should be given.

7.5. BUILT-IN MODELS OF MICROSTRIP LINE DISCONTINUITIES (DS)

One the most powerful features of MICAP are the built-in models of typical microstrip line discontinuities (DS) frequently used in practice. MICAP handles the discontinuities concerning opened, gap in microstrip, step in width, right bend, compensated right bend, right bend with change in width, symmetric and asymmetric T-junctions [7, 9, 10, 11].

In general it is assumed that the discontinuities are determined by previously extracted microstrip lines (TL). It must be underlined that the reference planes (T and Y) of the discontinuities must be retained as shown in the relevant figures in Table 5 and Table 6. The dimensions of the microstrip lines creating the discontinuity are determined with respect to these reference planes. The above approach is particularly convenient for the user, because it is easy to introduce any of the built-in discontinuities. The general record of the group DS is as follows:

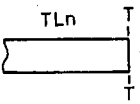
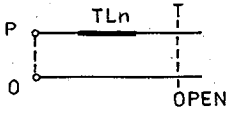
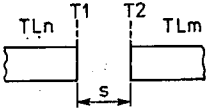
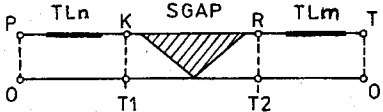
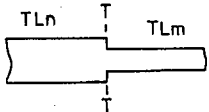
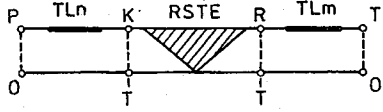
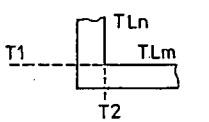
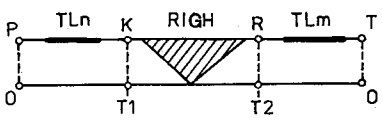
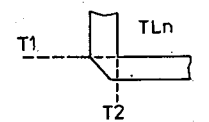
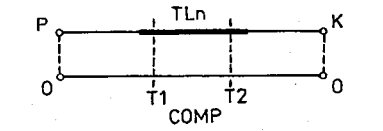
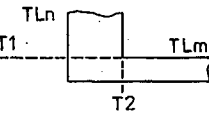
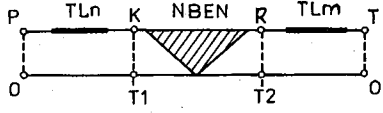
$$\text{DS}k, \text{ XXXX}, \text{ TL} = m \times n, \text{ X} = v;$$

where DS, TL, X are the identifiers of discontinuity, microstrip lines creating the discontinuity and additional parameter respectively; XXXX is the discontinuity code given in Table 5 and Table 6; m, x, n are the numbers of the microstrip lines creating the discontinuity and v is the value of the parameter. The detailed records dealing with particular discontinuities as well as the reference planes are given in Table 5 and in Table 6.

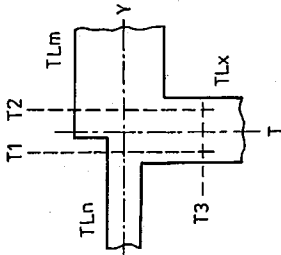
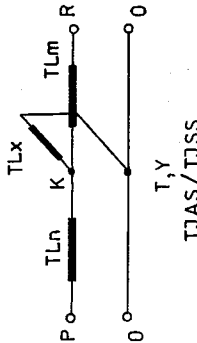
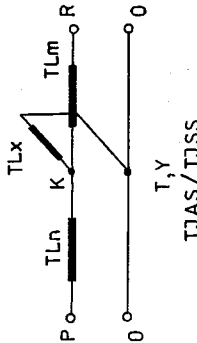
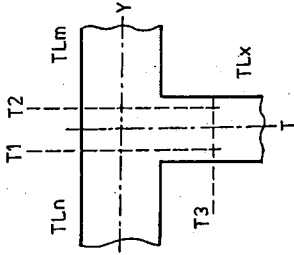
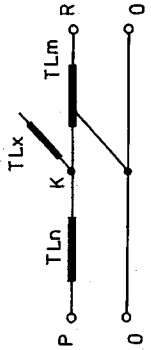
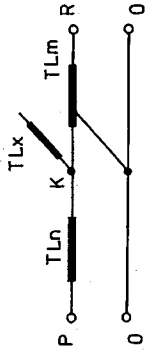
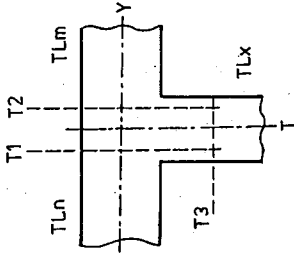
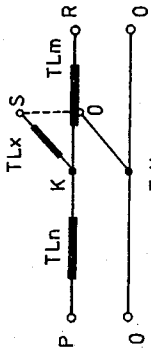
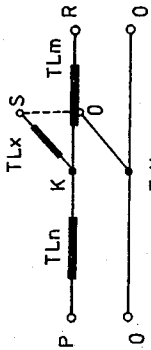
In the case of T-junctions there are introduced six separate cases depending on the symmetry and the type of the microstrip lines TL_x perpendicular to the other two TL_n, TL_m , as shown in Table 6. If any of the microstrip creating the discontinuity is modified then these changes are also automatically taken into the discontinuity equivalent model calculations. In this case the records dealing with group DS are not to be altered.

Table 5

Built-in Models of Microstrip Discontinuities of MICAP

Discontinuity type			Reference planes in MICAP	Discontinuity record form
1	Opened end			DSk, OPEN, $TL=n;$
2	Gap			DSk, SGAP, $TL=n\ m, X=s$
3	Step in width			DSk, RSTE, $TL=n\ m$
4	Right bend			DSk, RIGH, $TL=n\ m;$
5	Compensated right bend			DSk, COMP, $TL=n;$
6	Right bend with change in width			DSk, NBEN, $TL=n;$

Built-in Models of T-junction Discontinuities of MICAP

T-junction type		TLx type	T-junction structures	Reference planes in MICAP	Discontinuity record form		
1	Asymmetric	Short-circuited stub			DSK, TJAS, TL=n x m;		
	Symmetric				DSK, TJSS, TL=n x m;		
3	Asymmetric	Opened stub				DSK, TJAO, TL=n x m;	
	Symmetric					DSK, TJSO, TL=n x m;	
5	Asymmetric	Transmission line					DSK, TJAT, TL=n x m;
	Symmetric						DSK, TJST, TL=n x m;

8. INPUT AND OUTPUT PORTS

MICAP recognizes the analyzed microwave circuit to be a two-port as shown in Fig. 3, where its input ($P'_1 - K'_1$) and output ($P'_2 - K'_2$) ports are connected to the signal source port ($P_1 - K_1$) and load port ($P_2 - K_2$) respectively.

The record dealing with external ports is of the form:

$$\boxed{Z(P_i - K_i);} \quad \text{or} \quad \boxed{Z(P'_1 - K'_1, P'_2 - K'_2);}$$

for source, load and analyzed microwave circuit respectively.

9. MICROWAVE AMPLIFIER ANALYSIS USING MICAP

The microwave amplifier, shown in Fig. 7a, assigned to operation in the S-band has been analyzed using MICAP. The simplified equivalent circuit of the amplifier intended for small-signal frequency analysis is shown in Fig. 7b. Initially blocking

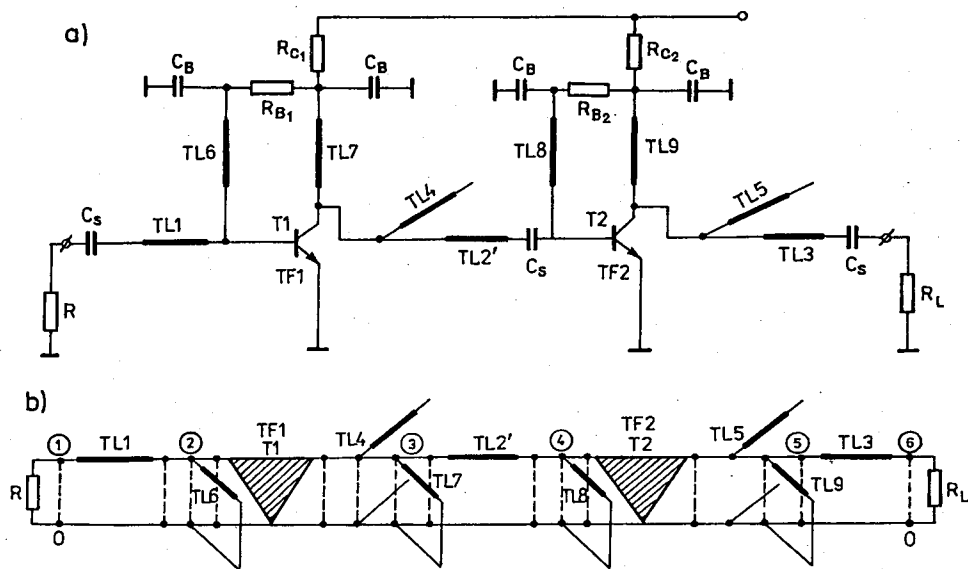


Fig. 7. S-band microwave amplifier (a) and its simplified small-signal equivalent circuit (b).

and coupling capacitors are short-circuited and microstrip discontinuities and bias circuitry are not taken into account. Three-terminal elements TF1 and TF2 represent microwave bipolar transistors described by their [S] parameters given at chosen dc operating point. The matching microstrip lines (TL1, TL2', TL3, TL4, TL5) and these of the bias circuitry (TL6, TL7, TL8, TL9) are determined by alumina substrate parameters $\epsilon_r = 9.95$, $h = 1.5$ mm and metal strip thickness $t = 10$ μm . The length and

Table 7

MICAP Input Data of the Microwave Amplifier in Figure 8b

MICAP

THE MICROWAVE AMPLIFIER - DISCONTINUITIES
INCLUDED

HEIGHT 1.5

THICK 10.

EPS.R 9.95

DATA

SOURCE

R1,(1-0)W=50. ; Z(1-0) ;

END

LOAD

R1,(1-0)W=50. ; Z(1-0) ;

END

MICNET

R1,(9-10)W=26.E3; R2,(10-0)W=300.;

R3,(13-14)W=33.E3; R4,(14-0)W=100.;

C1,(8-1)W=150E-12; C2,(12-18)W=150E-12;

C3,(16-17)W=150E-12; C4,(9-0)W=1.E-9;

C5,(10-0)W=1.E-9; C6,(13-0)W=1.E-9;

C7,(14-0)W=1.E-9;

TL1,NLPT,(1-0,2-0)W=5. , L=12.96 ;

TL2,NLPT,(18-0,4-0)W=1.4 , L=4.05 ;

TL11,NLPT,(11-0,12-0)W=1.4 , L=2.0 ;

TL3,NLPT,(15-0,16-0)W=0.8 , L=17.61 ;

TL4,NLPR,(11-0)W=1.4 , L=4.88 ;

TL5,NLPR,(15-0)W=0.8 , L=28.91 ;

TL6,NLPT,(2-0,9-0)W=0.2 , L=18.3 ;

TL7,NLPT,(3-0,10-0)W=0.2 , L=18.3 ;

TL8,NLPT,(4-0,13-0)W=0.2 , L=18.3 ;

TL9,NLPT,(15-0,14-0)W=0.2 , L=18.3 ;

TL10,NLPT,(3-0,11-0)W=1.4 , L=2.4 ;

TL12,NLPT,(5-0,15-0)W=1.4 , L=3.0 ;

TL13,NLPT,(17-0,6-0)W=1.4 , L=16. ;

TL14,NLPT,(7-0,8-0)W=1.4 , L=4.1 ;

DS1,OPEN,TL=4 ;

DS2,OPEN,TL=5 ;

DS3,TJST,TL=4 10 11 ;

DS4,TJAO,TL=12 5 3 ;

DS5,COMP,TL=2 ;

DS6,COMP,TL=3 ;

DS7,COMP,TL=5 ;

DS8,COMP,TL=5 ;

DS9,COMP,TL=5 ;

TF1,S3MF,(2-3-0)RN=50. ;

F=1.7E9

0.608 166. 0.123 77. 2.09 65. 0.442 -66.

TF2,S3MF,(4-5-0)RN=50. ;

F=1.7E9

0.608 161. 0.136 81. 2.37 61. 0.385 -66

Z(7-0,6-0) ;

END

FREQ

IFR: 1.7E9 ;

END

OUTPUT

GAMI, GAMO, KPSK ;

END

ANAL

STOP

width of relevant microstrips are given in MICAP input list in Table 7 (TL2' consists of TL2 and TL11). The signal source and load impedances (R_s and R_L) are 50Ω resistors. The frequency range considered is $1.5 \text{ Hz} \div 1.9 \text{ GHz}$ with a linear step.

The desired circuit functions are input and output reflection coefficients Γ_{in} , Γ_{out} (GAMI, GAMO) and transducer power gain G_t (KPSK). The results are given in Fig. 9. (dashed lines). The amplifier lay-out is presented in Fig. 8a, where the

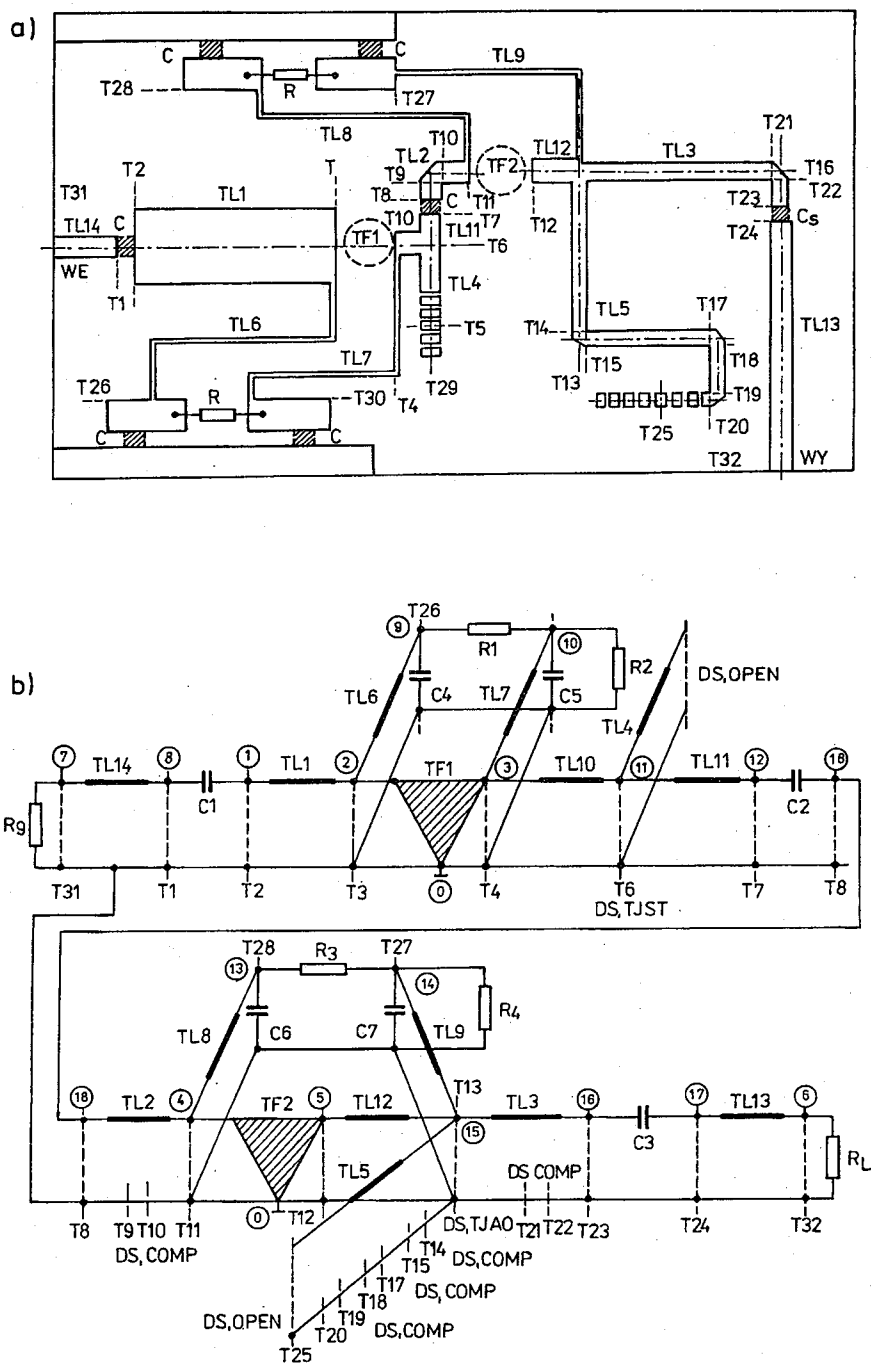


Fig. 8. The microwave amplifier lay-out (a), and its equivalent circuit (b)

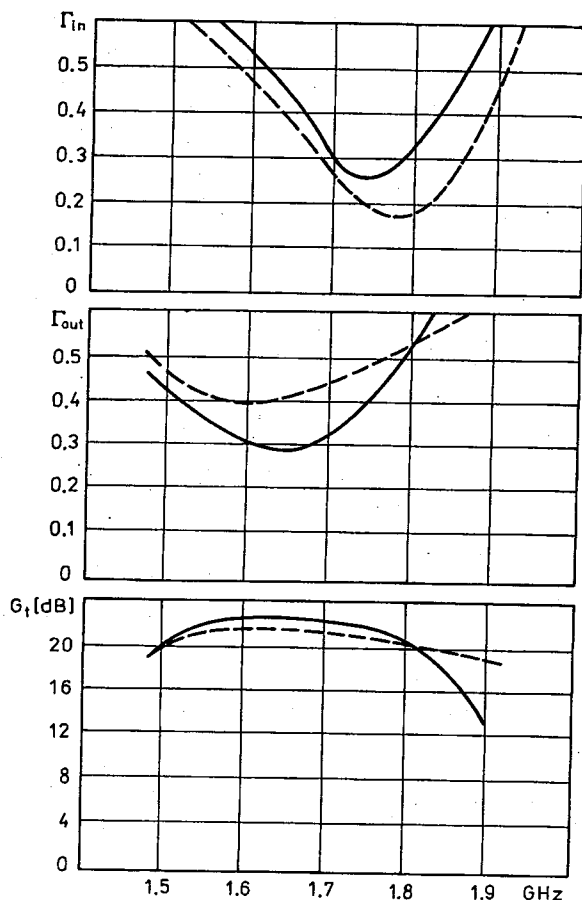


Fig. 9. The circuit functions Γ_{in} , Γ_{out} and G_t of the amplifier modeled in Fig. 7b (dashed line) and in Fig. 8b (solid line).

reference planes (T) are marked in order to introduce the majority of microstrip discontinuities. The equivalent circuit of the microwave amplifier prepared on the basis of its lay-out is shown in Fig. 8b. There are 14 transmission lines recognized, where TL2' is represented by TL2 and TL11; TL10 and TL12 are contact pads of the transistors TF1 and TF2 collectors; TL13, TL14 deal with the input and output of the amplifier. Furthermore 9 discontinuities are considered of the following type: **OPEN** (TL4, TL5), **COMP** (TL2, TL3, TL5), **TJST** (TL4, TL10, TL11) and **TJAO** (TL12, TL5, TL3). The others were omitted here because they have negligible influence on calculation accuracy. In the case of **TJST**, the microstrip TL10 is perpendicular to TL4 and TL11 and the effects of discontinuity are considered in the T6 reference plane. Similarly, for **TJSO** the perpendicular microstrip is TL5 and T13 is the affected plane.

The microstrip geometry, values of the coupling C_1 , C_2 , C_3 and blocking C_4 , C_5 , C_6 capacitors, bias circuitry resistors R_1 , R_2 , R_3 , R_4 and the records of the discontinuities are given in Table 7. Furthermore for clarity the data of TF1 and TF2 at only one frequency point $f=1.7$ GHz is inserted there. The results of the calculations for one frequency point $f=1.7$ GHz are shown in Fig. 9. (solid lines). The microwave amplifier example described above illustrates the abilities of MICAP and its usefulness in the design of hybrid MICs.

10. CONCLUSIONS

The MICAP program is intended for the analysis and design of linear microwave integrated circuits. Its algorithms are based on the nodal-admittance Y matrix together with sparse-matrix technique. Such an approach results in accuracy and efficiency of calculations and savings for the computer memory as well. MICAP handles lumped (R , G , C , L , S), distributed parameter (TL) and multiport (DF , TF) elements.

Built-in models of bipolar and unipolar microwave transistors (MT) are also accessible. Furthermore MICAP handles built-in models of microstrip line and its discontinuities (DS) based on their geometry. Therefore MICAP is a particularly useful tool for the analysis and design of linear hybrid MICs, where technology, circuit lay-out and functions are to be recognized simultaneously. One of the very handy and useful features of MICAP is its input language specifically oriented for use by microwave circuit designers. It contains many diagnostic messages designed to assist the user in correcting input errors. Furthermore MICAP also holds diagnostic tests dealing with the accuracy of the calculations. Because of the applied general nodal circuit description and sparse-matrix technique MICAP is capable of effective analysis of a wide variety of linear microwave circuits, particularly hybrid MICs.

ACKNOWLEDGMENT

The author wishes to thank Dr A. Sawicki for his helpful comments and suggestions as well as supporting some computer procedures dealing with microstrip discontinuities.

REFERENCES

1. R.D. Berry: *An Optimal Ordering of Electronic Circuit Equations for Sparse Matrix Solution*. IEEE Trans. Circuit Theory, vol. CT-18, pp.139-145, Jan. 1971
2. L.O. Chua, Lin Pen-Min: *Computer-Aided Analysis of Electronic Circuits*. Algorithms and Computational Techniques, Prentice-Hall, New Jersey, 1975
3. *Compact Engineering Microwave Software*, Palo Alto, USA
4. W.R. Curtice: *GaAs MESFET Modeling and Nonlinear CAD*. IEEE Trans. Microwave Theory Techn., vol. MPT-36, pp.220-230, February 1988

5. J.A. Dobrowolski: *Introduction to Computer Methods for Microwave Circuit Analysis and Design*. Dedham, Artech House, 1991
6. *EEsof Microwave Software*. Westlake Village, USA
7. R. Garg, I.J. Bahl: *Microstrip Discontinuities*. Int. J. Electronics, vol.45, pp.81—87, No 1, 1978
8. W.J. Getsinger: *Microstrip Dispersion Model*. IEEE Trans. Microwave Theory and Techn., vol. MTT—21, pp.34—39, 1973
9. K.C. Gupta, R. Garg, I.J. Bahl: *Microstrip Lines and Slotlines*. Dedham, Mass; Artech House, 1979
10. K.C. Gupta, R. Garg, R. Chadha: *Computer-Aided Design of Microwave Circuits*. Dedham, Mass: Artech House, 1981
11. E. Hammerstad: *Equations for Microstrip Circuit Design*. Proc. 8th European Microwave Conf., pp.268—273, Hamburg 1975
12. E. Hammerstad, O. Jensen: *Accurate Models for Microstrip Computer-Aided Design*. 1980 MTT-S Int. Microwave Symposium Digest, Washington, pp.407—409
13. Hewlett-Packard Co.: *CAD/CAE Microwave Circuit Design System*. Microwave Journal, pp.220—222, November 1987
14. R.K. Hoffman: *Handbook of Microwave Integrated Circuits*, Dedham, Mass: Artech House, 1987
15. K.S. Kundert, A. Sangiovanni-Vincentelli: *Simulation of Nonlinear Circuits in the Frequency Domain*. IEEE Trans. Computer Aided-Design, vol. CAD-5, No 4, pp.521—535, October 1986
16. R. Litwin, M. Suski: *Microwave Technique*. WNT, Warszawa, 1972 (in Polish)
17. H.M. Markowitz: *The Elimination Form of the Inverse and its Application to Linear Programming*. Management Sci., vol.3, pp.255—269, April 1957
18. V.A. Monaco: *Computer-Aided Analysis of Microwave Circuits*. IEEE Trans. Microwave Theory and Techn., vol. MTT-22, pp.249—263, March 1974
19. Z. Musiałowski, A. Sawicki, A. Zygmunt: *Computer-Aided Frequency Analysis and Design of Microwave Circuits with Regard to Specific Properties of Microstrip Line*. Reprt No I28/SPR-015/87, Institute of Telecommunication and Acoustics, Wrocław 1987
20. Z. Musiałowski: *The Frequency-Domain Computer-Analysis Program of the Hybrid Microwave Integrated Circuits*. Proc. of the 10th National Conf. Circuit Theory and Electronic Circuits, Sobieszewo, Oct. 1987, pp.719—724, (in Polish)
21. A. Ralston: *A First Course in Numerical Analysis*. McGraw-Hill, NY 1965
22. V. Rizzoli, C. Cecchetti, A. Lipparini, A. Neri: *User-Oriented Software Package for the Analysis and Optimization of Nonlinear Microwave Circuits*. IEE Proceedings, vol.133, Part H, No 5, pp.385—391, October 1986
23. S. Rosłonec: *Algorithms for Computer-Aided Design of Linear Microwave Circuits*. Dedham, Mass: Artech House, 1990
24. E. Sanchez-Sinencio, T.N. Trick, CADMIC — *Computer-Aided Design of Microwave Integrated Circuits*. IEEE Trans. Microwave Theory and Techn., vol. MTT—22, pp.309—316, March 1974
25. M.I. Sobhy, A.K. Jastrzębski: *Direct Integration Methods of Non-Linear Microwave Circuits*. Proc. 15th European Microwave Conference, pp.1110—1118, Paris 1985
26. B. Smólski: *High-frequency Transistor Amplifiers*. Warszawa, WNT, 1992, (in Polish).

Z. MUSIAŁOWSKI

**MICAP — PROGRAM KOMPUTEROWY DO ANALIZY I PROJEKTOWANIA LINIOWYCH
SCALONYCH UKŁADÓW MIKROFALOWYCH****Streszczenie**

W pracy przedstawiono komputerowy program **MICAP** przeznaczony do analizy i projektowania liniowych mikrofalowych układów scalonych. Analiza w dziedzinie częstotliwości bazuje na metodzie węzłowej w powiązaniu z metodą rzadkiej macierzy. **MICAP** analizuje układy z elementami rozłożonymi, skupionymi i aktywnymi o dowolnej topologii oraz dysponuje rezystorami, konduktorami, induktorami, kondensatorami, źródłami prądowymi sterowanymi napięciowo i elementami wielowrotowymi, takimi jak tranzystory, opisanymi ich parametrami rozproszenia lub admitancyjnymi. Ponadto dysponuje on wbudowanymi modelami mikrofalowych tranzystorów bipolarnych i unipolarnych oraz mikropaskowymi liniami transmisyjnymi i wbudowanymi modelami kilku często używanych nieciągłości niesymetrycznej linii paskowej. Efekty związane z dyspersją i grubością paska są również zawarte. **MICAP** jest dlatego szczególnie użyteczny w analizie i projektowaniu liniowych hybrydowych układów scalonych, gdzie technologia, rozkład elementów i własności sieci są jednocześnie rozważane.

A new method of PAS cell measurements

TOMASZ STARECKI

Instytut Podstaw Elektroniki, Politechnika Warszawska

Received 1993.01.14

Authorized 1993.05.26

A photoacoustic measurement system based on concept of virtual instruments is presented. Replacing hardware operations by software reduces size and cost of the arrangement and improves its metrological properties e.g. chopping frequency range, accuracy, time of measurement. Implementation of additional operation like filtering or further data processing is easy while in traditional arrangement requires some additional hardware or is not possible.

INTRODUCTION

In first photoacoustic instruments only non-resonant cells have been used [1, 2]. But it has been found in 1970's that sensitivity of the method can be significantly increased by the use of resonant cells [3, 4]. Since then a lot of different types of resonant cells have been invented [2, 5, 6] and nowadays nonresonant cells are not very often used. However if a resonant cell is to be used, its resonance curve must be determined very precisely. It is of great importance because even a small deviation from the resonance frequency decreases the output signal from the chamber. The higher Q factor is, the stronger is the decrease.

If a lot of number of measurements should be performed and a high accuracy should be obtained, the best solution is to use an automatic measurement system. In recent years some photoacoustic experiments carried out by means of such systems have been reported [7–11]. However in all of that solutions the changes of the setup structure in comparison to traditional experimental arrangement are rather small. Main parts of the setup e.g. chopper, lock-in amplifier are used in the same way as in traditional arrangements and only some interfaces to the computer are added. It is the goal of this work to realize an automatic measurement system based on concept of virtual instruments [12], that is to design such a system in which hardware would be replaced by software as far as possible.

1. STRUCTURE OF THE SYSTEM

Structure of the system is shown in Fig. 1. The main elements of the system are described below.

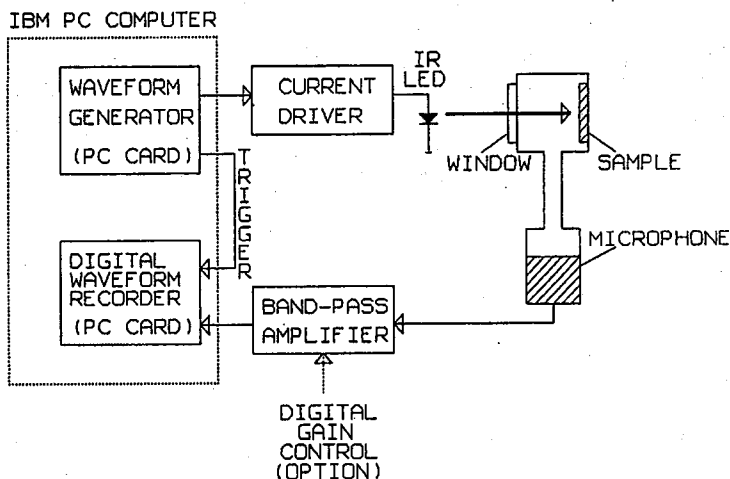


Fig. 1. Block diagram of an experimental arrangement for photoacoustic measurements. Structure of the system based on concept of virtual instruments

A. SOURCE OF RADIATION

In most experimental arrangements laser or high power arc lamp is used as a source of radiation [2, 5, 6, 13]. Modulation of the emitted light can be obtained by means of electro-optic or acousto-optic modulator, though mechanical chopper is still the most popular choice [2, 5, 6, 13]. Application of a light emitting diode as a source of modulated light has been also reported [14]. The main drawback of the last solution is small output light power. However LED diode is small, cheap and its light intensity can be easily modulated with the frequency of $0\text{ Hz} \div 1\text{ MHz}$. Moreover modulating signal can be of any shape (square, sine, triangle etc.) and stability of the light modulation frequency is equal to stability of the modulating signal frequency. Besides in some applications commercial IR LEDs working in the $0.78 \div 1.0\text{ }\mu\text{m}$ range can be used. Output light power that can be obtained from such IR LED diodes has typical value of $1 \div 100\text{ mW}$. In this work a LD271 (Litronix) diode was used. LD271 is a commercial IR LED which can emit up to 16 mW of an average light power at the wavelength of $0.95\text{ }\mu\text{m}$. The diode was controlled by a waveform generator card buffered with a simple current driver (Fig. 1).

B. SIGNAL RECORDER

Pressure changes in the photoacoustic cell used in this work were detected by an electret microphone TONSIL — MEO 55. Signal from the microphone was increased by the gain of about 10^4 in a simple band-pass amplifier built of two op-amps (NE5534A, LM318). In this way the level of the signal was matched to the input voltage level of a waveform recorder card. The card used in experiments had 8-bit resolution and 10 MHz maximum sampling rate. The signal was recorded with the resolution of 100 points a period, hence it was possible to record the signals which repetition rate were up to 100 kHz. Such high frequencies are much above the upper bandwidth limit of the microphone, but can be useful if a piezoelectric detector is used [15, 16]. Recording of a full waveform not only its amplitude and/or phase gives more information about the effects which occur in chamber. Hence applying a waveform recorder instead of a lock-in detector should improve metrological properties of the system. Small voltage resolution of the recorder can be a disadvantage, but can be easily compensated by means of an amplifier with a digital gain control. Precise self-calibration of the circuit can be done in configuration shown in Fig. 2.

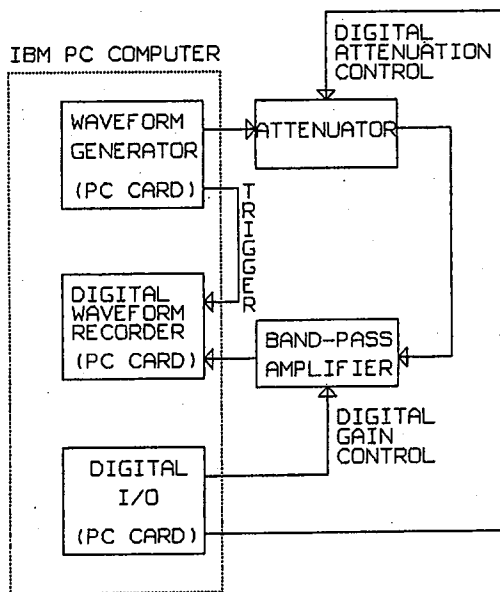


Fig. 2. Block diagram of a self-calibration arrangement for the system from fig. 1

C. SOFTWARE

The system was tested by means of a simple program for photoacoustic cell measurements. A flow chart of the program shown in Fig. 3 describes the principle of operation of the system.

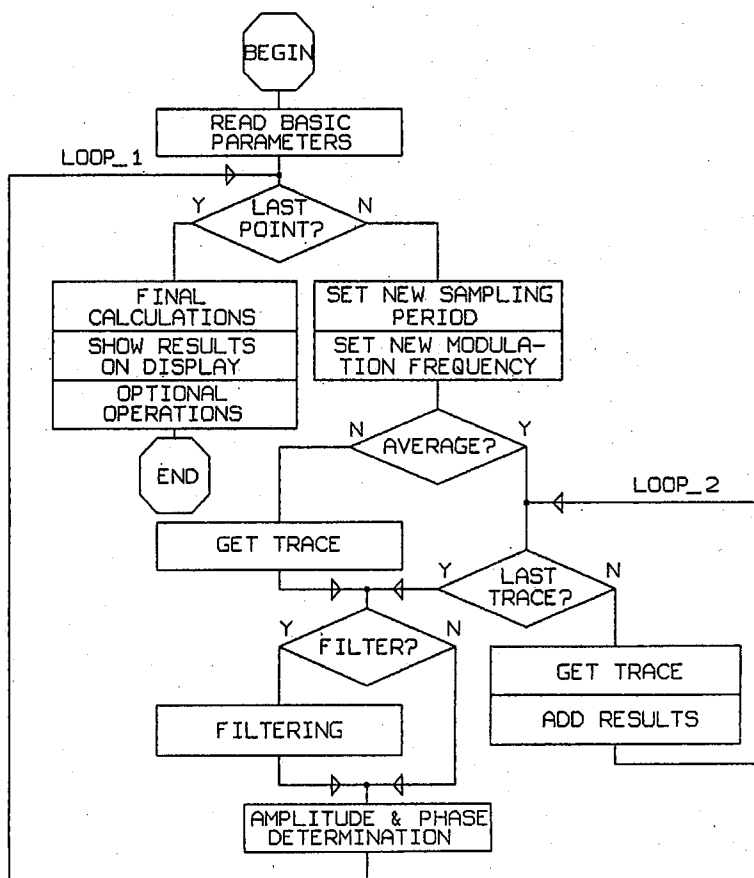


Fig. 3. Flow chart of a program for photoacoustic cell measurements performed in the setup from fig. 1

At first the program reads some basic parameters of the measurement. Then for each test point the same loop (LOOP_1) is performed. At first a modulation frequency in the generator and a new sampling rate in recorder are set. The ratio of the frequencies is always constant and equal 1:100, which makes each period of the signal consisted of 100 samples. But more important is that constant ratio of the frequencies simplifies subsequent signal processing, because all the parameters of digital filtering, averaging etc. are constant and independent of the frequency. In the next step one trace of the signal is recorded. The waveform recorder is triggered by the waveform generator (Fig. 1) in order to obtain full synchronizaton between the modulating and recorded signals. If an „average” option is switched on, the operation of signal recording and summing the samples is repeated again and again according to the set number of averages (LOOP_2). Setting of the „filter” switch determines whether digital filtering of the signal proceeds or not. Values of the amplitude and phase of the signal are determined by the least-square method. When

the last iteration in LOOP_1 is finished some additional calculations can be performed and the results of the measurements are presented on the computer screen. Certainly the results of any mentioned operation can be at any time written on a disk for further data processing or displayed on the screen in userdepending format. The contents of the screen can be at any time copied on a printer.

2. RESULTS AND DISCUSSION

A typical photoacoustic response recorded in the system is shown in Fig. 4a. The SNR of the observed signal is not very high but can be significantly increased by passing the signal through a three-pole low-pass digital elliptic filter (Fig. 4b).

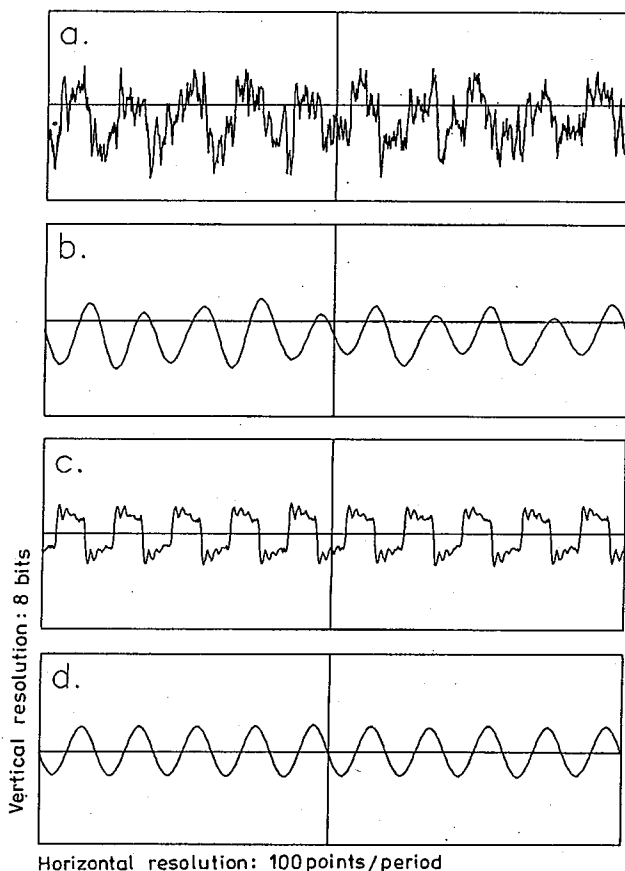


Fig. 4. Comparison of a typical photoacoustic response before and after digital signal processing.
(a) Single trace of the recorded signal. (b) Strong noise reduction as a result of filtering. (c) Original shape of the signal obtained by averaging of 100 traces. (d) Nearly pure sine wave observed after filtering of the averaged signal

However this method eliminates all the harmonics of the signal. The original shape of the microphone response can be reproduced by signal averaging (Fig. 4c). Averaged signal passed through the mentioned filter gives nearly pure sine wave (Fig. 4d). It should be noticed, that filtering changes the signal phase. However the value of the phase deviation is known and constant, hence it can be easily corrected in final calculations.

The system presented in this paper was used for determination of resonance curves of a Helmholtz resonator. The cell was equipped with a carbon black sample. The resonance curve obtained in the measurements is presented in Fig. 5.

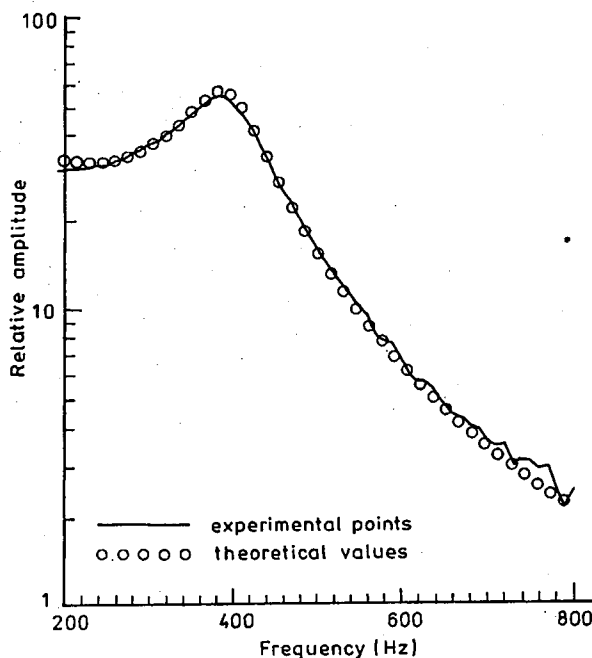


Fig. 5. Photoacoustic signal amplitude versus modulation frequency

The amplitude of the recorded signal is related to the value of the least significant bit of the waveform recorder. As can be seen, the experimental values (solid line) are close to the theoretical ones (small circles) obtained from the calculations of the extended Helmholtz resonator model [17, 18]. Noise reduction due to signal averaging and very good properties of the least-square method lead to good agreement between the theoretical and experimental values even in the range of small amplitudes.

Concept of virtual instruments applied in photoacoustic experiments gives promising results. The size and cost of equipment are considerably reduced. Additional functions e.g. temperature control, light power determination can be easily implemented. Due to high upper limit of the frequency modulation both microphone and piezoelectric detectors can be used.

REFERENCES

1. D.W. Hill, T. Powell: *Non-Dispersive Infra-Red Gas Analysis in Science, Medicine and Industry*. Plenum Press, New York, 1968
2. *Optoacoustic Spectroscopy and Detection*, edited by Y. -H. Pao, Academic Press, New York, 1977
3. C.F. Dewey, R.D. Kamm, C.E. Hackett: *Acoustic amplifier for detection of atmospheric pollutants*, Appl. Phys. Lett. 1973, vol.23, pp.633÷635
4. P.D. Goldan, K. Goto: *An acoustically resonant system for detection of low-level infrared absorption in atmospheric pollutants*. J. Appl. Phys. 1974, vol.45, pp.4350÷4355
5. A. Rosencwaig: *Photoacoustics and Photoacoustic Spectroscopy*, Wiley, New York, 1980
6. V.P. Zharov, V.S. Letokhov: *Laser Optoacoustic Spectroscopy*, Springer Series of Optical Sciences, Vol.37, edited by T. Tamir, Berlin Heidelberg, Springer-Verlag, 1986
7. A. Karbach, J. Roper, P. Hess: *Computer-controlled performance of photoacoustic resonance experiments*. Rev. Sci. Instrum. 1984, vol.55, pp.892÷895
8. P. Korpiun, W. Herrmann, A. Kindermann, M. Rothmeyer, B. Buchner: *Sorption of water investigated with the photoacoustic effect*. Can. J. Phys. 1986, vol.64, pp.1042÷1048
9. J.T. Dodgson, A. Mandelis, C. Andreetta: *Optical-absorption coefficient measurements in solids and liquids using correlation photoacoustic spectroscopy*. Can. J. Phys., 1986, vol.64, pp.1074÷1080
10. H.S. Lee, D.D. Lee: *Applicability of the Helmholtz resonator to photoacoustic difference spectroscopy*. Appl. Opt. 1988, vol.27, pp.10÷11
11. J. Roper, A. Neubrand, P. Hess: *Study of intracavity absorption and laser behaviour for C_2H_4 , C_2H_6 , CH_3Cl , and CH_4/N_2 by photoacoustic resonance spectroscopy*. J. Appl. Phys. 1988, vol.64, pp.2838÷2846
12. W. Nowakowski: *Pomiary Automatyka Kontrola — Tendencje rozwojowe w dziedzinie testowania i kontroli jakości wyrobów CAT-CAQ. Komputerowe systemy pomiarowe; wirtualne przyrządy pomiarowe*. Pomiary Automatyka Kontrola 1991, vol.37, s.15÷18
13. G.A. West, J.J. Barrett, D.R. Siebert, K.V. Reddy: *Photoacoustic spectroscopy*. Rev. Sci. Instr. vol.54 (1983), pp.797÷817
14. J.W. Chey, P. Sultan, H.J. Gerritsen: *Resonant photoacoustic detection of methane in nitrogen using a room temperature infrared light emitting diode*. Appl. Opt. 1987, vol.26, pp.3192÷3194
15. M.M. Farrow, R.K. Burnham, M. Auzanneau, S.L. Olsen, N. Purdie, E.M. Eyring: *Piezoelectric detection of photoacoustic signals*, Appl. Opt. 1978, vol.17, pp.1093÷1098
16. J.-M. Heritier, J.E. Fouquet, A.E. Siegman: *Photoacoustic cell using elliptical acoustic focusing*. appl. Opt. 1982, vol.21, pp.90÷93
17. O. Nordhaus, J. Pelzl: *Frequency dependence of resonant photoacoustic cells: The extended Helmholtz resonator*. Appl. Phys. 1981, vol.25, pp.221÷229
18. J. Pelzl, K. Klein, O. Nordhaus: *Extended Helmholtz resonator in low-temperature photoacoustic spectroscopy*. Appl. Opt. 1982, vol.21, pp.94—99

T. STARECKI

NOWA METODA POMIARU KOMÓR FOTOAKUSTYCZNYCH

Streszczenie

W artykule przedstawiono fotoakustyczny system pomiarowy skonstruowany w oparciu o ideę przyrządów wirtualnych. Zastąpienie operacji sprzętowych oprogramowaniem radykalnie zmniejsza koszt i rozmiary układu pomiarowego oraz poprawia jego parametry metrologiczne np. zakres częstotliwości modulacji, dokładność, czas pomiaru. Łatwa jest też implementacja dodatkowych operacji jak np. filtracji czy dalszej obróbki wyników. W przypadku tradycyjnych układów pomiarowych operacje te wymagały stosowania dodatkowego sprzętu lub nie były w ogóle możliwe.

Metoda uzmienniania stanów początkowych w analizie układów liniowych o okresowej odpowiedzi stanu ustalonego

ANDRZEJ LEŚNICKI

Wydział Elektroniki, Politechnika Gdańska

Otrzymano 1993.06.07

Autoryzowano do druku 1993.09.30

W pracy adaptowano metodę uzmienniania stanów początkowych (ang. shooting method) do analizy układów liniowych o okresowej odpowiedzi stanu ustalonego. Jest to metoda komputerowa pozwalająca wyznaczyć w efektywny sposób okresową odpowiedź stanu ustalonego układu liniowego. Metoda jest efektywna, gdyż jeden okres odpowiedzi stanu ustalonego jest wyznaczany od razu, bez potrzeby analizy czasowej układu w przedziale, w którym trwają procesy przejściowe. W układach zawierających podukłady o małym tłumieniu (np. układy prostownicze, wzmacniacze rezonansowe, powielacze częstotliwości) procesy przejściowe trwają przez dziesiątki i setki okresów nadmiernie wydłużając czas analizy w przypadku, gdy interesuje nas wyłącznie stan ustalony. Metoda w wersji przedstawionej w niniejszej pracy może być zastosowana do analizy tylko układów liniowych lub tych układów nieliniowych, w których część nieliniową można zamodelować stosownym źródłem napięciowym lub prądowym. Takie modelowanie jest na ogół możliwe z zadowalającą dokładnością w układach prostowniczych, wzmacniaczach klasy C, powielaczach częstotliwości, itp. W pracy zilustrowano działanie metody na przykładach analizy wzmacniacza klasy C i podwajacza częstotliwości 21/42 MHz. Metoda uzmienniania stanów początkowych jest metodą analizy w dziedzinie czasu i jest alternatywna względem analizy w dziedzinie częstotliwości.

1. WPROWADZENIE

W komputerowej analizie układów elektronicznych jedną z częściej wykonywanych analiz jest analiza w dziedzinie czasu. W jej wyniku otrzymuje się pełną odpowiedź układu z zanikającymi procesami przejściowymi, aż do osiągnięcia stanu ustalonego. Są układy, dla których jest mniej istotna znajomość procesów przejś-

ciowych, a praktycznego znaczenia nabiera możliwość szybkiego wyznaczenia okresowej odpowiedzi stanu ustalonego. Uniwersalne programy komputerowej analizy układów elektronicznych, takie jak PSPICE, NAP2, MCAP, wyznaczają pełną odpowiedź układu. Niepotrzebnie jest przeprowadzana analiza w długim przedziale czasu zanim zanikną procesy przejściowe (trwa to dziesiątki a nawet setki okresów) i osiągnięty będzie okresowy stan ustalony. W takich układach jak filtry, wzmacniacze rezonansowe, wzmacniacze klasy C, powielacze częstotliwości, prostowniki, interesuje użytkownika kształt odpowiedzi okresowej w stanie ustalonym. Przykładowo wystarcza znajomość jednego okresu odpowiedzi do przeprowadzenia analizy zniekształceń metodą szeregu Fouriera.

Istnieje wiele metod pozwalających wyznaczyć okresową odpowiedź układu *nieliniowego* w stanie ustalonym z pominięciem, lub znaczną redukcją czasu analizy przypadającego na przedział, w którym trwają procesy przejściowe:

- (a) Metoda szeregów Volterra [1]
- (b) Metoda równowagi harmonicznych [2]
- (c) Metoda równania całkowego [3]
- (d) Metoda uźmienniania stanów początkowych (ang. shooting method) [4, 5, 6].

W niniejszej pracy adaptowano metodę uźmienniania stanów początkowych dla przypadku analizy układu *liniowego*, *skupionego*, *stacjonarnego*. Metoda uźmienniania stanów początkowych ma na tle innych metod tę zaletę użytkową, że pozwala wykorzystać dobrze opracowane i rozpowszechnione podprogramy całkowania równań różniczkowych opisujących układ elektroniczny. Stosując inne metody trzeba opracowywać od podstaw stosowne programy komputerowe. Jest to metoda analizy w dziedzinie czasu i jest alternatywna względem analizy w dziedzinie częstotliwości.

Metoda uźmienniania stanów początkowych ma szereg odmian [4, 5, 6, 7]. W niniejszej pracy jest prezentowana jeszcze jedna odmiana metody polegająca zasadniczo na adaptowaniu metody do analizy układów liniowych. Nie wyklucza to jednak całkowicie możliwości stosowania jej do analizy w określonych sytuacjach układów nieliniowych. W pewnych układach ich część nieliniowa może być z dostateczną dokładnością zamodelowana jako źródło napięciowe lub prądowe. Przykładowo w tranzystorowym wzmacniaczu klasy C lub powielaczu częstotliwości wyjście tranzystora modeluje się źródłem prądowym obciążonego przebiegu sinusoidalnego o określonym kącie odcięcia θ . Podobnie w prostowniku diodowym część prostownicza układu przed filtrem wygładzającym może być przedstawiona jako źródło napięciowe przebiegu sinusoidalnego wyprostowanego jednopółkowo (prostownik jednopółkowy) lub dwupółkowo (prostownik dwupółkowy).

2. METODA ANALIZY

Zostanie wyprowadzona metoda analizy polegająca na adaptowaniu metody uźmienniania zmiennych stanu dla przypadku analizy układu liniowego.

Niech układ liniowy będzie opisany macierzowym równaniem stanu o postaci normalnej

$$\frac{d\bar{x}(t)}{dt} = \bar{A}\bar{x}(t) + \bar{B}\bar{w}(t), \quad \bar{x}(t_0) = \bar{x}_0, \quad (1)$$

gdzie $\bar{x}(t)$ jest n -wymiarowym wektorem stanu, zaś $\bar{w}(t)$ jest wektorem pobudzeń. Rozwiązanie równania (1) jest postaci

$$\bar{x}(t) = \bar{\Phi}(t-t_0)\bar{x}_0 + \int_{t_0}^t \bar{\Phi}(t-\tau)\bar{B}\bar{w}(\tau) d\tau, \quad (2)$$

gdzie

$$\bar{\Phi}(t-t_0) = e^{\bar{A}(t-t_0)} \quad (3)$$

jest macierzą tranzycji (przenoszenia) stanów. Rozwiązanie dane wzorem (2) zawiera pełną odpowiedź układu, łącznie z procesami przejściowymi. Stan ustalony może być wyznaczony z tego wzoru jako stan po dostatecznie długim czasie analizy

$$\bar{x}_{ss}(t) = \lim_{t \rightarrow \infty} \left[\bar{\Phi}(t-t_0)\bar{x}_0 + \int_{t_0}^t \bar{\Phi}(t-\tau)\bar{B}\bar{w}(\tau) d\tau \right]. \quad (4)$$

Jeżeli wektor pobudzeń jest T -okresowy

$$\bar{w}(t) = \bar{w}(t+T) \quad (5)$$

to istnieje rozwiązanie okresowe równania (1)

$$\bar{x}(t) = \bar{x}(t+T). \quad (6)$$

Metoda uzmienniania stanów początkowych polega na poszukiwaniu takiego stanu początkowego $\bar{x}_0^*$, począwszy od którego rozwiązanie równania stanu (1) jest rozwiązaniem okresowym stanu ustalonego (6), z pominięciem procesów przejściowych.

Uzmienniony stan początkowy $\bar{x}_0^*$ zostanie wyznaczony z warunku okresowości (6) zapisanego dla dowolnej chwili początkowej t_0 (najczęściej $t_0 = 0$)

$$\bar{x}(t_0) = \bar{x}_0 = (\bar{x}_0 + T). \quad (7)$$

Wartość $\bar{x}(t_0 + T)$ może być wyznaczona ze wzoru (2) i wtedy warunek okresowości (7) przyjmie postać równania

$$\bar{x}_0 = \bar{\Phi}(T)\bar{x}_0 + \int_{t_0}^{t_0+T} \bar{\Phi}(t-\tau)\bar{B}\bar{w}(\tau) d\tau \quad (8)$$

o poszukiwanym rozwiązaniu

$$\bar{x}_0^* = [\bar{I} - \bar{\Phi}(T)]^{-1} \int_{t_0}^{t_0+T} \bar{\Phi}(t-\tau) \bar{B} \bar{w}(\tau) d\tau, \quad (9)$$

gdzie $\bar{I}$ — macierz jednostkowa o wymiarach $n \times n$. Całkowanie równania stanu (1) począwszy od stanu początkowego $\bar{x}_0^*$ określonego wzorem (9) daje od razu odpowiedź okresową stanu ustalonego, z pominięciem procesów przejściowych. Wyprowadzony wzór (9) można zinterpretować jako wynik zastosowania pojedynczego kroku iteracyjnego algorytmu Newtona z metody Colana Tricka [5] dla układów nieliniowych, w którym wykorzystano uproszczenia wynikające z założenia liniowości układu.

Sposób wykorzystania wzoru (9) w obliczeniach numerycznych zależy od przyjętej metody analizy układu elektronicznego. Jeżeli układ jest analizowany metodą zmiennych stanu, to jest znana macierz $\bar{A}$ i wartość macierzy tranzykcji stanów $\bar{\Phi}(T)$ oblicza się bezpośrednio ze wzoru (3)

$$\bar{\Phi}(T) = e^{\bar{A}T}. \quad (10)$$

Wystarczy jednorazowa analiza układu w przedziale czasu $[t_0, t_0+T]$ dla zerowego stanu początkowego $\bar{x}_0(t_0) = \bar{x}_0 = \bar{0}$, aby zgodnie ze wzorem (2) wyznaczyć wartość całki

$$\bar{x}(t_0+T) \Big|_{\bar{x}_0=0} = \int_{t_0}^{t_0+T} \bar{\Phi}(t-\tau) \bar{w}(\tau) d\tau \quad (11)$$

potrzebną do obliczenia $\bar{x}_0^*$ z zależności (9).

W uniwersalnych programach analizy komputerowej układów elektronicznych (np. PSPICE, NAP2, MCAP) jest powszechnie stosowana analiza uogólnioną metodą napięć węzłowych [8]. W tym przypadku nie jest znana macierz $\bar{A}$, i aby wyznaczyć $\bar{\Phi}(T)$ należy wykonać n -krotnie analizę układu w przedziale czasu $[t_0, t_0+T]$, przy zerowym pobudzeniu $\bar{w}(t) = \bar{0}$ i jednostkowych stanach początkowych

$$\bar{x}_0^{(1)} = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ \vdots \\ 0 \end{bmatrix}, \quad \bar{x}_0^{(2)} = \begin{bmatrix} 0 \\ 1 \\ \vdots \\ \vdots \\ 0 \end{bmatrix}, \quad \dots, \quad \bar{x}_0^{(n)} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ \vdots \\ 1 \end{bmatrix}. \quad (12)$$

Przy tych warunkach będzie zgodnie ze wzorem (2)

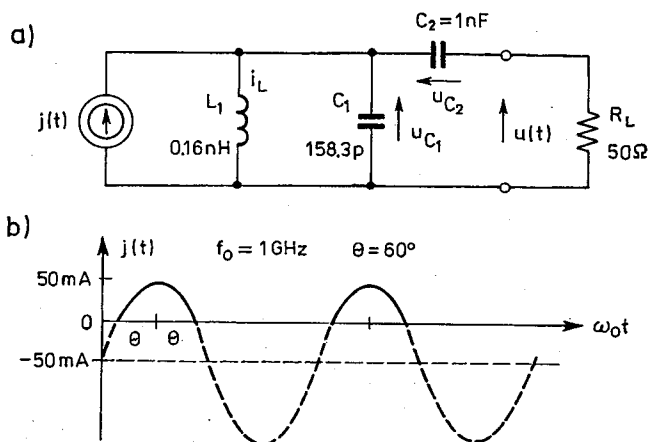
$$\left[\bar{x}^{(1)}(t_0+T), \bar{x}^{(2)}(t_0+T), \dots, \bar{x}^{(n)}(t_0+T) \right] = \bar{X}(t_0+T) = \bar{\Phi}(T) \bar{I} = \bar{\Phi}(T) \quad (13)$$

Tak więc, aby było możliwe wyznaczenie stanu początkowego $\bar{x}_0^*$ ze wzoru (9), należy przeprowadzić analizę układu w przedziale czasu $[t_0, t_0+T]$ w sumie $(n+1)$ -krotnie. Jeden raz przy zerowych warunkach początkowych $x_0 = \bar{0}$, aby wyznaczyć całkę (11), i n razy przy jednostkowych stnach początkowych i zerowym pobudzeniu $w(t) = \bar{0}$, aby wyznaczyć wartość macierzy tranzycji stanów (13).

3. PRZYKŁADY OBLICZENIOWE

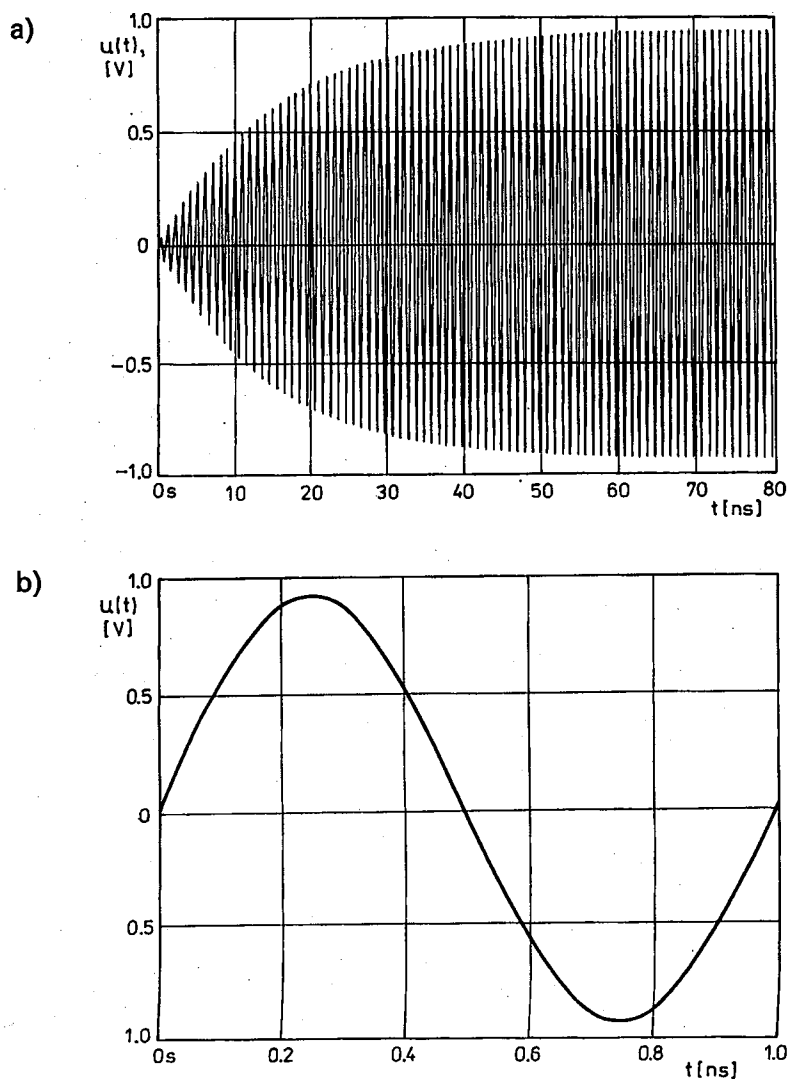
Dwa przykłady obliczeniowe ilustrują sposób i efekty wykorzystania wyprowadzonej metody analizy.

W pierwszym przykładzie będzie analizowany układ wzmacniacza klasy C o częstotliwości środkowej $f_0 = 1$ GHz (rys. 1). Wzmacniacz klasy C jest układem nieliniowym, ale część nieliniowa układu służy tylko do przetworzenia pobudzenia harmonicznego na pobudzenie z kątem odcięcia θ . Dlatego z punktu widzenia ob-



Rys. 1. Wzmacniacz klasy C: a) obwód wyjściowy wzmacniacza, b) pobudzenie prądowe obwodu wyjściowego

ciążenia R_L układ może być przedstawiony jako liniowy obwód wyjściowy pobudzony przebiegiem prądowym o kącie odcięcia θ . Analiza układu począwszy od zerowego stanu początkowego, aż do osiągnięcia stanu ustalonego (z dokładnością 1%) przebiega przez około 80 okresów (rys. 2a) i trwa 1060,96s na komputerze typu IBM/PC/486DX. Analiza układu metodą uzmienniania stanów początkowych pozwala wyznaczyć okresową odpowiedź stanu ustalonego (rys. 2b) w ciągu 30,81s (tj. 34,4 razy szybciej). W tym przypadku analiza była przeprowadzona metodą zmien-



Rys. 2. Napięcie wyjściowe wzmacniacza klasy C: a) pełna odpowiedź układu; b) odpowiedź w stanie ustalonym

nych stanu za pomocą własnego programu napisanego w języku TURBO PASCAL, w którym całkowanie równań stanu o postaci normalnej odbywa się niejawną metodą Eulera ze względną tolerancją 0,00003. Macierz tranzycji stanów wyznaczona ze wzoru (10) w układzie jednostek pochodnym od SI; V, A, nH, nF, ns, GHz

$$\bar{\Phi}(T) = e^{\bar{A}T} = \begin{bmatrix} 0,9358 & 0 & 0,002 \\ 0 & 0,9802 & -0,0001 \\ -0,0002 & 0,0008 & 0,9388 \end{bmatrix} \quad (14)$$

Wartości uziemionych stanów początkowych wyznacza się w praktyce rozwiązując układ równań liniowych (8), a nie ze wzoru (9) wymagającego odwrócenia macierzy, co zajęłoby ok. 3 razy więcej czasu obliczeń. Wyniki są następujące

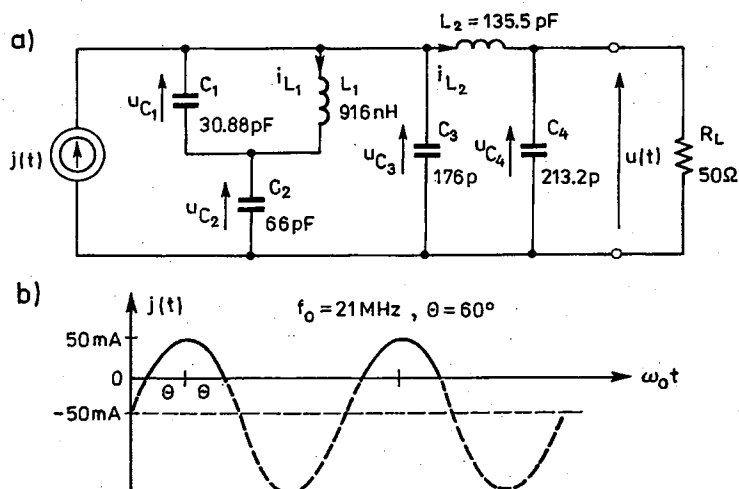
$$u_{C1}(0) = 0,0196V, \quad u_{C2}(0) = 0,0029V, \quad i_{L1}(0) = -0,9053A \quad (15)$$

Rozwiązanie wyznaczone począwszy od tego stanu początkowego jest rozwiązaniem okresowym stanu ustalonego (rys. 2b). Analiza Fouriera wyznaczonego rozwiązania okresowego pozwala określić amplitudy harmonicznych: $U_1 = 0,9267V$, $U_2 = 0,0068V$, $U_3 = 0,0030V$; oraz współczynnik zawartości harmonicznych $h = 0,83\%$.

Ten sam układ analizowano tą samą metodą, ale za pomocą uniwersalnego programu PSPICE. Trzykrotna analiza układu przy zerowym pobudzeniu i jednostkowych stanach początkowych daje według wzoru (13) macierz tranzycji stanów

$$\bar{\Phi}(T) = \begin{bmatrix} 0,9391 & -0,0001 & 0,0007 \\ 0 & 0,9802 & -0,0001 \\ -0,0007 & 0,0008 & 0,9389 \end{bmatrix} \quad (16)$$

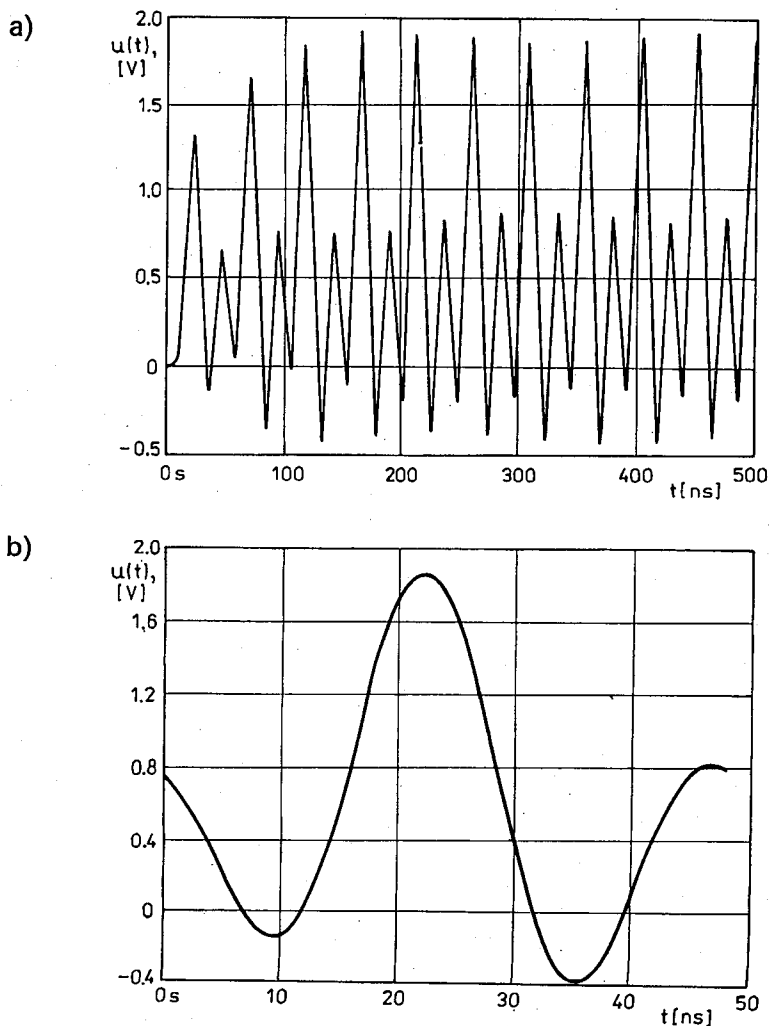
nie różniącą się praktycznie od poprzedniego wyniku (14) i prowadzącą do takich jak poprzednio wartości stanów początkowych (15). Sumaryczny czas analizy



Rys. 3. Podwajacz częstotliwości 21/42 MHz: a) obwód wyjściowy b) pobudzenie prądowe obwodu wyjściowego

metodą uziemienniania stanów początkowych wyniósł 6 s, podczas gdy analiza układu od zerowego stanu początkowego do osiągnięcia stanu ustalonego trwała 56 s.

W drugim przykładzie obliczeniowym będzie analizowany układ podwajacza częstotliwości 21 MHz/42 MHz (rys. 3). Układ podwajacza częstotliwości został przedstawiony z punktu widzenia obciążenia R_L jako liniowy obwód wyjściowy pobudzony przebiegiem prądowym o kącie odcięcia θ . Zadaniem obwodu wyjściowego jest przepuszczenie drugiej harmonicznej przy jednoczesnym stłumieniu pierwszej, trzeciej i wyższych harmonicznych. Analiza układu począwszy od zerowego stanu początkowego, aż do osiągnięcia stanu ustalonego (rys. 4a) przebiega przez około 10 okresów i trwa 281 s. Analiza układu metodą uziemienniania stanów



Rys. 4. Napięcie wyjściowe podwajacza częstotliwości 21/42 MHz: a) pełna odpowiedź układu, b) odpowiedź w stanie ustalonym

początkowych pozwala wyznaczyć okresową odpowiedź stanu ustalonego (rys. 4b) w ciągu 47,2 s (tj. 6 razy szybciej). Analiza była przeprowadzona metodą zmiennych stanu, znana była macierz $\bar{A}$ i macierz tranzycji stanów wyznaczono ze wzoru (10). Wartości uzmiennionych stanów początkowych obliczone ze wzoru (9) są następujące: $u_{C2}(0) = 0,5524$ V, $u_{C3}(0) = -0,3883$ V, $u_{C4}(0) = 0,7454$ V, $i_{L1}(0) = -0,0034$ A, $i_{L2}(0) = 0,0057$ A (napięcie u_{C1} nie było zmienną stanu, jest ono uzależnione od napięć u_{C2} , u_{C3}). Analiza Fouriera wyznaczonego sygnału okresowego wyjściowego pokazanego na rys. 4b pozwala określić składową stałą i amplitudy harmonicznych: $U_0 = 0,5423$ V, $U_1 = 0,5141$ V, $U_2 = 0,7830$ V, $U_3 = 0,0294$ V, $U_4 = 0,0064$ V. Analiza tego układu metodą uzmienniania stanów początkowych, ale za pomocą programu PSPICE trwała 11 s, podczas gdy analiza układu od zerowego stanu początkowego do osiągnięcia stanu ustalonego trwała 18 s. W tym przykładzie efektywność metody uzmienniania stanów początkowych jest mniejsza niż w poprzednim przykładzie, gdyż procesy przejściowe trwają w układzie znacznie krócej. Dodatkowo, jeśli macierz tranzycji stanów jest wyznaczana według wzoru (13), to efektywność metody maleje proporcjonalnie do liczby zmiennych stanu, co wiąże się z nakładem czasu na n-krotną analizę układu przy jednostkowych stanach początkowych.

4. ZAKOŃCZENIE

Przedstawiono efektywną metodę analizy układów liniowych o okresowej odpowiedzi stanu ustalonego. Metoda ta jest adaptacją metody uzmienniania stanów początkowych, znanej dla układów nieliniowych, na przypadek układów liniowych, skupionych, stacjonarnych. Wyprowadzono wzory pozwalające obliczyć takie wartości zmiennych stanu, począwszy od których całkowanie równań różniczkowych opisujących układ elektroniczny daje poszukiwane rozwiązanie okresowe stanu ustalonego. Mimo że metodę opracowano dla układów liniowych, to może ona być stosowana do analizy tych układów nieliniowych, w których część nieliniowa daje się przedstawić jako źródło napięciowe lub prądowe. Układami tego typu są prostowniki, wzmacniacze klasy C, powielacze częstotliwości. Zamieszczone dwa przykłady obliczeniowe wskazują na efektywność metody analizy. Pozwala ona skrócić czas analizy od kilku do kilkudziesięciu razy w porównaniu z wyznaczaniem okresowego stanu ustalonego poprzez całkowanie równań układu tak długo, aż zanikną procesy przejściowe.

BIBLIOGRAFIA

1. L.O. Chua, Y.S. Tang: *Nonlinear oscillation via Volterra series*. IEEE Trans. on CS, Vol. CAS-29, No 2, pp. 150–168, Feb. 1982
2. M.S. Nakhla, J. Vlach: *A piece-wise harmonic balance technique for determination of periodic response of nonlinear systems*. IEEE Trans. on CS, Vol. CAS-23, No 2, pp. 85–91, Feb. 1976
3. D.R. Frey, O. Norman: *An integral equation approach to the periodic steady-state problem in nonlinear circuits*. IEEE Trans. on CS, Vol. 39, No 9, pp. 744–755, Sept. 1992

4. T.J. Aprille, T.N. Trick: *Steady-state analysis of nonlinear circuits with periodic inputs*. Proc. of the IEEE, Vol. 60, No 1, pp. 108–114, Jan. 1972
5. F.R. Colon, T.N. Trick: *Fast periodic steady-state analysis for large-signal electronic circuits*. IEEE J. Solid-State Circuits, Vol. SC-8, No 4, pp. 260–269, Aug. 1973
5. M.S. Nakhla, F.H. Branin: *Determining the periodic response of nonlinear systems by a gradient method*. Int. Jour. of Circuit Theory and Appl., Vol. 5, No 3, pp. 255–273, Jul. 1977
6. S. Skelboe: *Computation of the periodic steady-state response of nonlinear networks by extrapolation methods*. IEEE Trans. on CS, Vol. CAS-27, No 3, pp. 161–175, March 1980
7. C.W. Ho, A.E. Ruehli, P.A. Brennan: *The modified nodal approach to network analysis*. IEEE Trans. on CS, Vol. CAS-22, No 6, pp. 504–509, June 1975

A. LEŚNICKI

THE SHOOTING METHOD IN ANALYSIS OF LINEAR CIRCUITS WITH PERIODIC STEADY-STATE RESPONSES

S u m m a r y

The shooting method has been adapted for analysing of linear circuits with periodic steady-state responses. It is a computer method for effective calculating of the periodic steady-state responses of linear circuits. The method is effective because the one period of the steady-state response is calculated directly, with no analysis necessary in the time where transient processes last. These transient processes last for tens and hundreds of periods in circuits with lightly damped subcircuits, such as rectifiers, high-Q amplifiers, frequency multipliers. It unnecessarily extends the time of analysis in the case we are interested only in the periodic steady-state. In the version presented in this paper, the method can be used only for analysing of linear circuits and these nonlinear circuits where nonlinear parts can be modeled by proper voltage or current sources. Such a modeling is usually possible with an excepted exactness for rectifiers, class C amplifiers, frequency multipliers, etc. In the paper, the effectiveness of the method is demonstrated on two examples. The class C amplifier and the frequency doubler 21/42 MHz have been analysed. The class C amplifier and the frequency doubler 21/42MHz have been analysed. The shooting method is a time domain method and is an alternative method in relation to frequency domain methods.

53.087.92
621.383.813

Problemy zapewnienia jakości w procesie wytwarzania drutów. Badania nieniszczące z wykorzystaniem metody prądów wirowych

ANNA LEWIŃSKA-ROMICKA

*Centrum Uczelniano-Przemysłowe Metrologii i Systemów Pomiarowych,
Politechnika Warszawska*

Otrzymano 1993.06.02

Autoryzowano do druku 1993.

Opisano metodykę oceny jakości obiektów metalowych — na przykładzie kontroli drutów wolframowych i molibdenowych. Przedstawiono nowe kierunki, jakie pojawiły się w badaniach nieniszczących opartych na wykorzystaniu zjawiska prądów wirowych w odniesieniu do różnych sposobów podejścia do wstępnie przetworzonych w układach defektoskopów sygnałów przetworników wiropądowych.

1. WSTĘP

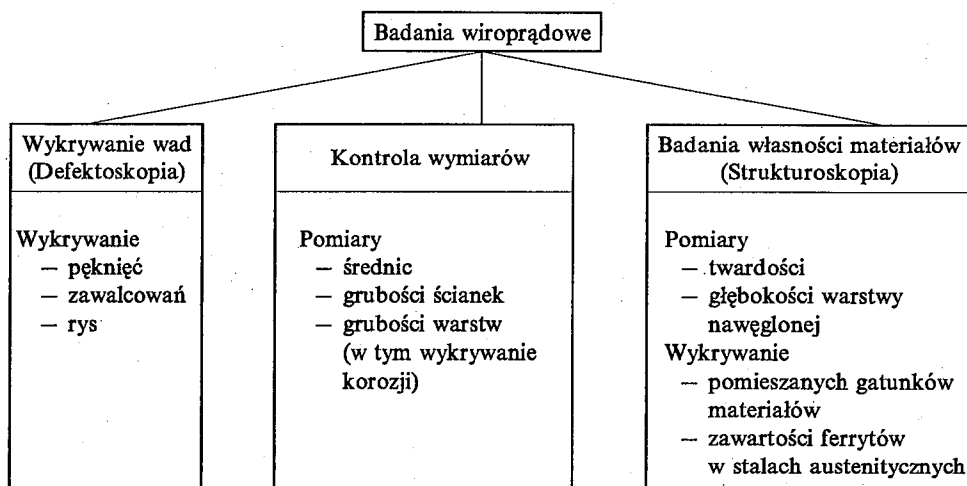
Procesy zapewnienia jakości mają na celu takie ustalenie warunków produkcji, procesów i serwisu, aby spełnione być mogły oczekiwania klientów, takie, jak zgodność z dokumentacją, przydatność do określonego celu, brak uszkodzeń, minimalny wymagany zakres napraw oraz stosowność wartości produktu do jego ceny. Istnieje wobec powyższego potrzeba opracowywania systemów monitorowania zarówno procesów wytwarzania, jak i jakości [16] otrzymywanego produktu. Jednym z szeregu ogniw w procesach zapewnienia jakości są badania nieniszczące. Ich wdrażanie umożliwia wykrycie uszkodzeń linii produkcyjnych i/lub defektów produktów we wczesnych fazach wytwarzania, a także opracowanie przez personel lub system diagnozy i podjęcie odpowiednich kroków dla poprawy zaistniałego stanu rzeczy. Często niewystarczająca lub zła jakość produktu ma swoją przyczynę w uszkodzeniach linii wytwarzania.

Wymagania jakościowe w stosunku do produktów wytwarzanych przez przemysł metalowy powodują, że konieczne staje się wdrażanie do procesów wytwarzania różnych technik badań nieniszczących. Do technik tych należą: radiograficzna (wykorzystanie promieniowania rentgenowskiego i promieniowania izotopowego),

elektromagnetyczna (wykorzystanie zjawiska prądów wirowych), magnetyczna, ultradźwiękowa i penetracyjna. Każda z wymienionych technik oznacza się specyficznymi cechami, związanymi z możliwościami zastosowań ale i z ograniczeniami tychże [6, 9, 11 ÷ 15, 18 ÷ 24, 28, 29, 32, 33 ÷ 37, 39].

Systemy zapewnienia jakości produktu i prawidłowego funkcjonowania zakładu wytwarzającego muszą spełniać wymaganie odnośnie krótkiego czasu uzyskiwania odpowiednich informacji w związku z pracą w czasie rzeczywistym. Wybór wielkości, jakie mają być mierzone w procesie zapewnienia jakości, zależy od obiektu wytwarzanego. Podczas wytwarzania drutów, prądów i rur do oceny procesów i produktów są odpowiednie wielkości uzyskiwane podczas badań nieniszczących, takich, jak wiroprowowe, magnetyczne oraz ultradźwiękowe.

Możliwości zastosowania techniki wiroprowowej ilustrują rys. 1 i rys. 2.



Rys. 1. Przegląd zastosowań metody prądów wirowych w badaniach nieniszczących

Typ wady	Badania wiroprowowe	Inne techniki badań
powierzchniowa	+	penetracyjna magnetyczna ultradźwiękowa spadku potencjału
podpowierzchniowa	+	magnetyczna ultradźwiękowa
wewnętrzna	– (z wyjątkiem obiektów cienkościennych)	ultradźwiękowa radiograficzna

Rys. 2. Ilustracja do opisu zastosowań metody prądów wirowych – na tle innych metod badań nieniszczących

Metoda prądów wirowych może być zastosowana do defektoskopii obiektów (wykrywania wad), strukturokopii (oceny struktury i wykrywania pomieszanych gatunków materiałów) oraz kontroli wymiarów (rys. 1). Poszczególne czynniki, jakie mają być wykrywane, wpływają na zmiany bądź przewodności elektrycznej, bądź przenikalności magnetycznej obiektów, względnie jednej i drugiej równocześnie lub też — na wypełnienie przetwornika przez obiekt lub oddalenie przetwornika od obiektu.

Technika wiroprowadowa umożliwia bezkontaktowe badanie różnorodnych produktów, w szczególności drutów. Jej zasadniczym ograniczeniem jest stosowalność jedynie do elementów przewodzących prąd elektryczny (metalowych). Umożliwia przeprowadzanie, z dużą szybkością, kontroli zarówno półproduktów, jak i produktów gotowych.

Szczególną cechą rozplywu w obiekcie badanym prądów wirowych jest ich naskórkowość wynikająca z ograniczonej głębokości wnikania pola magnetycznego (pkt. 5.3). Wykrywanie wad wewnętrznych możliwe jest jedynie w obiektach o grubości porównywalnej z głębokością penetracji (rys. 2) [42].

Przedmiotem niniejszego artykułu są zagadnienia defektoskopii drutów.

2. CHARAKTERYSTYKA I WADY OBIEKTÓW BADANYCH

Jako obiekty badane zostały wybrane druty wykonane z materiałów nieferromagnetycznych. Wybrano druty wolframowe i molibdenowe, które ze względu na swoje właściwości strukturalne i zastosowania powinny podlegać 100% kontroli nieniszczącej. Szczególnie przydatna do ich kontroli jest metoda oparta na wykorzystaniu zjawiska prądów wirowych.

Niżej zostanie podana charakterystyka struktury i wad drutów wolframowych i molibdenowych. Druty te wytwarza się na drodze metalurgii proszkowej. I tak na przykład chemicznie czysty tlenek wolframowy, po wprowadzeniu doń domieszek w postaci krzemianu potasu i związków glinu, poddaje się redukcji wodorem. Otrzymany w ten sposób proszek wolframowy prasuje się następnie w sztaby o przekroju kwadratowym. Sztaby te są spiekane w temperaturze około 3000 stopni Celsjusza — w celu otrzymania litego metalu, nadającego się do obróbki plastycznej. Sposób otrzymywania drutów molibdenowych jest następujący. Wzbogacony koncentrat molibdenitu, to jest rudy składającej się z siarczku molibdenu praży się w piecach hutniczych w celu otrzymania bezwodnika molibdenowego, który następnie oczyszcza się z domieszek. W tym celu rozpuszcza się go w amoniaku i krystalizuje w postaci paramolibdenianu amonu. Jest to półfabrykat do produkcji molibdenu metalicznego. Półfabrykat ten przesiewa się, suszy się i ponownie przesiewa. Dalsze operacje są analogiczne, jak w przypadku wolframu. Zaznaczyć trzeba, że proszek molibdenowy łatwo absorbuje tlen.

Następujące po sobie kolejno operacje kucia nagrzaných sztab wolframowych, jak też molibdenowych, a dalej młotkowania prętów oraz przeciągania drutów,

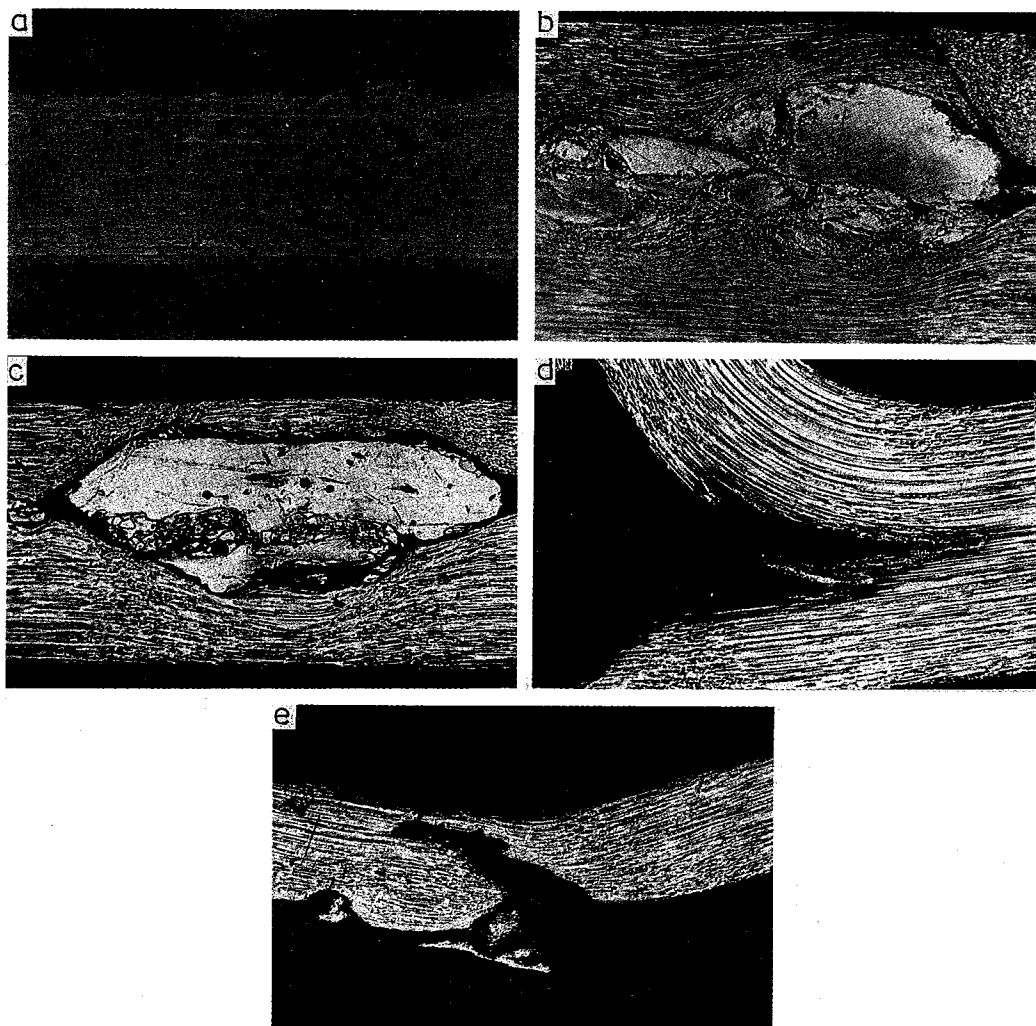
mające na celu stopniowe zmniejszenie ich przekroju, powodują powstawanie włóknistej struktury. Struktura taka jest charakterystyczna dla otrzymywanych drutów.

W toku ciągnięcia kryształy ulegają stopniowemu wydłużeniu, co nadaje drutom silnie zaznaczoną teksturę w całym przekroju. Tego typu tekstura wraz z małą spójnością międzykrystaliczną powodują, że druty bardzo często zawierają różnego rodzaju nieciągłości takie, jak pęknięcia, rozwarstwienia, a także — niezgrzane pustki oraz wtrącenia. Często wady tego typu przebiegają wzdłuż płaszczyzn międzykrystalicznych — w kierunku osi drutów [19].

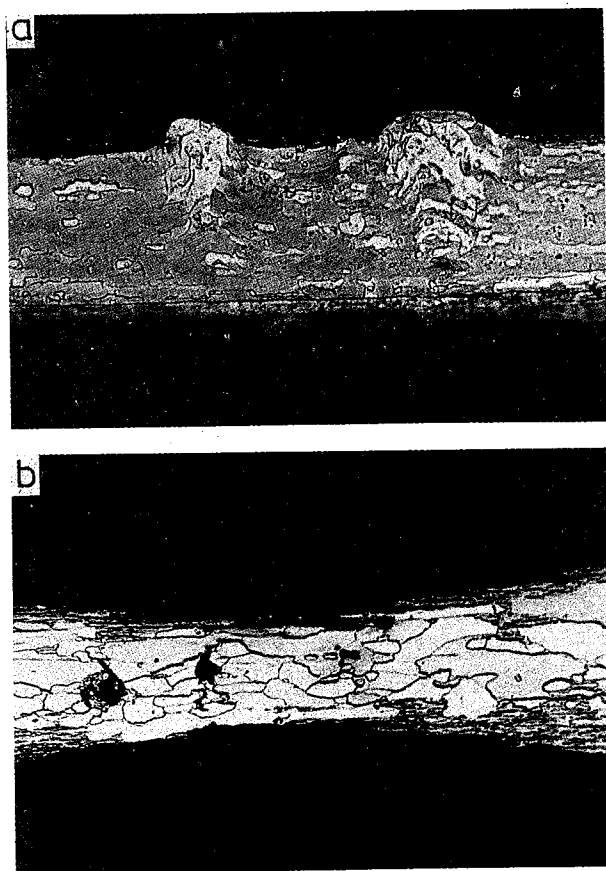
Pod wpływem ciągnięcia, jak już zaznaczono, kryształy metalu przekształcają się w cienkie włókna, których orientacja pokrywa się z osią drutów, co zostanie zilustrowane na zamieszczonych zglądach metalograficznych. Włókna wydłużają się, a ich średnica maleje wraz ze zmniejszaniem średnicy drutów. Jednocześnie ze wzrostem stopnia deformacji ziaren zmniejsza się porowatość drutów. Przy wyżarzaniu (w temperaturach 1100—2300 stopni Celsjusza) zachodzi rekrytalizacja pierwotna. Włókna wolframu ulegają rozpadowi i tworzą się drobne ziarna równoosiowe. W temperaturach 2400—2500 stopni Celsjusza zachodzi rekrytalizacja wtórna. Małe ziarna zrastają się tworząc duże kryształy.

Czysty wolfram, ze względu na wysoką temperaturę topnienia oraz dobre właściwości wytrzymałościowe w wysokich temperaturach, znalazł w postaci drutów zastosowanie głównie do produkcji żarników do lamp żarowych, katod oraz grzejników do lamp elektronowych dużej oraz małej mocy, do wyrobu elementów grzejnych do pieców oporowych oraz elektrod do lamp rentgenowskich (katody i antykatody), a także do świec zapłonowych i w postaci prętów — jako elektrody do spawania i zgrzewania. Czysty molibden w postaci drutów znalazł zastosowanie przede wszystkim w technice wysokiej próżni i w radiotechnice do produkcji elektrod do lamp elektronowych, do wyrobu prętów podtrzymujących włókna żarowe w lampach żarowych oraz elementów grzejnych w piecach oporowych, itp. Piece molibdenowe znalazły zastosowanie do spiekania wyprasek molibdenowych i wolframowych w atmosferze wodorowej. Z molibdenu wykonuje się również elektrody do topienia szkła.

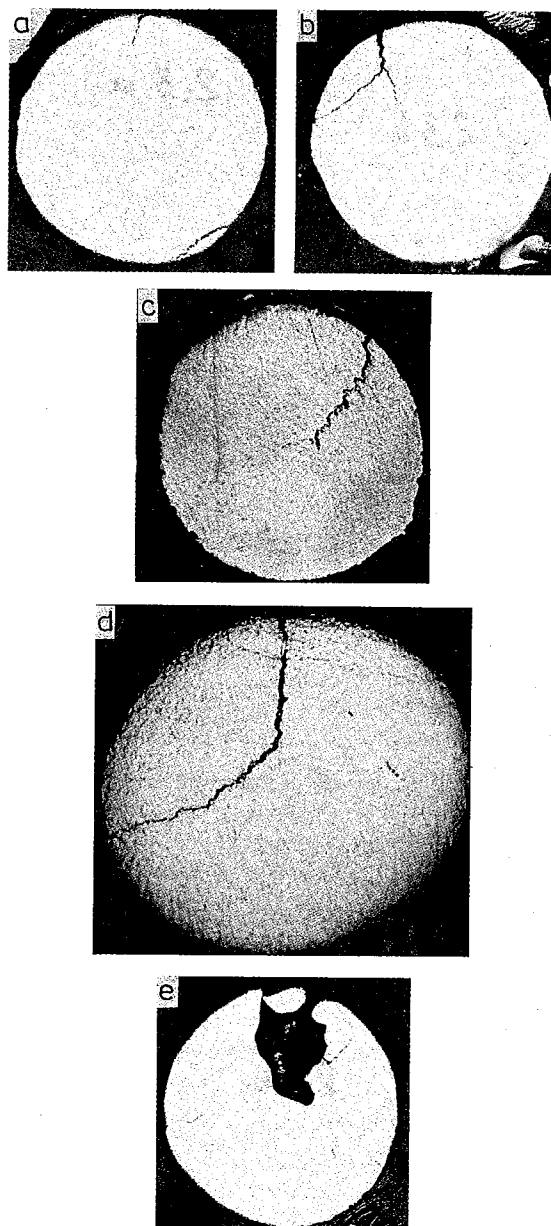
Dla zilustrowania celowości implementowania nieniszczącej, wiroprowodowej kontroli drutów wolframowych i molibdenowych przedstawione zostaną przykładowe fotogramy zglądów metalograficznych wzdłużnych (nietrawionych i trawionych) drutów molibdenowych o średnicach od 0,05 mm do 1,02 mm. Zaobserwowane wady zostały wykryte podczas badań defektoskopowych z zastosowaniem metody indukcyjnej opartej na wykorzystaniu zjawiska prądów wirowych. Wady te to przede wszystkim nieciągłości makrostruktury objawiające się w postaci pustek (na przykład fot. 3a), wtrąceń (fot. 3b, c), rozwarstwień (fot. 3d) lub pęknięć (fot. 3e). Inną kategorią wad zaobserwowanych w drutach molibdenowych są wady o charakterze deformacji geometrycznych. Należą do nich na przykład, pokazane na fot. 4a, wycieki oraz przedstawiona na fot. 4b znaczna nieciągłość średnicy drutów.



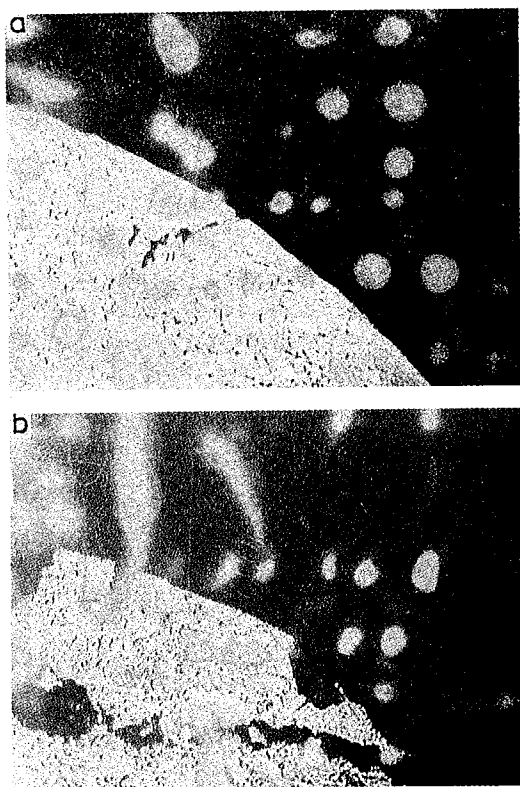
Rys. 3. Fotogramy zglądów metalograficznych wzdłużnych drutów molibdenowych o średnicach: 0,05 mm (a), 1,02 mm (b), 0,3 mm (c,d) i 0,33 mm (e)



Rys. 4. Fotogramy zglądów metalograficznych wzdłużnych drutów molibdenowych o średnicach: 0,05 mm (a) i 0,2 mm (b)



Rys. 5. Fotogramy zglądów metalograficznych poprzecznych drutów wolframowych średnicy 0,2 mm



Rys. 6. Fotogramy zgładów metalograficznych poprzecznych prętów wolframowych średnicy 2,6 mm

Dla porównania zostanie przedstawionych kilka przykładów wad drutów i prętów wolframowych. Na fot. 5 pokazano fotogramy zgładów metalograficznych poprzecznych drutów o średnicach 0,2 mm, 0,28 mm i 0,4 mm. Pokazane na fot. 5 nieciągłości drutów wolframowych są bardzo dla nich typowe [19]. Są to pęknięcia (fot. 5a, c), rozwarstwienia (fot. 5b, d) oraz pustki (fot. 5e).

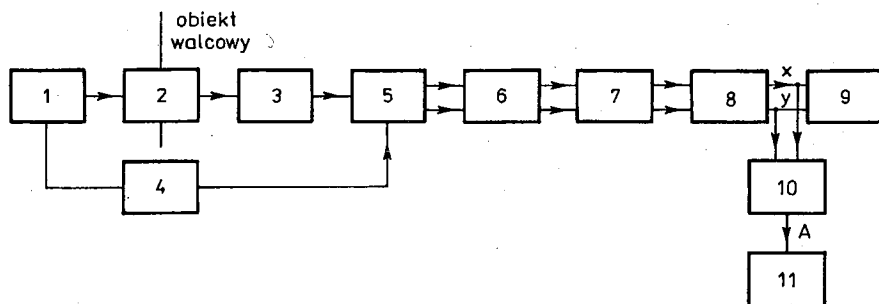
Na fot. 6a, b przedstawiono przykładowe fotogramy zgładów metalograficznych poprzecznych prętów wolframowych (średnicy 2,6 mm) odznaczających się odpowiednio pustkami oraz rozwarstwieniem, pęknięciem i ubytkiem powierzchniowym. Wystąpienie jakiegokolwiek z pokazanych tu wad powinno eliminować pręty, czy też druty z dalszej przeróbki oraz zastosowań.

Głównymi przyczynami występowania wad drutów — w postaci nieciągłości makrostruktury — są wady sztab i prętów młotkowanych, a na przykład nieciągłości parametrów geometrycznych — zły stan ciągaideł.

3. DEFEKTOSKOPY WIOPRĄDOWE II GENERACJI

Na rys. 7 przedstawiono schemat blokowy defektoskopu II-ej generacji [18, 25, 26, 28, 41]. Jako przykładowy obiekt badany podano tu obiekt walcowy, który może być kontrolowany przy użyciu przetwornika wioprądowego przelotowego. Nie mniej układ defektoskopu jest ogólny – mogą być w nim także analizowane sygnały przetworników stykowych, stosowanych do kontroli obiektów o różnych kształtach.

Przedstawiony na rys. 7 układ defektoskopu II generacji różni się od defektoskopu konwencjonalnego tym, że w układzie prostowania fazoczułego uzyskiwana jest informacja odnośnie dwu składowych napięcia wyjściowego przetwornika:



Rys. 7. Schemat blokowy defektoskopu wioprądowego II generacji 1 – generator, 2 – przetwornik wioprądowy, 3 – układy kompensacji składowych napięcia przetwornika, 4 – przesuwnik fazy, 5 – układ prostowania fazoczułego, 6 – wzmacniacz, 7 – filtry, 8 – układ nastawiania fazy (obrotu) trajektorii, 9 – monitor wad, 10 – układ obliczania amplitudy sygnału, 11 – układ selekcji sygnałów

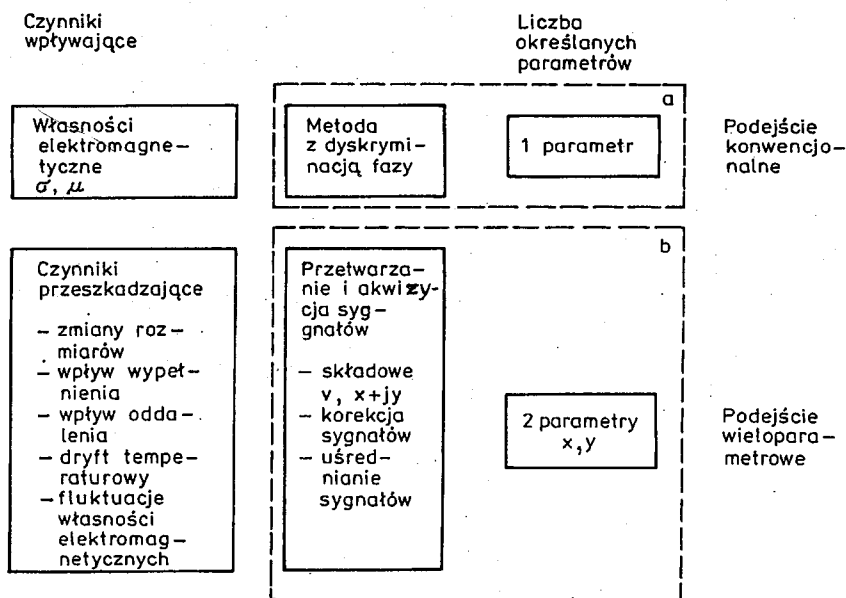
składowej będącej w fazie z prądem zasilającym przetwornik („kierunek X”) oraz – składowej przesuniętej w fazie o kąt $\pi/2$ („kierunek Y”) podczas, gdy w defektoskopach konwencjonalnych uzyskiwana jest informacja o jednej składowej.

Dane co do składowych powyższych sygnałów mogą być uzyskane poprzez rozkład tych sygnałów w szereg Fouriera i wyodrębnienie pierwszych wyrazów rozwinięcia (a_1 , b_1) – zależność (3.1)

$$u(t) = 0,5 a_0 + a_1 \cos \omega t + a_2 \cos 2\omega t + a_3 \cos 3\omega t + \dots + a_n \cos n\omega t + \\ + b_1 \sin \omega t + b_2 \sin 2\omega t + b_3 \sin 3\omega t + \dots + b_n \sin n\omega t$$

przy czym:

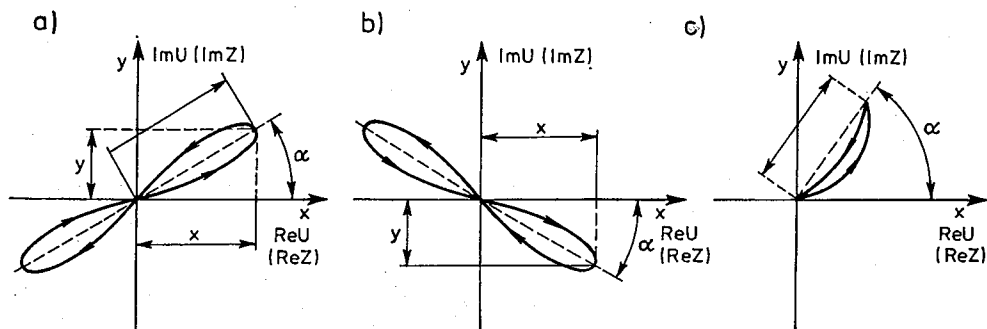
$$a_1 = (2/T) \int_0^T u(t) \cos \omega t \, dt \\ b_1 = (2/T) \int_0^T u(t) \sin \omega t \, dt, \quad (3.1)$$



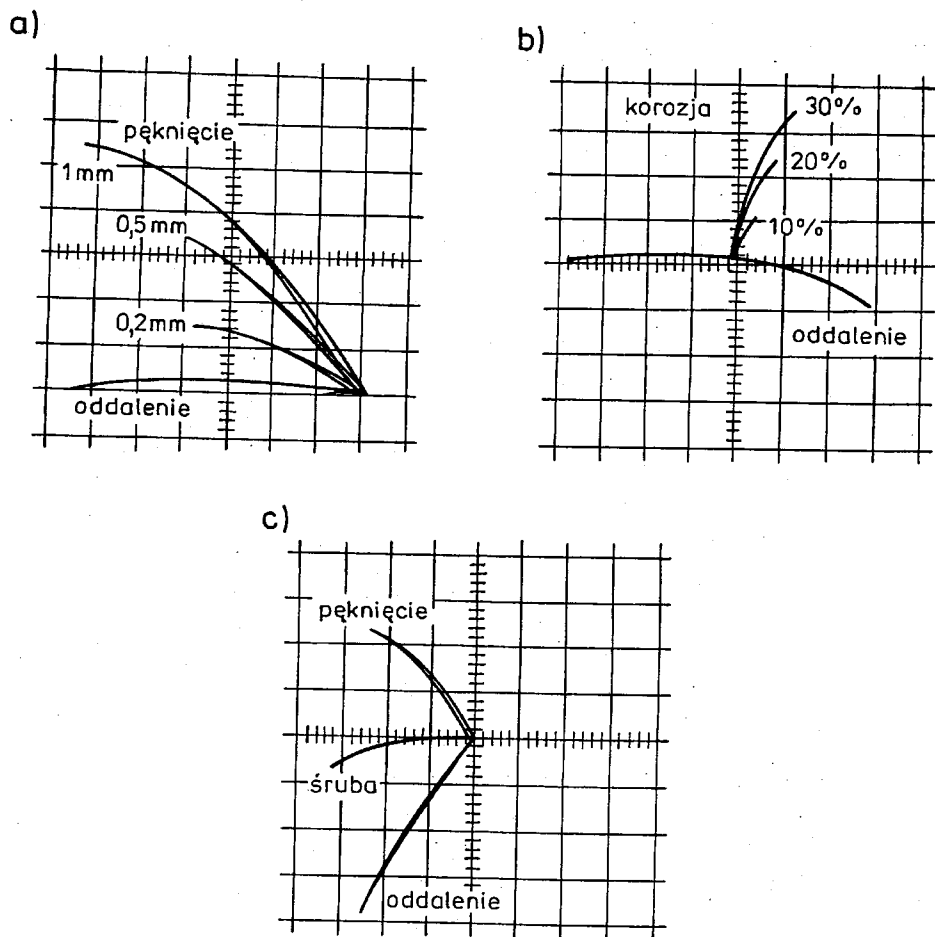
Rys. 8. Ilustracja do opisu sposobów podejścia do sygnałów przetworników wiroprowodowych w defektoskopach konwencjonalnych i defektoskopach II generacji

Dane odnośnie tych składowych, przy zastosowaniu odpowiedniego oprogramowania, mogą być wykorzystane do rekonstrukcji, przedstawianych na płaszczyźnie zmiennej zespolonej, trajektorii zmian sygnałów przetworników wiroprowodowych [18, 25, 42].

Przykładowe trajektorie zmian sygnałów przetwornika wiroprowodowego przelotowego zawierającego obiekt walcowy pokazano na rys. 9, a dla przetwornika stykowego zastosowanego do kontroli podzespołów samolotów – na rys. 10.



Rys. 9. Trajektorie zmian sygnałów przetwornika wiroprowodowego przelotowego; a, b – różnicowego, c – bezwzględnego



Rys. 10. Trajektorie zmian sygnałów przetwornika stykowego bezwzględnie zastosowanego do kontroli warstw aluminium

Na rys. 8 przedstawiono schemat funkcjonalny obrazujący różnice pomiędzy podejściem jednoparametrowym, konwencjonalnym (a), a podejściem wieloparametrowym (b), które może być zastosowane w defektoskopie II generacji.

W układzie przedstawionym na rys. 7 obliczana jest także amplituda sygnału A. W zależności od stopnia rozbudowy układu defektoskopu pracą urządzeń wykonawczych (dołączonych do defektoskopu) mogą sterować sygnały odpowiadające składowym x i y oraz amplitudzie A sygnału. W defektoskopach konwencjonalnych uzyskiwana była jedynie informacja odnośnie amplitudy sygnału wyjściowego przetwornika lub w najlepszym przypadku – informacja co do amplitudy sygnału, o określonej fazie (rys. 10) – w celu zmniejszenia wpływu czynników przeszkadzających [18, 19, 42]. W tym ostatnim przypadku umożliwiałoby to uzyskanie

informacji na przykład o pęknięciu obiektów — przy eliminacji przeszkadzających sygnałów wywołanych przez zmiany geometrycznych parametrów obiektów. Przy prezentacji trajektorii zmian sygnałów przetworników wiroprowodowych możliwe jest uzyskanie informacji odnośnie wszystkich rodzajów wad występujących w obiektach.

4. TRAJEKTORIE ZMIAN SYGNAŁÓW PRZETWORNIKÓW WIROPROWODOWYCH

Na rys. 9 i 10 podano przykładowe trajektorie zmian sygnałów przetworników wiroprowodowych przelotowych i stykowych [18, 25, 26, 28].

Trajektorie zmian sygnałów przetwornika różnicowego (rys. 9a, b) mają dwa ramiona, które odpowiadają przejściu wady przez dwa przeciwsobnie połączone uzwojenia przetwornika. Sygnał odpowiadający środkowi wady jest równy zeru (środek układu współrzędnych). Sygnał wywołany niesymetrią przetwornika jest skompensowany. Odpowiednio trajektoria zmian sygnału przetwornika bezwzględnego ma jedno ramię (rys. 9c). Trajektorie o kształcie, jak na rys. 9a, mogą być otrzymane na przykład dla nieciągłości rur od strony ścianki zewnętrznej (pęknięcia, przetop częściowy itp.); trajektorie, jak na rys. 9b — dla nieciągłości położonych od strony ścianki wewnętrznej. Informację o położeniu wady w zależności od odległości od powierzchni obiektu zawiera faza sygnałów przetworników, co związane jest z rozkładem pola magnetycznego wewnątrz obiektu [zależność (4.1)] [18, 28]

$$B = B_0 \exp(-d\sqrt{\pi f \mu \gamma}) \sin(2\pi ft - d\sqrt{\pi f \mu \gamma}) \quad (4.2)$$

gdzie: B — indukcja magnetyczna w dowolnej odległości d od powierzchni obiektu.

B_0 — indukcja na powierzchni obiektu,

f — częstotliwość pracy przetwornika,

μ — przenikalność magnetyczna obiektu,

γ — przewodność elektryczna właściwa materiału obiektu.

Natomiast trajektorie o kształcie pokazanym na rys. 9 odpowiadają wadom w postaci zmian średnicy (grubości ścianki) rur. Trajektorie pokazane na rys. 9 uzyskano przy wykorzystaniu układu opisanego w [18, 25].

Na rys. 10 przedstawiono trajektorie zmian sygnałów przetwornika wiroprowodowego stykowego bezwzględnego, zastosowanego do kontroli obiektów o różnych kształtach. Trajektorie z rys. 10a dotyczą przypadku badania bloku aluminiowego z pęknięciami o głębokościach 0,2 mm, 0,5 mm i 1,0 mm. Trajektorie z rys. 10b i rys. 10c dotyczą przypadku kontroli dwu połączonych ze sobą warstw aluminium. Trajektorie z rys. 10b zostały otrzymane dla odpowiednio 10%, 20% i 30% korozji wierzchniej warstwy pomiędzy tą warstwą a warstwą głębszą. Jedna z trajektorii z rys. 10c dotyczy przypadku, gdy w spodniej warstwie wystąpiło pęknięcie; drugą otrzymano przy umieszczeniu przetwornika na śrubie łączącej warstwy. Trajektorie

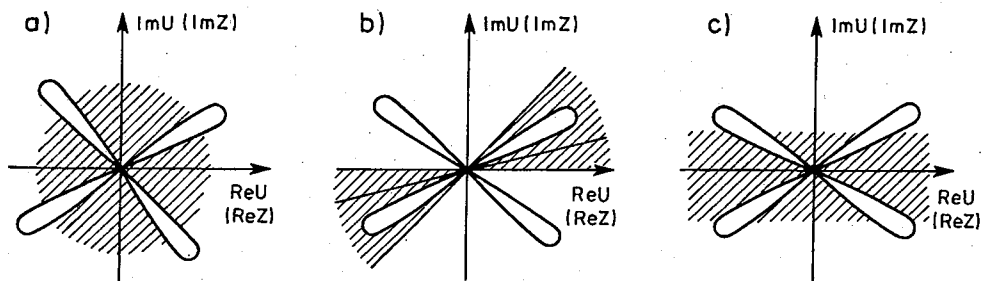
z rys. 10 otrzymano przy wykorzystaniu defektoskopu Elotest B2 niemieckiej firmy Rohmann GmbH.

Jak widać z rys. 9 i rys. 10 czynniki związane z różnymi obiektami, takie, jak obecność pęknięć, oddalenie przetwornika od powierzchni obiektu (na przykład obecność warstwy pokrycia) i zmniejszenie grubości materiału charakteryzują się odmiennymi pod względem kształtu i rozpiętości ramion trajektoriami. Tak więc przy zapewnieniu odpowiednich środków można dążyć do rozróżniania niektórych wad obiektów – na podstawie trajektorii. Znane są na podstawie dostępnej literatury metody oceny sygnałów przetworników wiropędowych, przedstawianych w postaci trajektorii, poprzez analizę współczynników rozwinięcia trajektorii w szereg Fouriera [47], analizę sieci neuronowych [17, 46] oraz bezpośrednią analizę kształtu trajektorii [18, 25].

Trajektorie zmian sygnałów przetworników wiropędowych uzyskiwane są dla stanu dynamicznego – podczas przesuwu odcinka z wadą poprzez przetwornik (przelotowy) lub podczas przesuwania przetwornika (stykowego) nad miejscem obiektu z wadą. Informacja o długości wady ukryta jest w kolejnych próbkach sygnału przetwornika, składających się na trajektorię. Rozpiętość ramion trajektorii odpowiada amplitudzie sygnału wywołanego przez wadę.

Obecnie stosowane są trzy sposoby podejścia do informacji zawartej w trajektoriach zmian sygnałów przetworników wiropędowych (rys. 11): analiza amplitudowa, analiza fazoczuła (sektorowa) i analiza składowej urojonej sygnałów.

Pierwszy ze sposobów podejścia do sygnałów przetworników polega na analizie ich amplitudy (rys. 11). Użytkownikowi przekazywana jest wówczas informacja, (niezależnie od fazy sygnałów) odnośnie wszystkich wad występujących w obiekcie – bez rozróżnienia ich rodzajów. Możliwe jest w tym przypadku popełnienie pomyłek odnośnie jakości obiektu poprzez klasyfikację na podstawie amplitudy sygnałów miejsc niewadliwych jako wadliwych. Może mieć to miejsce na przykład podczas wykrywania pęknięć – przy braku pokrycia galwanicznego na części obiektu lub przy zmianie sprzężenia przetwornika z obiektem – na przykład przy przekoszeniu przetwornika względem ustalonego w stosunku do obiektu położenia lub w przypadku niecentrycznego lub nieosiowego prowadzenia obiektu poprzez przetwornik.



Rys. 11. Ilustracja do opisu sposobów podejścia do sygnałów przetworników wiropędowych a – analiza amplitudowa, b – analiza fazoczuła (sektorowa), c – analiza składowej urojonej sygnałów

W celu wyeliminowania sygnałów otrzymywanych przy występowaniu „pseudo-defektów” zaproponowana została analiza korekcyjna sygnałów przetworników, polegająca na tym, że analizowane są sygnały dwu jednakowych przetworników oddalonych od siebie o określoną stałą odległość [42]. Dla poprawy stosunku sygnałów do zakłóceń, wywołanych na przykład przez wady powierzchni drutów, może być zastosowana funkcja autokorelacji [38, 41].

Drugi ze sposobów podejścia do sygnałów przetworników wiroprowadowych polega na analizie amplitudy sygnałów o określonej fazie (rys. 11), dla których trajektorie znajdują się w określonym sektorze płaszczyzny zmiennej zespolonej. Przyjmowane jest przy tym założenie, że sygnały wywołane przez czynniki zakłócające mają inne położenia fazowe, co jednak nie jest zawsze spełnione [18, 25].

Trzeci sposób podejścia do sygnałów przetworników wiroprowadowych polega na analizie składowej urojonej tych sygnałów (rys. 11c). Przyjmowane jest przy tym założenie, że trajektorie mogą być elektronicznie „obracane” na płaszczyźnie zmiennej zespolonej, tak, aby trajektorie otrzymywane dla czynników zakłócających przyjmowały położenie poziome. Ten sposób podejścia jest przydatny na przykład przy kontroli podzespołów samolotów (relacje fazowe defektoskopu ustala się w ten sposób, aby położenie poziome przyjmowały sygnały wywołane oddaleniem przetwornika od obiektu). W tym przypadku na podstawie analizy składowej urojonej sygnałów przetwornika jako wada może zostać na przykład sklasyfikowana krawędź obiektu.

Niżej przedstawiony zostanie opis systemu umożliwiającego otrzymywanie i analizę sygnałów przetworników wiroprowadowych, co jest bardzo istotne z punktu widzenia potrzeb kontroli potokowej, — na podstawie zależności sygnałów przetworników wiroprowadowych od czasu (długości próbek). Odpowiada to, przy wzięciu pod uwagę wyżej opisanych sposobów podejścia do sygnałów przetworników, analizie amplitudowej. Do układów wykonawczych systemu podawana jest informacja o wszystkich wadach obiektu, zarówno o nieciągłościach o charakterze pęknięć, jak i o charakterze zmian geometrii drutów. System przeznaczony jest do kontroli drutów wykonanych z materiałów nieferromagnetycznych.

5. UKŁAD DO OTRZYMYWANIA I ANALIZY SYGNAŁÓW PRZETWORNIKÓW WIROPROWADOWYCH

5.1. ZAŁOŻENIA DO UKŁADU

Najbardziej istotnymi założeniami do układu są:

1. Główne urządzenia i podzespoły układu: defektoskop wiroprowadowy, komputer wraz z urządzeniami peryferyjnymi. Przesuw obiektów — w linii technologicznej wytwarzania.
2. Generator wewnętrzny układu — o częstotliwości dobranej do obiektu.
3. Jednoczęstotliwościowe zasilanie przetworników wiroprowadowych.

4. Rekonstrukcja przebiegu zmian napięcia przetworników wiroprowadowych w funkcji długości obiektów (analiza amplitudowa).

5. Statystyczna analiza napięć wyjściowych przetworników

5.2. METODYKA OCENY JAKOŚCI DRUTÓW

Na podstawie sygnałów przetworników — przedstawianych w zależności od długości obiektów [40] przy zaimplementowaniu odpowiedniej analizy statystycznej, odróżniane mogą być długie wady od wad incydentalnie pojawiających się. Na podstawie tak przedstawianych sygnałów mogą być określone przyczyny powstawania wad obiektów (w ciągu technologicznym). Do oceny jakości drutów może być zastosowana niżej opisana metodyka rozróżniania wad drutów — na podstawie, przedstawianych w zależności od długości, sygnałów wyjściowych przetworników wiroprowadowych.

Sygnały wyjściowe przetworników podlegają dyskryminacji amplitudowej. Na wyjściu odpowiednich dyskryminatorów uzyskiwana jest informacja odośnie występowania następujących zdarzeń elementarnych:

- Istnienia sygnałów o poziomie poniżej wybranego określonego poziomu progowego. Niech zdarzenia te będą określone jako C_i .
- Istnienia sygnałów o poziomie zawierającym się pomiędzy dwoma poziomami progowymi A i B. Zdarzenia te określimy jako A_i .
- Istnienia sygnałów o poziomie powyżej wybranego progu B. Zdarzenia te określimy jako B_i .

Jeśli w badanym obiekcie występują incydentalnie niewielkie wady powierzchniowe, wówczas występują jedynie zdarzenia C_i i obiekt może być zakwalifikowany jako „dobry” — na całej swojej długości. W przypadku, gdy w badanym obiekcie występują na przykład niezbyt głębokie pęknięcia, to pojawiają się zdarzenia A_i oraz ewentualnie C_i — przy zbiegu z poprzednim przypadkiem, lub — odpowiednio występują zdarzenia A_i i B_i . W przypadku istnienia w obiekcie długich wad możemy mieć do czynienia z sytuacją, gdy wystąpi duże natężenie bardzo małych sygnałów i duże natężenie dużych i bardzo dużych sygnałów, tj. mogą wystąpić zdarzenia

$$C \cap K.$$

przy czym K jest pewną wartością progową ujmującą ilościowe wystąpienie razem zdarzeń A_i oraz C_i ($A_i \cap C_i$).

Wzrost intensywności występowania bardzo małych sygnałów (zdarzenia C_i) bez przekroczenia przez duże sygnały określonego progu K , może z dużym prawdopodobieństwem być przypisany złemu stanowi powierzchni obiektu, co ma miejsce w przypadku wystąpienia niewielkich wklęsłości i chropowatości. W takim przypadku wystąpią zdarzenia

$$C_i \cap K,$$

Z kolei pojedyncze defekty obiektu (pęknięcia, wtrącenia itp.) mogą być diagnozowane, jeśli wystąpią poszczególne zdarzenia A_i oraz B_i — bez jednoczesnego

wzrostu intensywności występowania najmniejszych sygnałów C lub coraz to większych sygnałów. Wówczas wystąpią zdarzenia

$$(A_i \cup B_i) \cap C_i \cap K,$$

W powyżej opisany sposób mogą być uzyskane korelacje pomiędzy poszczególnymi typami wad (defekty krótkie lub długie) — a odpowiadającymi im sygnałami przetworników wiropływowych. Jeśli suma sygnałów, które w ramach eksperymentu zakwalifikowane zostaną jako odpowiadające wadom pojedynczym przekroczy określoną wartość dla całej partii obiektu, to wówczas

$$\sum_{\substack{\text{partia} \\ \text{obektu}}} [(A_i \cup B_i) \cap C_i \cap K] \geq S_{\max}.$$

W praktyce oznacza to nieprawidłowo prowadzony proces technologiczny wytwarzania obiektu.

Natomiast w przypadku na przykład zużycia się niektórych elementów ciągu technologicznego może występować chropowatość powierzchni obiektu. W tym przypadku będziemy mieć do czynienia z dużą intensywnością najmniejszych sygnałów (zdarzenia C_i) przy braku występowania zdarzeń odpowiadających długim wadom i licznym wadom pojedynczym. Wówczas zostaną zarejestrowane zdarzenia

$$C \cap K$$

Może tu mieć miejsce przypadek okresowego występowania wad obiektu, które mogą czynić go nieprzydatnym do zastosowań.

5.3. DEFEKTOSKOP WIROPŁYWOWY DO KONTROLI DRUTÓW WYKONANYCH Z MATERIAŁÓW NIEFERROMAGNETYCZNYCH

Częstotliwość pracy defektoskopu

W tablicy 5.1 podano wartości częstotliwości granicznych, przyjętych częstotliwości pracy defektoskopu oraz stosunku częstotliwości pracy defektoskopu do częstotliwości granicznych. Wartości częstotliwości granicznych obliczane są według zależności (5.1) [15, 28]

$$f_g = 5066 \cdot 10^8 / \sigma \mu_r D^2 \quad (5.1)$$

gdzie: σ — przewodność elektryczna właściwa materiału obiektu, w S/m; dla wolframu i molibdenu $\sigma = 18,2$ MS/m,

μ_r — przenikalność magnetyczna względna obiektu (dla wolframu $\mu_r = 1$)

D — średnica obiektu, w mm.

Częstotliwości pracy przetworników odnosi się do wartości częstotliwości granicznych, przyjmując, że dla uzyskania dużej czułości wykrywania nieciągłości stosunek f/f_g powinien wynosić od około dwu do około kilkunastu [15, 19, 28].

Tablica 5.1

Parametry kontroli wiroprowdowej drutów wolframowych
średnic od 0,2 mm do 1,0 mm

D [mm]	f_g [kHz]	f [MHz]	f/f_g
0,2	696	2	2,9
0,3	309	2	6,5
0,35	227	2	8,8
0,4	174	1	5,7
0,5	111	1	9,0
0,6	77	1	12,9
0,7	57	1	17,6
0,8	44	0,5	11,5
0,9	34	0,5	14,6
1,0	28	0,5	18,0

Maksimum czułości wykrywania wad przy amplitudowej detekcji sygnałów przetworników występuje dla $f/f_g = 6,25$.

Defektoskop przeznaczony jest do kontroli drutów o średnicach od 0,2 do 1,0 mm. Przy wzięciu pod uwagę powyższego przyjęto trzy częstotliwości pracy defektoskopu wynoszące 2 MHz, 1 MHz i 0,5 MHz – do kontroli drutów o średnicach w zakresie odpowiednio 0,2 mm – 0,35 mm, 0,35 mm – 0,70 mm i 0,70 mm – 1,00 mm (tabl. 5.1).

Tablica 5.2

Indukcja magnetyczna wewnątrz drutów wolframowych w stosunku do indukcji na powierzchni drutów dla różnych częstotliwości pracy przetwornika

d [mm]	f [MHz]			B/B ₀		
	0,1	0,25	0,5	1,0	2,0	5,0
0,01	0,97	0,96	0,94	0,91	0,89	0,83
0,05	0,88	0,81	0,74	0,66	0,55	0,39
0,117	0,73	0,61	0,50	0,37	0,25	0,11
0,2	0,59	0,43	0,30	0,18	0,09	0,023
0,3	0,45	0,28	0,17	0,08	0,03	0,003
0,37	0,37	0,21	0,11	0,04	0,01	0,001

Dane zamieszczone w tablicach 5.2 i 5.3 ilustrują zagadnienie penetracji prądów wirowych we wnętrzu badanych drutów. W tablicy 5.2 podano wartości indukcji magnetycznej wewnątrz drutów w stosunku do wartości na powierzchni, obliczone według zależności (4.1) dla sześciu częstotliwości pracy przetwornika, w tym – dla przyjętych częstotliwości. Im większa jest częstotliwość pracy przetwornika, tym szybciej zanika wewnątrz drutów pole magnetyczne.

W badaniach wiroprowdowych przydatne jest obliczenie głębokości wnikania pola magnetycznego (prądów wirowych). Przyjęta została definicja opisana zależnością (5.2), zgodnie, z którą za głębokość wnikania przyjmuje się tę głębokość penetracji,

przy której indukcja magnetyczna maleje e -krotnie w stosunku do swojej wartości na powierzchni obiektu [15, 18, 19, 28].

$$\delta = 5 \cdot 10^5 / \sqrt{f \mu_r \gamma} \quad (5.2)$$

gdzie: f – częstotliwość pracy przetwornika, Hz.

Tablica 5.3

Głębokość wnikania pola magnetycznego
i prądów wirowych dla różnych częstotliwości
pracy przetwornika

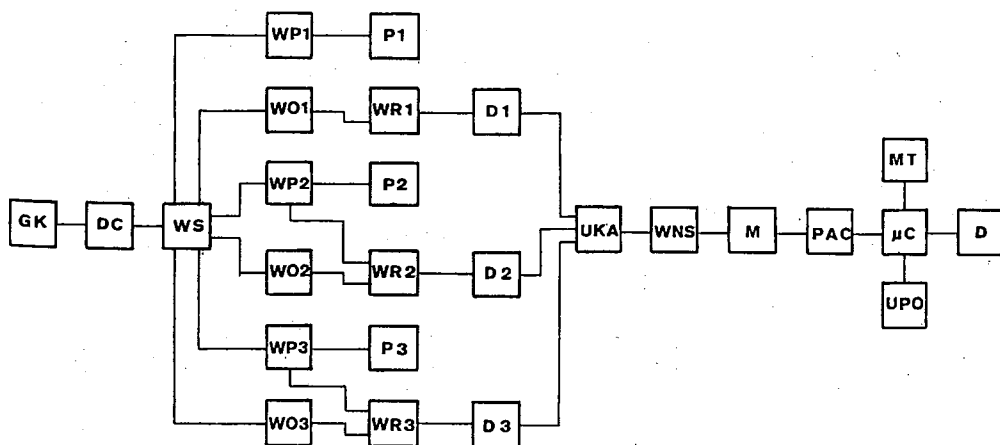
f [MHz]	δ [mm]
0,1	0,3729
0,25	0,2358
0,5	0,1667
1,0	0,1170
2,0	0,0830
5,0	0,0530

W tablicy 5.3 podano głębokość wnikania do obiektów, wykonanych z wolframu i molibdenu, dla sześciu częstotliwości pracy przetwornika, jak w tablicy 5.1. Jak wynika z zamieszczonych danych przy przyjętych częstotliwościach pracy przetworników wykrywane będą przede wszystkim wady powierzchniowe i podpowierzchniowe drutów.

5.4 UKŁAD DEFECTOSKOPU

Na rys. 12 przedstawiono schemat blokowy defektoskopu wiroprowadowego przeznaczonego do kontroli drutów wykonanych z materiałów nieferromagnetycznych. Przetworzone w układzie defektoskopu sygnały zawierające informację o stanie jakościowym badanych drutów podawane są do komputera. Defektoskop został opracowany w Centrum Uczelniano-Przemysłowym Metrologii i Systemów Pomiarowych Politechniki Warszawskiej, autorzy defektoskopu A. Lewińska-Romicka, S. Piskorski, J. Papis.

Układ defektoskopu (rys. 12) zawiera następujące grupy układów elektronicznych: generator napięć wysokiej częstotliwości, dzielnik częstotliwości wzmacniacze selektywne, wzmacniacze pomiarowe, wzmacniacze odniesienia, wzmacniacze różnicowe, detektory, układ komutacji analogowej, oraz wzmacniacz napięcia stałego. Układy elektroniczne defektoskopu mogą pracować z trzema wiroprowadowymi przetwornikami przelotowymi P1, P2 i P3, dzięki zastosowaniu których możliwa jest kontrola drutów o trzech zakresach średnic. Generator GK umożliwia otrzymanie napięć wysokiej częstotliwości. Generator ten jest stabilizowany przetwornikiem kwarcowym o częstotliwości 5 MHz. W związku z przyjętym założeniem, że częstotliwości napięć, które zasilac mają przetworniki wiroprowadowe P1, P2 i P3 w związku

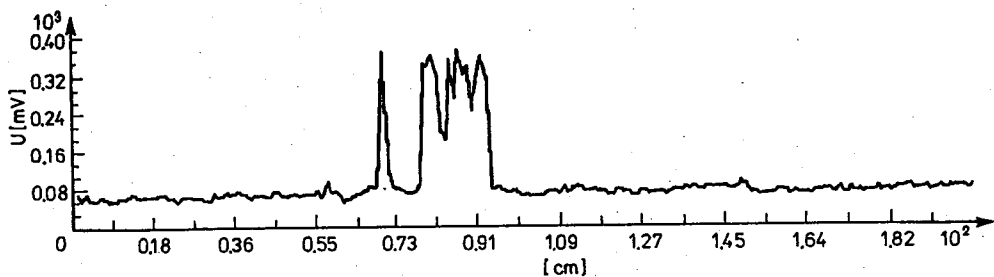


Rys. 12. Schemat blokowy defektoskopu do kontroli drutów wolframowych i molibdenowych GK – generator, DC – dzielnik częstotliwości, WS – wzmacniacz selektywny, WP, WO1÷WP3 – wzmacniacze pomiarowe, P1÷P3 – przetworniki wiropądowe odpowiednio do kontroli drutów o średnicach w zakresie 0,2 mm–0,35 mm, 0,35 mm–0,70 mm i 0,70 mm–1,00 mm, WR1÷WR3 – wzmacniacze różnicowe, D1÷D3 – detektory wartości maksymalnej, UKA – układ komutacji analogowej, WNS – wzmacniacz napięcia stałego, M – multiplexer, PAC – przetwornik analogowo-cyfrowy, μC – komputer, MT – monitor, UPO – układ przetwarzania sygnałów, D – drukarka

z kontrolą drutów o wymienionych wyżej średnicach i przewodności właściwej 18,2 MS/m, wynosić mają 2 MHz, 1 MHz i 0,5 MHz napięcia pobierane z wyjścia generator GK podawane są do dzielnika częstotliwości DC. Następnie sygnały te wzmacniane są w układzie wzmacniacza selektywnego WS. Sygnały z wyjścia wzmacniacza selektywnego WS podawane są do trzech analogicznych do siebie zespołów układów zawierających odpowiednio wzmacniacz pomiarowy WP1 i przetwornik wiropądowy P1 oraz wzmacniacz odniesienia WO1 i wzmacniacz różnicowy WR1 a także detektor D1 (lub WP2, P2, WO2, WR2 i D2 albo też WP3, P3, WO3, WR3 i D3). Sygnały napięcia stałego pobierane z wyżej wymienionych torów defektoskopu podlegają komutacji w układzie komutacji analogowej UKA, a dalej podawane są do wzmacniacza napięcia stałego WNS. Sygnały te dla celów analizy komputerowej podawane są do komputera μC poprzez kartę zawierającą multiplexer M i przetwornik analogowo-cyfrowy P A/C. Sygnały przetworników wiropądowych, przetworzone w układach defektoskopu, prezentowane są na ekranie monitora MT oraz zapisywane są przy użyciu drukarki D.

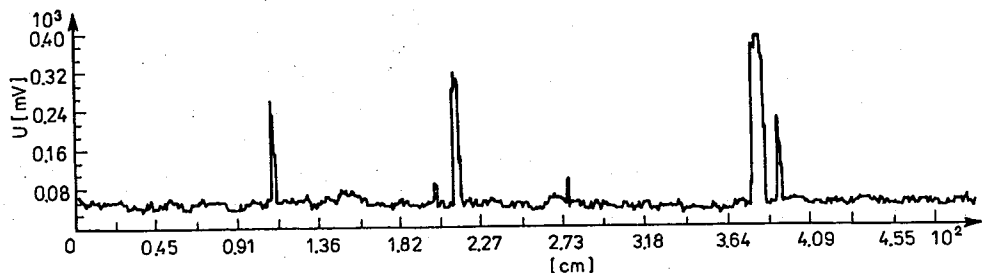
Akwizycja sygnałów zawierających informację o wadach drutów pozwala na zaimplementowanie dalszego przetwarzania tych sygnałów w układzie UPO, zgodnie z metodyką przedstawioną w p. 5.2.

Na rys. 13 i rys. 14 przedstawiono przykładowe wyniki badania (w zależności od długości) drutu wolframowego o średnicy 0,33 mm. Wynik ten otrzymano przy wykorzystaniu opisanego defektoskopu. Zarejestrowano wystąpienie w drutach wad o długości od około kilkudziesięciu milimetrów do około dwustu milimetrów.



Rys. 13. Sygnał wyjściowy defektoskopu wiroprowodowego DWK-101 dla odcinka drutu wolframowego średnicy 0,33 mm

analiza przedstawianych w zależności od długości drutu sygnałów wyjściowych defektoskopu pozwala na przeprowadzanie analizy statystycznej, na przykład z użyciem pakietu Statgraphics, który umożliwi opracowywanie histogramów odpowiednio próbkowanych danych [38]. Histogramy takie stanowią charakterystyki jakości obiektu.



Rys. 14. Sygnał wyjściowy defektoskopu wiroprowodowego DWK-101 dla odcinka drutu wolframowego średnicy 0,33 mm

Obecnie w szeregu firm zagranicznych prowadzone są prace nad określeniem zespołów cech charakteryzujących sygnały przetworników wiroprowodowych. W dalszej kolejności posłuży to do tworzenia baz wiedzy i budowy systemów ekspertowych [1 ÷ 5, 7, 8, 10 ÷ 14, 27, 30, 31, 43 ÷ 45].

PODSUMOWANIE

Przedstawiono problem, jakim jest nieniszcząca analiza stanu jakościowego wybranych obiektów, wykonanych z materiałów nieferromagnetycznych, to jest drutów wolframowych i molibdenowych. Do analizy tej jest szczególnie przydatna metoda wykorzystująca zjawisko prądów wirowych.

Konieczność wdrażania nieniszczących badań wymienionych drutów wynika z ich dużej wadliwości, co zilustrowano poprzez podanie fotogramów szeregu zglądów metalograficznych drutów o różnych średnicach.

Opisano metody uzyskiwania sygnałów przetworników wiroprowodowych i analizy tych sygnałów przy diagnostyce różnych obiektów.

Opisano współpracujący z komputerem defektoskop wiroprowodowy przeznaczony do badania drutów o średnicach od 0,2 mm do 1,0 mm.

Do oceny jakości drutów wykorzystano analizę (po odpowiednim przetworzeniu) sygnałów przetworników przelotowych, przez które prowadzone są druty — w linii lub poza linią wytwarzania. Zaproponowano metodykę statystycznej oceny jakości drutów. Może ona być również wykorzystana do diagnozy stanu linii wytwarzania.

BIBLIOGRAFIA

1. A. Corazza, E. Milana, F.A. Zanardi, F.M. Ziprani: *A New Smart Eddy-Current System for on-Line Flaws Detection. Proceedings of the 12th World Conference on Non-Destructive Testing*. Amsterdam, Elsevier Science Publishers, Amsterdam — Oxford — New York — Tokyo 1989, t.1, s.352÷354
2. M. Cousin: *An on-Line Eddy-Current System for the Quality Evaluation of Hot Rods*. Materials Evaluation, 1985, 43, nr 13, s.1649÷1654
3. G. Delasarte, R. Levy: *Mascotte. An Analytical Model for Eddy Current Signals*. Proceedings of the 13th World Conference on Non-Destructive Testing. Sao-Paulo, Elsevier Science Publishers B.V., Amsterdam — London — New York — Tokyo 1992, s. 264÷268
4. R.O. Duda, P.E. Hart: *Pattern Classification and Scene Analysis*. Wiley and Sons, New York — London — Sydney — Toronto 1973
5. D.E. Dudgeon, R.M. Mersereau: *Multidimensional Digital Signal Processing*. Prentice Hall, Englewood Cliffs 1984
6. K. Dydyński, M. Dźwiarek, A. Lewińska-Romicka, S. Piskorski: *Defektoskop do badania elementów niemagnetycznych*. Patent nr 145 699, 1984
7. B. Georgel, R. Zorgati: *Extraction: A System for Automatic Eddy Current Diagnosis of Steam Generator Tubes in Nuclear Power Plants*. Proceedings of the 13th World Conference on Non-Destructive Testing. Sao-Paulo, Elsevier Science Publishers B.V., Amsterdam — London — New York — Tokyo 1992, s. 278÷282
8. A. Golke, M. Tietze: *Computergestützte Wirbelstromprüfung an Drahten, Stangen und Rohren*. Materialprüfung, 1989, 31, Nr 11—12, s. 363÷366
9. W.J. McGonnagle: *Nondestructive Testing*, McGraw-Hill Book Company. New York, Toronto, London 1961
10. A. Grabner: *An Integrated Condition Monitoring System Increases Productivity in Wire Rolling Lines*. Institut Dr Förster. Report No 299, 1991
11. A. Grabner: *A Project for Integrated Condition Monitoring Systems to Improve Product Quality and Plant Availability in Wire Production*. Institut Dr Förster. Report No 294, 1991
12. A. Grabner: *Integrierte Monitoring Systeme zur Sicherung von Produktqualität und Anlagenverfügbarkeit*. Institut Dr Förster. Instituts-Bericht Nr 292, 1991
13. A. Grabner: *Integriertes Condition Monitoring System zur Sicherung von Produktqualität und Anlagenverfügbarkeit für Stab- und Drahtwalzstrassen*. Institut Dr Förster. Instituts Bericht Nr 296, 1991
14. K. Grotz, B. Lutz: *Electromagnetic Multiparameter Determination of Material Characteristics*. Institut Dr Förster. Report No 279, 1991

15. H. Heptner, H. Stroppe: *Magnetyczne i indukcyjne badania metali*. Wyd. Śląsk, Katowice 1972
16. A. Kiliński: *Jakość*. WNT, Warszawa, 1979
17. H. Komatsu, Y. Matsumoto, K. Aoki, F. Nakayasu, M. Hashimoto, K. Miya: *Basic Study on ECT Data Evaluation Method with Neural Network*. Proceedings of the 13th World Conference on Non-Destructive Testing, Sao-Paulo, Elsevier Science Publishers B.V., Amsterdam — London — New York — Tokyo 1992, s. 322÷326
18. A. Lewińska-Romicka: *Badania nieniszczące rur metalowych metodą prądów wirowych*. PWN, Warszawa 1991
19. A. Lewińska-Romicka: *Badania nieniszczące drutów i prętów metalowych metodą prądów wirowych*. Centrum Uczelniano-Przemysłowe Metrologii i Systemów Pomiarowych Politechniki Warszawskiej, Warszawa 1985
20. A. Lewińska-Romicka: *Defektoskopia wiroporządowa rur nieferromagnetycznych spawanych*. „Kwartalnik Elektroniki i Telekomunikacji” 1990, 36, nr 1 s. 175÷206
21. A. Lewińska-Romicka: *Devices for Eddy-Currents Testing of Cracks in Steel Wires and Bars*. Proceedings of the 3rd European Conference on Nondestructive Testing, Firenze 1984, t. 4, s. 100÷109
22. A. Lewińska-Romicka: *Eddy Current Method in Bar and Wire Inspection*. Proceedings of the XI World Congress IMEKO, Houston 1988, t. 4, s. 415—422
23. A. Lewińska-Romicka: *Nieniszczące wykrywanie defektów drutów i prętów niemagnetycznych i magnetycznych metodą prądów wirowych*. Aparatura Naukowa i Dydaktyczna, 1984, 10, nr 3, s. 11÷16
24. A. Lewińska-Romicka: *Problemy nieniszczących badań rur niemagnetycznych metodą prądów wirowych*. Rozprawy Elektrotechniczne 1986, 32, nr 3, s. 801÷812
25. A. Lewińska-Romicka: *Trajektorie napięcia przetworników wiroporządowych samoporównawczych*. Kwartalnik Elektroniki i Telekomunikacji 1990, 36, nr 3, s. 741÷761
26. A. Lewińska-Romicka, S. Piskorski, J. Papis: *Defektoskop wiroporządowy o dwuwymiarowej prezentacji sygnałów przetworników do kontroli rur nieferromagnetycznych*. Kwartalnik Elektroniki i Telekomunikacji 1991, 37, nr 3—4 s. 785÷794
27. B. Lutz, D. Sy: *Integration of Mathematical Statistical Evaluation Methods in Automatic Quality Control Leads to New Application Areas of Magneto-Inductive Tests*. Proceedings of the 12th World Conference on Non-Destructive Testing. Amsterdam, Elsevier Science Publishers, Amsterdam — Oxford — New York — Tokyo 1989, t. 2, s. 1324—1327; również Institut Dr. Förster. Report No 142, 1989
28. *Nondestructive Testing Handbook. Vol. Four Electromagnetic Testing*. Published by the American Society for Nondestructive Testing. Wyd. II 1985
29. Z. Pawłowski: *Badania nieniszczące*. Poradnik. Wyd. Ośrodka Doskonalenia Kadr SIMP, Warszawa 1988
30. M.A. Rothe: *Reliable and Predictive Quality and Process Control of NFE Wire and Rod Using Latest Eddy-Current Technology*. Institut Dr. Förster. Report No. 287, 1991
31. A. Schlicht, A. Zhirabok: *The Integrated Expert System for Nondestructive Testing in Quality Control Systems*. Proceedings of the 13th World Conference on Non-Destructive Testing. Sao-Paulo, Elsevier Science Publishers B.V., Amsterdam — London — New York — Tokyo 1992, s. 201÷205
32. W. Solnica, K. Dydyński, A. Lewińska-Romicka: *Niektóre problemy nieniszczących badań drutów niemagnetycznych przy użyciu przetworników wykorzystujących zjawisko prądów wirowych*. Pomiary, Automatyka, Kontrola, 1975, XXI, nr 2, s. 74÷77
33. W. Solnica, A. Lewińska-Romicka: *System nieniszczących badań i sposoby rejestracji wyników kontroli drutów niemagnetycznych*. Rozprawy Elektroniczne, 1977, 23, nr 3, s. 535÷552
34. W. Solnica, A. Lewińska-Romicka, K. Dydyński, M. Owsiany: *Układ cyfrowej rejestracji wad do defektoskopów*. Patent Nr 103 992, 1980

35. W. Solnica, A. Lewińska-Romicka, K. Dydyński, M. Owsiany: *Wykrywanie wad w drutach wolframowych*. Elektronika, 1976, 17, nr 3, s. 114÷117
36. W. Solnica, A. Lewińska-Romicka, K. Dydyński, S. Piskorski, M. Owsiany: *Defektoskopy wiroprądowe do nieniszczącego wykrywania i rejestracji pęknięć drutów wolframowych i molibdenowych*. Mechanik 1981, LIV, nr 5÷6, s. 271÷273
37. W. Solnica, A. Lewińska-Romicka, M. Owsiany: *Układ do wykrywania wad drutów z materiałów przewodzących prąd elektryczny*. Patent nr 118 366, 1984
38. W. Solnica, A. Lewińska-Romicka, J. Papis, S. Piskorski: *Metody analizy sygnałów przetworników wiroprądowych*. Zesz. Nauk. Pol. Śląskiej. Zesz. 111, ss. 81—101
39. W. Solnica, A. Lewińska-Romicka, S. Piskorski: *Urządzenie do nieniszczących badań jakości i znakowania wad materiałowych w drutach niemagnetycznych*. Rozprawy Elektrotechniczne 1980, 26, nr 3, s. 609÷615
40. W. Solnica, A. Lewińska-Romicka, S. Piskorski, J. Papis: *Analiza sygnałów z defektoskopu wiroprądowego wspomagana komputerowo*. Metrologia i Systemy Pomiarowe, 1990, nr 5, s. 121÷135
41. W. Solnica, A. Lewińska-Romicka, S. Piskorski, J. Papis: *Raporty tematu VII.03 Problemy badań nieniszczących obiektów metalowych metodą prądów wirowych*, CPBP 02.20 „Wybrane zagadnienia podstawowych problemów współczesnej metrologii oraz technologii i konstrukcji systemów i urządzeń pomiarowo-kontrolnych”. Etapy I÷V (1986÷1990)
42. K. Stöwer: *The Non-Destructive Testing of Materials by Electro-Magnetic Techniques*. Institut Dr. Förster. Report No 132
43. A. Śluzek: *Komputerowa analiza sygnałów*. WPW, Warszawa 1991
44. R. Tadeusiewicz, M. Flasiński: *Rozpoznawanie obrazów*. PWN, Warszawa 1991
45. M. Tietze: *Valuable Eddy-Current Inspection Utilising Latest Computer Technology*. British Journal of NDT, 1991, 33, Nr 5, s. 217÷220; również Institut dr Förster. Report No 288, 1991
46. L. Udpa, S.S. Udpa: *Eddy Current Defect Characterization Using Neural Networks*. Materials Evaluation, 1990, 48, March, s. 342÷347
47. C.T. Zahn, R.Z. Roskies: *Fourier Descriptors for Plane Closed Curves*. IEEE Transactions on Computers, 1972, March, s. 269÷281

A. LEWIŃSKA-ROMICKA

QUALITY ASSURANCE PROBLEMS IN WIRE PRODUCTION. EDDY CURRENT NON-DESTRUCTIVE TESTING

Summary

In the article quality estimation methodology exemplified to the tungsten and molybdenum testing is described. New trends appeared in eddy current non-destructive testing which rely on a novel approach to preprocessed in defectoscopes networks. Eddy current transducers signals are presented.

SPIS TREŚCI

T. A d a m s k i: Analiza szumu addytywnego wprowadzonego przez niestałość chwilową momentu próbkowania w konwerterach częstotliwości próbkowania	447
J. W o j a s: Analiza fotoemisji elektronów	467
E. K o r n a t o w s k i: FFT jako algorytm postprocessingu zadań syntezy pola	477
M. Ż ó ł t o w s k i: O liniowej adaptacyjnej filtracji ze stałym krokiem adaptacji	483
M. Ż ó ł t o w s k i: Krótkoczasowa zbieżność algorytmów adaptacyjnej filtracji w terioobwodowo- wopodobnym ujęciu	503
Z. W r ó b e l: D-równoważność struktur układów realizujących funkcje wymierne	523
Z. W r ó b e l: Synteza struktur kanonicznych czwórników składających się z dwóch rodzajów elementów	541
Z. M u s i a ł o w s k i: MICAP — program komputerowy do analizy i projektowania liniowych scalonych układów mikrofalowych	561
T. S t a r e c k i: Nowa metoda pomiaru komór fotoakustycznych	584
A. L e ś n i c k i: Metoda uzmienniania stanów początkowych w analizie układów liniowych o okresowej odpowiedzi stanu ustalonego	593
A. L e w i Ń s k a—R o m i c k a: Problemy zapewnienia jakości w procesie wytwarzania drutów. Badania nieniszczące z wykorzystaniem prądów wirowych	603

CONTENTS

T. A d a m s k i: Analysis of additive noise introduced by sampling time jitter in sampling frequency converters	447
J. W o j a s: Analysis of photoemission of electrons	467
E. K o r n a t o w s k i: FFT as the algorithm of postprocessong of an assignment of the field synthesis	477
M. Ż ó ł t o w s k i: On linear adaptive filtering with constant step of adaptation	483
M. Ż ó ł t o w s k i: Circuits theory oriented approach to the short term convergence of adaptive digital filtering algorithms	503
Z. W r ó b e l: D-Equivalence structure circuits realizing rational function	523
Z. W r ó b e l: Topological synthesis of canonical structures of bigraphs two-port circuits consisting of two kinds of elements.	541
Z. M u s i a ł o w s k i: MICAP — the computer program for the analysis and design of linear microwave integrated circuits	561
T. S t a r e c k i: A new method of PAS cell measurements	584
A. L e ś n i c k i: The shooting method in analysis of linear circuits with periodic steady—state responses	593
A. L e w i Ń s k a—R o m i c k a: Quality assurance problems in wire production. Eddy current non-destructive testing	603

Computer Assisted Mechanics and Engineering Sciences (CAMES)

to be published quarterly starting in 1994

Editor: M. Kleiber, Warsaw

Associate Editor: H.A. Mang, Vienna

Scope of the Journal

Computer Assisted Mechanics and Engineering Sciences (CAMES) is a refereed international journal, published quarterly, providing an international forum and an authoritative source of information in the field of computational mechanics and related problems of applied science and engineering.

The specific objective of the journal is to support researchers and practitioners based in Central Europe by offering them a means facilitating (a) access to newest research results by leading experts in the field (b) publishing their own contributions and (c) dissemination of information relevant to the scope of the journal.

Papers published in the journal will fall largely into three main categories. The first will contain state-of-the-art reviews with the emphasis on providing the Central European readership with a guidance on important research directions as observed in the current world literature on computer assisted mechanics and engineering sciences.

The second category will contain contributions presenting new developments in the broadly understood field of computational mechanics including solid and structural mechanics, multi-body system dynamics, fluid dynamics, constitutive modeling, structural control and optimization, transport phenomena, heat transfer, etc. Variational formulations and numerical algorithms related to implementation of the finite and boundary element methods, finite difference method, hybrid numerical methods and other methodologies of computational mechanics will clearly be the core areas covered by the journal.

The third category will contain articles describing novel applications of computational techniques in engineering practice; areas of interest will include mechanical, aerospace, civil, naval, chemical and architectural engineering as well as software development.

The journal will also publish book reviews and informations on activities of the Central European Association of Computational Mechanics.

Subscription and sale of single issues is managed by the Editorial Office (for 1994).

Price of single issue: 25 USD (in Poland: 40 000 zł)

Address: Editorial Office, CAMES

Polish Academy of Sciences

Institute of Fundamental Technological Research

ul. Świątokrzyska 21, PL 00-049 Warsaw, Poland

Our Bankers: Bank Przemysłowo-Handlowy Kraków, XIV O/M Warszawa,
32-0007-1906.

Computer Assisted Mechanics and Engineering Sciences (CAMES)

kwartalnik — wydawany w języku angielskim od zeszytu pierwszego 1994 roku

Redaktor naczelny: Prof. Michał Kleiber

Adres Redakcji:

Instytut Podstawowych Problemów Techniki PAN
ul. Świętokrzyska 21, 00-049 Warszawa

Cena zeszytu pierwszego: 40 000 zł. —

Zamówienia prosimy składać pod podanym adresem. Cena zeszytów: drugiego, trzeciego i czwartego po 40 000 zł. Zamówienia należy składać w jednostkach kolportażowych „RUCH” S.A., właściwych dla miejsca zamieszkania lub siedziby prenumeratora w terminach podanych w warunkach prenumeraty lub w IPPT PAN.

Zakres czasopisma

Computer Assisted Mechanics and Engineering Sciences (CAMES) jest międzynarodowym czasopismem publikowanym kwartalnie, stanowiącym międzynarodowe forum i autorytatywne źródło informacji w dziedzinie metod komputerowych w mechanice i zbliżonych zagadnień w naukach stosowanych i inżynierskich.

Szczególnym celem czasopisma jest wspomaganie naukowców, studentów i inżynierów-praktyków w Europie Środkowej poprzez: (a) zapewnienie dostępu do najnowszych wyników badań przodujących specjalistów, (b) umożliwienie publikowania własnych prac oraz (c) upowszechnianie wiedzy w dziedzinach należących do zakresu czasopisma.

W czasopiśmie publikowane są trzy zasadnicze kategorie prac. Pierwsza obejmuje prace przeglądowe przybliżające czytelnikom środkowoeuropejskim ważne kierunki rozwoju obserwowane we współczesnej literaturze z dziedziny zastosowań metod komputerowych w mechanice i naukach inżynierskich.

Do drugiej kategorii należą prace prezentujące nowe osiągnięcia w szeroko pojmowanych metodach komputerowych w różnych działach mechaniki, takich jak mechanika ciała stałego, mechanika konstrukcji, dynamika układów wielu ciał, mechanika płynów, modelowanie konstytutywne, zagadnienia sterowania i optymalizacji konstrukcji, zjawiska transportu, przepływu ciepła itp. Centralne miejsce w tematyce czasopisma zajmują sformułowania wariacyjne i algorytmy numeryczne tworzone pod kątem zastosowań metody elementów skończonych i brzegowych, metody różnic skończonych, metod hybrydowych oraz innych metod komputerowych mechaniki.

Do trzeciej kategorii zaliczyć można artykuły opisujące nowatorskie zastosowania technik komputerowych w praktyce inżynierskiej, szczególnie w zagadnieniach mechaniki konstrukcji i budowy maszyn, oraz tworzenia oprogramowania komputerowego.

W czasopiśmie publikowane są ponadto recenzje książek i informacje na temat działalności Środkowoeuropejskiego Stowarzyszenia Metod Komputerowych Mechaniki.

Subscription for external subscribers:

The promotional subscription price in 1994 is \$ 80 including postage for institutions. Subscriptions should be sent too the publisher, Polish Scientific Publishers PWN Ltd, Journal Division, Miodowa 10, 00-251 Warsaw, POLAND, fax (48) (22) 26 09 50, (48) (22) 26 71 63, with a copy by fast to Electronics and Telecommunications Quarterly 00-665 Warsaw, ul. Nowowiejska 15/19, p. 470. Subscription is occupted after showing a checque or transfer documents. Our bank account is as follows:

Bank account: PBK VIII O/Warszawa nr 370028-1052-139-112. At subscriber's request this journal will be air mailed at additional postage to 50% of gross price to European countries and 65% overseas.

Kwartalnik Elektroniki i Telekomunikacji

WARUNKI PRENUMERATY

Wpłaty na prenumeratę przyjmowane są na okresy kwartalne:

na teren kraju

- jednostki kolportażowe „RUCH” S.A. i urzędy pocztowe oddawcze, właściwe dla miejsca zamieszkania lub siedziby prenumeratora, oraz doręczyciele w miejscowościach, gdzie dostęp do urzędu jest utrudniony,
- od osób lub instytucji, zamieszkających lub mieszkających się w miejscowościach, w których nie ma jednostek kolportażowych „RUCH”, wpłaty należy wносить do „RUCHU” S.A. Oddział Warszawa, 00-958 Warszawa, ul. Towarowa 28. Konto: PBK XIII Oddział Warszawa nr 370044-1195-139-11. „RUCH” S.A. zapewni dostawę pod wskazanym adresem pocztą zwykłą w ramach opłaconej prenumeraty.

na zagranicę

- „RUCH” S.A. Oddział Warszawa, 00-958 Warszawa, ul. Towarowa 28. Konto: PBK XIII Oddział Warszawa nr 370044-1195-139-11. Dostawa odbywa się pocztą zwykłą w ramach opłaconej prenumeraty, z wyjątkiem zlecenia dostawy pocztą lotniczą, której koszt w pełni pokrywa zleceniodawca.

Prenumerata ze zleceniem dostawy za granicę jest o 100 % wyższa od krajowej.

Dostawa zamówionej prasy następuje:

- przez jednostki kolportażowe „RUCH” S.A. — w sposób uzgodniony z zamawiającym,
- prenumerata pocztowa — pod wskazanym adresem, w ramach opłaconej prenumeraty.

Terminy przyjmowania przez „RUCH” S.A. wpłat na prenumeratę krajową i zagraniczną oraz przez Pocztcę Polską (tylko prenumerata krajowa):

<u>„RUCH” S.A.</u>		<u>Poczta Polska</u>	
— do 20 XI	na I kw. roku następnego	— do 25 XI	na I kw. roku następnego
— do 20 II	na II kw.	— do 25 II	na II kw.
— do 20 V	na III kw.	— do 25 V	na III kw.
— do 20 VIII	na IV kw.	— do 25 VIII	na IV kw.

Bieżące numery można nabyć w Księgarni Wydawnictwa Naukowego PWN, ul. Miodowa 10, 00-251 Warszawa. Również można je nabyć, a także zamówić (przesyłka za zaliczeniem pocztowym) we Wzorcowni Ośrodka Rozpowszechniania Wydawnictw Naukowych PAN, Pałac Kultury i Nauki, 00-901 Warszawa.

UWAGA!

Poczynając od 1993 r. prenumeratę naszego czasopisma prowadzi również Zakład Kolportażu Wydawnictwa SIGMA-NOT. Zamówienia na prenumeratę należy składać w terminach kwartalnych. Wpłata na II półr. wynosi 60 000 zł

Warunkiem przyjęcia prenumeraty jest wpłata na konto wartości zamówionych egzemplarzy. Wpłat na prenumeratę należy dokonywać na blankietach do wpłat na rachunki bankowe na konto:

Zakład Kolportażu Wydawnictwa SIGMA-NOT Sp. z o.o.
00-950 Warszawa nr PBK III O/Warszawa 370015-1573-139-11

Zakład Kolportażu oferuje również możliwość zamawiania prenumeraty ze zleceniem wysyłki za granicę. Jej cena jest dwukrotnie wyższa i ceny prenumeraty normalnej, a zlecający powinien podać dokładny adres odbiorcy za granicą. Dodatkowych informacji udziela:

Zakład Kolportażu Wydawnictwa SIGMA-NOT Sp. z o.o.
00-716 Warszawa, skr. poczt. 1004, ul. Bartycka 20.
Telefony: 40-30-86, 40-35-89, 40-00-21 w. 249, 295, 299