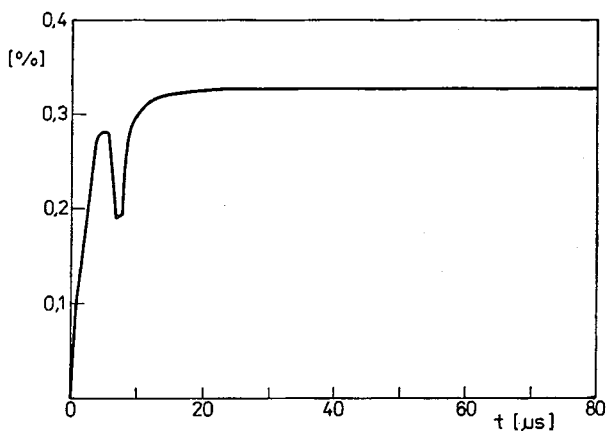


W krzemowych przyrządach półprzewodnikowych gęstości prądu są rzędu $100 - 2500 \text{ A/cm}^2$ [17–19], odpowiada im współczynnik α rzędu $0,1\%$, a co za tym idzie nieizotermiczne efekty termoelektryczne mogą być być w bilansie procesów rozpraszania i transportu energii w tych przyrządach potraktowane jako pomijalnie małe. Inaczej wygląda sytuacja w pozostałych materiałach półprzewodnikowych, dla uzyskania $\alpha = 10\%$ wymagana jest np. w GaAs gęstość prądu w granicach $80 - 13 \text{ kA/cm}^2$ a dla InSb w granicach $20 - 4.7 \text{ kA/cm}^2$. Są to gęstości prądu



Rys. 8. Względne obniżenie oszacowania maksymalnej temperatury w analizowanej strukturze lasera wywołane uwzględnieniem w modelu cieplnym nieizotermicznych efektów termoelektrycznych

występujące realnie w niektórych przyrządach optoelektrycznych i mikrofalowych wykonywanych z tych materiałów [20–23], a co za tym idzie można oczekiwać w tych przyrządach dostrzegalnego wpływu nieizotermicznych efektów termoelektrycznych na ich własności.

W celu zbadania tego wpływu oraz ocenienia na ile jest miarodajne szacowanie istotności nieizotermicznych efektów termoelektrycznych w oparciu o współczynnik α wykonano obliczenia testowe dla modelu cieplnego realnego przyrządu półprzewodnikowego, w którym występują duże gęstości prądu. Wybrano do tego celu opisane w [20,21] struktury powierzchniowych diod laserowych wykonanych z GaAs, w których robocze gęstości prądu sięgają 50 kA/cm^2 . Opis opracowanego dla tych struktur jednowymiarowego modelu cieplnego wraz z podstawowymi danymi liczbowymi przyjętymi do obliczeń zawiera Dodatek II. Model ten wykorzystano do wyznaczenia odpowiedzi termicznej na prostokątny impuls prądu przyjmując typową dla tego typu laserów wartość gęstości prądu w obszarze aktywnym, równą 32 kA/cm^2 . Obliczenia przeprowadzono z uwzględnieniem i bez uwzględnienia nieizotermicznych efektów termoelektrycznych a ich wynikiem były zmiany rozkładu temperatury wzdłuż linii przepływu prądu pozwalające na wyznaczenie zmian w czasie maksymalnej temperatury w strukturze. Na rys. 8 przedstawiono procentową różnicę w oszacowaniu maksymalnej temperatury w strukturze obliczonej z uwzględnieniem i bez uwzględnienia nieizotermicznych efektów termoelektrycznych. Zależność ta

posiada charakterystyczny pik „efektywności” nieizotermicznych efektów termoelektrycznych dla krótkich czasów trwania impulsu prądu. Jest on związany z tym, że odprowadzanie ciepła w tych efektach jest bezinercyjne, podczas gdy odprowadzanie ciepła na drodze klasycznego przepływu obarczone jest pewną inercją. Uzyskane procentowe zmniejszenie maksymalnej temperatury w wyniku uwzględnienia nieizotermicznych efektów termoelektrycznych jest niewielkie i nie przekracza 0.4% ale też odpowiada mu niewielka wartość parametru α gdyż rozważane struktury są strukturami o niewielkiej grubości warstwy GaAs odprowadzającej ciepło z obszaru aktywnego do ciepłowod. Wielkość α jest w nich rzędu 0.3%, a więc jest niższa od oszacowanych zmian procentowych maksymalnej temperatury.

PODSUMOWANIE

Praca jest poświęcona ocenie wkładu do bilansu energetycznego przyrządu półprzewodnikowego zaproponowanych w [1] źródeł ciepła związanych z izotermicznymi i nieizotermicznymi efektami termoelektrycznymi. Dla każdej z tych dwóch grup źródeł ciepła przeprowadzono odpowiednie obliczenia testowe mające na celu określenie warunków, w których te źródła mogą mieć istotny wpływ na własności cieplne przyrządów półprzewodnikowych. Jako struktury testowe wykorzystano w pracy strukturę diody idealnej p-n, diody prostowniczej p-v-n oraz diody laserowej.

Izotermicznym efektem termoelektrycznym jest poświęcony rozdz. 2. Przedstawione tam wyniki obliczeń testowych dla idealnej diody p-n, wykazują, że efekty te mogą w pewnych przypadkach uzyskać decydujące znaczenie dla własności cieplnych takiej struktury i np. w wyniku pochłaniania ciepła przy kontaktach moc rozpraszana wewnątrz struktury może nawet dwukrotnie przewyższać moc elektryczną dostarczaną do przyrządu. Obliczenia przeprowadzone dla struktury p-v-n, odpowiadającej rzeczywistej diodzie prostowniczej BYP 401, nie wykazały już tak wielkich wewnętrznych przepływów energii związanych z izotermicznymi efektami termoelektrycznymi, jednak ich wkład do bilansu energetycznego struktury był istotnie duży ($\geq 20\%$) dla całego zakresu gęstości prądów jakie mogą w niej wystąpić w typowych warunkach pracy.

W rozdz. 3 dokonano oceny wpływu nieizotermicznych efektów termoelektrycznych na cieplne własności przyrządów półprzewodnikowych. Ponieważ efekty te są konkurencyjne w stosunku do dyfuzyjnego przepływu ciepła, wyprowadzono w pracy zależność na procentowy współczynnik α określający stosunek mocy rozpraszanej w nieizotermicznych efektach termoelektrycznych w wyniku przepływu prądu przez próbkę o grubości w , w której występuje gradient temperatury, do mocy przenoszonej przez tą próbkę w wyniku dyfuzyjnego przepływu ciepła w warunkach cieplnego stanu ustalonego. Dla podstawowych materiałów półprzewodnikowych zależność ta jest przedstawiona graficznie na rys. 7 i pozwala na oszacowanie warunków, w których efekty te mogłyby wpływać w sposób istotny na cieplny stan ustalony w przyrządach wykonanych z tych materiałów. Dla typowych grubości pastylek półprzewodnikowych, wymaga to np. w krzemie gęstości prądu rzędu setek kA/cm^2 ,

czyli znacznie przekraczających najwyższe spotykane gęstości występujące w przyrządach wykonanych z tego materiału. Tak więc, nieizotermiczne efekty termoelektryczne mogą być w tych przyrządach zaniedbane. Inaczej wygląda sytuacja w optoelektronicznych przyrządach wykonanych np. GaAs czy InSb. Występujące w nich robocze gęstości prądu sięgają 100 kA/cm^2 , czemu może odpowiadać wartość współczynnika α do ułamka do kilkunastu procent, w zależności od grubości pastylki półprzewodnikowej i rodzaju materiału. Obliczenia odpowiedzi cieplnej na prostokątny impuls prądu przeprowadzone dla modelowej struktury symulującej przepływ ciepła w laserze powierzchniowym wykazały, że np. przy współczynniku $\alpha = 0.3\%$ zmiana maksymalnej temperatury w laserze wywołana nieizotermicznymi efektami termoelektrycznymi może być rzędu 0.4% .

DODATEK I: ELEKTRYCZNY MODEL STRUKTURY p-v-n

W Dodatku został przedstawiony elektryczny model struktury p-v-n dla stanu przewodzenia opracowany w oparciu o zaproponowaną w [25] metodę polegającą na wyodrębnieniu w rozważanej strukturze elementarnych obszarów, dla których rozkłady gęstości nośników można wyznaczyć w stosunkowo prosty sposób. W modelu tym przyjęto założenie, że domieszkowanie poszczególnych warstw jest jednorodne a złącza mają charakter skokowy, co pozwala na wyodrębnienie w rozważanej strukturze trzech elementarnych obszarów. Stanowią je, przedstawione na rys. I-1, dwa skrajne obszary emiterowe o niskim poziomie wstrzykiwania oraz obszar środkowy, „*zalaną*” przez nośniki nadmiarowe.

Rozkład nośników nadmiarowych w obszarach niskiego poziomu wstrzykiwania jest określony przez znane wzory (I-1a) i (I-1b):

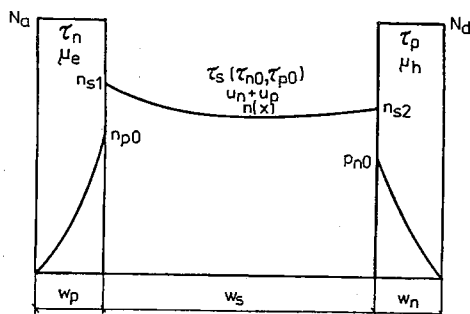
$$\Delta n = \Delta n_0 \left(\cos h \frac{x}{L_e} - \cot h \frac{w_p}{L_e} \sin h \frac{x}{L_e} \right) \quad (\text{I-1a})$$

$$\Delta p = \Delta p_0 \left(\cos h \frac{x}{L_h} - \cot h \frac{w_n}{L_h} \sin h \frac{x}{L_h} \right) \quad (\text{I-1b})$$

gdzie: $x=0$ odpowiada granicy ładunku przestrzennego, w_n i w_p – grubości odpowiednich warstw a L_e i L_h – stałe dyfuzji.

W obszarach tych panują warunki niskiego poziomu wstrzykiwania, dla których ruchliwość i czas życia są wielkościami stałymi, określonymi przez domieszkowanie oraz właściwości materiału wyjściowego, z którego wykonano diodę. Do wyznaczenia granicznych gęstości nośników mniejszościowych, wymaganych we wzorach (I-1), wykorzystano w modelu znane wzory (I-2) [18] łączące te gęstości z brzegowymi gęstościami nośników nadmiarowych w obszarze środkowym.

$$n_{p0} = \frac{n_{s1}^2}{N_a} \quad (\text{I-2a}) \quad p_{n0} = \frac{n_{s2}^2}{N_d} \quad (\text{I-2b})$$



Rys. I-1. Rozkład domieszkowania i gęstości nośników w stanie przewodzenia modelowanej struktury p-v-n

Do wyznaczenia rozkładu nośników w obszarze „zatopionym” wykorzystano wzór (I-3) podany przez Schlangenotto i Gerlacha [26]. Graniczne wartości gęstości nośników n_{s1} i n_{s2} występujące w tym wzorze zostały wyznaczone przy pomocy wzorów (I-4)–(I-7) wyprowadzonych przez Berza [27] oraz wzorów (I-8)–(I-9) wyprowadzonych w oparciu o prace [18] i [28].

$$p(x) = n(x) = \frac{n_{s1} \sinh \frac{w_s - x}{L} + n_{s2} \sinh \frac{x}{L}}{\sinh \frac{w_s}{L}} \quad (\text{I-3})$$

$$n_{s1} = \frac{D}{2h_1L} \left(-\cot h \frac{w_s}{L} + \left(\cot h^2 \frac{w_s}{L} + \frac{J}{J_1} \right)^{0.5} \right) \quad (\text{I-4})$$

$$n_{s2} = \frac{D}{2h_2L} \left(-\cot h \frac{w_s}{L} + \left(\cot h^2 \frac{w_s}{L} + \frac{J}{J_2} + \frac{4h_2n_{s1}}{D \sinh(w_s/L)} \right)^{0.5} \right); \quad (\text{I-5})$$

$$J_1 = \frac{1 + b}{4b} \frac{qD}{h_1\tau_s} \quad (\text{I-6}) \quad J_2 = \frac{1 + b}{4b} \frac{qD}{h_2\tau_s} \quad (\text{I-7})$$

$$h_1 = \frac{D_e}{L_e N_a} \cot h \frac{w_p}{L_e} \quad (\text{I-8}) \quad h_2 = \frac{D_h}{L_h N_d} \cot h \frac{w_n}{L_h} \quad (\text{I-9})$$

gdzie: $L = \sqrt{D\tau_s}$

$b = \mu_n/\mu_p$

D – ambipolarna stała dyfuzji w obszarze „zatopionym”

J – gęstość prądu płynącego przez diodę

D_e i L_e – stała dyfuzji i droga dyfuzji w warstwie p

D_h i L_h – stała dyfuzji i droga dyfuzji w warstwie n

Równania (I-3)–(I-9) zostały wyprowadzone przy założeniu stałości czasu życia τ_s i ambipolarnej stałej dyfuzji D w obszarze „zatopionym”. Założenie takie nie może być utrzymane przy obliczeniach dla dużego zakresu zmian gęstości prądu J , kiedy obie te wielkości zmieniają się o kilka rzędów. Dlatego układ równań

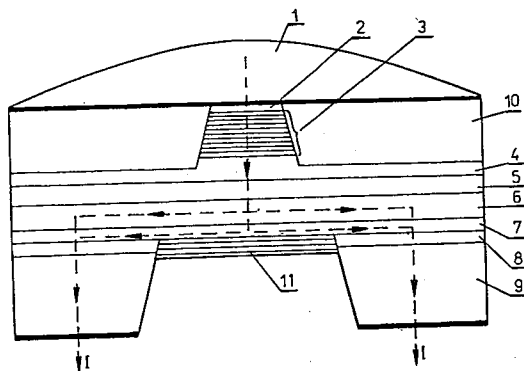
opisujących model został uzupełniony o równania (I-10) i (I-11), uzależniające te stałe od aktualnej wartości średniej gęstości nośników w obszarze „zatopionym”:

$$\tau_s = \frac{n_{sr}}{R(n_{sr})} \quad (\text{I-10}) \quad D = \frac{2kT(\mu_n + \mu_p)_{n=m_{sr}}}{q(1 - b^2)} \quad (\text{I-11})$$

Do wyznaczenia aktualnej wartości szybkości rekombinacji R wykorzystano w modelu zależność (4), natomiast do wyznaczenia aktualnej wartości sumy $\mu_n + \mu_p$ wykorzystano dane doświadczalne podane przez Dannhäusera [24].

DODATEK II: JEDNOWYMIAROWY MODEL ROZPŁYWU CIEPŁA W DIODZIE LASEROWEJ

Punktem wyjścia dla modelu są opisane w [20,21] struktury powierzchniowych diod laserowych, w których robocze gęstości prądu mogą przekraczać 50 kA/cm².



Rys. II-1. Poprzeczny przekrój analizowanej struktury diody laserowej

Tabela II-1

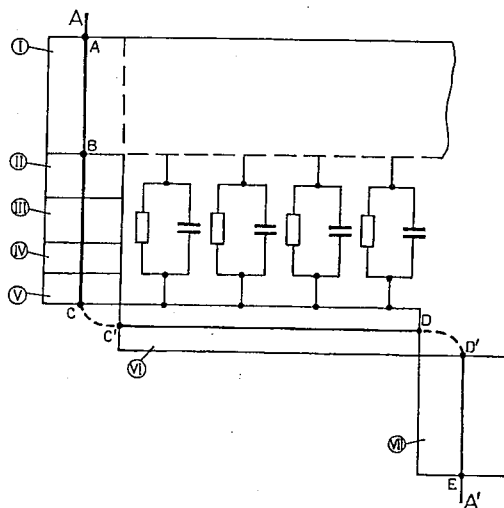
Dane o podstawowych warstwach uwzględnionych w modelu rozptyłu ciepła

nr	materiał	w [μm]	λ [W/mK]	c [J/gK]	r _o [Ωm]
2	p ⁺ GaAs	0,06	44	0,318	7e-5
3	p: Al _{0,1} Ga _{0,9} As/ Al _{0,7} Ga _{0,3} As		12,9	0,397	4. 4e-4
4	p: Al _{0,33} Ga _{0,67} As	0,3	12,2	0,358	9. 4e-4
5	n: GaAs	2	44	0,318	2. 2e-5
6	n: Al _{0,56} Ga _{0,44} As	1	10,9	0,385	6. 2e-4
7	n: Al _{0,3} Ga _{0,7} As	0,2	12,2	0,358	5. 8e-5
9	n: GaAs	90	44	0,318	2. 2e-5

Rys. II-1 przedstawia poprzeczny przekrój struktury a dane o poszczególnych warstwach są zebrane w Tabeli II-1. Linia przerywaną zaznaczono na rysunku drogę przepływu prądu. Jest to struktura o geometrii walcowatej, do której prąd

wpływa przez centralnie umieszczony kontakt anodowy. Prąd ten ma charakter czysto większościowy i przepływa kolejno przez warstwy 1,2,3 w kierunku obszaru aktywnego 4, w którym następuje wymiana nośników prądu w procesie rekombinacji. Następnie zmienia on kierunek i już jako prąd elektronowy płynie w poprzek struktury przez warstwy 5,6 by, po osiągnięciu warstwy 7, ponownie zmienić kierunek i poprzez warstwę 7 osiągnąć kontakt katodowy.

W rozważanej diodzie przepływ prądu i związana z nim generacja ciepła mają charakter trójwymiarowy, tak więc odpowiadający jej model cieplny powinien w ogólnym przypadku być także modelem trójwymiarowym. Biorąc jednak pod uwagę, że przeznaczeniem tego modelu ma być umożliwienie prześledzenia wpływu nieizotermicznych efektów termoelektrycznych na własności modelowanej diody, postanowiono wykorzystać przy jego tworzeniu fakt, że występowanie nieizotermicznych efektów termoelektrycznych jak i lokalizacja źródeł ciepła związanych ze stratami omowymi i rekombinacją są ściśle związane z przepływem prądu, a kanały przepływu prądu są w rozważanej strukturze wyraźnie wyodrębnione. Pozwoliło to na wykorzystanie do zamodelowania własności cieplnych diody modelu, w którym równanie przewodnictwa cieplnego jest rozwiązywane jednowymiarowo wzdłuż linii przepływu prądu a efekt trójwymiarowości jest uzyskiwany poprzez uwzględnienie przy tworzeniu algorytmu obliczeń symetrii cząstkowych rozważanego problemu



Rys. II-2. Schemat przepływu ciepła w modelowej strukturze diody

oraz poprzez wprowadzenie dodatkowych źródeł ciepła symulujących przepływy ciepła w kierunkach prostopadłych do analizowanej linii modelu jednowymiarowego.

Użyty w modelu algorytm obliczeń został stworzony w oparciu o przedstawiony na rys. II-2 schemat przepływu ciepła, który odwzorowuje procesy przepływu i generacji ciepła w modelowej strukturze. Obejmuje on 7, ponumerowanych liczbami rzymskimi warstw reprezentujących poszczególne fragmenty struktury lasera, przez

które przechodzi linia A-A' reprezentująca model jednowymiarowy. Przerywane fragmenty tej linii odpowiadają obszarom nieciągłości modelu jednowymiarowego, a ograniczające je punkty, np. C i C', są traktowane w modelu jako jeden punkt, dla którego jest wprowadzany odpowiedni warunek graniczny.

— WARSTWA I — reprezentuje miedzianą elektrodę anody, która nie jest uwzględniona na rys. II—1. Pełni ona rolę ciepłowodę odprowadzającego ciepło do obudowy. Z uwagi na jej prawie punktowy kontakt z warstwą 2 przyjęto w modelu, że ciepło wypływając z punktu B rozprzestrzenia się w niej równomiernie we wszystkich kierunkach. Oznacza to przyjęcie założenia o symetrii sferycznej przepływu ciepła, co pozwala na opisanie tego przepływu jednowymiarowym równaniem przepływu ciepła we współrzędnych sferycznych (II—1a), przy założeniu, że $r=0$ odpowiada punktowi B na rys. II—2. W modelu wykorzystano do zamodelowania warstwy I dyskretną postać tego równania (II—1b) uwzględniając jednocześnie, że kontakt pomiędzy warstwami w punkcie B nie jest punktowy ale ma skończony wymiar.

$$\lambda \left(\frac{2}{r} \frac{\delta T}{\delta r} + \frac{\delta^2 T}{\delta^2 r} \right) - c\rho \frac{\delta T}{\delta t} = -g(r, t) \quad (\text{II—1a})$$

$$T_i^{j+1} = T_{i-1}^j \left(1 + \frac{h}{r_i} \right) \alpha + T_{i+1}^j \left(1 - \frac{h}{r_i} \right) \alpha + T_i^j (1 - 2\alpha) + g_i^j, \quad (\text{II—1b})$$

gdzie: i, j — odpowiednio, przestrzenna i czasowa zmienna dyskretna, h — krok przestrzenny, k — krok czasowy a $\alpha = \lambda h^2 / (c\rho k)$;

— WARSTWA II — odpowiada spoiwu indowemu łączącemu pastylkę półprzewodnikową z ciepłowodem. Dla warstwy tej, podobnie jak dla następujących po niej warstw III, IV i V, przyjęto, że gęstość prądu, strumienia cieplnego oraz rozpraszania ciepła w wyniku strat omowych i efektów termoelektrycznych są stałe w całym przekroju poprzecznym a kierunek przepływu prądu i ciepła jest zgodny z linią A—A'. Tak więc przepływ ciepła w tej warstwie może być opisany jednowymiarowo, odpowiednio, równaniami (II—2a) i (II—2b).

$$\lambda \frac{\delta^2 T}{\delta^2 x} - c\rho \frac{\delta T}{\delta t} = -g(x, t) \quad (\text{II—2a})$$

$$T_i^{j+1} = (T_{i-1}^j + T_{i+1}^j) \alpha + T_i^j (1 - 2\alpha) + g_i^j; \quad (\text{II—2b})$$

— WARSTWA III — odpowiada warstwie 3 na rys. II—1 stanowiącej zwierciadło warstwowe. Jest ono utworzone z 25 segmentów, z których każdy składa się z pary cienkich warstw o istotnie różnych parametrach cieplnych i elektrycznych. Dla zamodelowania tej warstwy wprowadzono w modelu jednorodną warstwę zastępczą o tak dobranych parametrach cieplnych i elektrycznych aby odwzorowywała ona wiernie przepływ prądu oraz przepływ i generację ciepła w rzeczywistej warstwie.

— WARSTWA IV — odpowiada warstwie 4 na rys. I—1.

— WARSTWA V — odpowiada warstwie 5 na rys. I—1. Jest to obszar „aktywny”, w którym następuje zmiana charakteru prądu z elektronowego na

dziurowy w wyniku rekombinacji. W warstwie tej, obok strat omowych, występują straty rekombinacyjne związane z tym, że rekombinacji promienistej towarzyszy rekombinacja fononowa i rekombinacja Augera. Wielkość tych strat jest równa $(1-\eta)RW_g$, gdzie η — sprawność kwantowa lasera.

— WARSTWA VI — jest warstwą zastępczą dla warstw 6 i 7 z rys. II—1 a jej parametry elektryczne i cieplne dobrano tak, aby mogła ona wiernie symulować przepływ ciepła wzdłuż linii A—A'. Stanowi ona pierścień, poprzez który prąd i ciepło płyną równomiernie od jego centrum ku zewnętrznemu brzegowi. Pozwala to na przyjęcie za podstawę do opisu przepływu ciepła równania jednowymiarowego przepływu ciepła we współrzędnych walcowych (II—3a), któremu odpowiada dyskretne równanie (II—3b). Kierunek A—A' jest w tej warstwie najważniejszy z punktu widzenia przepływu prądu i rozkładu generacji ciepła, jednak podstawowy strumień ciepła płynie w niej nierównoległe a prostopadle do linii A—A', odprowadzając ciepło do ciepłowod. Aby ten fakt uwzględnić wprowadzono do modelu, przedstawione na rys. II—2, dyskretne elementy R C symulujące rezystancję i pojemność cieplną warstw oddzielających warstwę VI od ciepłowod. Odpowiada im w równaniach (II—3) funkcja gęstości rozpraszania mocy g_{th} , której wartość w każdym kroku czasowym jest określana m.in. w oparciu o aktualny rozkład temperatury w ciepłowodzie.

$$\lambda \left(\frac{1}{r} \frac{\delta T}{\delta r} + \frac{\delta^2 T}{\delta^2 r} \right) - c\rho \frac{\delta T}{\delta t} = -g(r, t) - g_{th}(r, t) \quad (\text{II—3a})$$

$$T_i^{j+1} = T_{j-1}^j \left(1 + \frac{h}{2r_i} \right) \alpha + T_{i+1}^j \left(1 - \frac{h}{2r_i} \right) \alpha + T_i^j (1 - 2\alpha) + g_i^j + g_{th_i}^j \quad (\text{II—3b})$$

— WARSTWA VII — odpowiada warstwie 8 i 9 na rys. II—1. Z uwagi na stosunkowo najmniejsze znaczenie tej warstwy dla procesu odprowadzania ciepła, przyjęto w modelu, że stanowi ona pierścień o stosunkowo dużej powierzchni przekroju poprzecznego, przez który ciepło płynie wzdłuż linii A—A' równoległej do tworzącej pierścienia. Założeniu temu odpowiada równanie przepływu ciepła (II—2), w którym uwzględniono jedynie zmiany gęstości generacji ciepła wynikające z rozprzestrzeniania się prądu w warstwie 7.

Warunki brzegowe sformułowano w modelu następująco:

— na zewnętrznych powierzchniach skrajnych warstw, temperatura jest stała np. równa 0.

— na granicy warstw, pomiędzy którymi zachowana jest ciągłość, jest spełniony warunek ciągłości gęstości strumienia ciepła

— na granicy warstw, pomiędzy którymi nie jest zachowana ciągłość modelu, jest zachowana zasada, że prąd i strumień ciepła opuszczający jedną warstwę jest równy prądowi i strumieniowi ciepła wnikającemu do warstwy następnej.

Bezpośrednimi wynikami symulacji komputerowej są w tym modelu chwilowe rozkłady temperatury wzdłuż linii A—A', w oparciu o które można wyznaczyć zmiany maksymalnej temperatury w laserze.

BIBLIOGRAFIA

1. Z. Lisik: *Modelowanie generacji ciepła w przyrządach półprzewodnikowych — Część II: Funkcje rozpraszania ciepła*, Kwart. Elektr. i Tel. 1991, z. 1, ss
2. A. F. Ioffe, A. S. Silbans, E. K. Iordansvili, T. S. Stavickaja: *Termoelektriceskoe ochlazhenie*, Izd. Akad. Nauk SSSR, Moskva—Leningrad, 1956
3. A. E. Middleton, W. W. Scanlon: *Measurement of the thermoelectric power of germanium at temperatures above 78*, K. Phys. Rev., 1953, vol. 91, nr 2, ss. 219—232
4. C. Harrington: *Theory of thermoelectric power of semiconductors*, Phys. Rev., 1954, vol. 96, nr 5, ss. 1163—1187
5. P. J. Price: *Theory of transport effects in semiconductors: Thermoelectricity*, Phys. Rev., 1956, vol. 104, nr 5, ss. 1223—1239
6. M. C. Steele, D. Rosi: *Thermal conductivity and thermoelectric power of germanium—silicon alloys*, J. Appl. Phys., 1958, vol. 29, nr 11, ss. 1517—1520
7. P. I. Baranskij: *Obemno-gradientnyj efekt Thomsona*, Fiz. Tver. Tela, 1960, vol. 11, nr 3, ss. 445—457
8. A. F. Ioffe: *Termoelektriceskie generatory*, Izv. Akad. Nauk SSSR Ser. Fiz., 1956, vol. 20, nr 1, ss. 76—80
9. A. W. Van Herwaarden, P. M. Sarro: *Thermal sensors based on the Seebeck effect*, Sensors & Actuators, 1986, vol. 10, nr ss. 321—346
10. R. N. Hall: *An analysis of the performance of thermoelectric devices made from long-lifetime semiconductors*, Solid-St. Electron., 1961, vol. 2, ss. 113—122
11. E. S. Rittner: *On the Peltier heat pump*, J. Appl. Phys., 1959, vol. 30, nr 5, ss. 702—707
12. L. S. Stilbanc, E. K. Pordanisvili, T. S. Stavickaja: *Termoelektriceskoje ochlazhenie*, Izv. Akad. Nauk SSSR Ser. Fiz., 1956, vol. 20, nr 1, ss. 82—88
13. T. S. Silday: *Performance of composite Peltier Junctions of Bi_2Te_3* , J. Appl. Phys., 1957, vol. 28, nr 9, ss. 1035—1042
14. A. Abete, M. Parvis: *Prova convenzionale per le caratteristiche dei moduli a semiconduttori termoelettrici*, Alta Freq., 1986, vol. 55, nr 1, ss. 63—71
15. C. M. Cortes, R. G. Hunsperger: *Effects of contact resistance and dopant concentration in metal-semiconductor thermoelectric coolers*, IEEE Trans. Electron Devices, 1980, vol. ED—27, nr 3, ss. 521—525
16. Z. Lisik: *Elektrotermiczny model przewodzącej diody p-n*, VII Kraj. Konf. „Teoria Obwodów i Układy Elektroniczne”, Kazimierz Dolny 23—26.X.1984, ss. 241—5
17. J. H. Bowen, M. E. Breese, V. A. Mikenas, M. Weiss, Liu Shing-Gong, H. Sobol: *Analytic and experimental techniques for evaluating transient thermal characteristics of TRAPATT diodes*, IEEE Trans. Electron Devices, 1974, vol. ED—21, nr 8, ss. 480—487
18. A. Herlet: *The forward characteristic of silicon power rectifiers at high current density*, Solid-St. Electron., 1968, vol. 11, ss. 717—742
19. A. Silard, M. Bodea, M. Luca: *Predicting the surge capability of power thyristors*, Electron. Lett., 1980, vol. 16, nr 9, ss. 325—327
20. D. Botez, L. M. Zinkiewicz, T. J. Roth, L. J. Mawst, G. Peterson: *Low-threshold-current-density vertical-cavity surface-emitting AlGaAs/GaAs diode lasers*, IEEE Photonics Technology Letters, 1989, vol. 1, nr 8, ss. 205—208
21. M. Sakamoto, D. F. Welch, J. G. Endriz, D. R. Scifres, W. Streifer: *75 W cw monolithic laser diode arrays*, Appl. Phys. Lett., 1989, vol. 54, ss. 2299—2301
22. R. C. Goodfellow, A. C. Carter, G. J. Res, R. Davis: *Radiance saturation in small-area GaInAsP/InP and GaAlAs/GaAs*, LED's, IEEE Trans. Electron Devices, 1981, vol. ED—28, nr 4, ss. 365—371
23. H. Sada, Y. Motegi, K. Iga: *GaInAsP/InP surface emitting injection laser with short cavity length*, IEEE J. Quant. Electron., 1983, vol. QE—19, nr 6, ss. 1035—1041
24. F. Danhäuser: *Die Abhängigkeit der Trägerbeweglichkeit in Silizium von der Konzentration der freien Ladungsträger — I*, Solid-St. Electron., 1972, vol. 15, nr 12, ss. 1371—1376

25. Z. Lisik: *Computer thermal model of thyristor for dynamic states*, ECCTD'80 Warszawa 1980, ss. 343–348
26. H. Schlängenotto, W. Gerlach: *On the effective carrier lifetime in p-s-n rectifiers at high injection levels*, Solid-St. Electron., 1969, vol. 12, nr 4, ss. 267–275
27. F. Berz: *A simplified theory of the p-i-n diode*, Solid-St. Electron., 1977, vol. 20, nr 8, ss. 709–714
28. J. Burtscher, F. Dannhäuser, J. Krausse: *Die Rekombination in thyristoren und Gleichrichtern aus Silizium: Ihr Einfluss auf die Durchlasskennlinie und das Freierzeitverhalten*, Solid-St. Electron., 1975, vol. 18, nr 1, ss. 35–83

Z. LISIK

MODELLING OF HEAT GENERATION IN SEMICONDUCTOR DEVICES

PART III: THE CONTRIBUTION OF THERMOELECTRIC EFFECTS TO THE ENERGY BALLANCE OF SEMICONDUCTOR DEVICE

S u m m a r y

In Part III, the contribution of heat sources associated with isothermal and non-isothermal thermoelectric effects to the energy balance of a semiconductor device has been evaluated. For each group of the heat sources testing calculations have been performed in order to determine the conditions under which these sources essentially affects semiconductor device properties. As testing structures as an ideal p-n diode, a rectifier p-v-n diode and a laser diode have been used.

Skalowanie i korekcja zniekształceń geometrycznych w systemie cyfrowego przetwarzania obrazów z kamerą widikonową

ANDRZEJ MATERKA, WŁODZIMIERZ GÓRNIŚIEWICZ, MARIUSZ KUKUŁA

Instytut Elektroniki, Politechnika Łódzka

Otrzymano 1990.04.03

Autoryzowano do druku 1990.07.31

W pracy rozważono zagadnienia skalowania i korekcji zniekształceń geometrycznych w systemie cyfrowego przetwarzania obrazów z kamerą widikonową. Na bazie twierdzenia o warunkach jednoznacznego odwzorowanie odległości punktu od obiektu na wartość sygnału wizyjnego opracowano metodę skalowania systemu z błędem średniokwadratowym kilkakrotnie mniejszym od odstępów próbkowania obrazu. Własności metody potwierdzono symulacją komputerową i pomiarami współczynnika skalowania w rzeczywistym systemie. Zbadano możliwość korekcji zniekształceń geometrycznych obrazu z wykorzystaniem trzech różnych funkcji matematycznych do opisu zniekształceń. Wykazano, że zastosowanie interpolacji dwuliniowej pozwala uzyskać 6-krotne skrócenie czasu obliczeń w porównaniu do aproksymacji wielomianami stopnia 3, przy podobnym, około dwukrotnym zmniejszeniu rozrzutu wyników pomiaru wybranych figur geometrycznych w polu obrazu. Uzyskane wyniki mogą być wykorzystane w ilościowych badaniach biomedycznych i przemysłowej kontroli jakości, a także w robotyce.

1. WSTĘP

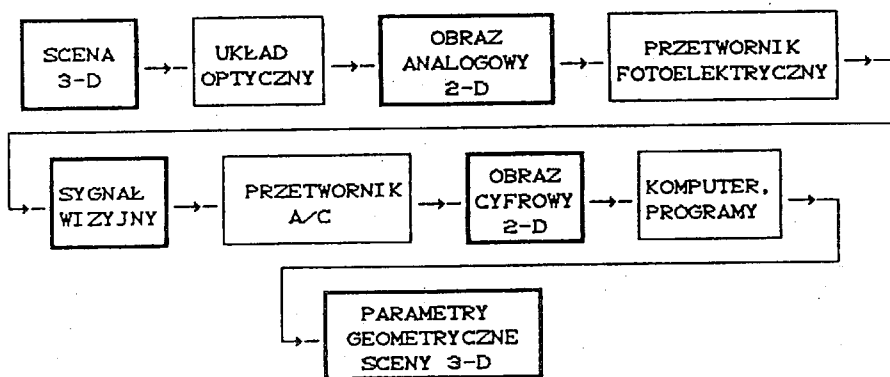
Możliwość pomiaru parametrów geometrycznych ciał fizycznych (figur 2—D i 3—D) za pomocą systemów cyfrowego przetwarzania obrazów jest praktycznie niezastąpiona w wielu zastosowaniach, jak diagnostyka medyczna, przemysłowa kontrola jakości i robotyka. Często nie istnieje inna realistyczna metoda przeprowadzenia określonych badań ilościowych, np. w analizie obrazów angiograficznych [15] lub biometrii naczyń krwionośnych [4], w których to przypadkach niezbędna jest ocena pojemności pojedynczych żył i tętnic o średnicy ułamka milimetra [11].

Podobnie w potomorfologii, system obrazowy daje możliwość obiektywnej oceny wpływu substancji toksycznej na tkanki żywe, poprzez cyfrową analizę obrazów mikroskopowych, z użyciem odpowiednio przygotowanych preparatów [5].

W każdym z zastosowań ważna jest identyfikacja głównych źródeł błędów pomiaru. Może ona posłużyć do optymalizacji metod pomiarowych lub aparatury

pomiarowej w celu przeprowadzenia zamierzonych badań ilościowych przy możliwie najniższym koszcie.

Ogólny schemat procesów przetwarzania informacji obrazowej przedstawiono na rys. 1. Rozkład jakości badanej sceny 3-D jest rzutowany, przez układ optyczny, na płaszczyznę fotoczułą przetwornika fotoelektrycznego, np. w kamerze telewizyjnej. Powstaje w ten sposób analogowy obraz 2-D, będący sygnałem wejściowym przetwornika. Rozkład jasności tego obrazu jest przetwarzany na sygnał wizyjny, który jest próbkowany i zapamiętywany jako dwuwymiarowy obraz cyfrowy. Jest to w istocie zbiór liczb wykorzystywany w obliczeniach mających na celu ocenę parametrów geometrycznych sceny 3-D. Programy tych obliczeń powinny uwzględniać zniekształcenia informacji obrazowej zachodzące w poszczególnych elementach systemu. Są wśród nich zniekształcenia systemowe wynikające z dyskretyzacji obrazu, są też zniekształcenia wywołane zjawiskami niepożądanymi, jak zakłócenia szumowe, olśnienie i skończona rozdzielczość przetwornika fotoelektrycznego, nieliniowe odchylenie strumienia elektronów w widikonie i wiele innych [3]. Zagadnienia redukcji wpływu czynników zniekształcających na błąd pomiaru parametrów geometrycznych figur płaskich oraz objętości walców obrotowych o zmiennej średnicy rozważano w pracach [11,13], w których zaproponowano między innymi nieobciążone estymatory długości obwodu i pola powierzchni figur. W [11,13] wykazano, że wyniki pomiarów, np. pola powierzchni koła w systemie z kamerą widikonową, w dużym stopniu zależą od położenia modelu koła w polu widzenia kamery. Jest to efektem zniekształceń geometrycznych obrazu związanych z niedoskonałością układu optycznego, a także z niezamierzoną nieliniowością procesów odchylenia strumienia elektronów. Opracowanie metody cyfrowej korekcji tych zniekształceń jest jednym z celów niniejszej pracy.



Rys. 1. Schemat przetwarzania informacji obrazowej

Koszt systemu obrazowego silnie zależy od wielkości pamięci obrazu cyfrowego, wynikłej z przyjętych odstępów próbkowania. System o rozdzielczości 256×256 punktów zawiera pamięć o pojemności 64 kB, która jest czterokrotnie większa w systemie o rozdzielczości 512×512 punktów. Jednocześnie, przy zwiększonej

rozdzielczości wzrasta czas obliczeń przetwarzania informacji obrazowej. W związku z tym, celowe jest poszukiwanie metod pomiaru parametrów geometrycznych zadaną dokładnością przy możliwie małej rozdzielczości systemu. Metody takie, zaaplikowane w systemach dużej rozdzielczości, pozwalają na zwiększenie wymiarów bezwzględnych analizowanej sceny 3-D, a tym samym jednoczesną analizę większej liczby figur (obiektów), co ma istotne znaczenie w badaniach statystycznych. Poszukując metod pomiaru z dokładnością lepszą niż wynikałoby to z odstępu próbkowania (subpixel accuracy) zaproponowano w niniejszym artykule metodę dokładnego skalowania systemu obrazkowego, wykorzystującą własności matematyczne funkcji odwzorowującej odpowiedź impulsową przetwornika fotoelektrycznego.

2. MODEL MATEMATYCZNY OBRAZU

Niech Oxy będzie układem współrzędnych prostokątnych dwuwymiarowego obrazu analogowego na wejściu przetwornika fotoelektrycznego. Oznaczmy przez $f(x, y)$ jasność obrazu w punkcie (x, y) . Jasności tej odpowiada sygnał wizyjny $v(x+dx, y+dy)$. Funkcje $dx=dx(x,y)$ i $dy=dy(x,y)$ reprezentują zniekształcenia geometryczne obrazu powstające w przetworniku fotoelektrycznym [17]. Jeśli proces przetwarzania energii promienistej na sygnał elektryczny jest liniowy i nie zależy od przesunięcia na płaszczyźnie współrzędnych Oxy [14] to sygnał wizyjny jest określony zależnością

$$v(x + dx, y + dy) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(\xi, \eta) h(x - \xi, y - \eta) d\xi d\eta + v_N, \quad (1)$$

w której $h(x,y)$ jest funkcją opisującą odpowiedź impulsową przetwornika fotoelektrycznego (funkcja dyspersji punktu, ang. point spread function), v_N reprezentuje zakłócenia szumowe, natomiast ξ i η są zmiennymi całkowania. Odpowiedź impulsowa reprezentuje połączone zjawiska skończonej rozdzielczości układu optycznego (w wyniku dyfrakcji światła i różnego rodzaju aberacji) i skończonych wymiarów obszaru, w którym następuje faktyczny proces przetwarzania obrazu optycznego na sygnał wizyjny. Obszar ten jest określony rozmiarami plamki świetlnej w skanerze elektromechanicznym [9] albo rozkładem gęstości nośników ładunku w wiązce elektronów w widikonie [3].

W każdym fizycznie realizowanym przetworniku fotoelektrycznym obszar przetwarzania ma powierzchnię większą od zera. Gwarantuje to odpowiedni stosunek mocy sygnału do mocy szumu na wyjściu przetwornika. Wynika stąd, że kształt odpowiedzi impulsowej jest różny od idealnego przypadku delty Diraca i przyjmuje niezerowe wartości w pewnym obszarze $\mathcal{D} \subset \mathbb{R}^2$:

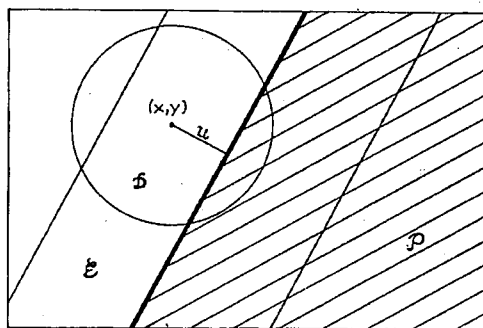
$$h(x, y) \neq 0, \quad (x, y) \in \mathcal{D}, \quad \text{mes}(\mathcal{D}) \neq 0 \quad (2)$$

Tym samym, jak wynika z (1) i (2) każdy przetwornik wprowadza zniekształcenia liniowe, zwane zniekształceniami apertury [19]. Charakteryzują się one rozmyciem linii brzegowej obiektów w obrazie.

Sygnał wizyjny $v(x,y)$ jest próbkowany i kwantowany za pomocą przetwornika analogowo-cyfrowego (por. rys. 1.). Zakładamy, że czas przetwarzania jest pomijalnie mały w porównaniu z odstępem próbkowania w dziedzinie czasu. Przyjmujemy także jednakowe odstępki na płaszczyźnie obrazu, odpowiednio w kierunku x i w kierunku y , $\Delta_x = \Delta_y$. Ponadto zakładamy, że liczba bitów przetwornika A/C jest na tyle duża, by można było zaniedbać błędy kwantowania sygnału wizyjnego. Tym samym obraz cyfrowy jest dany zbiorem próbek $v_{ij} = v(x_0 + i\Delta x, y_0 + j\Delta y)$, gdzie (x_0, y_0) jest punktem początkowym, natomiast i, j są liczbami całkowitymi $i=0,1,\dots,N-1$, $j=0,1,\dots,M-1$, przy czym zwykle $N=M$.

Założmy teraz, że przetwornik fotoelektryczny nie wprowadza zniekształceń geometrycznych, oraz że zakłócenia szumowe można zaniedbać:

$$dx = dy = 0, \quad v_N = 0. \quad (3)$$



Rys. 2. Równomiernie oświetlona półpłaszczyzna P z obszarem, przejściowym $\mathcal{E}$

Wpływ niezerowych wartości dx , dy , i v_N na błędy pomiaru zostanie zbadany w dalszej części pracy.

Niech sygnał wejściowy odpowiada równomiernie oświetlonej półpłaszczyźnie $\mathcal{P}$:

$$f(x, y) = \begin{cases} 1, & (x,y) \in \mathcal{P} \\ 0, & (x,y) \notin \mathcal{P} \end{cases} \quad (4)$$

Symbolem $\mathcal{E}$ na rys. 2 oznaczono obszar przejściowy w otoczeniu brzegu półpłaszczyzny $\mathcal{P}$. Jest on zbiorem takich punktów (x,y) , że przesunięty o wektor $[x,y]$ obszar $\mathcal{D}$ (dziedzina odpowiedzi impulsowej przetwornika) ma niezerową część wspólną z obiektem $\mathcal{P}$. W przypadku obszaru $\mathcal{D}$ w kształcie koła o promieniu δ , obszar $\mathcal{E}$ jest pasem o szerokości 2δ , z linią środkową na brzegu półpłaszczyzny. Oznaczmy przez u odległość euklidesową punktu (x,y) od brzegu półpłaszczyzny, przy czym $u < 0$ dla punktów (x,y) leżących poza $\mathcal{P}$ i $u > 0$ dla $(x,y) \in \mathcal{P}$. Dla punktów brzegowych $u=0$. Udowodnimy teraz następujące twierdzenie [11]

Jeżeli 1) Odpowiedź impulsowa $h(x,y)$ przetwornika fotoelektrycznego jest funkcją obrotowo symetryczną

$$\wedge h(\rho, \theta) = h(\rho) \quad (5)$$

$$\theta \in [0, 2\pi)$$

gdzie: $\rho = \sqrt{x^2 + y^2}$, $\theta = \arctan(x/y)$, $(x,y) \in \mathcal{D}$;
2) Dla każdego $(x,y) \in \mathcal{D}$ całka

$$g(y) = \int_0^\infty h(x, y) dx \quad (6)$$

jest ciągłą i dodatkową funkcją zmiennej y , to odległość u dowolnego punktu $(x,y) \in \mathcal{E}$ od brzegu równomiernie oświetlonej półpłaszczyzny jest przyporządkowana wzajemnie jednoznacznie sygnałowi wizyjnemu $u(x,y)$ w tym punkcie.

DOWÓD [12]

Zdefiniujmy pomocniczy układ współrzędnych $0\xi\eta$, z początkiem w punkcie (x,y) według rys. 2. W wyniku symetrii obrotowej funkcji $h(x,y)$ można przyjąć dowolną orientację tego układu na płaszczyźnie. Załóżmy, że oś 0ξ jest równoległa do brzegu półpłaszczyzny, a oś 0η prostopadła do niego, z dodatnim kierunkiem 0η w kierunku od półpłaszczyzny. Sygnał wizyjny w punkcie (x,y) jest zgodnie z (1), (2), (3) dany wzorem

$$u = \int_{-\infty}^u \left[\int_{-\infty}^\infty h(-\xi, -\eta) d\xi \right] d\eta \quad (7)$$

co można z uwzględnieniem 1) i 2) zapisać jako

$$u = 2 \int_{-\infty}^u g(\eta) d\eta = \varphi(u). \quad (8)$$

Ponieważ zgodnie z 2) $g(\eta)$ jest ciągłą funkcją η w każdym punkcie $\eta=u$, to funkcja $\varphi(u)$ posiada pochodną w tym punkcie [2]:

$$\frac{d\varphi(u)}{du} = g(u). \quad (9)$$

Z drugiej strony, ponieważ $g(u)$ jest ograniczona to $\varphi(u)$ jest ciągłą. Jednocześnie $g(u) > 0$, czyli $v = \varphi(u)$ jest funkcją rosnącą. Istnieje zatem funkcja odwrotna

$$u = g^{-1}(v). \quad (10)$$

która również jest rosnącą funkcją zmiennej v [2]. Tym samym, dla przyjętych założeń 1) i 2) istnieje wzajemna jednoznaczność między odległością punktu od brzegu $u(x,y)$ i wartością sygnału wizyjnego $v(x,y)$, c.b.d.o.

Pomiary dowodzą [12], że założenia 1) i 2) powyższego twierdzenia są w przybliżeniu spełnione w przypadku kamer widikonowych. Łatwo wykazać, że w takich przypadkach progowanie obrazu, z progiem równym średniej arytmetycznej jasności tła i jasności obiektu, odtwarza kształt obiektu mimo rozmycia wywołanego skończoną zdolnością rozdzielczą przetwornika fotoelektrycznego. Własność tę wykorzystano w [12] do opracowania metod pomiaru pola powierzchni i długości obwodu figur płaskich oraz objętości brył o przekroju kołowym z dużą dokładnością.

Założenie 1) nie jest na ogół spełnione w przypadku przetworników CCD. W przypadkach tego typu obszar D , ogólnie zdefiniowany przez (2), jest w przybliżeniu prostokątem i odpowiada elementarnym komórkom akumulacji ładunku, w „studniach potencjału” pod metalizacją elektrod sterujących [7]. W takim przypadku wartość sygnału wizyjnego odpowiadającego danej komórce zależy nie tylko od odległości środka symetrii tej komórki od brzegu obiektu (oświetlonej półpłaszczyzny) ale też od kąta między brzegiem obiektu a osią symetrii rozważanego obszaru D . Prowadzi to do określonych i łatwych do zaobserwowania zniekształceń linii brzegowej obiektów w obrazach rejestrowanych za pomocą kamer CCD. Badanie takich zniekształceń będzie przedmiotem dalszych prac.

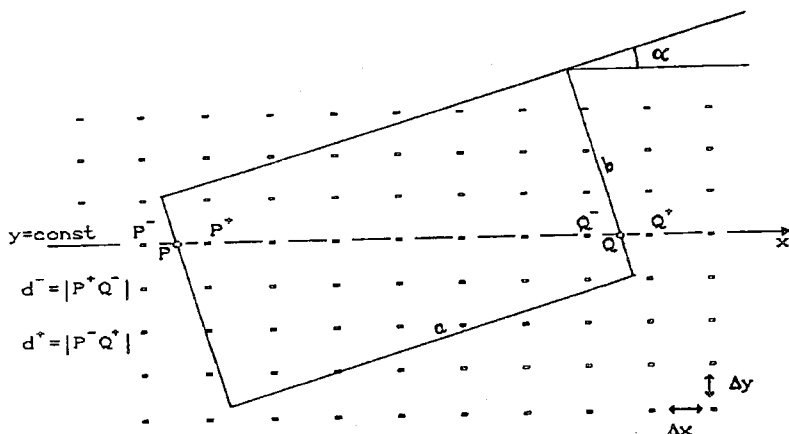
Założenia 2) nie spełniają przetworniki fotoelektryczne wykazujące cechy sieci neuronowej z hamowaniem obocznym. Odwzorowanie $v = \varphi(u)$ dla takich sieci nie jest wzajemnie jednoznaczne, co można sprawdzić doświadczalnie, za pomocą cienkich elektrod do pomiaru odpowiedzi żywych neuronów [10]. Jest to jedną z przyczyn złudzenia znanego jako pasy Macha [10], które można interpretować w kategoriach niejednoznacznego odwzorowania odległości receptorów światła od linii zmiany jasności obrazu, na intensywność odpowiedzi receptorów (np. w oku człowieka).

W pracy niniejszej założymy, że warunki 1) i 2) powyższego twierdzenia są spełnione. Wykorzystamy to do wyprowadzenia metody skalowania systemu przetwarzania obrazów z błędem mniejszym, co do modułu, od odstępów próbkowania Δx .

3. SKALOWANIE SYSTEMU PRZETWARZANIA OBRAZÓW

Załóżmy, że wzorcem do skalowania jest prostokąt o znanej długości boków i jasności I_0 większej od jasności tła obrazu I_b . Dłuższy bok prostokąta tworzy z linią $y = \text{const}$ kąt α . Kropkami zaznaczono na rys. 3 punkty próbkowania obrazu. Odległość między krótszymi bokami wzorca $|PQ|$ jest dla $\alpha \simeq 0$ w przybliżeniu równa długości boku a . W ogólnym przypadku dowolnego, losowego położenia wzorca w płaszczyźnie obrazu punkty P i Q nie pokrywają się z punktami próbkowania.

W przypadku skalowania „ręcznego”, operator, za pomocą znacznika sterowanego manipulatorem kulkowym, zaznacza te dwa punkty dyskretnej siatki obrazu,



Rys. 3. Prostokątny wzorec do skalowania na półpłaszczyźnie obrazu (a, b – boki prostokąta)

które według jego subiektywnej oceny znajdują się najbliżej nieznanymi punktów P oraz Q. Przy pewnej wprawie, na podstawie porównania jasności punktów, np. Q^- i Q^+ , można, w nawiązaniu do kształtu funkcji $v = \varphi(u)$ dla danego przetwornika obrazowego, ocenić czy punkt Q znajduje się bliżej punktu Q^- czy też Q^+ . Podobne rozumowanie można przeprowadzić dla punktów P^- , P, P^+ w pobliżu drugiego końca wzorca. Można uważać, że dla populacji wzorców o długości z przedziału $(d^-, d^- + \Delta x)$ długości d wzorca zostanie przypisana długość $d^- = |P^+ Q^-|$. Wartość oczekiwania błędu skalowania jest zatem różna od zera (przy założeniu równomiernego rozkładu prawdopodobieństwa)

$$\bar{\varepsilon}_1 = \frac{1}{\Delta x} \int_{d^-}^{d^- + \Delta x} (\xi - d^-) d\xi = \frac{\Delta x}{2}, \quad (11)$$

przy czym błąd maksymalny wynosi $\varepsilon_M = \Delta x$. Z kolei, wariancja błędu jest dana całką

$$\alpha_1^2 = \frac{1}{\Delta x} \int_{d^-}^{d^- + \Delta x} \left(\xi - d^- - \frac{\Delta x}{2} \right)^2 d\xi = \frac{(\Delta x)^2}{12}, \quad (12)$$

czemu odpowiada odchylenie średnie kwadratowe błędu skalowania

$$\alpha_1 = \frac{\Delta x}{2\sqrt{3}}. \quad (13)$$

W przypadku automatycznego skalowania z wykorzystaniem obrazu binarnego wzorca, obraz podlega segmentacji przez progowanie z progiem równym średniej arytmetycznej jasności tła i jasności obszaru reprezentującego wzorec. Odpowiedni program komputerowy służy do automatycznego wyznaczania punktów krańcowych wzorca, wzdłuż linii $y = \text{const}$. Są nimi punkty P^+ i Q^- dla populacji wzorców o długości z przedziału (d^-, d^+) . Jeśli każdej długości $d \in (d^-, d^+)$ przypisać jednakowe

prawdopodobieństwo zdarzenia, że długość wzorca jest równa d , błąd średni skalowania wyniesie

$$\bar{\varepsilon}_2 = \Delta x \quad (14)$$

z odchyleniem średnim kwadratowym

$$\sigma_2 = \frac{\Delta x}{\sqrt{3}} \quad (15)$$

Porównanie wzorów (14) z (11) oraz odpowiednio (15) z (13) wskazuje na możliwość redukcji błędu skalowania poprzez wykorzystanie informacji o rozkładzie jasności obrazu w pobliżu krawędzi wzorca. Rozkład ten jest opisany funkcją $v = \varphi(u)$, która jednoznacznie przetwarza odległość punktu od krawędzi na wartość sygnału wizyjnego, o ile założenia podanego wyżej twierdzenia są spełnione.

Załóżmy, że odpowiedź impulsowa przetwornika fotoelektrycznego może być przybliżona aperturą kołową o promieniu δ :

$$h(x, y) = \begin{cases} (\pi\delta^2)^{-1}, & x^2 + y^2 \leq \delta^2 \\ 0, & x^2 + y^2 > \delta^2, \end{cases} \quad (16)$$

której odpowiada następująca funkcja φ

$$v = \varphi(u) = \frac{1}{2} + \frac{1}{\pi} \left[\arcsin(u/\delta) + (u/\delta) \sqrt{1 - (u/\delta)^2} \right]. \quad (17)$$

W przypadku kamery widikonowej TPK 162 funkcja ta przybliża rzeczywiste rozmycie krawędzi obiektów w obrazach z odchyleniem średnim kwadratowym 0.71 przy gradacji jasności 128-stopniowej [12], nie gorzej od przypadku, w którym (16) zastąpiono funkcją Gaussa zmiennej $u = \sqrt{x^2 + y^2}$ [8]. Według proponowanej metody skalowania, w celu wyznaczenia współrzędnej x_p punktu P należy do ciągu próbek obrazu wzdłuż linii $y = \text{const}$ dopasować, metodą najmniejszych kwadratów, funkcję

$$v_i = v_b + \Delta v \varphi \left(\frac{x_0 + i\Delta x - x_p}{\delta} \right) \quad (18)$$

gdzie $\varphi(\cdot)$ jest dana wzorem (17), natomiast $\{v_b, \Delta v, x_p, \delta\}$ jest zbiorem wyznaczanych parametrów. Mają one następującą interpretację:

v_b – jasność tła

Δv – różnica między wartością sygnału wizyjnego w obszarze wzorca i w obszarze tła (kontrast wzorca)

x_p – współrzędna x punktu P

δ – promień kołowego obszaru D będącego dziedziną odpowiedzi impulsowej przetwornika

Zbiór ten jest wyznaczany przez minimalizację funkcji

$$f(v_b, \Delta v, x_p, \delta) = \sum_{i=i_p-W}^{i_p+W} (v_i - v_b)^2 \quad (19)$$

w przestrzeni parametrów $\{v_b, \Delta v, x_p, \delta\}$ przy czym v_i jest wartością i -tej próbki obrazu wzdłuż linii $y = \text{const}$ w ciągu złożonym z $(2W+1)$ elementów, natomiast i_p jest numerem próbki leżącej w pobliżu punktu P, a v_i jest dane wzorem (18). Podobnie można wyznaczyć współrzędną x_Q punktu Q określającego położenie prawej krawędzi wzorca. Współczynnik skali jest obliczany jako

$$k = \frac{a \cos \alpha}{x_Q - x_p} \simeq \frac{a}{x_Q - x_p}. \quad (20)$$

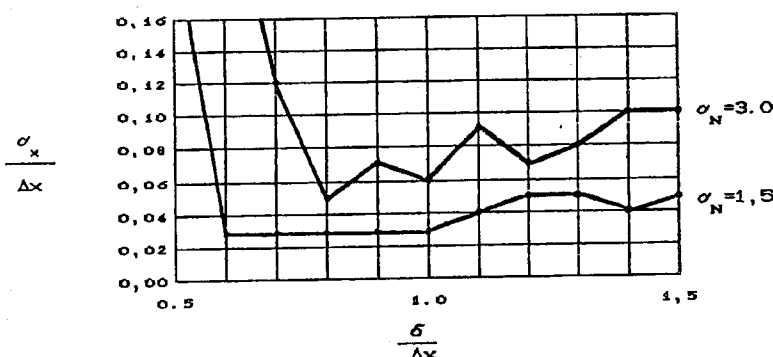
Jest oczywiste, że proponowana procedura skalowania byłaby obciążona znacznym błędem w przypadku, w którym promień δ apertury przetwornika byłby dużo mniejszy od odstepu próbkowania Δx . Wówczas „rozmycie” brzegów albo nie występowałoby w ogóle, albo jedna tylko próbka w pobliżu brzegu wzorca miałaby jasność różną od v_0 oraz v_b , a wtedy przypadkowe zakłócenia miałyby duży wpływ na lokalizację położenia brzegu obiektu. Ale w takim przypadku, jak można wykazać, częstotliwość próbkowania obrazu jest dużo mniejsza od częstotliwości wynikającej z twierdzenia Shannona [16], należy się więc liczyć z utratą informacji w wyniku próbkowania.

Dla oceny związków między rozdzielczością przetwornika fotoelektrycznego a dokładnością lokalizacji brzegów wzorca przeprowadzono symulację komputerową. Dla ustalonych wielkości $v_b = 32$, $\Delta v = 64$ oraz $x_0 + 10\Delta x - x_p = 0.5$ we wzorze (18) i kilku wartości δ z przedziału $<0,5 - 1,5> \Delta x$ obliczono wartości $2W+1=9$ kolejnych próbek sygnału wizyjnego według (18). Następnie, dla każdego ciągu 9 próbek odpowiadającego ustalonej wartości δ przeprowadzono 30 razy dopasowanie funkcji (18) do danego ciągu próbek, zaburzonego dodatkowo szumem addytywnym o rozkładzie Gaussa i ustalonej wariancji. Wyniki otrzymane dla każdego z 30 losowań szumu uśredniono, otrzymując między innymi wartość średnią $\bar{x}_p$ i odchylenie standardowe σ_x parametru x_p , w funkcji promienia δ i odchylenia standardowego szumu α_N . Pseudolosowe wartości zakłóceń o rozkładzie gaussowskim obliczano korzystając ze standardowego generatora liczb losowych i wzorów Boxa-Mullera [1], a minimalizację funkcji celu realizowano bezgradientową metodą Daviesa-Swanna-Campeya, bez ograniczeń.

Jak wynika z obliczeń, wartość średnia położenia krawędzi jest równa wartości rzeczywistej $\pm 0,02\Delta x$ dla $\delta > 0,5\Delta x$ oraz $\alpha_N \in <0,4>$, w skali 128-stopniowej. Błąd estymacji położenia dla $\delta < 0,5\Delta x$ wynosi $\pm 0,15\Delta x$, z odchyleniem średnim kwadratowym rzędu $0,22\Delta x$ co odpowiada w przybliżeniu błędom $\bar{\epsilon}$ i σ_1 , wg (11) i (13). Zgodnie z przewidywaniami zatem, proponowana metoda nie daje korzyści dla $\delta < 0,5\Delta x$. Zależność odchylenia standardowego σ_x położenia brzegu wzorca od promienia apertury δ i wielkości zakłóceń σ_N przedstawiono na rys. 4.

Na uwagę zasługuje istnienie minimum błędu σ_x w funkcji promienia apertury przetwornika. Wzrost promienia δ prowadzi do zwiększenia szerokości obszaru przejściowego ξ (por. rys. 2), w którym następuje zmiana wartości sygnału wizyjnego od wartości odpowiadającej jasności tła do jasności wzorca. Jest to obszar rozmycia brzegów obiektu (ang. blurring), według funkcji (18), w którym pochodna funkcji

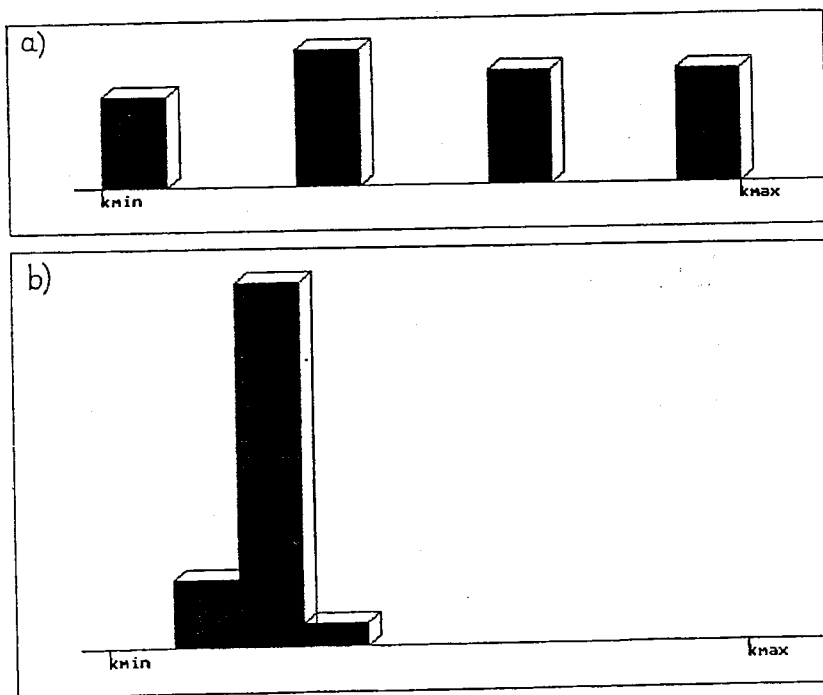
błądu (19) względem parametru x_p jest największa. Dla ustalonego odstepu próbkowania Δx , ze wzrostem δ rośnie zatem liczba tych próbek v_i , które są zlokalizowane w obszarze przejściowym. Przy ustalonej wariancji zakłóceń



Rys. 4. Odchylenie standardowe σ_x położenia brzegu wzorca w funkcji promienia apertury δ dla dwóch wartości odchylenia standardowego σ_N zakłóceń szumowych i kontraście wzorca $\Delta v = 64$

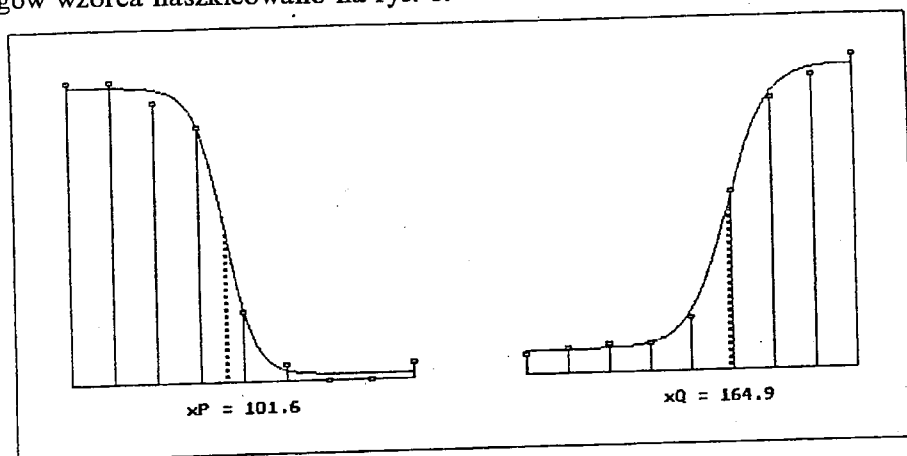
i wzrastającej liczebności próby w zakresie największej wrażliwości metody pomiarowej niepewność pomiaru maleje [21], co objawia się na rys. 4 malejącym w funkcji δ odchyleniem standardowym σ_x , dla $\delta < 0,7\Delta x$. Z drugiej strony, na mocy podanego wyżej twierdzenia, zakłócenie v_N sygnału wizyjnego w punkcie $(x,y) \in \xi$ może być interpretowane jako zakłócenie u_N położenia odległości tego punktu od brzegu obiektu. Dla dostatecznie małych wartości v_N można napisać: $u_N \approx v_N(d\varphi/du)^{-1}$. Ze wzrostem wartości δ maleje pochodna $d\varphi/du$, co prowadzi do wzrostu wariancji σ_x^2 położenia brzegu x_p przy stałej wariancji zakłóceń σ_N^2 . Wzrost ten jest widoczny na rys. 4, dla $\delta > \Delta x$. Jak wynika z wykresów przedstawionych na tym rysunku, dla typowej w praktyce wartości stosunku sygnału do szumu $\Delta v/\sigma_N \approx 40$ (32 dB), dyskutowana metoda skalowania pozwala na utrzymanie małej wartości średniokwadratowego błędu skalowania (dużo mniejszej od odstępów próbkowania Δx) z kamerą o małej rozdzielczości, z promieniem równoważnej apertury kołowej rzędu $\delta \approx \Delta x$.

Dla sprawdzenia praktycznej użyteczności proponowanej metody skalowania przeprowadzono pomiary współczynnika skali z użyciem wzorca o długości $a = 10$ cm i kontraście $\Delta v \approx 60$ jednostek w skali 128-stopniowej, w systemie o rozdzielczości 256×256 punktów z kamerą widikonową TPK 162, zmodyfikowaną dla uzyskania synchronizmu procesów odchyłania [5]. Zastępczy promień apertury kołowej (idealizowanej odpowiedzi impulsowej) przetwornika zmierzony dla tego typu kamery wynosi $\delta \approx 0,8\Delta x$, a odchylenie standardowe szumów $\sigma_N = 2,3$ w skali 128-stopniowej [12]. Histogramy współczynnika skali wyznaczonego za pomocą manipulatora kulkowego i za pomocą proponowanej metody przedstawiono odpowiednio na rys. 5a i rys. 5b. Wartość średnia współczynnika skali (dla 20 losowych położów wzorca w polu obrazu) wynosi w pierwszym przypadku 0,160 cm/ Δx , a w drugim 0,158



Rys. 5. Histogramy wyników skalowania a) „ręcznego”, b) z użyciem proponowanej metody

$cm/\Delta x$. Jednocześnie, odchylenie standardowe współczynnika skali dla skalowania „ręcznego” wynosi $0,027\text{ cm}/\Delta x$ (17%) i zostaje zmniejszone do wartości $0,003\text{ cm}/\Delta x$ (2%) za pomocą opracowanej metody. Przykładowe wykresy funkcji (18) w pobliżu brzegów wzorca naszkicowano na rys. 6.



Rys. 6. Przykładowy kształt funkcji (18) aproksymującej rozkład jasności obrazu w pobliżu brzegów wzorca długości

4. KOREKCJA ZNIEKSZTAŁCEŃ GEOMETRYCZNYCH OBRAZU

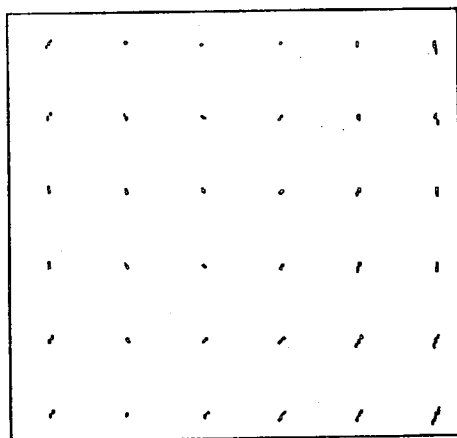
W wyniku nieliniowości układów elektronicznych i niedokładności wykonania cewek odchylających, równym odstępom chwil próbkowania sygnału wizyjnego nie odpowiadają równe odstępstwa w przestrzeni xy obrazu analogowego na wejściu przetwornika. W efekcie, wielkości dx i dy w (1) przyjmują niezerowe wartości, co oznacza istnienie zniekształceń geometrycznych obrazu. Podobne zniekształcenia mogą być wprowadzane przez układ optyczny [6], toteż zagadnienie ich korekcji jest ważne również w przypadku przetworników typu CCD.

Do identyfikacji zniekształceń wykorzystuje się obrazy wzorcowe w postaci regularnych wzorów geometrycznych. Rozważmy, dla przykładu wzorec złożony z ciemnych pasów poziomych i pionowych na jasnym tle, tworzących kratę [6]. Odległości między węzłami kraty wzorca są znane, natomiast położenie węzłów w obrazie cyfrowym może być zmierzone (np. pośrednio, za pomocą metody opisanej w poprzednim rozdziale) i użyte do obliczenia wielkości dx , dy w (1). W rezultacie pomiarów otrzymuje się K próbek funkcji $dx(x_i, y_i) = \hat{x} - x_i$ oraz K próbek funkcji $dy(x_i, y_i) = \hat{y} - y_i$ w punktach (x_i, y_i) , $i = 1, 2, \dots, K$, przy czym punkty $(\hat{x}_i, \hat{y}_i)$ są wyznaczone współrzędnymi zmierzonymi, natomiast (x_i, y_i) współrzędnymi skorygowanymi węzłów kraty.

Korekta zniekształceń geometrycznych polega na obliczeniu wartości dx i dy dla każdego punktu (x, y) obrazu skorygowanego i odpowiednim przesunięciu próbek obrazu z punktów $(\hat{x}, \hat{y})$ do położenia $(\hat{x} - dx, \hat{y} - dy)$. Na ogół wielkości dx i dy nie są liczbami całkowitymi, wobec czego wyznaczenie wartości próbek (jasności) obrazu skorygowanego wymaga dodatkowych obliczeń. Ponieważ powierzchnie $dx(x, y)$ oraz $dy(x, y)$ są dane w postaci zbioru ich K próbek, obliczenie skorygowanego położenia każdego z punktów obrazu wymaga przybliżenia wartości dx i dy funkcją ciągłą. Nakład obliczeń zależy w tym przypadku od wyboru funkcji dwóch zmiennych do opisu powierzchni korygujących. W pracy [6] zaproponowano zastosowanie do tego celu wielomianów Beziera [16]. Całkowity czas korekcji zniekształceń obrazu o wymiarach 480×512 punktów z wzorcem o 100 węzłach wymaga 4,5 min obliczeń z pomocą komputera MicroVax II [6]. Do rekonstrukcji jasności obrazu zastosowano w [6] prostą regułę „najbliższego sąsiada”, która prowadzi jednak do deformacji linii brzegowej obszarów obrazu.

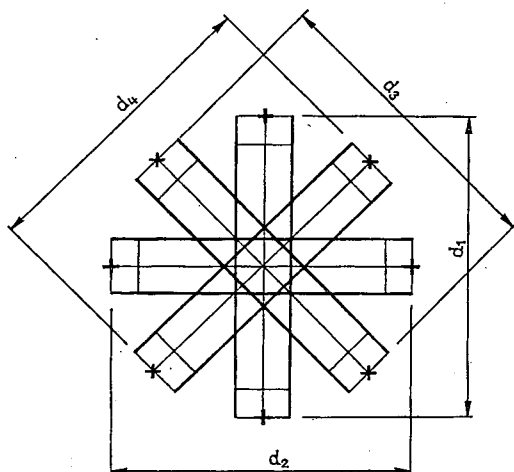
Jednym z celów niniejszej pracy jest opracowanie metody cyfrowej korekcji zniekształceń geometrycznych obrazu przy jak najkrótszym czasie obliczeń i małych zniekształceniach rozkładu jasności w pobliżu krawędzi obiektów. Do oceny wpływu zniekształceń geometrycznych na dokładność pomiarów przeprowadzono badania powtarzalności pomiaru parametrów modeli wybranych figur geometrycznych w polu obrazu. Zastosowano skalowanie za pomocą metody zaproponowanej w poprzednim rozdziale, co pozwoliło jednocześnie na ocenę jej praktycznej użyteczności.

Na rys. 7 przedstawiono wykresy wektorów $[dx, dy]$ w polu obrazu, otrzymane w wyniku identyfikacji zniekształceń geometrycznych kamery TPK 162 z obiektywem Televidon 2/10, Carl Zeiss Jena. Zmierzono je w punktach węzłowych obrazu wzorcowego, zawierającego w tym przypadku 36 kół o średnicy 2 cm, których środki



Rys. 7. Odcinki reprezentujące wektory $[dx, dy]$ zniekształceń geometrycznych dla kamery TPK 162

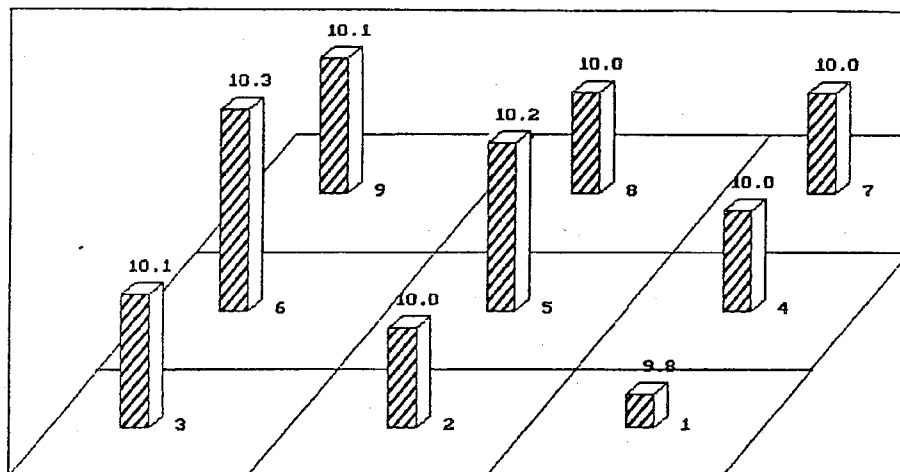
były oddalone od siebie o stałą wartość, równą 6 cm. Liczba punktów owych nie jest w przypadku badanej kamery krytyczna. Wybrana wartość $K=36$ dla obrazu o rozdzielczości 256×256 punktów jest dostatecznie duża dla odwzorowania stosunkowo gładkich zależności $dx(x,y)$ i $dy(x,y)$ i dostatecznie mała dla umożliwienia automatycznego pomiaru położenia środków kół z zadowalającą dokładnością [18]. Użycie wzorca zawierającego koła zamiast obrazu kraty upraszcza procedurę wyznaczania położenia punktów węzłowych, w tym przypadku środków kół.



Rys. 8. Obraz rozety do badania dokładności pomiarów odległości wzdłuż linii prostych

Dla oceny wpływu zilustrowanych na rys. 7 zniekształceń geometrycznych obrazu na dokładność pomiarów przeprowadzono pomiary długości obwodu i pola powierzchni oraz współczynników Fereta modelu koła średnicy 10 cm i wymiarów linowych d_1, d_2, d_3, d_4 rozety jak na rys. 8. Długość obwodu i pole powierzchni koła obliczano

metodami opracowanymi w [11,12], natomiast pomiary długości ramion rozety realizowano za pomocą manipulatora kulkowego. Obraz rejestrowany przez kamerę podzielono umownie na 9 kwadratowych części (pól), przy czym każda z figur przyjmowała 10 losowych położenia w każdym z pól, a system był skalowany za pomocą wzorca umieszczonego, również losowo, w polu środkowym. Do skalowania wykorzystano estymację średniokwadratową położenia brzegów wzorca, zaproponowaną w poprzednim rozdziale, średnica obszaru reprezentującego koło w obrazie cyfrowym wynosiła około $60\Delta x$, co odpowiada w przybliżeniu wielkości obiektów badanych w patamorfologii [5].



Rys. 9. Rozkład wartości średniej długości d_3 ramion rozety (bez korekcji zniekształceń geometrycznych)

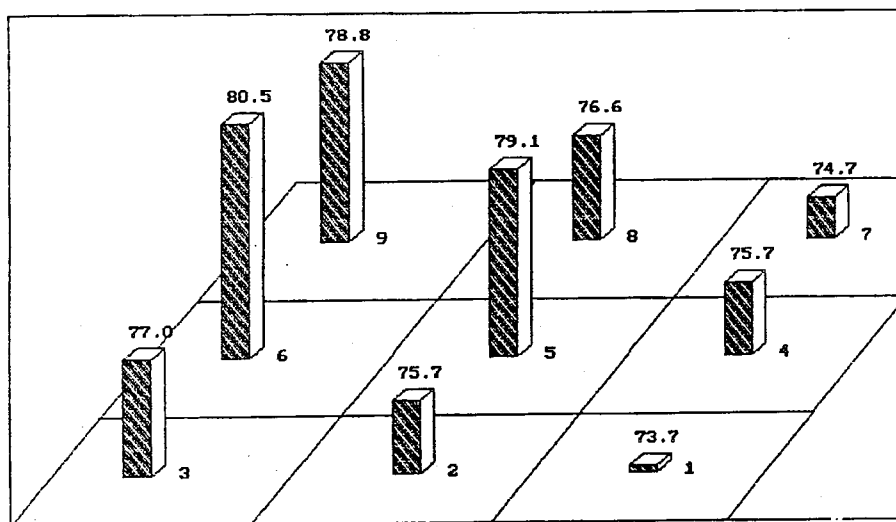
Rozrzut wyników pomiaru zilustrowano na wykresach długości d_3 i pola powierzchni A , przedstawionych odpowiednio na rys. 9 i rys. 10. Sumaryczne wyniki pomiarów dla obrazu bez korekcji zniekształceń geometrycznych zestawiono w Tabelicy 1. Symbolem $\bar{\cdot}$ oznaczono średnią arytmetyczną 90 niezależnych pomiarów dla każdego z rozważanych wielkości; s jest odchyleniem średnim kwadratowym, natomiast z — współczynnikiem zmienności, czyli stosunkiem $s/\bar{\cdot}$. Przez min i max oznaczono najmniejszą i największą wartość danej wielkości.

W monografii [17] zaproponowano zastosowanie aproksymacji średniokwadratowej wielomianem dwóch zmiennych x, y do opisu zależności $dx(x, y)$ oraz $dy(x, y)$:

$$dx(x, y) = a_0 + a_1x + a_2y + a_3x^2 + a_4xy + a_5y^2 \quad (21)$$

$$dy(x, y) = b_0 + b_1x + b_2y + b_3x^2 + b_4xy + b_5y^2 \quad (22)$$

W niniejszej pracy, do identyfikacji współczynników $\{a_i, i=0, \dots, 5\}$, $\{b_i, i=0, \dots, 5\}$ dla $K=36$ wykorzystano gradientową metodę minimalizacji bez ograniczeń, Fletchera-Powella. Czas obliczeń za pomocą komputera PC AT z zegarem 16 MHz wynosi w tym przypadku 3 s. Do obliczenia wartości próbek obrazu skorygowanego wykorzystano interpolację dwuliniową [16,20] w zbiorze czterech sąsiadujących ze



Rys. 10. Rozkład wartości średniej pola powierzchni A koła (bez korekcy zniekształceń geometrycznych)

Tabela 1

Wyniki pomiarów bez korekcy geometrycznej

<i>i</i>	<i>d1</i>	<i>d2</i>	<i>d3</i>	<i>d4</i>	<i>A</i>	<i>P</i>	<i>Fx</i>	<i>Fy</i>
<i>w</i>	10	9.9	10.1	9.8	76.9	31.2	9.8	9.7
<i>s</i>	0.2	0.2	0.2	0.1	2.2	0.5	0.2	0.2
<i>z</i> [%]	2.0	2.1	1.7	1.5	2.8	1.5	2.2	2.0
<i>w</i> min	9.6	9.4	9.7	9.3	72.5	30.2	9.4	9.3
<i>w</i> max	10.5	10.4	10.5	10.1	81.1	32.2	10.1	10.1

sobą punktów obrazu (rys. 11). Według tej metody, estymowana jasność próbki $\hat{z}$ położonej między czterema próbkami z_1, z_2, z_3, z_4 w unormowanej odległości α od pary $\{z_1, z_4\}$ i w odległości β od pary $\{z_1, z_2\}$, $0 \leq \alpha \leq 1$, $0 \leq \beta \leq 1$, jest dana wzorem

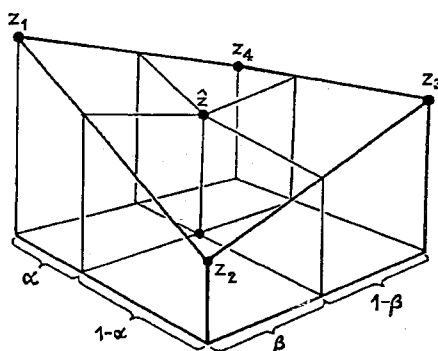
$$\hat{z} = z_1 \alpha \beta + z_2 (1 - \alpha) \beta + z_3 (1 - \alpha) (1 - \beta) + z_4 \alpha (1 - \beta). \quad (23)$$

Czas obliczeń dających w wyniku korekcję zniekształceń geometrycznych za pomocą wielomianu stopnia 2 (21), (22) z interpolacją jasności próbek obrazu (23) wynosi 4 min, dla komputera PC AT z zegarem 16 MHz i obrazu o wymiarach 256×256 punktów. Korekcja taka daje widoczne zmniejszenie współczynnika zmienności każdego z badanych parametrów, co przedstawiono w Tablicy 2.

Dalsze zmniejszenie rozrzutu wyników pomiarów otrzymano z zastosowaniem wielomianów stopnia trzeciego:

$$dx(x, y) = a_0 + a_1 x + a_2 y + a_3 x^2 + a_4 xy + a_5 y^2 + a_6 x^3 + a_7 x^2 y + a_8 xy^2 + a_9 y^3 \quad (24)$$

$$dy(x, y) = b_0 + b_1x + b_2y + b_3x^2 + b_4xy + b_5y^2 + b_6x^3 + b_7x^2y + b_8xy^2 + b_9y^3 \quad (25)$$



Rys. 11. Ilustracja metody interpolacji dwuliniowej

Tabela 2

Wyniki pomiarów z korekcją wielomianami stopnia 2

<i>i</i>	<i>d1</i>	<i>d2</i>	<i>d3</i>	<i>d4</i>	<i>A</i>	<i>P</i>	<i>Fx</i>	<i>Fy</i>
<i>w</i>	9.9	9.9	9.9	9.9	76.0	30.9	9.8	9.6
<i>s</i>	0.2	0.1	0.1	0.1	1.3	0.3	0.1	0.1
<i>z</i> [%]	1.7	1.2	1.3	1.0	1.7	1.0	0.9	1.4
<i>w</i> min	9.6	9.6	9.7	9.7	73.1	30.4	9.6	9.4
<i>w</i> max	10.2	10.2	10.2	10.2	78.7	32.0	9.9	9.9

W miejsce funkcji (21), (22). Liczba współczynników wzrasta tu do dwudziestu, a czas ich wyznaczania wynosi odpowiednio 15 s. Wzrasta również czas korekcji zniekształceń obrazu, który wynosi dla wielomianu stopnia trzeciego 6 min. Wyniki pomiarów otrzymane dla tego przypadku przedstawiono w Tablicy 3. Warto zwrócić uwagę na ponad dwukrotne zmniejszenie współczynnika zmienności dla długości obwodu *P* i pola powierzchni koła *A*.

W celu dalszego skrócenia czasu obliczeń, do wyznaczenia wartości $dx(x, y)$ oraz $dy(x, y)$ zastosowano interpolację (23) zamiast aproksymacji wielomianu. Wartości

Tabela 3

Wyniki pomiarów z korekcją wielomianami stopnia 3

<i>i</i>	<i>d1</i>	<i>d2</i>	<i>d3</i>	<i>d4</i>	<i>A</i>	<i>P</i>	<i>Fx</i>	<i>Fy</i>
<i>w</i>	9.9	10.0	9.9	9.9	77.4	31.2	9.8	9.7
<i>s</i>	0.1	0.1	0.1	0.1	0.9	0.2	0.1	0.1
<i>z</i> [%]	1.3	1.4	1.5	1.2	1.2	0.6	1.1	1.0
<i>w</i> min	9.6	9.5	9.6	9.7	75.5	30.8	9.5	9.5
<i>w</i> max	10.1	10.3	10.2	10.3	79.2	31.6	10.1	9.8

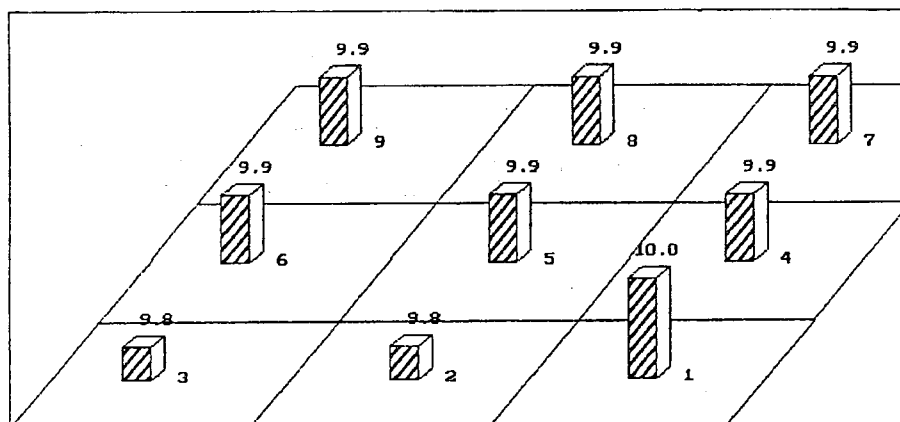
przesunąć dx , dy w przestrzeni między węzłami obliczano z pomocą (23) przy odpowiednim podstawieniu zmierzonych K wartości węzłowych dx i dy pod $z_i, i=1,2,3,4$. Wartości dx , dy na powierzchni obrazu nie objętej węzłami interpolacji (na brzegach obrazu) wyznaczono metodą ekstrapolacji, utrzymując stałe wartości dx , dy w kierunku od węzłów do brzegu obrazu. Korekcja zniekształceń geometrycznych za pomocą interpolacji dwuliniowej (z interpolacją dwuliniową jasności próbek obrazu jak w poprzednich przypadkach wielomianów stopnia 2 i 3) jest realizowana w znacznie krótszym czasie (około 50 s), przy dokładności pomiarów równie dobrej jak w przypadku korekcji wielomianem stopnia 3. Wyniki pomiarów obrazu skorygowanego za pomocą interpolacji dwuliniowej zestawiono w Tablicy 4. Na rys. 12 i rys. 13 przedstawiono wykresy ilustrujące rozrzut wartości średniej odpowiednio długości d_3 i pola powierzchni A , dla przypadku korekcji za pomocą interpolacji dwuliniowej.

Tabela 4

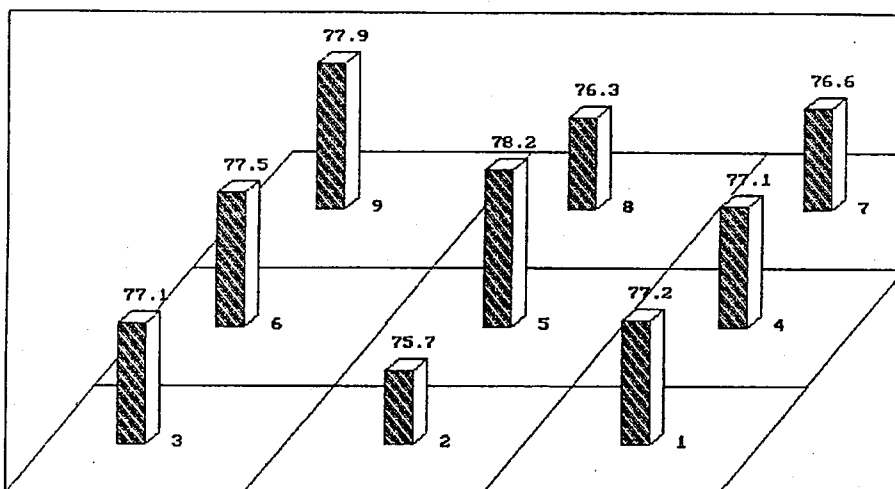
Wyniki pomiarów z korekcją przy użyciu interpolacji dwuliniowej

i	$d1$	$d2$	$d3$	$d4$	A	P	F_x	F_y
w	9.9	10.0	9.9	9.9	77.1	31.2	9.8	9.6
s	0.1	0.2	0.1	0.1	0.8	0.2	0.1	0.1
z [%]	1.4	1.5	1.1	1.2	1.1	0.6	1.2	1.4
w min	9.6	9.6	9.6	9.6	74.8	30.7	9.6	9.3
w max	10.3	10.3	10.2	10.2	78.6	31.5	10.1	10.0

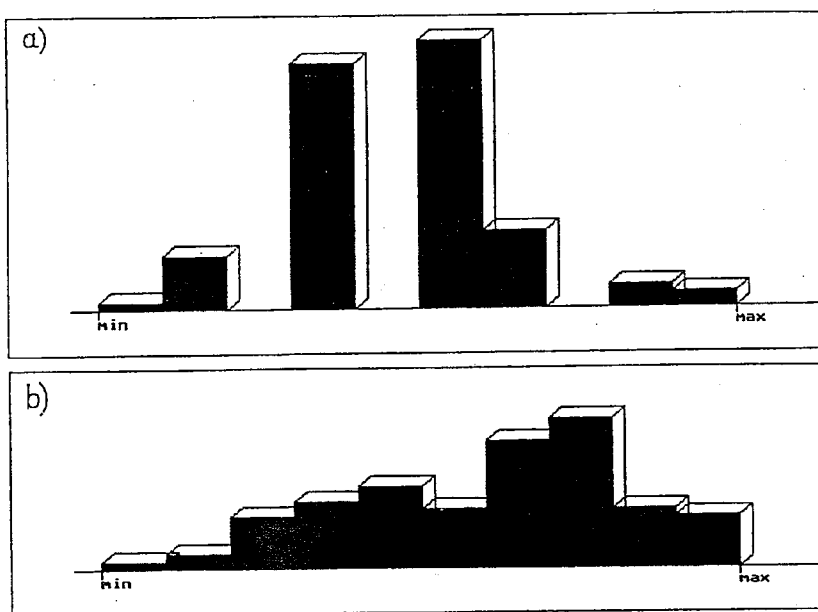
Na rys. 14a wykreślono histogram długości d_3 , mierzonej przez operatora za pomocą manipulatora kulkowego. Histogram ten wykazuje charakterystyczną nieciągłość niepustych przedziałów mierzonej wielkości. W tym przypadku, jak również w odniesieniu do wielkości d_1 , d_2 , d_4 , F_x i F_y niepewność pomiaru jest wyznaczana



Rys. 12. Rozkład wartości średniej długości d_3 ramion rozety (z korekcją zniekształceń metodą interpolacji dwuliniowej)



Rys. 13. Rozkład wartości średniej pola powierzchni A koła (z korekcją zniekształceń metodą interpolacji dwuliniowej)



Rys. 14. Histogramy a) długości d_3 ramion rozety b) pola powierzchni A koła (z korekcją zniekształceń geometrycznych metodą interpolacji dwuliniowej)

odstępem próbkowania obrazu Δx . Odmienne właściwości wykazuje histogram długości obwodu koła przedstawiony na rys. 14b. Gęstość prawdopodobieństwa tej wielkości, uzyskana z pomiarów, ma rozkład zbliżony do normalnego, dla dostatecznie dużego stosunku promienia koła do dostępu próbkowania, który w rozważanym

przypadku jest większy od 30. Podobne własności wykazują pomiary pola powierzchni A modelu koła.

PODSUMOWANIE

W artykule rozważono zagadnienie zwiększenia dokładności pomiaru parametrów figur geometrycznych za pomocą systemu cyfrowego przetwarzania obrazów. Przedstawiono twierdzenie o warunkach jednoznacznego odwzorowania odległości punktu od oświetlonego obiektu na wartość sygnału wizyjnego [12]. Korzystając z tego twierdzenia zaproponowano metodę skalowania systemu poprzez dopasowanie funkcji rozmycia brzegów wzorca długości do ciągu próbek obrazu w pobliżu krawędzi wzorca. Dowiedziono, za pomocą symulacji komputerowej i pomiarów w systemie z kamerą widikonową, że możliwe jest skalowanie systemu obrazowego z dokładnością o rząd wielkości lepszą od wyznaczonej odstępami próbkowania obrazu, dla typowego poziomu zakłóceń szumowych (subpixel accuracy). Opracowaną metodę skalowania zastosowano do wyznaczenia próbek funkcji korygujących $dx(x,y)$, $dy(x,y)$, przy korekcji zniekształceń geometrycznych obrazu. Zaproponowano wykorzystanie interpolacji dwuliniowej do odwzorowania funkcji korygujących i wykazano, że skuteczność takiej korekcji jest równie dobra (w sensie redukcji rozrzutu wyników pomiaru parametrów modeli wybranych figur geometrycznych) jak w przypadku opisu funkcji korygujących za pomocą wielomianu stopnia trzeciego, przy kilkukrotnie zmniejszonym czasie obliczeń. Uzyskane wyniki przedstawionej pracy mogą mieć zastosowanie w badaniach naukowych, diagnostyce medycznej, laboratoriach przemysłowych i robotyce — w tych działach gospodarki, w których wykorzystuje się ilościową analizę obrazów.

Planowane są dalsze prace, obejmujące badanie własności metrologicznych systemu przetwarzania obrazów z kamerą CCD.

Pracę wykonano częściowo w ramach CPBP 02.14 i CPBR 8.6

BIBLIOGRAFIA

1. G. Dalquist, A. Björck: *Metody numeryczne*. PWN, Warszawa 1983
2. G. M. Fichtenholz: *Rachunek różniczkowy i całkowy*, PWN, Warszawa, 1985
3. R. E. Flory: *Image Acquisition Technology*. Proc. IEEE, 73, 1985, 613–633
4. J. Gielecki: *Rozwój koła tętniczego mózgu u płodów bydła*. Rozprawa doktorska, Akademia Medyczna, Bydgoszcz, 1988
5. J. Gociłowski, A. Materka, M. Pałuba, A. Sitek, J. Stetkiewicz, P. Strumiłło, W. Zatorski: *Mikrokomputerowy system do przetwarzania obrazów medycznych dla komputerów IBM PC*, Archiwum Informatyki Teoretycznej i Stosowanej, 1/89
6. A. Goshtasby: *Correction of Image Deformation from Lens Distortion Using Bezier Patches*. Computer Vision, Graphics, and Image Processing, 47, 1989, 385–394

7. G. S. H o b s o n: *Przyrządy z przenoszeniem ładunku*. WNT, Warszawa, 1981
8. B. R. H u n t, J. R. B r e e d l o v e: *Scan and Display Consideration in Processing Images by Digital Computer*. IEEE Trans. Computers, 1975, 848–853
9. M. D. L e v i n e: *Vision in Man and Machine*, McGraw-Hill, New York, 1985
10. P. H. L i n d s a y, D. A. N o r m a n: *Przetwarzanie informacji u człowieka*. PWN, Warszawa, 1984
11. A. M a t e r k a: *Digital Image Analysis in Biometric Research of Blood–Vessel Systems: Selected Methods*, wykład zaproszony na seminarium „Computer Processing of Biological Images, Mathematical Methods, Systems, and Applications”. Międzynarodowe Centrum Biocybernetyki PAN, Mądralin, listopad 1989
12. A. M a t e r k a: *On Measurements of Geometric Figures Parameters Using Digital Image Processing System*. artykuł złożony w Computer Vision, Graphics, and Image Processing
13. A. M a t e r k a, P. S t r u m i ł o: *Wpływ wybranych zniekształceń sygnału wizyjnego na błąd pomiaru pola powierzchni figur płaskich w cyfrowych systemach analizy obrazu*. XII Konf. Teoria Obwodów i Układy Elektroniczne, Myczkowce 1989, 295–300
14. A. V. O p p e n h e i m, R. W. S c h a f e r: *Cyfrowe przetwarzanie sygnałów*, WKŁ, Warszawa, 1979
15. T. N. P a p p a s, J. S. L i m: *A New Method for Estimation of Coronary Artery Dimensions in Angiograms*. IEEE Trans. Acoust. Speech Signal Process. ASSP–36, 1988, 1501–1513
16. T. P a v l i d i s: *Grafika i przetwarzanie obrazów*. WNT, Warszawa, 1987
17. W. K. P r e t t: *Cifrowa obróbka obrazów*. Mir, Moskwa, 1982
18. S. M. T h o m a s, Y. T. C h a n: *A Simple Approach for the Estimation of Circular Arc Center and its Radius*. Computer Vision, Graphics, and Image Processing, 45, 1989, 362–370
19. J. W a ł c z y k: *Podstawy systemów telewizji użytkowej*. WKŁ, Warszawa, 1982
20. L. P. Y a r o s l a v s k y: *Digital Picture Processing*. Springer, Berlin 1979
21. H. S z y d ł o w s k i (red.): *Teoria pomiarów*. PWN, Warszawa, 1981

A. MATERKA, W. GÓRNICIEWICZ, M. KUKUŁA

CALIBRATION AND GEOMETRIC DISTORTION CORRECTION IN A DIGITAL IMAGE PROCESSING SYSTEM WITH A VIDICON CAMERA

S u m m a r y

Methods for calibration and geometric distortion correction in a digital image processing system with a vidicon-type camera are discussed in the paper. Based on a theorem about conditions of unique transformation between video signal value and point-to-object distance, a method for system calibration with a subpixel accuracy is proposed. The method performance is verified using computer simulation and measurements in a physical system. Three methods of image geometric distortion correction are investigated using different mathematical functions. It is shown, that bilinear interpolation gives 6 times shorter computational time, as compared to least-squares approximation using polynomials of order 3, with similar increase of geometric objects measurement accuracy in the image field. The results can be applied to quantitative biomedical investigations and industrial quality control, also in robotics.

Digital 2-D picture restoration and enhancement

NGUYEN KIM SACH

Instytut Radioelektroniki, Politechnika Warszawska

Received 1989.10.12

Authorized 1989.12.27

This paper presents operations of 2-D noisy picture restoration, by combination of histogram with median filter, and operations of picture enhancement, with gradient module (8), 3-time predictor (9), Laplace operator (10) and Robert operator (11). Combination of operators of picture restoration with picture enhancement is discussed. Experimental results, their mean-square error and their application are presented.

1. INTRODUCTION

Noise and interferences are very important in television. Noise is effect stochastic electric fluctuations, with are formed by granular structure of electronic elements or electric charges. There are numerous types of noise, e.g. thermal noise, shot noise, impulse noise,..., which are described in [4], [5], [9]. For the 2-D picture restoration linear or nonlinear filters are used. Nonlinear filters, e.g. rank filter and median filter, have high possibility of picture restoration [6], [7]. Digital 2-D picture restoration is estimation process of original picture. In opposite to it, 2-D picture enhancement has an aim to improve picture edges in suitable forms for direct observation or machine analysis. Picture enhancement is not the same as picture restoration. Observation of enhancing pictures is more comfortable than observation of original pictures. Fine details of pictures are described by high frequency of their spatial-frequency bands. High frequency band filters are useful to enhance picture edges. Different techniques to enhance picture edges have been well established, for instance, contrast regulation, variation of statistic distribution of gray levels in picture, edge detection [4], [5], [9]. Particularly for simultaneous picture restoration and picture enhancement, combination of picture restoration operators with picture enhancement operators can be used.

2. OPERATION OF 2-D NOISY PICTURE RESTORATION

By 2-D discrete picture restoration, process of Fourier transformation must be realized so, that discrete filtering is presented correctly by convolution (1):

$$\hat{g}(x, y) = g_2(x, y) * h(x, y) \quad (1)$$

where $\hat{g}_1(x, y)$ — estimation of continuous object $g_1(x, y)$
 $g_2(x, y)$ — continuous picture (observation)
 $h(x, y)$ — point spread function
 $*$ — operator of convolution.

and impulse response of continuous filter is modeled by discrete impulse response form of picture restoration filter [5]. Median filter (as a special case of rank filter) is the way of nonlinear signal processing, which is used reduce noise in picture without considerable interference and destruction of edges [2], [3]. For that purpose is used 2-D window of definite forms (e.g. square, rectangle, etc.), which shifted at the picture. Central pixel of the window is represented in median value of window pixel. Median filter removes pixels, whose values are considerably deviated away from remaining pixels in the surroundings. By this way all intensities of greatest values characterized by noise or interference can be removed. In pictures subsized by high numbers of details, median filter small window, e.g. 3×3 pixel window, is advantageous. In opposite, for the picture subsized by a small number of details of similar intensity, window dimension should be increased, to such an extent as 5×5 -pixel window or more. For median filter with various constant or variable dimensions, cascade may be used for picture restoration. The median filter reduces the impulse noise more efficiently than smoothing noise and can preserve considerably the edges [1], [2]. The 2-D median filter with window B in picture A is defined by (2)

$$y_{ij} = \underset{B}{med} \left(x_{ij} \right) \triangleq \underset{B}{med} \left[x_{i+p, j+q}; (p, q) \in B \right], \quad (2)$$

if the 2-D picture is described by the sequence $(x_{ij}; (i, j) \in N^2)$,

where $i, j = 1, 2, 3, \dots, N$

p, q — line and column element of window B,

x_{ij} — intensity value of pixel of input picture,

y_{ij} — intensity value of pixel of output picture.

Edge is a subpicture of the form of a line, by which all pixels at one side of which have common value l , and pixels at other side of which have the common value h (l and h are situated in definite partition of values). Difference between l and h determines the edge height (3)

$$d = l - h. \quad (3)$$

If the ensemble of values of pixels (of window B) is symmetrical to the point $(0, 0)$, i.e.

$$(l, h) \in B, \quad (-l, -h) \in B \quad \text{and} \quad (0, 0) \in B \quad (4)$$

then the median filter

$$y_{ij} = \text{med}[x_{i+l, j+h}, (l, h) \in B] \quad (5)$$

$$(i, j) \in N^2$$

will preserve the edge [2]. If we treat the edge K as difference of values between two pixels at all directions, conditions for edge preservation can be described by (6):

$$(l, h) \in B, \quad (l, h) \in K \quad \text{and} \quad d = l - h \neq 0 \quad (6)$$

3. DIGITAL PICTURE ENHANCEMENT

For the edge preservation and enhancement, order statistic filters (OS-filters), comparison and selection filters (CS-filters) and median filters can be used [8]. Median filters can preserve, but can not enhance edges. OS-filter is a modification of median filter. It can enhance picture edge and reduce impulse noise and some parts of additive white noise. On the contrary, CS-filter can reduce impulse noise and some parts of additive nonimpulse white noise (Gaussian) by enhancing edge. CS-filter can be described by (7):

$$y_k = \begin{cases} x_k^{(N+1-J)}, & \text{if } \mu_k \geq M_k \\ x_k^{(N+1+J)}, & \text{otherwise} \end{cases} \quad (7)$$

where

y_k — output of CS-filter

$x_k^{(i)}$ — i -th smallest sample in the window

$x_{k-N}, \dots, x_{k+N}$ — input values

$2N+1$ — dimension of window

J — parameter of repeat number k of filtering

μ_k, M_k — mean values and median of sample.

If $J=N$, then CS-filter receives min. or max. value of intensity in the window, and $y_k = x_k$. For $J=0$, CS-filter becomes median filter, and $y_k = x_k^{N+1}$. Because CS-filter is nonlinear, we can't describe it in common form. It can reduce well impulse and nonimpulse noise, if J is small, e.g. $J=1$. If J is large, CS-filters causes a good enhancement of edge. For edge enhancement different operators are used, e.g. Laplace's operator, Robert's operator [2], [8]. In our experiment the following operators for edge enhancement are used: gradient module (10), 3-time predictor (9), Laplace's operator (10) and Robert's operator (11), Fig. 1 [5]:

a. gradient module

$$y_{ijG} = |x_{ij} - x_{i+1, j+1}| + |x_{ij} - x_{i-1, j-1}| + |x_{ij} - x_{i+1, j-1}| + |x_{ij} - x_{i-1, j+1}| \quad (8)$$

b. 3-time predictor

$$y_{ij3} = (x_{i+1, j} - x_{ij}) + (x_{i+2, j} - x_{ij}) + (x_{i+3, j} - x_{ij}) \quad (9)$$

c. Laplace's operator

$$y_{ijL} = (x_{i+1,j} + x_{i-1,j} + x_{i,j+1} + x_{i,j-1}) - 4x_{ij} \quad (10)$$

d. Robert's operator

$$y_{ijR} = |x_{ij} - x_{i+1,j+1}| + |x_{i+1,j} - x_{i,j+1}| \quad (11)$$

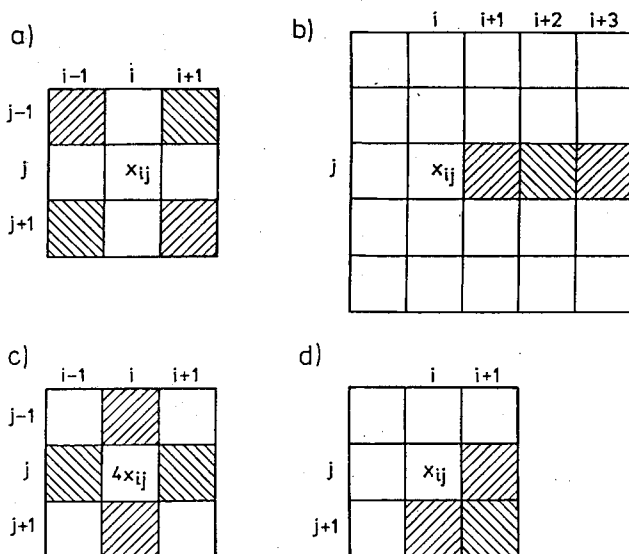


Fig. 1. Description of edge enhancement in the windows

4. EXPERIMENT OF PICTURE RESTORATION AND ENHANCEMENT

In the present experiment we use a combination of median filter (devised by histogram technique) and operators of edge enhancement in (8)...(11). At the beginning, the program (represented by the language C) realizes operation of picture restoration, and then operation of edge enhancement. Through these operations different types of noise are treated, viz. noise of uniform distribution, Poisson distribution and impulse noise, and different types of test pictures, e.g. pictures of small and large details and TV test pictures. Results are satisfactory at different levels. In Fig. 2, some results of operations of median filter and gradient module (a), 3-time predictor (b), Laplace's operator (c) and Robert's operator (d) for noise of uniform distribution and gaussian distribution are presented.

Table presents computing values of mean-square errors in the Fig. 2.

10% noise



original picture

a.



b.



c.



d.



Fig. 2. Some results of realizing operations for picture restoration and enhancement in 10% noise of Gaussian distribution

T a b l e
Mean-square error NMSE for 256×256 pixel picture
restoration and enhancement

No.	picture in fig. 2	NMSE	notes
1	a	0.016	noise of the Gaussian distribution
2	b	0.022	
3	c	0.018	
4	d	0.017	

CONCLUSION

For the aim of restoration of noisy pictures, the median filter with combination of histogram techniques is used [7]. While for edge enhancement, the gradient module and 3-time predictor, which are adaptive median filters, Laplace's operator and Robert's operator are useful (point 3). Combining the operator of picture restoration and operators of edge enhancement noise in noisy pictures can be reduced and at the same time edges of picture can be corrected (point. 4). Experimental results are satisfactory.

Another direction of this work is development of presented methods considering their application possibility to TV picture restoration in the real time by partitioning each picture to subpictures and parallel subpictures processed by means of 64×64 -pixel subpicture processing modules [5]. Digital 2-d picture restoration and enhancement can be applied to TV purposes at the TV centers and at the relay TV stations or at the satellite TV stations.

REFERENCES

1. W. K. P r a t t: *Digital image processing*. New York, 1978
2. T. S. H u a n g: *Two-dimensional digital signal processing*. Springer Verlag, Berlin 1981
3. G. S t a n k e, W. O s t e n: *Rangordnungsoperatoren auf einem parallelen Bildprozessor*, Bild und Ton, No 3/1989, p. 84—88
4. N. K. S a c h: *Szumy i odtwarzanie obrazów telewizyjnych metodami cyfrowymi*, Przegląd techniki. Radio i Telewizja, No 2/1989, p. 12—22
5. N. K. S a c h: *Poprawa jakości obrazów telewizyjnych metodami cyfrowymi*. (Monografie, Technical University Warsaw)
6. N. K. S a c h: *Zmniejszanie zakłóceń obrazów za pomocą filtru medianowego*. Rozprawy Elektrotechniczne, No 3/1990 (on printing)
7. N. K. S a c h: *Odtwarzanie cyfrowe obrazów przy pomocy modyfikującego filtru medianowego*. Krajowe Sympozjum telekomunikacji KST'89, Bydgoszcz, Poland. 6—8 August 1989, Tom D, p. 170—182
8. Y. H. L e e, A. T. F a m: *An edge gradient enhancement adaptive order statistic filter*. IEEE Trans. on acoustics, speech, and signal processing vol. ASSP—35, May, No. 5/1987, p. 680—695
9. N. K. S a c h: *Cyfrowe odtwarzanie zasumionych obrazów telewizyjnych*. Rozprawy Elektroniczne, 1990, No. 3—4

N. KIM SACH

**CYFROWE ODTWARZANIE I UWYDATNIENIE KRAWĘDZI DWUWYMIAROWYCH
OBRAZÓW ZASZUMIONYCH****S t r e s z c z e n i e**

W artykule omówiono operację odtwarzania dwuwymiarowych obrazów zaszumionych (tj. kombinacja histogramu i filtru medianowego) i operację uwydatnienia krawędzi obrazu (tj. moduł gradientowy (8), predyktor trzykrokowy (9), operator Laplace'a (10) i operator Roberta (11)). Opisano połączenie operacji odtwarzania i uwydatnienia krawędzi obrazu. Zaprezentowano wyniki przeprowadzonych eksperymentów, wartości błędów średnio-kwadratowych oraz ich zastosowania.

Kwantyzery adaptacyjne w modulacji DPCM

KINDA ABOU-KASSEM, PRZEMYSŁAW DYMARSKI

Instytut Telekomunikacji, Politechnika Warszawska

Otrzymano 1989.11.22

Autoryzowano do druku 1990.07.12

Opisano metodę projektowania kwantyzatorów równomiernych i nierównomiernych z adaptacją dla zadanych właściwości statystycznych kwantowanego sygnału, liczby poziomów kwantyzacji i szybkości adaptacji. Zaprojektowane kwantyzery przetestowano metodą symulacji w układzie DPCM z adaptacyjnym predyktorem. Symulacje potwierdziły prawidłowość proponowanej metody obliczania kwantyzatorów i pozwoliły na optymalizację niektórych parametrów, przede wszystkim szybkości adaptacji

1. WSTĘP

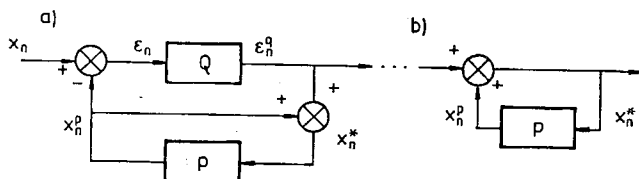
W kodowaniu i transmisji cyfrowej sygnału mowy do niedawna dominowały systemy oparte na próbkowaniu i kwantyzacji (modulacja PCM). Obecnie coraz większe znaczenie zyskują systemy oparte na próbkowaniu, kwantyzacji i predykcji (modulacja DPCM), umożliwiające zmniejszenie przepływności binarnej przy zachowaniu dobrej jakości sygnału mowy. Przykładem jest system DPCM dla transmisji sygnału telefonicznego kanałem o przepływności 32 kbit/s [9].

W systemie DPCM stosuje się kwantyzery adaptacyjne, równomierne i nierównomierne. Algorytm projektowania kwantyzatorów równomiernych z adaptacją, przy założeniu że kwantowany sygnał ma charakter gaussowski, podał Jayant [1]. Przyjmując bliższy rzeczywistości rozkład prawdopodobieństwa próbek sygnału mowy (np. rozkład Laplace'a), otrzymuje się niegaussowski rozkład kwantowanego sygnału [10]. Wymaga to przeprowadzenia modyfikacji metody Jayanta projektowania kwantyzatorów, co zostało opisane w p. 3. Proponowany algorytm nadaje się również do projektowania kwantyzatorów nierównomiernych przy dowolnej szybkości adaptacji. Wykorzystano go do zaprojektowania szeregu kwantyzatorów współpracujących z układami DPCM o szybkości transmisji 16, 24 i 32 kbit/s (p. 4).

2. OPIS UKŁADU DPCM

2.1 SCHEMAT UKŁADU I JEGO WŁAŚCIWOŚCI

Ogólny schemat blokowy kodera i dekodera DPCM pokazano na rys. 1.



Rys. 1. Układ kodera (a) i dekodera (b) systemu DPCM: Q — kwantyzator, P — predyktor

Kwantowaniu podlegają próbki sygnału różnicowego ε_n , powstającego w wyniku odjęcia sygnału aproksymującego (sygnału predykcji) x_n^p od sygnału mowy x_n (n jest numerem próbki). Poza szczególnymi przypadkami (modulacja delta [8]) próbkowanie przebiega z częstotliwością 8000 próbek/s. Sygnał aproksymujący x_n^p powstaje w predyktorze P, w wyniku cyfrowego przetwarzania poprzednich próbek sygnału wyjściowego x_{n-1}^* , x_{n-2}^* , ..., x_{n-p}^* . Przy braku przekłamań transmisji w koderze i dekodrze generowany jest ten sam sygnał predykcji x_n^p — wówczas spełnione jest równanie:

$$x_n^* - x_n = (\varepsilon_n^q + x_n^p) - (\varepsilon_n + x_n^p - \varepsilon_n + x_n^p) = \varepsilon_n^q - \varepsilon_n = e_n \quad (1)$$

Oznacza to, że sygnał wyjściowy dekodera różni się od sygnału wejściowego kodera jedynie o błąd kwantyzacji $e_n = \varepsilon_n^q - \varepsilon_n$.

Miarą jakości sygnału wyjściowego $x_n^* = x_n + e_n$ jest stosunek mocy składowej użytecznej x_n do mocy zniekształceń e_n :

$$\text{SNR} = \frac{(\sigma_x)^2}{(\sigma_e)^2} = \frac{(\sigma_x)^2}{(\sigma_\varepsilon)^2} \cdot \frac{(\sigma_\varepsilon)^2}{(\sigma_e)^2} = G_p \cdot \text{SNR}_q, \quad (2)$$

gdzie: $G_p = \left(\frac{\sigma_x}{\sigma_\varepsilon}\right)^2$ — zysk predykcji równy stosunkowi mocy sygnału mowy do mocy sygnału różnicowego, czyli błędu predykcji,

$\text{SNR}_q = \left(\frac{\sigma_\varepsilon}{\sigma_e}\right)^2$ — stosunek mocy sygnału kwantowanego do mocy szumu kwantowania.

Dla uzyskania możliwie najlepszej jakości sygnału wyjściowego, predyktor projektuje się na maksimum G_p , a kwantyzator — na maksimum SNR_q .

2.2 PREDYKTOR

Zadaniem predyktora jest wytwarzanie sygnału aproksymującego x_n^p na podstawie poprzednich próbek sygnału wyjściowego x_{n-1}^* , ..., x_{n-p}^* w taki sposób, aby moc sygnału różnicowego ε_n (błędu predykcji) była jak najmniejsza. Predyktory

działają najczęściej na zasadzie przekształceń liniowych — najprostszy jest predyktor transwersalny opisany równaniem

$$x_n^p = \sum_{i=1}^p a_i \cdot x_{n-i}^* = \bar{a}^T \bar{x}(n), \quad (3)$$

gdzie $\bar{a}^T = [a_1, a_2, \dots, a_p]$ — transponowany wektor współczynników predykcji, p — rząd predyktora, $\bar{x}(n) = [x_{n-1}^*, x_{n-2}^*, \dots, x_{n-p}^*]^T$ — wektor poprzednich próbek sygnału wyjściowego.

Przy stosowaniu predyktora o stałych współczynnikach można uzyskać zysk predykcji G_p rzędu 6 dB; większy zysk predykcji, rzędu kilkunastu dB, można uzyskać stosując adaptację predyktora. Metody adaptacji predyktora można podzielić na dwie grupy: są to metody adaptacji „w przód” (zwane blokowymi) i metody adaptacji „w tył” (zwane sekwencyjnymi). Adaptacja „w przód” polega na zgromadzeniu pewnego bloku próbek sygnału mowy (obejmującego 10–20 ms), wyznaczeniu optymalnych współczynników predykcji (najczęściej metodą Levinsona-Durbina [5]), po czym następuje zakodowanie całego bloku w układzie DPCM.

Wadą tych metod adaptacji jest konieczność transmitowania do odbiornika, co 10–20 ms, całego wektora współczynników predykcji. Wady tej nie mają metody sekwencyjne (adaptacja „w tył” [7]), w których predyktor jest wyznaczany jedynie w oparciu o transmitowany do odbiornika, skwantowany sygnał błędu predykcji ε_n^q .

W koderze i dekoderze przechowuje się aktualną wartość wektora współczynników predykcji $\bar{a}(n)$, którą modyfikuje się z częstotliwością próbkowania:

$$\bar{a}(n+1) = \bar{a}(n) + \Delta\bar{a}(n), \quad (4)$$

gdzie $\Delta\bar{a}(n)$: poprawka, wprowadzona w chwili n . Najprostszym algorytmem sekwencyjnej adaptacji jest algorytm stochastycznego gradientu, wykorzystujący następujący wektor poprawki:

$$\begin{aligned} \Delta\bar{a}(n) &= -\frac{\alpha}{2} \cdot \text{grad}_{\bar{a}} (\varepsilon_n^q)^2 = -\frac{\alpha}{2} \cdot \text{grad}_{\bar{a}} (x_n^* - x_n^p)^2 = -\frac{\alpha}{2} \cdot \text{grad}_{\bar{a}} [x_n^* - \\ &\quad - \bar{a}^T(n) \cdot \bar{x}(n)]^2 = \alpha \cdot [x_n^* - \bar{a}^T(n) \cdot \bar{x}(n)] \cdot \bar{x}(n) = \alpha \cdot \varepsilon_n^q \cdot \bar{x}(n), \end{aligned} \quad (5)$$

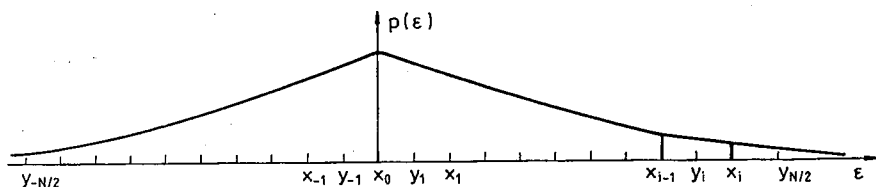
gdzie: α — współczynnik charakteryzujący szybkość adaptacji.

Poprawka skierowana jest w stronę przeciwną do gradientu $(\varepsilon_n^q)^2$, a więc w kierunku zmniejszania się mocy sygnału błędu predykcji. W celu uniezależnienia wielkości poprawki $\Delta\bar{a}(n)$ od poziomu sygnału mowy, należy współczynnik α uczynić odwrotnie proporcjonalnym do mocy sygnału mowy (jest to tzw algorytm stochastycznego gradientu z normalizacją).

2.3 KWANTYZER

Zadaniem kwantyzera jest zaokrąglenie wartości próbki sygnału wejściowego ε_n do jednej z N wartości (poziomów kwantyzacji), co umożliwia zakodowanie próbki ε_n^q w słowie binarnym o długości $B = \log_2 N$ bitów.

Kwantyzator jest w pełni określony, gdy znane są poziomy kwantyzacji $y_{-N/2}, \dots, y_{-1}, y_1, \dots, y_{N/2}$ oraz granice przedziałów kwantyzacji $x_{-N/2}, \dots, x_{-1}, x_0, x_1, \dots, x_{N/2}$; rys. 2 ($x_{-N/2} = -\infty, x_{N/2} = \infty$)



Rys. 2. Zasada działania kwantyzera: y_i – i -ty poziom kwantyzacji, od x_{i-1} do x_i rozciąga się i -ty przedział kwantyzacji (próbki leżące w tym przedziale będą zaokrąglone do wartości y_i), $p(\varepsilon)$ – gęstość prawdopodobieństwa wartości próbki

Ze względu na zerową wartość średnią sygnału, zapewnia się symetrię przedziałów i poziomów kwantyzacji: $x_0 = 0$, $x_i = -x_{-i}$, $y_i = -y_{-i}$. Mając dany rozkład prawdopodobieństwa wartości próbki $p(\varepsilon)$, można znaleźć moc szumu kwantyzacji ze wzoru [5]:

$$\sigma_e^2 = 2 \cdot \sum_{i=1}^{N/2} \int_{x_{i-1}}^{x_i} (\varepsilon - y_i)^2 \cdot p(\varepsilon) d\varepsilon. \quad (6)$$

Podstawiając moc kwantowanego sygnału $\sigma_s^2 = 2 \cdot \int_0^\infty \varepsilon^2 p(\varepsilon) d\varepsilon$ można określić stosunek mocy sygnału do szumu kwantyzacji $\text{SNR}_q = \sigma_s^2 / \sigma_e^2$. W przypadku kwantyzera równomiernego, w którym kolejne punkty $x_0, y_1, x_1, y_2, x_2, y_3, \dots$ tworzą sieć o jednakowych odstępach Δ , wartość Δ i liczba poziomów kwantyzacji N stanowią pełny opis kwantyzera.

Dla danego sygnału ($p(\varepsilon)$ znane) i liczby poziomów kwantyzacji N można zawsze znaleźć wartość Δ_{opt} , minimalizującą moc szumu kwantowania σ_e^2 , a co za tym idzie, maksymalizującą SNR_q .

W przypadku kwantyzera nierównomiernego, minimalizację σ_e^2 – wzór (6) przeprowadza się ze względu na wszystkie wartości $y_1, \dots, y_{N/2}$ i $x_1, \dots, x_{N/2-1}$ [5]. Otrzymuje się że w ten sposób kwantyzator optymalny, „dopasowany” do danego sygnału.

Ażeby zapewnić prawidłową pracę kwantyzera dla sygnałów o różnym poziomie (tzn. w warunkach gdy σ_e^2 , a więc także $p(\varepsilon)$ zależy od czasu), stosuje się adaptację kwantyzera, uzależniając wielkości y_i, x_i od czasu.

Adaptacja jest niezbędna w kodowaniu mowy, gdy liczba poziomów kwantyzacji jest mała ($N=4-16$).

Podobnie jak w przypadku adaptacji predyktora, istnieją metody adaptacji kwantyzera „w przód” i „w tył”. Adaptacja „w przód” polega na określeniu optymalnych parametrów kwantyzera dla pewnego fragmentu sygnału wejściowego, a następnie skwantowaniu tego fragmentu sygnału. Tego typu adaptacją nie będzie-

my się zajmować, gdyż wymaga ona transmisji parametrów kwantyzera do odbiornika, a ponadto nie może być stosowana w układach DPCM, w których sygnał wejściowy kwantyzera nie może być z góry określony, ze względu na zamkniętą pętlę sprzężenia zwrotnego (rys. 1).

Adaptacja „w tył” polega na korekcji parametrów kwantyzera jedynie w oparciu o wartości transmitowanych próbek ε_n^q . Większość stosowanych algorytmów adaptacji „w tył” wywodzi się z koncepcji Jayanta [1] opracowanej dla kwantyzera równomiernego. W algorytmie Jayanta parametr Δ , określający odległości poziomów kwantyzacji, zmienia się z częstotliwością próbkowania wg wzoru:

$$\Delta(n+1) = \Delta(n) \cdot M[|I(n)|] \quad (7)$$

gdzie n — indeks próbki, $I(n)$ — numer poziomu kwantowania wykorzystywany w chwili n , $M(1) \dots M(N/2)$ — współczynniki adaptacji.

Siatka poziomów kwantyzacji „rozciąga się” lub „kurczy” w zależności od wykorzystywanych poziomów kwantowania (niskie poziomy kwantowania implikują kurczenie się siatki, wysokie poziomy — jej rozciąganie).

Według tej zasady działa szereg kwantyzatorów równomiernych i nierównomiernych [1],[2],[4],[5],[9], jednak w literaturze brak jest jasnych wskazówek co do wyboru parametrów, w szczególności wartości $M[|I(n)|]$.

W przypadku systemu DPCM nie jest również jasne, na jakim modelu probabilistycznym $p(\varepsilon)$ sygnału wejściowego kwantyzera należy oprzeć się w procesie projektowania. Zagadnieniu optymalizacji parametrów kwantyzatorów adaptacyjnych równomiernych i nierównomiernych będzie poświęcona następna część artykułu.

3. KWANTYZER Z ADAPATACJĄ

3.1 KWANTOWANIE SYGNAŁU STACJONARNEGO

Adaptację kwantyzera stosuje się ze względu na niestacjonarność sygnału mowy, polegającą głównie na zmianach poziomu (wariancji σ_x^2) w czasie. Nie ulega jednak wątpliwości, że kwantyzator adaptacyjny winien prawidłowo przetwarzać również sygnały stacjonarne — gwarantuje to prawidłowe kwantowanie odcinków mowy o charakterze quasi-stacjonarnym (np. przeciągle wypowiedzianych samogłosek). Dla sygnału stacjonarnego lepsze wyniki daje jednak kwantyzator bez adaptacji o odpowiednio dobranych parametrach. Od kwantyzera adaptacyjnego można jedynie wymagać, aby działał w tych warunkach w sposób zbliżony do kwantyzera bez adaptacji.

Aby prawidłowo zaprojektować kwantyzator bez adaptacji (w przypadku kwantyzera równomiernego wystarczy obliczyć Δ_{opt} , w przypadku kwantyzera nierównomiernego wszystkie poziomy i przedziały kwantyzacji) wystarczy znać rozkład prawdopodobieństwa próbek sygnału wejściowego $p(\varepsilon)$. Jayant [1] posługiwał się rozkładem gaussowskim. Rozkład ten niebyt dokładnie odzwierciedla właściwości sygnału mowy — krótkie fragmenty mowy (głoski, pojedyncze wyrazy) lepiej opisuje

rozkład eksponencjalny Laplace'a [5]. Modelując sygnał mowy skorelowanym procesem stochastycznym o rozkładzie eksponencjalnym musimy wziąć pod uwagę, że sygnał wejściowy kwantyzera w układzie DPCM (tzn. sygnał błędu predykcji) może mieć inny rozkład prawdopodobieństwa. Wyjaśnieniu tego zagadnienia jest częściowo poświęcona praca [10]. Po wykonaniu szeregu badań testem chi-kwadrat okazało się, że sygnał błędu predykcji może być opisany rozkładem prawdopodobieństwa, będącym kombinacją liniową dwóch znanych rozkładów: eksponencjalnego i gaussowskiego.

Zawartość tych dwóch składników zależy od korelacji kolejnych próbek sygnału eksponencjalnego modelującego mowę, np. dla współczynnika korelacji $C=0.9$, $p(\varepsilon)$ jest w 70% rozkładem gaussowskim w 30% eksponencjalnym, co można zapisać wzorem:

$$p(\varepsilon) = 0.7 \frac{1}{\sqrt{2 \cdot \Pi \cdot \sigma_x^2}} \exp\left(-\frac{x^2}{2 \cdot \sigma_x^2}\right) + 0.3 \frac{1}{\sqrt{2 \cdot \sigma_x^2}} \exp\left(-\sqrt{2} \frac{|x|}{\sigma_x}\right). \quad (8)$$

W dalszych rozważaniach będziemy posługiwać się gęstością prawdopodobieństwa opisaną wzorem (8).

Dla sygnału o jednostkowej mocy ($\sigma_\varepsilon^2=1$) można znaleźć optymalną siatkę poziomów i przedziałów kwantyzacji opisaną parametrami $\bar{y}_1 \dots \bar{y}_{N/2}$, $\bar{x}_1 \dots \bar{x}_{N/2-1}$. Dla kwantyzera równomiernego wystarczy wyliczyć $\bar{\Delta}_{opt}$ na minimum σ_ε^2 – wzór (6) – podstawiając poziomy kwantowania $\bar{y}_i = (2i-1) \bar{\Delta}_{opt}$ i granice przedziałów kwantowania $\bar{x}_i = 2i \bar{\Delta}_{opt}$. Algorytm wyznaczania optymalnych parametrów kwantyzera nierównomiernego był wielokrotnie publikowany – por. [5].

Dla sygnału o tym samym rozkładzie $p(\varepsilon)$ ale innej mocy $\sigma_\varepsilon^2 \neq 1$ optymalne parametry $y_1 \dots y_{N/2}$, $x_1 \dots x_{N/2-1}$ mogą być wyznaczone z zależności $y_i = \bar{y}_i \cdot \sigma_\varepsilon$, $x_i = \bar{x}_i \cdot \sigma_\varepsilon$, gdzie σ_ε jest wartością skuteczną sygnału wejściowego kwantyzera (dispersja rozkładu $p(\varepsilon)$). Dla kwantyzera równomiernego wystarczy wyznaczenie $\Delta_{opt} = \bar{\Delta}_{opt} \cdot \sigma_\varepsilon$ i odtworzenie parametrów y_i i x_i ze wzorów $y_i = (2i-1) \Delta_{opt}$, $x_i = 2i \Delta_{opt}$. W zależności od σ_ε optymalna siatka poziomów i przedziałów kwantyzacji „rozciąga się” lub „kurczy się”.

W przypadku gdy σ_ε^2 nie jest znane, możemy posłużyć się jedynie estymatorem wartości skutecznej sygnału σ_ε . Do estymacji możemy użyć te wielkości, które są znane zarówno po stronie nadawczej jak i odbiorczej układu DPCM – są numerami poziomów kwantyzacji $I(n)$ w kolejnych chwilach czasowych n . W najprostszym przypadku otrzymuje się estymator:

$$\hat{\sigma}_\varepsilon(n+1) = \hat{\sigma}_\varepsilon(n) \cdot M[|I(n)|], \quad (9)$$

gdzie $M(k)$ – współczynniki adaptacji (tzw. mnożniki), zależne od numeru poziomu kwantyzacji. Odpowiednie poziomy kwantyzacji w chwili $n+1$ obliczymy z zależności $y_1(n+1) = \bar{y}_1 \hat{\sigma}_\varepsilon(n+1)$, a granice przedziałów kwantowania ze wzoru $x_i(n+1) = \bar{x}_i \hat{\sigma}_\varepsilon(n+1)$.

W szczególności, dla kwantyzera równomiernego obowiązują zależności $\Delta(n+1) = \Delta_{opt} \hat{\sigma}_\varepsilon(n+1)$, $\Delta(n) = \bar{\Delta}_{opt} \hat{\sigma}_\varepsilon(n)$, co po podstawieniu do (9) i pomnożeniu

stronami przez $\bar{\Delta}_{opt}$ daje wzór (7). Oznacza to, że algorytm adaptacji opisany wzorem (9) jest uogólnieniem wzoru Jayanta [1] na przypadek kwantyzera nierównomiernego.

Działanie kwantyzera adaptacyjnego zależy w głównej mierze od wyboru parametrów adaptacji $M(1) \dots M(N/2)$. Swoboda wyboru tych współczynników jest ograniczona pewnym warunkiem, wynikającym z wymagania stabilnej pracy w procesie kwantowania sygnału stacjonarnego.

Załóżmy, że mamy sygnał o jednostkowej mocy ($\sigma_\varepsilon^2 = 1$) i rozpoczynamy kwantyzację z optymalnymi parametrami $\bar{x}_i, \bar{y}_i$ (oznacza to, że $\hat{\sigma}_\varepsilon(0) = 1$). Po dostatecznie długim czasie ($n \rightarrow \infty$) estymator (9) winien oscylować wokół wartości $\hat{\sigma}_\varepsilon(n) \approx 1$, co oznacza, że

$$\prod_{k=1}^n M[|I(k)|] \approx 1. \quad (10)$$

Jeżeli przez $P(i)$ oznaczyć prawdopodobieństwo wystąpienia i -tego poziomu kwantowania, równe prawdopodobieństwu wystąpienia $(-i)$ -tego poziomu kwantowania, to można zauważyć, że iloczyn (10) składa się z $2n$ $P(1)$ czynników o wartości $M(1)$, $2n$ $P(2)$ czynników o wartości $M(2)$ itd.

Po uwzględnieniu tego faktu we wzorze (10) i podniesieniu stronami do potęgi $\frac{1}{2 \cdot n}$ otrzymuje się opublikowany przez Jayanta [1] warunek równowagi:

$$M(1)^{P(1)} \cdot M(2)^{P(2)} \dots M(N/2)^{P(N/2)} \approx 1 \quad (11)$$

W procesie kwantowania sygnału stacjonarnego o jednostkowej mocy kwantyzator z adaptacją winien zachowywać się podobnie jak kwantyzator bez adaptacji o optymalnych parametrach $\bar{x}_i, \bar{y}_i$. Oznacza to, że średnie wartości $x_i(n), y_i(n)$ winny być zbliżone do $\bar{x}_i, \bar{y}_i$ (dla kwantyzera równomiernego wystarczy, że średnia wartość $\Delta(n)$ jest bliska $\bar{\Delta}_{opt}$ [6]). Ponadto prawdopodobieństwa wystąpienia poszczególnych poziomów kwantyzacji winny być zbliżone do odpowiednich prawdopodobieństw wyznaczonych dla kwantyzera bez adaptacji, to jest:

$$P(i) = \int_{\bar{x}_{i-1}}^{\bar{x}_i} p(\varepsilon) d\varepsilon, \quad (12)$$

Prawdopodobieństwa wyznaczone ze wzoru (12) należy podstawić do warunku równowagi (11).

3.2 WYZNACZENIE WSPÓŁCZYNNIKÓW ADAPTACJI

Współczynniki adaptacji $M(1) \dots M(N/2)$ muszą być wybrane w ten sposób, aby zapewnić „kurczenie się” siatki poziomów kwantyzacji przy zmniejszaniu się wartości skutecznej sygnału σ_ε , oraz „rozciąganie się” siatki w warunkach rosnącego σ_ε . W związku z tym wartości $M(i)$ powinny być mniejsze od 1 dla niskich poziomów kwantyzacji i większe od 1 dla wysokich poziomów. Ponadto winien być spełniony warunek (11). Istnieje nieskończenie wiele wartości współczynników adaptacji speł-

nających w/w warunki. Wartości współczynników adaptacji można wyznaczyć biorąc pod uwagę, że są one jednocześnie współczynnikami estymatora wartości skutecznej sygnału (wzór 9). Weźmy pod uwagę następujący estymator wariancji sygnału wejściowego kwantyzera:

$$\hat{\sigma}_\varepsilon^2(n+1) = \beta \cdot \varepsilon_n^2 + (1 - \beta) \cdot \hat{\sigma}_\varepsilon^2(n) \quad (13)$$

gdzie $\beta < 1$ jest miarą szybkości reakcji estymatora na zmiany wariancji sygnału. Wartość próbkki ε_n nie jest znana po stronie odbiorczej układu DPCM, wobec czego należy wyrazić ją przez wartość skwantowaną ε_n^q :

$$\varepsilon_n = \varepsilon_n^q - e_n = y_{I(n)} - e_n = \bar{y}_{I(n)} \hat{\sigma}_\varepsilon(n) - e_n \quad (14)$$

Po uwzględnieniu tego wyrażenia w (13) otrzymuje się:

$$[\hat{\sigma}_\varepsilon(n+1)]^2 = \beta \cdot [\bar{y}_{I(n)}]^2 \cdot [\hat{\sigma}_\varepsilon(n)]^2 - 2 \cdot \beta \cdot \bar{y}_{I(n)} \cdot \hat{\sigma}_\varepsilon(n) \cdot e_n + \beta \cdot e_n^2 + (1 - \beta) \cdot [\hat{\sigma}_\varepsilon(n)]^2 = [\hat{\sigma}_\varepsilon(n)]^2 \cdot \left\{ (1 - \beta) + \beta \cdot [\bar{y}_{I(n)}]^2 - \frac{2 \cdot \beta \cdot \bar{y}_{I(n)} \cdot e_n}{\hat{\sigma}_\varepsilon(n)} + \frac{\beta \cdot e_n^2}{(\hat{\sigma}_\varepsilon(n))^2} \right\} \quad (15)$$

Po porównaniu z (9) okazuje się, że wyrażenie w nawiasie klamrowym jest kwadratem współczynnika adaptacji $M[I(n)]$, czyli

$$M(i) \approx \sqrt{(1 - \beta) + \beta \cdot (\bar{y}_i)^2}. \quad (16)$$

We wzorze (16) pominięto składniki zawierające małą wartość błędu kwantyzacji $e(n)$. Jest to uzasadnione w przypadku, gdy kwantowanie odbywa się bez przesterowania. W warunkach przesterowania, gdy $\varepsilon_n > y_{N/2}$ przyjmowany jest poziom kwantowania $I(n) = \frac{N}{2}$ (gdy $\varepsilon_n < y_{-N/2}$ jest to poziom $I(n) = -\frac{N}{2}$). Wówczas błąd

$e_n = \varepsilon_n^q - \varepsilon_n$ może osiągnąć dużą wartość bezwzględną, przy czym iloczyn $e_n \cdot \varepsilon_n^q = e_n \cdot \bar{y}_{I(n)} \cdot \hat{\sigma}_\varepsilon(n)$ jest ujemny. Oznacza to, że oba pominięte składniki zawierające e_n są dodatnie i mogą mieć dużą wartość. Stąd wypływa wniosek, że wzór (16) szacuje wartość $M(N/2)$ z niedomiarem — należy więc dodać poprawkę, której wielkość trzeba dobrać eksperymentalnie. Po podstawieniu dla kwantyzera równomiernego $\bar{y}_i = (2i - 1) \bar{\Delta}_{opt}$ oraz $\beta = 0,5$ wzór (16) staje się identyczny z formułą Jayanta [1]. Oznacza to, że wzór (16) jest uogólnieniem formuły Jayanta na przypadek kwantyzera nierównomiernego, oraz dowolnej szybkości adaptacji (β).

4. WYNIKI OBLICZEŃ I SYMULACJI

4.1 PROJEKTOWANIE KWANTYZERA Z ADAPTACJĄ

W pierwszym etapie obliczeń, posługując się rozkładem prawdopodobieństwa próbek błędu predykcji (8) i zakładając jednostkową wariancję sygnału, wyznaczono parametry kwantyzera równomiernego ($\bar{\Delta}_{opt}$) lub nierównomiernego ($\bar{x}_i, \bar{y}_i$) bez adaptacji, podobnie jak w [5]. Wyznaczono również prawdopodobieństwa $P(i)$ ze

wzoru (12) oraz stosunek mocy sygnału wejściowego kwantyzera do szumu kwantyzacji $SNR_q = \left[\frac{\sigma_\varepsilon}{\sigma_e} \right]^2$ ze wzoru 6. Np. dla kwantyzera równomiernego o $N=8$ poziomach kwantyzacji otrzymano $SNR_q = 13.07$ dB.

Zakładając szybkość adaptacji β można wyliczyć współczynniki adaptacji $M(1) \dots M(N/2)$ ze wzoru (16). Dla 8-poziomowego kwantyzera równomiernego i $\beta=0.5$ wyznaczono $M(1)=0,74$, $M(2)=0,97$, $M(3)=1,32$, $M(4)=1,71$.

Wartość $M(N/2)$ wyznaczona jest z niedomiarem i należy ją skorygować. Próby wyznaczenia $M(N/2)$ z warunku (11) nie dały zadowalających rezultatów. Lepsze wyniki dała metoda symulacyjna. Eksperymenty polegały na generowaniu stacjonarnego sygnału ε_n o rozkładzie prawdopodobieństwa (8) i jednostkowej wariancji $\sigma_\varepsilon^2=1$.

Sygnał ten był kwantowany przy użyciu zasymulowanego kwantyzera adaptacyjnego (przyjmowano warunek początkowy $\hat{\sigma}_\varepsilon(0)=1$, co zapewniało start z optymalnymi wartościami $\bar{x}_i$, $\bar{y}_i$).

Po przeprowadzeniu serii obliczeń otrzymywano $SNR_q = \sum_n \varepsilon_n^2 / \sum_n \varepsilon_n^2$ w funkcji $M(N/2)$ i wybierano współczynnik $M(N/2)$ ze względu na maksimum SNR_q . W omawianym przykładzie skorygowano $M(4)$ z wartości 1.71 do wartości 2.2. Jednocześnie kontrolowano średnie wartości poziomów kwantyzacji w każdej symulacji (wystarcza obliczenie średniej wartości $\hat{\sigma}_\varepsilon(n)$). Np. wykonując kwantyzację 5000 próbek sygnału przy użyciu 8-poziomowego kwantyzera równomiernego o podanych wyżej współczynnikach adaptacji otrzymano średnią wartość estymatora wartości skutecznej $\hat{\sigma}_\varepsilon(n)$ równą 0.93, co uznano za zadowalające.

Po ostatecznym wyznaczeniu współczynników adaptacji można sprawdzić warunek równowagi (11). We wszystkich analizowanych przypadkach otrzymywano wartości lewej strony wyrażenia (11) z przedziału 0.97 – 1. Spełnienie w zadowalającym stopniu warunku równowagi oraz utrzymywanie się średnich wartości poziomów kwantyzacji zbliżonych do $\bar{y}_i$ świadczy o tym, że zaprojektowany kwantyzator adaptacyjny zachowuje się w procesie kwantowania sygnału stacjonarnego w podobny sposób, jak kwantyzator nieadaptacyjny o optymalnych parametrach.

Otwartą pozostaje sprawa wyboru szybkości adaptacji β . Im większa wartość β , tym szybsza reakcja kwantyzera na zmiany mocy sygnału, lecz gorsze wyniki w procesie kwantowania sygnału stacjonarnego. Ilustracją są wyniki symulacji zamieszczone w tabl. 1 i 2. Eksperymenty polegały na kwantowaniu $M=10, 100, 1000$ lub 5000 próbek stacjonarnego sygnału ε_n o rozkładzie (8) i jednostkowej wariancji $\sigma_\varepsilon^2=1$. Warunki początkowe kwantyzera tj. wartość $\hat{\sigma}_\varepsilon(0)$ ustalającą początkową siatkę poziomów i przedziałów kwantyzacji $y_i(0)=\bar{y}_i \cdot \hat{\sigma}_\varepsilon(0)$, $x_i(0)=\bar{x}_i \cdot \hat{\sigma}_\varepsilon(0)$ wybierano w pięciu wariantach: na optymalnym $\hat{\sigma}_\varepsilon(0)=1$, w odległości ∓ 10 dB od poziomu optymalnego: $\hat{\sigma}_\varepsilon(0)=\frac{1}{\sqrt{10}}$ i $\sqrt{10}$, oraz w odległości ∓ 20

dB od poziomu optymalnego: $\hat{\sigma}_\varepsilon(0)=0,1$ i 10. W tabelach zamieszczono wartości

$SNR_q = \sum_{i=0}^{N-1} \varepsilon_n^2 / \sum_{i=0}^{N-1} e_n^2$, czyli stosunek sygnału do szumu kwantyzacji. Porównując tabl. 1 i 2 można zauważyć, że większa szybkość adaptacji ($\beta=0,5$) gwarantuje lepsze wyniki dla krótkich fragmentów sygnału ($M=10-100$ próbek), natomiast mniejsza szybkość adaptacji ($\beta=0,2$) zapewnia o 1 dB lepsze SNR_q dla długich fragmentów sygnału ($M=1000-5000$ próbek). Porównując tablicę 2 i 3 można

T a b l i c a 1

SNR_q [dB] w procesie kwantowania M próbek sygnału stacjonarnego, przy użyciu kwantyzera równomiernego o $N=8$ poziomach kwantowania, szybkość adaptacji $\beta=0,5$

$\hat{\sigma}_\varepsilon(0) \backslash M$	10	100	1000	5000
0,1	5,9	11,6	10,6	10,9
$\frac{1}{\sqrt{10}}$	9,0	12,4	10,5	10,9
1	13,6	12,9	10,5	10,9
$\sqrt{10}$	8,1	12,2	10,5	10,8
10	-3,7	7,7	9,9	10,7

T a b l i c a 2

SNR_q [dB] jak w tabl. 1 — szybkość adaptacji $\beta=0,2$

$\hat{\sigma}_\varepsilon(0) \backslash M$	10	100	1000	5000
0,1	6,4	11,3	11,3	11,8
$\frac{1}{\sqrt{10}}$	9,3	13,2	11,3	11,8
1	14,2	13,4	11,7	11,9
$\sqrt{10}$	4,5	11,9	11,4	11,9
10	-6,8	5,1	10,1	11,6

stwierdzić, że kwantyzator nierównomierny zachowuje się lepiej od kwantyzera równomiernego o tych samych parametrach (liczba poziomów kwantowania, szybkość adaptacji) — zarówno w procesie kwantowania krótkich jak i długich fragmentów sygnału.

T a b l i c a 3

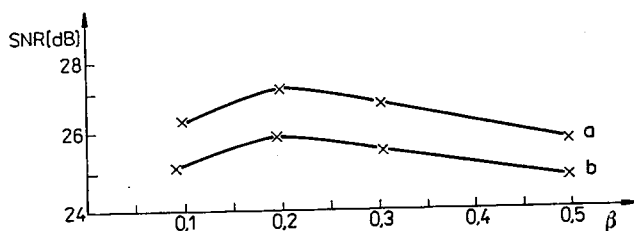
SNR_q [dB] w procesie kwantowania M próbek sygnału stacjonarnego przy użyciu kwantyzera nierównomiernego o $N=8$ poziomach kwantowania — szybkość adaptacji $\beta=0,2$

$\hat{\sigma}_e(0) \backslash M$	10	100	1000	5000
0,1	10,0	11,5	11,7	12,6
$\frac{1}{\sqrt{10}}$	8,1	12,8	11,9	12,7
1	14,6	14,1	12,0	12,8
$\sqrt{10}$	10,4	14,0	11,9	12,7
10	-0,8	9,1	11,3	12,5

4.2 SYMULACJA UKŁADU DPCM

Badania symulacyjne z wykorzystaniem sygnału stacjonarnego nie stanowią wystarczającej podstawy do określenia wszystkich parametrów kwantyzera — przede wszystkim szybkości adaptacji β . Można to uczynić jedynie wykorzystując wyniki symulacji całego kodera DPCM wraz z kwantyzерem i predyktorem (rys. 1) w przeprowadzonych badaniach symulacyjnych stosowano predyktor transwersalny opisany wzorem (3) o $p=8$ współczynnikach predykcji. Adaptacja współczynników predykcji przebiegała wg zasady stochastycznego gradientu (wzory 4 i 5) ze współczynnikiem $\alpha=0,1$. Przetwarzano frazę sygnału mowy o długości ok. 4 sekund.

Badano kwantyzery równomierne i nierównomierne zaprojektowane wg zasad podanych w p. 4.1. Na rys. 3 podano $SNR = \left[\frac{\sigma_x}{\sigma_e} \right]^2$ w funkcji szybkości adaptacji β dla kwantyzera równomiernego i nierównomiernego o $N=16$ poziomach. Optymalny współczynnik adaptacji β ma wartość 0,2, lecz nie jest to wartość krytyczna. Zmiany szybkości adaptacji w przedziale 0,1–0,5 nie obniżają SNR o więcej niż

Rys. 3. SNR kodera DPCM w funkcji szybkości adaptacji kwantyzera β ($N=16$ poziomów):

a — kwantyzator nierównomierny, b — kwantyzator równomierny

1 dB. Kwantyzer nierównomierny zachowuje ponad 1 dB przewagę nad kwantyzерem równomiernym.

Po ustaleniu optymalnej szybkości adaptacji $\beta=0,2$ wyznaczono SNR w funkcji liczby poziomów kwantyzacji N (tabl. 4).

T a b l i c a 4

Stosunek mocy sygnału do zniekształceń SNR [dB] dla układu DPCM z kwantyzерem o N poziomach kwantyzacji

N	4	8	16
kwantyzер			
równomierny	14,4	20,44	25,53
nierównomierny	14,3	20,92	26,67

ZAKOŃCZENIE

Proponowana metoda projektowania równomiernych i nierównomiernych kwantyzерów adaptacyjnych, dla dowolnego sygnału wejściowego i dowolnej szybkości adaptacji, okazała się skutecznym narzędziem praktycznym. Zaprojektowane kwantyzery umożliwiają prawidłową pracę układów DPCM o szybkości transmisji 24 i 32 kbit/s — przy szybkości transmisji 16 kbit/s szum kwantyzacji jest nazbyt dokuczliwy (tabl. 4). Zaprojektowane układy DPCM sprawdzono symulacyjnie — planuje się ich realizację w czasie rzeczywistym na procesorze sygnałowym TMS 32010.

BIBLIOGRAFIA

1. N. S. J a y a n t: *Adaptive quantization with a one-word memory*. B. S. T. J. vol 52, No 7, Sept 1973
2. P. C u m m i s k e y, N. S. J a y a n t, J. L. F l a n a g a n: *Adaptive quantization in DPCM coding of speech*. B. S. T. J. vol 52, No 7, Sept 1973
3. M. D. P a e z, T. H. G l i s s o n: *Minimum mean squared error quantization in speech PCM and DPCM systems*. IEEE Tran. on Communications, April 1972
4. J. D. G o o d m a n, R. M. W i l k i n s o n: *A robust adaptive quantizer*. IEEE Trans. on Communications. Vol COM-23, Nov 1975
5. N. S. J a y a n t, P. N o l l: *Digital coding of waveforms — principles and applications to speech and video*, Prentice-Hall 1984
6. D. J. G o o d m a n, A. G e r s h o: *Theory of an adaptive quantizer*. IEEE on Communications, vol. COM-22, No. 8, August 1974
7. R. C. R e i n i n g e r, J. D. G i b s o n: *Backward adaptive lattice and transversal predictors in ADPCM*, IEEE. Trans. on Communications, vol. 33, Jan. 1985
8. P. B u b l e w i c z: *Modulacja delta z adaptacją sylabową sterowaną obwiednią sygnału*, Referaty Instytut Telekomunikacji WP, Z. 64, 1979
9. CCITT: *Recomendation G. 721: 32 kbit/s adaptive differential pulse code modulation (ADPCM)*, August 1986
10. K. A b o u K a s s e m: *Optymalizacja kwantyzera błędu predykcji*. Rozprawa doktorska, Inst. Telekomunikacji Politechniki Warszawskiej, 1990

K. A. KASSEM, P. DYMARSKI

ADAPTIVE QUANTIZERS IN DPCM CODERS

S u m m a r y

Design method of the adaptive uniform and nonuniform quantizers has been described, given the statistical properties of the signal, the number of quantization levels and the speed of adaptation. The designed quantizers have been tested in a simulated DPCM coder with adaptive predictor. These simulations validated the proposed design method and permitted to optimize some parameters of the quantizer, e.g. the speed of adaptation.

Model czujnika fotoelektrycznego jako fotodetektora w układzie otwartego łącza optycznego

ANNA CYSEWSKA-SOBUSIAK

Instytut Elektroniki i Telekomunikacji, Politechnika Poznańska

Otrzymano 1990.05.12

Autoryzowano do druku 1990.07.14

W artykule przedstawiono zmodyfikowany model czujnika fotoelektrycznego o strukturze złączowej p-n lub p-i-n. Omówiono wyniki analizy właściwości statycznych i dynamicznych modelu oraz przyjęte zasady jego weryfikacji eksperymentalnej. Rozpatrzono oddziaływanie napięcia polaryzacji, temperatury, poziomu oświetlenia, charakteru i wartości parametrów obciążenia. W przedstawionych zaleceniach aplikacyjnych uwzględniono wnioski wynikające z przeprowadzonych badań reprezentatywnych egzemplarzy czujników.

1. SFORMUŁOWANIE PROBLEMU

Przyjęte określenie „otwarte łącze optyczne” dotyczy rodzaju kanału przesyłu i przetwarzania informacji o wielkości nieelektrycznej charakteryzującej obiekt badany. Nośnikiem informacji jest zdeterminowane promieniowanie świetlne oddziałujące na obiekt, a przetwarzanie odbywa się w wydzielonej optycznie przestrzeni pomiarowej. Elementy optyczne kanału nie są czujnikami wielkości mierzonej.

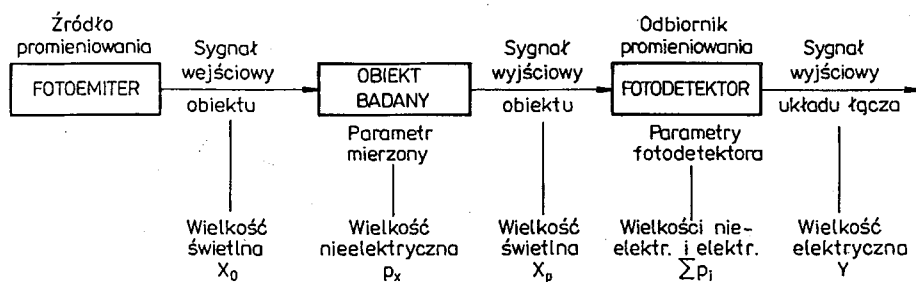
W układzie łącza można przeprowadzać pomiary laboratoryjne i techniczne wielkości świetlnych oraz przetwarzanych na wielkości świetlne, o zróżnicowanej dokładności i przeznaczeniu [10]. Ze względu na rodzaj badanych wielkości wyróżniono cztery grupy pomiarów:

a) pomiary fotometryczne związane z oceną warunków widzenia (np. pomiar natężenia oświetlenia),

b) pomiary związane z testowaniem, wzorcowaniem, selekcjonowaniem elementów optoelektronicznych pod względem ich właściwości metrologicznych (np. pomiar światłości fotoemitera),

c) pomiary wielkości nieelektrycznych przetwarzanych w obiekcie badanym na wielkości świetlne (np. spektrofotometryczne pomiary stężenia substancji absorbujących promieniowanie świetlne),

d) pomiary związane z identyfikacją przemieszczających się obiektów (np. pomiar gabarytów poruszającego się przedmiotu).



Rys. 1. Główne ogniwa w przykładowej strukturze otwartego łąca optycznego

Na rys. 1 przedstawiono uproszczoną strukturę przykładowego układu łąca, wyróżniając jego zasadnicze ogniwa, w kolejności przyczynowej zachodzących zjawisk. Otwarte łące optyczne tworzą w tym przypadku: źródło promieniowania (fotoemiter) — transiluminowany obiekt badany — odbiornik promieniowania (fotodetektor). Realizowany pomiar można zaliczyć do badań objętych w/w grupą c).

Miarą wiarygodności przetwarzania w układzie jest całkowity błąd pomiaru. Składowe tego błędu powstają na poszczególnych etapach przetwarzania. Uwzględniając oznaczenia z rys. 1, równanie przetwarzania dla bloku fotodetektora można zapisać w ogólnej postaci:

$$Y = f_2(X_p, \sum_{i=1}^n P_i), \quad (1)$$

gdzie $X_p = f_1(X_0, p_x)$.

Niektóre z parametrów fotodetektora mają znaczący wpływ na wiarygodność pomiarów.

W zależności od rodzaju pomiarów, fotodetektor zawiera pojedynczy czujnik fotoelektryczny lub zespół takich elementów. Obecnie powszechnie stosuje się krzemowe czujniki złączowe. W artykule przedstawiono zmodyfikowany model matematyczny czujnika fotoelektrycznego. Analizę stanów modelu zweryfikowano eksperymentalnie. Badania przeprowadzono dla reprezentatywnych egzemplarzy czujników. Analizując wpływ czynników oddziałujących na kształtowanie elektrycznej odpowiedzi czujnika (prąd wyjściowy I_0) na wymuszenie świetlne (wielkość X_p), uwzględniono m.in. obciążenie zewnętrzne. Ustalono, że może być celowy dobór nie tylko odpowiedniej wartości, ale i charakteru impedancji obciążenia, czego w badaniach poprzedzających nie brano pod uwagę.

2. ZAŁOŻENIA DO BUDOWY MODELU

Przedmiotem opracowania jest model czujnika fotoelektrycznego o strukturze złączowej typu p-n lub p-i-n. Zaliczając do tej grupy czujniki wykorzystujące wewnętrzny efekt fotoelektryczny, można w ramach niej wyróżnić: fotoogniwa,

fotodiody konwencjonalne jedno- i dwuzłączowe oraz fotodiody p-i-n i lawinowe. Czujniki te różnią się parametrami metrologicznymi określającymi ich przydatność aplikacyjną [3, 5, 6], ale można dla nich przyjąć uogólnioną postać modelu. Z uwagi na strukturę czujników, charakter zjawisk fizycznych i złożoność funkcji opisujących ich parametry wewnętrzne, celowe jest przyjęcie modelu w postaci schematu zastępczego i równań opisujących zależności między elementami składowymi schematu.

Model — rozsądnie złożony — ma być fizycznie uzasadniony i przydatny dla celów pomiarowych. Przyjęto, że winien on spełnić następujące wymagania:

- a) odzwierciedlenie istotnych zjawisk fizycznych zachodzących pod wpływem promieniowania, w warunkach statycznych i dynamicznych,
- b) wprowadzenie elementów możliwych do pomiaru lub obliczenia z możliwie małymi błędami,
- c) uwzględnienie w analizie możliwości nieliniowości charakterystyk statycznych elementów składowych,
- d) możliwość modyfikacji schematu oraz równań matematycznych tak, by dysponować wariantami będącymi podstawą do opracowania szczegółowych zaleceń aplikacyjnych.

3. UOGÓLNIONA POSTAĆ SCHEMATU ZASTĘPCZEGO CZUJNIKA

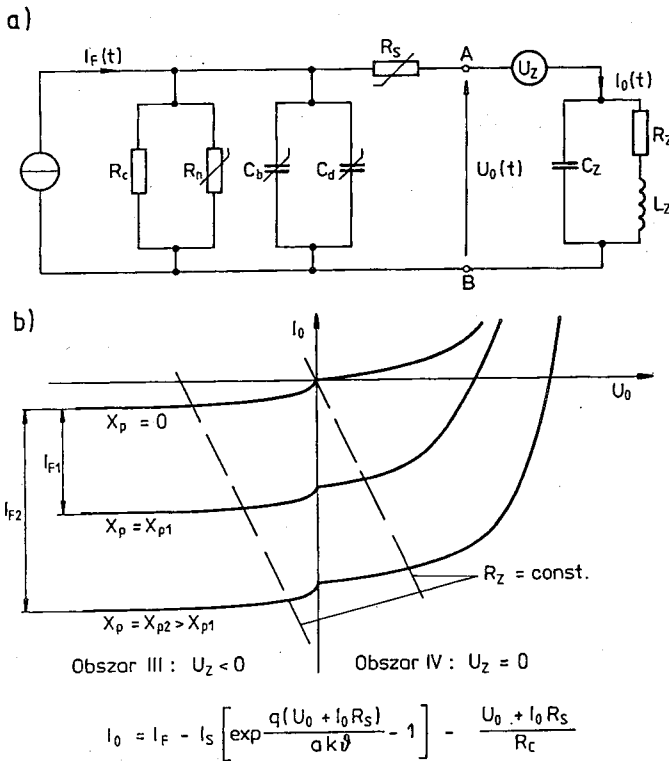
Proponowaną uogólnioną postać schematu zastępczego oświetlonego i obciążonego czujnika fotoelektrycznego o strukturze złączowej, przedstawiono na rys. 2a.

Czujnik stanowi źródło prądowe obciążone trzema gałęziami dedukującymi o dokładności i czasie trwania procesu przetwarzania wymuszenia świetnego X_p na wyjściowy prąd mierzony I_0 . Źródło prądowe reprezentuje fotoelektryczny prąd źródłowy I_F . Parametrami wewnętrznymi są pojemności C_b i C_d oraz rezystancje R_c , R_n i R_s . Do zacisków zewnętrznych A—B dołączona jest impedancja obciążenia o parametrach R_z , L_z , C_z oraz źródło napięcia U_z polaryzacji zewnętrznej. W praktycznych realizacjach układów z czujnikami złączowymi występuje dwojaki charakter polaryzacji i związany z tym rodzaj działania:

— brak polaryzacji zewnętrznej ($U_z = 0$) i działanie oparte na wykorzystaniu zjawiska fotowoltaicznego, wywołującego napięcie fotoelektryczne polaryzujące złącze w kierunku przewodzenia,

— polaryzacja zewnętrzna w kierunku zaporowym ($U_z < 0$), w wyniku której występuje silne pole elektryczne na warstwie ładunku przestrzennego złącza, a działanie opiera się na wykorzystaniu zjawiska fotoprzewodnościowego. Sens fizyczny pozostałych elementów schematu jest następujący:

Źródło prądowe określa proces generacji nośników prądu fotoelektrycznego I_F pod wpływem promieniowania świetlnego. Prąd ten wspomaga ciemny prąd nasycenia, generowany termicznie. W stanie ustalonym, wartości wielkości świetlnej $X_p = \text{const.}$ odpowiada wartość $I_F = \text{const.}$ Dla promieniowania świetlnego o określonym rozkładzie gęstości monochromatycznej, zależność tę można zapisać w postaci charakterystyki źródła:



Rys. 2. a) Uogólniony schemat zastępczy złączowego czujnika fotoelektrycznego, b) Przebieg statycznej charakterystyki prądowo-napięciowej

$$I_F = S \cdot X_p. \quad (2)$$

Parametr S oznacza bezwzględną czułość ogólną czujnika fotoelektrycznego.

Pojemność wewnętrzna uwidacznia się w warunkach zmiennej w czasie wielkości świetlnej $X_p(t)$, a określa ją suma:

$$C = C_b + C_d. \quad (3)$$

Składowe C_b i C_d zależą od właściwości fizycznych czujnika związanych z jego strukturą oraz od parametrów zewnętrznych. Można je określić w postaci następujących wyrażeń:

pojemność bariery C_b :

$$C_b = b \cdot \sqrt{\frac{1}{U_b - U_{pol}}}; \quad b = A_z \cdot \sqrt{\frac{\varepsilon}{2 \cdot \mu \cdot \rho}}; \quad (4)$$

pojemność dyfuzyjna C_d :

$$C_d = G_z \cdot \tau; \quad G_z = \frac{q \cdot I_s}{a \cdot k \cdot \theta} \cdot \exp \left(-\frac{U_b - U_{pol}}{a \cdot k \cdot \theta} \right), \quad (5)$$

gdzie: U_b — napięcie kontaktowe,

U_{pot} — napięcie polaryzujące złącze, o wartości i znaku ustalonym przez oświetlenie oraz wsteczne napięcie zewnętrzne,

A_z — powierzchnia złącza,

ϵ, ρ — stała dielektryczna oraz rezystywność materiału,

μ, τ — ruchliwość oraz czas życia nośników mniejszościowych,

I_s — ciemny prąd nasycenia,

a — współczynnik zależny od właściwości materiału ($a = 1 \div 2$),

k — stała Boltzmanna,

q — ładunek elementarny,

ϑ — temperatura otoczenia.

W stanie ciemnym pojemność C_d jest bliska zeru — główną składową pojemności czujnika jest pojemność bariery C_b . Przy oświetlaniu czujnika rosną wartości obu składowych, przy czym stopień ich udziału zależy od rodzaju polaryzacji złącza. Przy braku polaryzacji zewnętrznej główne znaczenie ma pojemność dyfuzyjna C_d , której wartość rośnie ze wzrostem poziomu oświetlania. Przy polaryzacji zewnętrznej napięciem wstecznym, pojemność C odpowiada pojemności bariery C_b o wartości tym mniejszej, im wyższa wartość tego napięcia.

Wartość pojemności czujnika uzależniona jest także od temperatury — jej wzrost przyczynia się do wzrostu wartości C . Szczególnie istotny jest wpływ temperatury na wartość prądu ciemnego, od którego bezpośrednio zależy wartość pojemności dyfuzyjnej określonej wyrażeniem (5). Zależność $I_s(\vartheta)$ można zapisać w postaci:

$$I_s = I_{s\vartheta_0} \cdot \exp k\vartheta \cdot (\vartheta - \vartheta_0), \quad (6)$$

gdzie: ϑ_0 — temperatura odniesienia, najczęściej równa $+25^\circ\text{C}$,

$k\vartheta$ — współczynnik temperatury o wartości $(0,06 \div 0,10)$.

Ponadto, ponieważ prąd ciemny wspomaga prąd fotoelektryczny generowany światłem, może przy wzroście temperatury otoczenia zostać błędnie zinterpretowany przez układ jako sygnał pochodzący od zmiany wartości wielkości świetlnej X_p .

Niekorzystne oddziaływanie prądu ciemnego na pojemność zaznacza się dla czujników pracujących przy polaryzacji w kierunku zaporowym.

Rezystancje wewnętrzne czujnika to rezystancja szeregową R_s i równoległą

$$R_r = \frac{R_c \cdot R_n}{R_c + R_n}.$$

Element R_s odpowiada rezystancji środkowej warstwy w strukturze czujnika. W złączach typu p-n warstwa przejściowa jest cienka — rezystancja szeregową ma mniejszą wartość aniżeli w złączach p-i-n z wysokooporową warstwą samoistną. Wyrażenie analityczne określające rezystancję R_s i uwzględniające wpływ parametrów wewnętrznych czujnika, można zapisać w postaci:

$$R_s = \frac{w}{q \cdot \mu \cdot N \cdot A_z}, \quad (7)$$

gdzie: w — grubość warstwy przejściowej,

N – koncentracja nośników większościowych.

Zewnętrzna polaryzacja napięciem zaporowym sprawia, że nawet znaczna wartość R_s nie wywiera tak ujemnego wpływu na pracę fotodiody p-i-n jak niewielka wartość R_s na pracę niespolaryzowanego zewnętrznie fotoogniwa o strukturze p-n. Fotoogniwo jest przetwornikiem energii a sprawność przetwarzania określa wyrażenie:

$$\eta = \frac{I_0 \cdot U_0}{W_F} \cdot 100\%, \text{ gdzie } W_F \text{ jest mocą promieniowania świetlnego. Występująca}$$

rezystancja R_s powoduje spadek wartości napięcia na zaciskach wyjściowych fotoogniwa i pogorszenie jego sprawności. Niekorzystny wpływ wywiera wzrost temperatury, który przyczynia się m.in. do zmniejszenia ruchliwości nośników, co prowadzi do niepożądanego wzrostu wartości rezystancji R_s . Ponadto, ze wzrostem temperatury maleje wartość napięcia fotoelektrycznego generowanego światłem, a tym samym zmniejsza się sprawność fotoogniwa.

Elementy R_c i R_n tworzą rezystancję równoległą bocznikującą źródło prądu I_F i ilustrują drogę upływu prądu w obrębie struktury czujnika. Rezystancja R_c ma największą wartość w stanie ciemnym i nieliniowo maleje ze wzrostem oświetlenia czujnika. Zależność tę można interpretować jako bocznikowanie w coraz większym stopniu stałej rezystancji $R_c = \text{const}$ nieliniową rezystancją R_n odzwierciedlającą diodową strukturę czujnika. Tę nieliniową składową rezystancji równoległej określa wyrażenie:

$$R_n = \frac{U_{pol}}{I_s \cdot \exp \frac{q U_{pol}}{a \cdot k \cdot T} - I_s} \quad (8)$$

W stanie ciemnym, polaryzacja w kierunku zaporowym sprawia, że rezystancja R_c osiąga bardzo duże wartości wskutek rosnącej z napięciem wstecznym składowej R_n . Polaryzacja w kierunku przewodzenia napięciem generowanym światłem powoduje zmniejszenie R_n i znaczny upływ prądu wewnątrz struktury czujnika.

Obciążenie zewnętrzne czujnika ma w warunkach statycznych charakter rezystancyjny (rezystancja R_z). Na podstawie bilansu prądów można zapisać równanie statycznej charakterystyki przetwarzania w postaciach:

$$I_0 = I_F \cdot \frac{1}{1 + \frac{R_s + R_z}{R_r}} \quad (9)$$

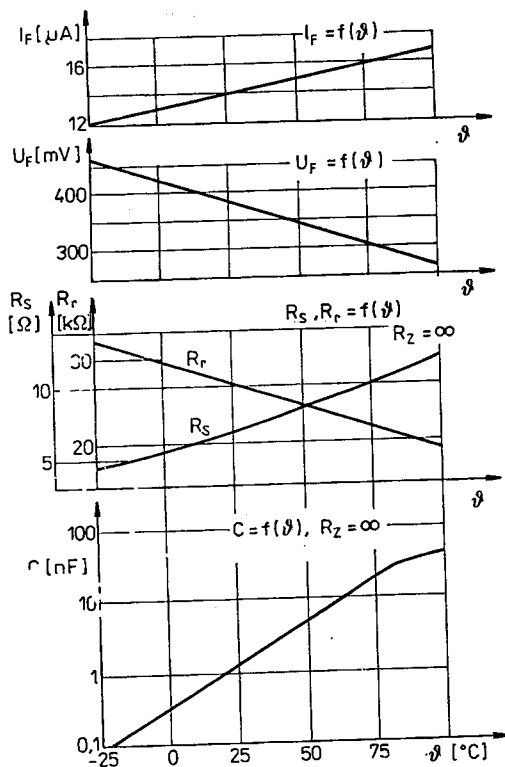
oraz po uwzględnieniu wyrażenia (2):

$$I_o = S \cdot X_p \cdot \frac{1}{1 + \frac{R_s + R_z}{R_r}} \quad (10)$$

W celu uzyskania liniowego przebiegu statycznej charakterystyki przetwarzania, konieczne jest spełnienie warunku:

$$R_s + R_z \ll R_r. \quad (11)$$

W określonych warunkach oświetlenia i obciążenia czujnika może wystąpić niekorzystne oddziaływanie temperatury. Z jej wzrostem zwiększa się wartość prądu ciemnego I_s , co powoduje wzrost rezystancji R_s i zmniejszenie rezystancji R_r . Wskutek tego może wystąpić niespełnienie warunku (11).



Rys. 3. Zależność parametrów I_F , U_F , R_s , R_r , C od temperatury dla przykładowej fotodiody krzemowej pracującej jako fotoogniwo (niespolaryzowana zewnętrznie i oświetlona z natężeniem $E = 1000 \text{ lx}$)

Na rys. 3 zestawiono wykresy ilustrujące zależność od temperatury wybranych parametrów przykładowego czujnika — fotodiody pracującej w warunkach $U_z = 0$, a więc jako fotoogniwo.

Przy zastosowaniu czujnika do pomiarów wielkości świetlnej o szybko zmieniających się wartościach, mogą odgrywać rolę wielkości L_z i C_z . Indukcyjność i pojemność w obwodzie obciążenia czujnika mogą być wprowadzone przez przyrządy służące do pomiarów elektrycznych parametrów wyjściowych jak również wielkości te mogą mieć charakter pasożytniczy, uwidaczniający się przy wyższych częstotliwościach modulacji wartości wielkości świetlnej ($f_{xp} > 1000 \text{ Hz}$). Jak wykazano dalej, istnieją także warunki, w których obecność indukcyjności w gałęzi obciążenia jest pożądana w celu poprawy właściwości dynamicznych układu.

W rzeczywistych układach łącza optycznego, warunki pracy czujnika fotoelektrycznego zależą w znacznym stopniu od sposobu jego polaryzacji napięciowej. Parametr ten posłużył jako punkt wyjścia do analizy wariantów modelu czujników należących do omawianej grupy fotodetektorów.

Uwzględniając sens fizyczny elementów składowych uogólnionego schematu z rys. 2, zbadano wpływ parametrów wewnętrznych i zewnętrznych na przebieg charakterystyk czujnika: statystycznej prądowo-napięciowej oraz dynamicznej odpowiedzi na zmienne w czasie wymuszenie świetlne.

Przyjęto założenie, że powierzchnia światłoczuła ma znaczne rozmiary; szczególnie duże są one w przypadku fotoogniw. Zgodnie z tym założeniem, pominięto w rozważaniach pojemność C_z , przyjmując warunek $C_z \ll C_b + C_d$. W gałęzi obciążenia występuje zatem *impedancja o parametrach R_z, L_z* .

4. WARIANTY MODELU ZŁĄCZOWEGO CZUJNIKA FOTOELEKTRYCZNEGO

4.1. SCHEMATY ZASTĘPCZE I RÓWNANIA CHARAKTERYSTYK STATYCZNYCH

Równanie ogólne statycznej charakterystyki prądowo-napięciowej przedstawionej na rys. 2b, można zapisać w postaci:

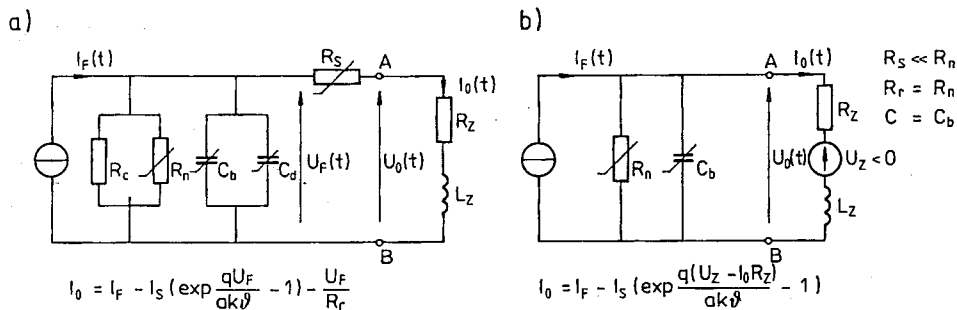
$$I_0 = I_F - I_S \left(\exp \frac{U_{pol}}{a \cdot k \cdot \vartheta} - 1 \right) - \frac{U_{pol}}{R_c} \quad (12)$$

Dla napięcia $U_z = 0$ (IV ćwiartka układu współrzędnych), napięcie polaryzacji U_{pol} określa napięcie fotoelektryczne U_F generowane światłem, co można zapisać jako:

$$U_{pol_{IV}} = U_F = I_0 (R_s + R_z). \quad (13)$$

Dla napięcia $U_z < 0$ (III ćwiartka układu współrzędnych), napięcie U_{pol} wyraża równanie:

$$U_{pol_{III}} = U_z - I_0 (R_s + R_z). \quad (14)$$



Rys. 4. Warianty modelu czujnika: a) dla IV ćwiartki układu współrzędnych charakterystyki $I_0(U_0): U_z = 0$, b) dla III ćwiartki układu współrzędnych charakterystyki $I_0(U_0): U_z < 0$

Uwzględniając wnioski wynikające z przedstawionego wpływu rodzaju polaryzacji na wartości parametrów wewnętrznych czujnika, można zestawić dwa warianty schematu zastępczego i odpowiadających im równań charakterystyk (12). Warianty te przedstawiono na rys. 4.

Wyrażenie ogólne (9) można zapisać odpowiednio w postaciach:

$$I_{0IV} = I_F \cdot \frac{1}{(R_s + R_z) \cdot \left(\frac{1}{R_n} + \frac{1}{R_c} \right) + 1} \quad (15)$$

oraz

$$I_{0III} = I_F \cdot \frac{1}{1 + \frac{R_z}{R_n}} \quad (16)$$

Prąd I_0 jest elektryczną odpowiedzią czujnika na wymuszenie świetlne X_p , które reprezentuje w schematach prąd fotoelektryczny I_F . Z punktu widzenia właściwości metrologicznych układu istotna jest wierność odtworzenia wielkości mierzonej X_p , a o jego przydatności do pomiarów dynamicznych decyduje przede wszystkim czas ustalenia odpowiedzi. Pierwsze z wymagań jest jednoznaczne z zapewnieniem możliwie małego błędu spowodowanego możliwością nieliniowego przebiegu charakterystyki przetwarzania określonej równaniem (10). Niezależnie od rodzaju czujnika pożądana jest w tym celu praca w warunkach obciążenia bliskich zwarcia.

W pomiarach dynamicznych, czas odpowiedzi układu limitowany jest bezwładnością czujnika. Poniżej przedstawiono wyniki analizy dotyczącej właściwości dynamicznych rozpatrywanego modelu.

4.2. ANALIZA WŁAŚCIWOŚCI DYNAMICZNYCH

Schematy zastępcze z rys. 4 mają charakter rozbudowanych członów oscylacyjnych. Ponieważ występują w nich elementy nieliniowe, analizę oparto na wprowadzonych założeniach upraszczających. Założenia te są fizycznie uzasadnione, zweryfikowane doświadczalnie [1, 2] i prowadzą do ustalenia zależności analitycznych rozsądnie złożonych. Przyjmując założenia bazowano na postulacie zapewnienia prawidłowości modelu pod względem jego sensu fizycznego i wymaganie to uczyniono priorytetowym w stosunku do rygorystycznego formalizmu.

Analizie poddano odpowiedź czujnika fotoelektrycznego na prostokątną zmianę wielkości świetlnej (impuls światła przerywanego). Przyjęto, że w ostatniej fazie odpowiedzi, a więc w bezpośrednim sąsiedztwie jej wartości ustalonej, parametry schematu zastępczego są stałe. Zgodnie z tym, dla skoku dodatniego wielkości X_p , parametry czujnika reprezentują ich wartości „jasne”, określone docelową wartością $I_F = \text{const.}$, a dla skoku ujemnego — wartości charakteryzujące czujnik w stanie ciemnym ($I_F = 0$).

Korzystając z metody opisu układów liniowych ustalono postać analityczną skokowej odpowiedzi $I_0(t)$, wprowadzając parametry charakterystyczne:

m — współczynnik obciążenia rezystancją,

ω_0 — pulsacja drgań nietłumionych,

β — współczynnik tłumienia drgań.

Parametry te można opisać przy pomocy odpowiednich wyrażeń dla wariantów a) i b) schematów z rys. 4:

$$m_{IV} = 1 + \frac{R_s + R_z}{R_n \cdot R_c} \cdot (R_n + R_c), \quad (17)$$

$$m_{III} = 1 + \frac{R_z}{R_n},$$

$$\omega_{0IV} = \frac{1}{\sqrt{L_z(C_b + C_d)}} \quad (18)$$

$$\omega_{0III} = \frac{1}{\sqrt{L_z \cdot C_b}},$$

$$\beta_{IV} = \frac{1}{2} \left(R_s + R_z \right) \cdot \sqrt{\frac{C_b + C_d}{L_z}} + \frac{1}{2} \cdot \frac{R_n + R_c}{R_n \cdot R_c} \cdot \sqrt{\frac{L_z}{C_b + C_d}}, \quad (19)$$

$$\beta_{III} = \frac{1}{2} R_z \sqrt{\frac{C_b}{L_z}} + \frac{1}{2 R_n} \cdot \sqrt{\frac{L_z}{C_b}}.$$

Skokowa odpowiedź czujnika może mieć charakter aperiodyczny lub periodyczny tłumiony.

Dla $\beta > \sqrt{m}$, aperiodyczną odpowiedź czujnika określa wyrażenie:

$$I_0(t) = I_F \cdot \frac{1}{m} \cdot \left[1 - \frac{1}{T_1 - T_2} (T_1 \cdot e^{-\frac{t}{T_1}} - T_2 \cdot e^{-\frac{t}{T_2}}) \right]. \quad (20)$$

Stałe czasowe T_1 i T_2 zależą od parametrów wewnętrznych i zewnętrznych czujnika; wartości ich są porównywalne lub jedna z nich przybiera wartość dominującą.

Dla $\beta = \sqrt{m}$ odpowiedź czujnika jest najszybszą z odpowiedzi aperiodycznych o stałej

czasowej $T = \frac{1}{\omega_0 \sqrt{m}}$ i określa ją wyrażenie:

$$I_0(t) = I_F \cdot \frac{1}{m} \cdot \left[1 - (1 + \omega_0 \sqrt{m} \cdot t) \cdot e^{-\omega_0 \sqrt{m} \cdot t} \right]. \quad (21)$$

Dla $0 < \beta < \sqrt{m}$, w odpowiedzi czujnika występują drgania tłumione o pulsacji $\omega = \omega_0 \cdot \sqrt{m - \beta^2}$ a przebieg $I_0(t)$ określa równanie:

$$I_0(t) = I_F \cdot \frac{1}{m} \cdot \left\{ 1 - e^{-\omega_0 \beta t} \cdot \left[\cos \omega_0 \sqrt{m - \beta^2} \cdot t + \frac{\beta}{\sqrt{m - \beta^2}} \cdot \sin \omega_0 \sqrt{m - \beta^2} \cdot t \right] \right\}. \quad (22)$$

Prowadzenie obliczeń modelu zgodnie z przyjętymi założeniami pozwala, bez konieczności rozwiązywania skomplikowanych równań nieliniowych, przewidzieć odpowiedź czujnika w zadanych warunkach oświetlania i obciążania.

5. WNIOSKI Z BADAŃ MODELU I PODSUMOWANIE WYNIKÓW ANALIZY

Przeprowadzone badania eksperymentalne rzeczywistych czujników fotoelektrycznych oraz analiza ich liniowych analogów potwierdziły słuszność przyjętych założeń upraszczających.

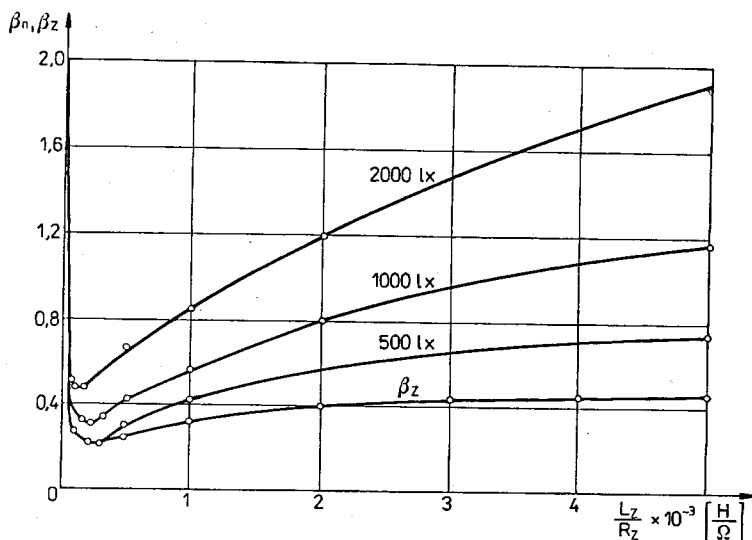
Badania i analiza dotyczą pomiarów w warunkach statycznych oraz dla światła przerywanego. Znajomość analitycznej postaci odpowiedzi skokowej czujnika umożliwia określenie jego odpowiedzi na dowolne wymuszenie $I_F(t)$. Wnioski można zatem uogólnić, rozszerzając je na inne przypadki promieniowania świetlnego o modulowanej wartości.

Niezależnie od typu złączeniowego czujnika fotoelektrycznego oraz postaci funkcji opisującej zmienność mierzonej wielkości świetlnej, wskazane jest zapewnienie w całym zakresie pomiarowym liniowości statycznej charakterystyki przetwarzania $I_0 = f(X_p, R_z)$, o możliwie dużym nachyleniu S .

W pomiarach wielkości szybko zmieniających się pożądanym jest ponadto krótki czas odpowiedzi. W omawianej grupie czujników osiągnięcie bardzo krótkich czasów odpowiedzi możliwe jest przy zastosowaniu fotodiody p-i-n. Inne czujniki mają większą pojemność, co ogranicza zakres ich zastosowań. Jednakże, fotoogniwa oraz fotodiody konwencjonalne cechuje wiele korzystnych parametrów metrologicznych [5, 7, 8]. Przeprowadzone badania wykazują, że poprzez zmianę charakteru obciążenia istnieje także możliwość poprawy ich właściwości dynamicznych.

Wprowadzenie do obciążenia czujnika indukcyjności może wpłynąć na zmianę jego czasu odpowiedzi. W zależności od wartości R_z , $\frac{L_z}{R_z}$ oraz poziomu oświetlania, obecność indukcyjności może:

- być szkodliwa, gdyż przyczynia się do wydłużenia czasu aperiodycznej odpowiedzi czujnika,
- nie powodować zmiany w charakterze przebiegu,
- być pożądana, gdyż zmniejszając $\beta \geq \sqrt{m}$ lub zwiększając $\beta < 1$, powoduje skrócenie czasu ustalania odpowiedzi.



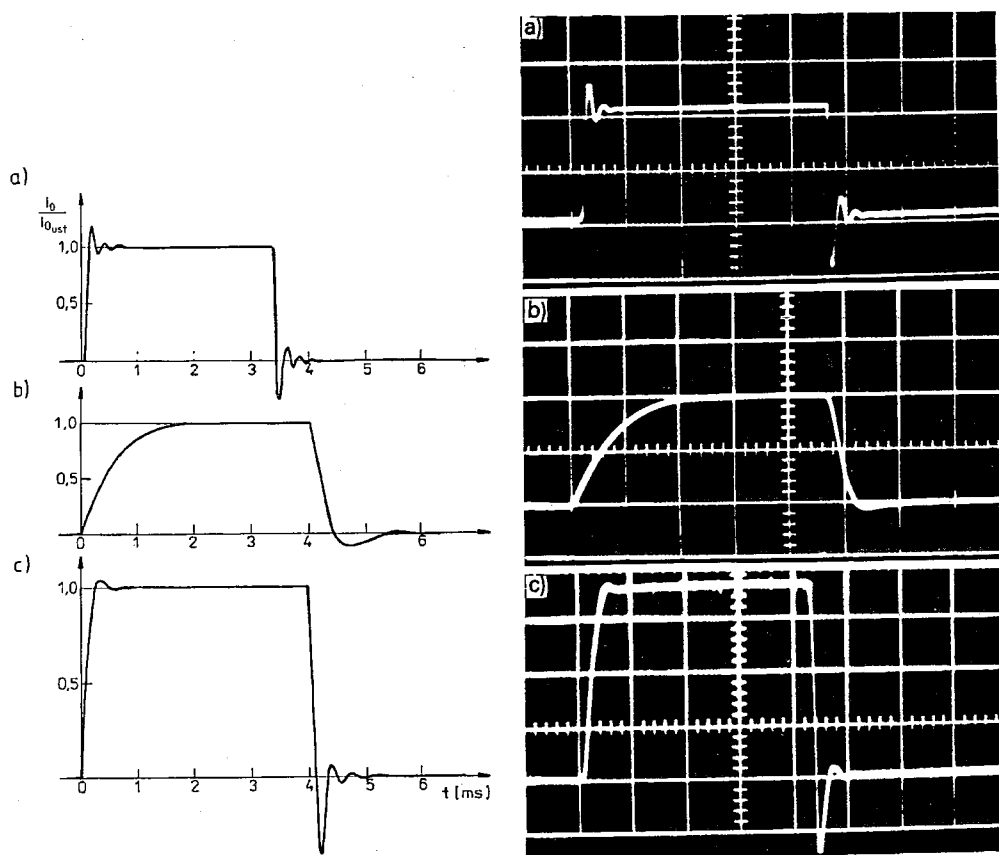
Rys. 5. Ilustracja wpływu indukcyjności obciążenia czujnika na ustalanie się jego odpowiedzi skokowej; β_n – współczynnik tłumienia przebiegu narastającego z oświetlaniem (X_p = natężenie oświetlenia E [lx]), β_z – współczynnik tłumienia przebiegu zanikającego z zaciemnianiem

Dla danego czujnika istnieją zalecane wartości $\frac{L_z}{R_z}$, przy których czas odpowiedzi jest najkrótszy. Na rys. 5 przedstawiono wykresy ilustrujące kształtowanie charakteru odpowiedzi $I_0(t)$ czujnika w zależności od zawartości indukcyjności w impedancji obciążenia. Wykresy te dotyczą badań przeprowadzonych dla przykładowego fotoogniwa krzemowego. W oparciu o wyniki badań tego czujnika, na rys. 6 zestawiono wybrane oscylogramy oraz wykresy przebiegów $I_0(t)$ otrzymane z obliczeń według wzorów (20) ÷ (22).

Poprawa właściwości dynamicznych umożliwia zastosowanie np. fotoogniw w układach sterujących i liczących lub fotodiod konwencjonalnych np. w układach współpracujących z laserowymi źródłami promieniowania.

Przedstawiony model złączowego czujnika fotoelektrycznego umożliwia analizę jego właściwości statycznych i dynamicznych z punktu widzenia przydatności do analizy wiarygodności przetwarzania wielkości w układzie otwartego łącza optycznego.

W rozdz. 1 scharakteryzowano grupy pomiarów, które można wykonywać w układzie tego typu. Przedmiotem obecnie prowadzonych badań są prace nad modelem otwartego łącza optycznego w systemie do wybranych nieinwazyjnych pomiarów medycznych [10]. Ten rodzaj pomiarów zalicza się do przedstawionej grupy c) a ogólną strukturę łącza ilustruje rys. 1. Parametry fotodetektora wywierają istotny wpływ na wiarygodność przetwarzania mierzonej wielkości. Można sformułować wniosek, że najbardziej przydatny w pomiarach byłby czujnik łączący zalety fotoogniwa i fotodiody. Przeprowadzona analiza i uzyskane dotychczas wyniki doświadczeń pozwalają na stwierdzenie, że jest to możliwe przy zastosowaniu fotodiody jednoz-



Rys. 6. Oscylogramy i wykresy przebiegów $I_0(t)$ dla fotoogniwa krzemowego oświetlonego prostokątnym impulsem żarowego światła przerywanego 0—E—0 i obciążonego impedancją R_z , L_z : a) $E=10000$ lx, $R_z=100\Omega$, $L_z=0,05$ H, b) $E=1000$ lx, $R_z=150\Omega$, $L_z=1,00$ H, c) $E=2000$ lx, $R_z=60\Omega$, $L_z=0,10$ H

łączowej, o dużej powierzchni światłoczułej, bez polaryzacji zewnętrznej, możliwie w stanie zwarcia, przy dostatecznie niskiej temperaturze.

BIBLIOGRAFIA

1. A. Cysewska: *Analiza stanów pracy selenowego ogniwa fotoelektrycznego obciążonego impedancją o charakterze indukcyjnym*. Praca doktorska, Politechnika Poznańska, 1978
2. A. Cysewska: *O dynamice układu z ogniwem fotoelektrycznym*. PAK, 1980 nr 3,
3. V. Härtel: *Optoelectronics. Theory and Practice*, Mc Graw Hill Book Company, USA 1978
4. M. H. Weik: *Fiber Optics and Lightwave Communications Standard Dictionary*. Van Nostrand Reinhold Company Inc., USA 1981
5. A. Pawlaczek: *Elementy i układy optoelektroniczne*, Warszawa: WKiŁ, 1984
6. A. Smoliński: *Optoelektronika światłowodowa*, Warszawa: WKiŁ, 1985

7. Z. O k r a j n i: *Szumy zespołu fotodioda-wzmacniacz w urządzeniach spektrofotometrycznych*, Probl. Techn. Med., 16/1, 1985
8. M. P o l o w c z y k: *Elementy i przyrządy półprzewodnikowe powszechnego zastosowania*, Warszawa: WKiŁ, 1986
9. M. K o n w i c k i, M. W ę g r z e c k i, R. G r o d e c k i: *BPYPO7 silicon measuring photocell*, 12 th Int. Symp. Techn. Comm. on Photon-Detectors, Varna, May 1986
10. A. C y s e w s k a - S o b u s i a k: *Dokładność pomiarów w systemie typu otwarte łącze optyczne*, Prace Naukowe IME Politechniki Wrocławskiej, nr 33, Seria Konferencje nr 15, 1989

A. CYSEWSKA-SOBUSIAK

MODEL OF THE PHOTOELECTRIC SENSOR AS THE PHOTODETECTOR IN THE SYSTEM OF THE OPEN OPTICAL COUPLER

In the paper the modified model for the photoelectric sensor of the junction structure type p-n or p-i-n has been presented. The results of the static and dynamic properties analysis and the assumption principles of the experimental verification for this model are described. Influence of the polarisation voltage, temperature, illumination level and the parameters of the load impedance are considered. In the presented application recommendations the conclusions for the realized investigations of the representative photodetectors have been taken.

Possibilities of digital TV picture in real time

NGUYEN KIM SACH

Instytut Radioelektroniki, Politechnika Warszawska

Received 1989.10.10

Authorized 1989.12.27

For the aim of TVpicture processing, picture is divided into subpictures and a module of subpicture processing is used. This paper presents the module of subpicture processing and it's use to picture processing in nonreal and real time.

1. INTRODUCTION

For digital TV picture processing in real time, the technique of very fast processing in real time is applied. It requires elevation of computing power in digital system of array processors. This system can be used for digital TV picture restoration, because its structure is suitable with tables of digital 2-D pictures and topology connection of elementary processors with computing structure in local operations. Array processors are considered as special processors.

Array processors belong to systems of processors SIMD (single instructure stream – multiple data stream) and MIMD (multiple instructure stream – multiple data stream) [1]...[9], by which each elementary processor is endowed in local memory and separated registers. Essential elements of the array processor are elementary array processors, input-output disposition, removed unit and they are coupled with computer. Due to progress in VLSI and UVLSI technology, computing power of array processors of small (16×16 pixels) or middle (32×32 , 64×64 pixels) arrays has become more and more large [4],[8],[10]...[13]. These processors are at present time recognized to belong to highly developed architectures in new generation of computers.

2. MODULE OF 64×64 PIXEL SUBPICTURE PROCESSING

If subpicture is presented by e.g. 64×64 pixels, then a whole active picture of 512×512 pixels contains 64 subpictures of 64×64 pixels. To process a 512×512 -pixel picture in real time, 64 array processors, which are in parallel activated, are utilized.

Each array processor presents a module of subpicture processing. In television, 25 pictures are performed per second (OIRT, CCIR), i.e. each picture persists in 40 ms, while an active picture only lasts for 29,9 ms. Therefore, the processing time of an active subpicture t_{spi} is:

$$t_{spi} \leq 29,9 \text{ ms}$$

and the processing time of a pixel is

$$t_{pix} \leq \begin{cases} 83 \text{ ns} & (\text{OIRT}) \\ 100 \text{ ns} & (\text{CCIR}) \\ 119 \text{ ns} & (\text{FCC}) \end{cases}$$

Table presents digital parameters of different TV standart and computers [14].

Parameter	CCIR 625 lin.	FCC 525 lin.	HDTV 1125 lin.	Computer
active lines	575	484	1050	-
format	$\frac{4}{3}$	$\frac{4}{3}$	$\frac{5}{3}; \frac{16}{9}$	$\frac{1}{1}$
active pixels /picture	640 × 431	634 × 338	1185 × 998	512 × 512
Whole number of active pix. /picture	0.27.10 ⁶	0.21.10 ⁶	1.18.10 ⁶	0.26.10 ⁶
dimension of pixel in horizon./mm	0.7	0.7	0.7	-

We can construct a module of subpicture processing by the use of a method of combination of histogram and median filter as described in the scheme shown in Fig. 1, 8-bit digital signals from the frame buffer pass the detector, which acts as a commutative unit transforming 8-bit signals to 1-bit signals prior to be counted.

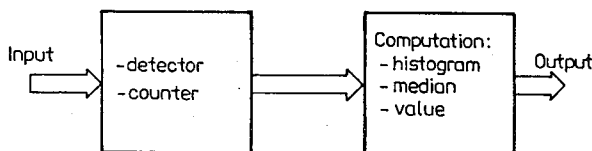


Fig. 1. Scheme of module of 64 × 64-pixel subpicture processing

The number of counters used will be 256. Each counter works as a register. Digital signals from the counter output are then processed into unit computing histograms, their threshold value and median value. Results are returned back to the frame buffer: If 5 × 5-pixel window is used, operation for combination of histogram and median filter (presented in [15]) can be made according to the following scheme, see Fig. 2.

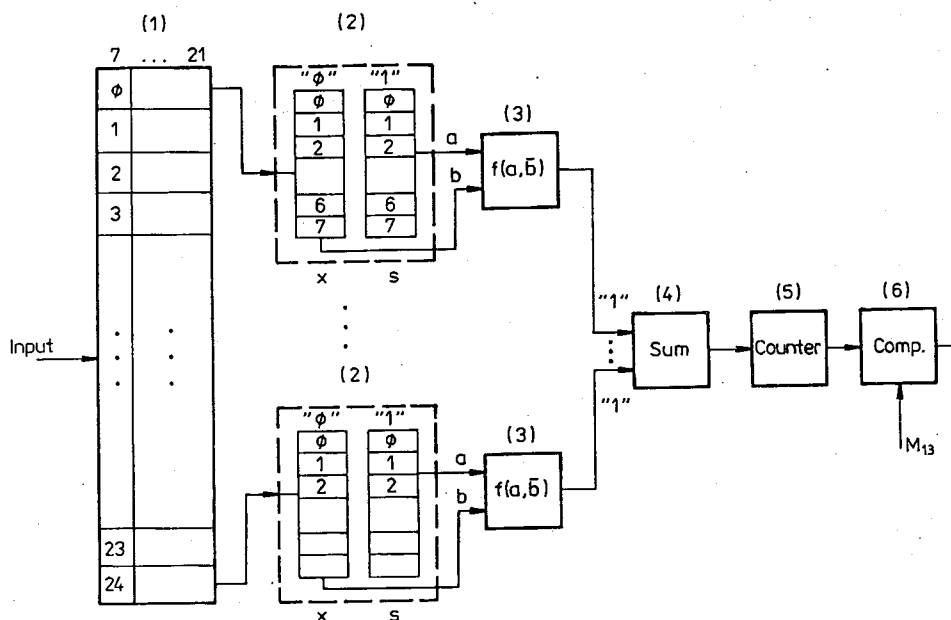


Fig. 2. Scheme of building procedure of histogram and median filter for a 5×5 -pixel window

Digital signals from shifted register (1) are led to registers (2) for arranging binary values of „0” and „1” before summing in unit (3). The blocks (2) and (3) contain 25 uniform units. Data from 25 sumators (3) are summed again in block (4) before reaching the counter (5). Block of comparator (6) selects median value M_{13} for 5×5 -pixel window. If the data from the counter are smaller than M_{13} and sum (4) is different from zero, the median is shifted to the left and 1 is added, and $S \leftarrow S \& X$. On the contrary, if data from the counter are larger than M_{13} and the sum is different from zero, the median is shifted to the left, but $S \leftarrow S \& \bar{X}$ and counter $\leftarrow$ (counter - sum). This scheme presented by the program in [15]) can be considered as a module of subpicture processing and is used for digital TV picture processing by sequential and parallel methods.

3. PROCESSING BY SEQUENTIAL METHOD

In our practical conditions, procedure of picture restoration and enhancement are carried out with sequential method [15]. Scheme in Fig. 3 presents blocks realised following this procedure. Each subpicture is processed in sequence and then stored in the frame buffer. The whole picture appears on the monitor screen at the end of the last subpicture processing.

Module of subpicture processing (block (2),(3),(4) — Fig. 2) repeats sequentially the processing cycle of subpictures. This method can be used for picture processing in nonreal time, because the time needed for processing of whole pictures is very

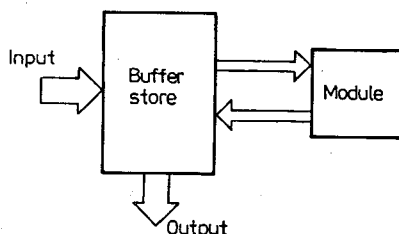


Fig. 3. Scheme of picture processing by sequence method using module of subpicture processing

large as compared to that used in television (Fig. 2). For digital TV picture processing in real time parallel method with parallel processor system can be used.

4. PROCESSING BY PARALLEL METHOD

Module of subpicture processing in Fig. 2 can be used for picture processing in television. Once the procedure of subpicture processing (i.e. picture restoration and enhancement) in real time (i.e. $t_{spi} < 29,9$ ms) is used for each 64×64 -pixel subpicture, the system of digital picture processing (e.g. 512×512 -pixel picture) will contain 64 processors. Such 64-processor system suits well to digital TV picture processing, Fig. 4. For that purpose, very fast array processors in the form of IC are used.

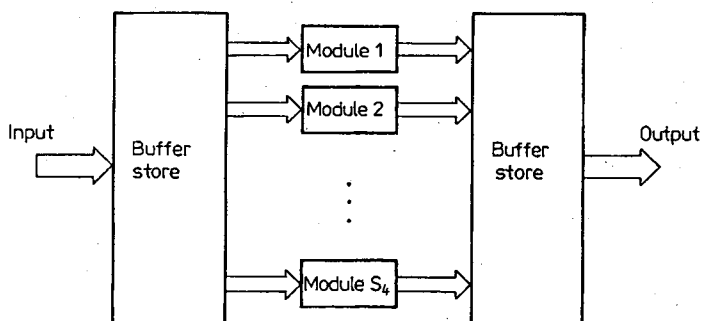


Fig. 4. Scheme of the 64-array processor system for TV picture processing

When processed subpictures are connected, distortions often appear at the edges of subpictures, because pixels in last windows from previous subpicture are desirable and black lines appear at the edges of subpictures. These black line can be considered as distortion. The distortions can be avoided by choosing a larger window and computing a smaller window from previous one. This means that at beginning the picture, more than two lines in horizontal and vertical direction should be computed. This method can also be used for sequential and parallel processing.

CONCLUSION

The present paper describes possibilities of sequential and parallel picture processing, which can be used for TV picture processing. The module of subpicture processing plays an important part. Traditional television is limited with 1-D signals. In the contrary, 2-D picture processed by computer methods is considerably suitable, because a 2-D picture can be seen as a table of data and the 2-D picture processing can be independent on standards and systems of colour television, which are now used in many countries and are very complicated. In the form of 64-array processor system, TV picture processing and procedures of picture restoration and enhancement can be developed with new procedures to increase its significance.

In brief, it is obvious that:

1. Utilizing the elaborated module of subpicture processing and system of parallel processors to improve the quality of TV pictures in real time is possible, if the array processors used are satisfactorily fast.
2. Improvement of quality of TV pictures is complicated and is not straightforward. Picture restoration and enhancement for television is, however, able to be successfully carried out, especially in the near future.
3. Digital 2-D picture restoration and enhancement can be applied to the TV purposes at the TV centers and at the relay TV stations or at the satellite TV stations.

REFERENCES

1. W. Cellary: *Systemy wielomikroprocesorowe, cz. I*, Informatyka, No. 11-12/1986, p. 1-3, 7
2. W. Iszkowski: *Architektury systemów n-mikroprocesorowych*, Informatyka, No. 7-8/1986, p. 12-15
3. J. Owczarczyk, M. Stolarski: *Architektura współbieżnych systemów przetwarzania obrazów*, Informatyka, No. 7-8/1986, p. 8-12
4. J. Owczarczyk, M. Stolarski: *Procesory tablicowe - Współbieżne systemy przetwarzania obrazów*, Informatyka, No. 11-12/1986, p. 8-11
5. G. Kaleta, J. Owczarczyk: *Wieloprocessorowa realizacja operacji przetwarzania obrazów*, Informatyka, No. 8/1988, p. 8-11
6. PIP-512/1024A. *Hardware manual: video digitizer board for IBM PC XT/AT*. Matrox Electronic Systems Limited, Canada, march 1986
7. D. Dreyer, K. Meyer, R. Berger: *Spezialprozessoren fuer die Echtzeit-Bildverarbeitung*. Teil I, II. Bild und Ton, No. 1/1989, p. 19-22; No. 2/1989, p. 46-50, 57
8. M. M. Sysło: *Maszyny i algorytmy równoległe, cz. I, II*. Informatyka, No. 10/1988, p. 5-7; No. 11-12/1988, p. 30-34
9. J. Deminiet: *Cm** - *Przykład komputera wieloprocessorowego o strukturze hierarchicznej*. Informatyka, No. 11-12/1986, p. 4-8, 30
10. A. Rosenfeld: *Parallel image processing using cellular arrays*, Computer, No. 1/1983, p. 14-20
11. P. Danielsson, S. Levialdi: *Computer architectures for pictorial information systems*, Computer, No. 11/1981, p. 53-67
12. K. Batchner: *Design of a massively parallel processor*. IEEE Trans. on Computer, Vol. C-29, No. 9/1980, p. 836-840

13. A. F i s h e r, H. K u n g: *Synchronizing large VLSI processor arrays*. IEEE Trans. on Computer, Vol. C-34, No. 8/1985, p. 734-740
14. M. C a l z i n i: *Hauptparameter von Kinofilm und Fernsehen sowie die Zone optimaler Sicht*. Bild und Ton, No. 3/1988, p. 85-88
15. N. K. S a c h: *Digital 2-D picture restoration and enhancement*, Rozprawy Elektrotechniczne, 1991, Vol. 2, z. 1-2

N. KIM SACH

MOŻLIWOŚCI POPRAWY JAKOŚCI CYFROWYCH OBRAZÓW TELEWIZYJNYCH W CZASIE RZECZYWISTYM

Streszczenie

W celu przetwarzania obrazów telewizyjnych można podzielić obraz na podobrazy i stosować moduł przetwarzania podobrazu. W artykule przedstawiono moduł przetwarzania podobrazu i jego zastosowanie do przetwarzania obrazów w czasie rzeczywistym i rzeczywistym.

Systolic circuits for fast median computation with application to sampling time jitter correctors

TOMASZ ADAMSKI

Instytut Podstaw Elektroniki, Politechnika Warszawska

Received 1990.02.23

Authorized 1990.06.20

In the paper systolic array circuits implementing sorting algorithm are presented. The input signal of these circuits is a sequence of N_0 real numbers, the output signal is nondecreasing ordered (i.e. sorted) input sequence. Along with ordering, each sorting algorithm chooses k -th (in the value) element of the ordered sequence then proposed arrays can be used to median and quantiles computations. Applications to sampling time jitter correctors reducing distortions caused by sampling time jitter in equivalent time sampling oscilloscopes are also described.

1. INTRODUCTION

The main aims of the paper are to propose fast systolic circuits implementing sorting algorithms and give some their applications in instrumentation. Specifically, we will show how to use described systolic arrays in sampling time jitter correctors which can be used to reduce distortion of the measured signal in equivalent time sampling oscilloscopes. Let N_0 denotes the length of the finite sequence $(a_n)_{n=1}^{N_0}$ of real numbers which are the input data to the sorting algorithm (or to implementing this algorithm circuit). Assume that $2|N_0$ (i.e. N_0 is even) and for simplicity a_n are integers from the interval $[0, m]$ where $m \in \mathbb{N}$. Sorting from formal point of view is equivalent to computing such a permutation $\pi: \{1, 2, \dots, N_0\} \rightarrow \{1, 2, \dots, N_0\}$ that $a_{\pi(1)} \leq a_{\pi(2)} \leq \dots \leq a_{\pi(N_0)}$.

The sequence $(a_{\pi(n)})_{n=1}^{N_0}$ is the output data. Sorting algorithm chooses k -th (in the value) element $a_{\pi(k)}$ of the ordered sequence for arbitrary $k \in \{1, 2, \dots, N_0\}$ then can be considered as a procedure for computing median, quantile or, in general, rank statistics. For fixed k algorithms of $a_{\pi(k)}$ computation are a little simpler than full sorting algorithms. A variety of sorting algorithms and their properties are detailed described in monographs [1] and [2]. Proposed below systolic circuits implement a parallelized modification of sorting algorithm called in [2] bubblesort procedure. In the simplified version used in the paper this procedure is the following:

```

procedure bubblesort (a);
  i, j: integer;
  a: array [0.. $N_0$ ] of real;
  for i=1 to  $N_0$  do
    for j=1 to  $N_0$  do
      if  $a[i] > a[i + 1]$  change positions  $a[i]$  and  $a[i + 1]$ .

```

Computational complexity of the above algorithm is proportional to N_0^2 . We can assume that N_0 is even, without loss generality because for N_0 odd, we can use a sorting circuit with $N_0 + 1$ inputs applying maximal number of the representation to one of them. On the other hand N_0 even gives natural symmetry and modularity to the circuit.

2. BASIC PROCESSING CELLS

Systolic arrays presented in the sequel consist of three kinds of elementary processing cells PC1, PC2, and PC3. The processing cell PC1 is a pair of multiplexers driven by a comparator and is described in Fig. 1. The PC2 cell, shown in Fig. 2 is one level of all proposed systolic circuits with parallel word processing and consists of N_0 PC1 cells and $2N_0$ registers. All registers have W -bits length equal to the length of processed words representing numbers. The processing cell PC3 (see Fig. 3) is a set of $N_0/2$ cascaded PC2 cells. The solution of PC1, PC2 and PC3 as circuits with parallel word processing is fast but rather expensive.

Besides, large number of pins (for large N_0) makes difficult on chip realization. More convenient for on chip implementation seems to be solution of the elementary processing cells as a serial word processing devices. The idea is similar to the distributed arithmetic concept used widely in DSP. The PC1, PC2 and PC3 cells with serial processing are shown respectively in Fig. 4, 5 and 3. The serial PC1 has a flip-flop FF reseted every W cycles of the clock. The first inequality in bit comparison indicating that the left side input is greater then the right one, set FF to 1. Two input multiplexers MUX choose right side input if controlled by $Q=0$ and left side input if $Q=1$. Circuits from Fig. 7 can be treated as modifications of the cell PC3 shown in Fig. 3.

Systolic circuits with serial word processing are not so fast as systolic arrays with paralel word processing (serial method is exactly W times slower) but for the same N_0 have W times less transistors and connecting pins.

3. SORTING CIRCUITS

Sorting circuits are systolic arrays using basic processing cells PC1, PC2 and PC3 described in previous section. We need $N_0/2$ cascaded PC2 cells or one PC3

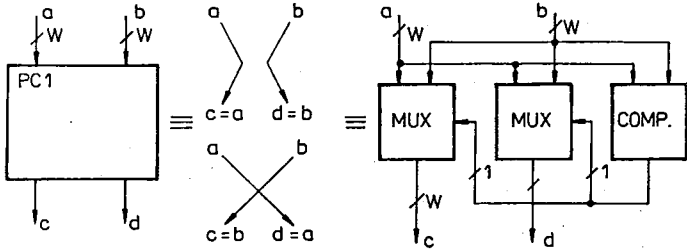


Fig. 1. The PC1 elementary cell

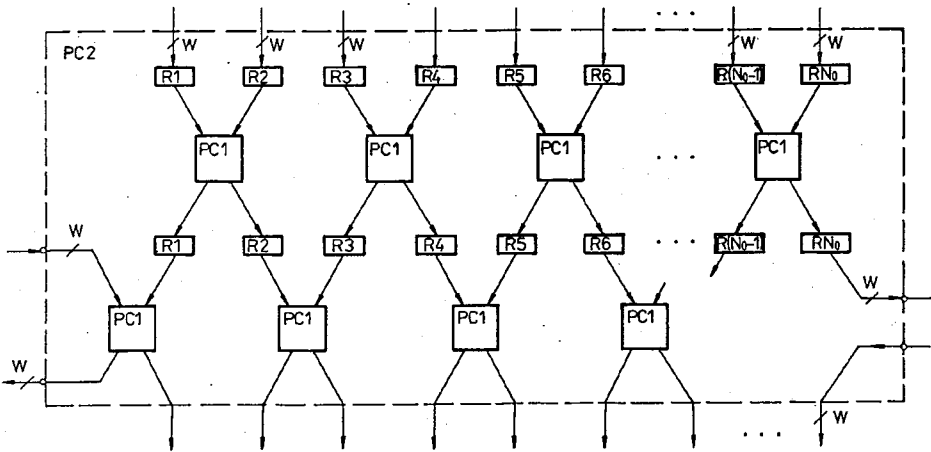


Fig. 2. The PC2 elementary cell-one level of the systolic array for sorting procedure with parallel word processing

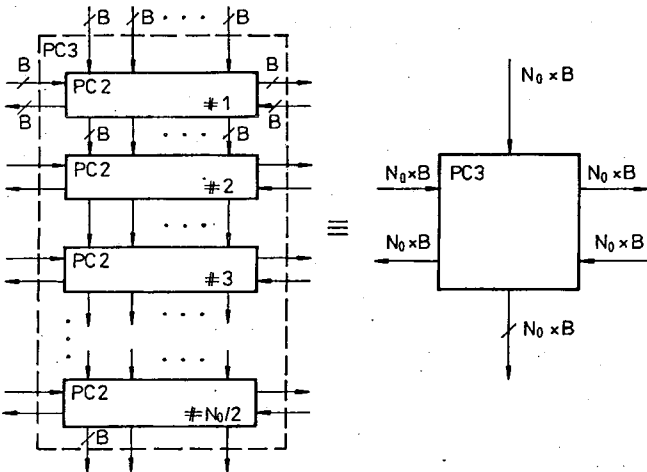


Fig. 3. The PC3 elementary cell, $B = W$ for parallel word processing and $B = 1$ for serial word processing

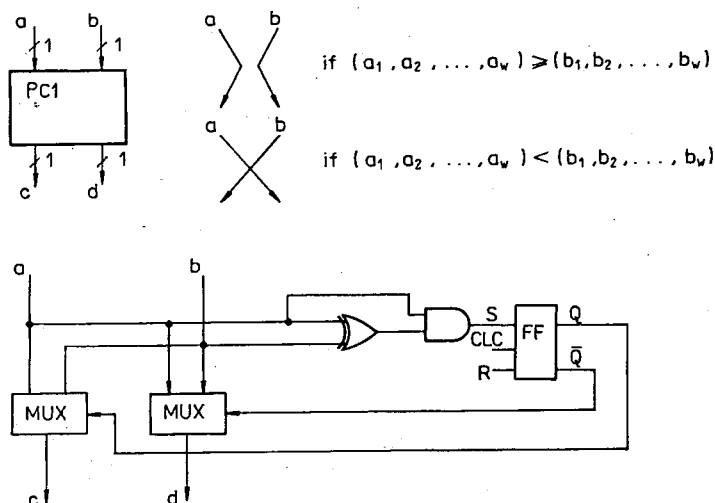


Fig. 4. The PC1 cell with serial word processing; $\leq$ is a lexicographic order, $a = (a_1, a_2, \dots, a_w)$; $b = (b_1, b_2, \dots, b_w)$; $c = (c_1, c_2, \dots, c_w)$; $d = (d_1, d_2, \dots, d_w)$

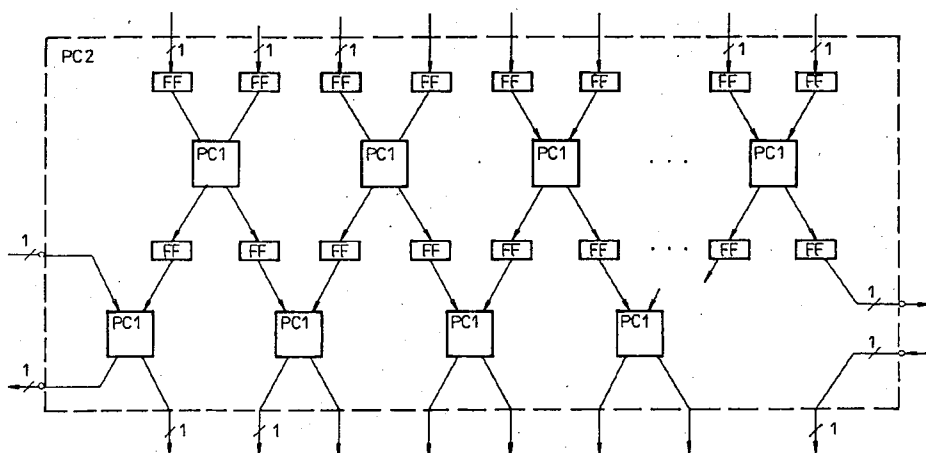


Fig. 5. The PC2 cell with serial word processing

appropriately connected (see circuits in Fig. 6 and 7) to sort N_0 binary numbers. From theoretical point of view these circuits implement lexicographic sorting of binary vectors. Circuits shown in Fig. 6 and 7 a), b) compute the output result (i.e. sorted sequence of numbers) in N_0 clock cycles in parallel case and WN_0 in serial solution. If the throughput of the circuit from Fig. 7a) is A then it is AM_0 and AN_0 appropriately for circuits shown in Fig. 7b) and 6. The number M_0 of levels (or more precisely cascaded PC2 cells) must be divisor of $N_0/2$. When $M_0 = N_0/2$ circuits from Fig. 6 and 7b) are equivalent. Shortcoming of the circuit from Fig. 7 is

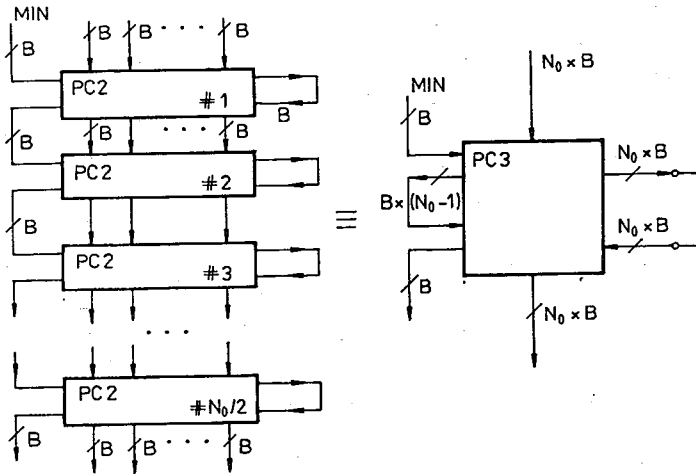


Fig. 6. The processing cell PC3 with basic N_0 level sorting connection, $B = W$ for parallel word processing and $B = 1$ for serial one

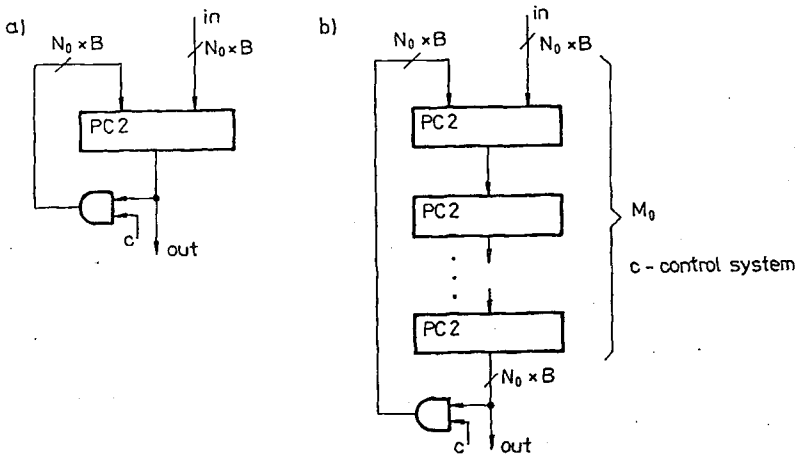


Fig. 7. One level a) and M_0 level connection b) of PC2 cells giving as a result sorting circuits $B = 1$ in serial word processing and $B = W$ for parallel one, $M|(N_0/2)$

more complicated synchronization of the data flow in comparison with the circuit from Fig. 6 but in the second concept we can easily exchange the throughput and complexity of the system. A method of extending the number N_0 of inputs is explained in the Fig. 8 and 9.

Described above circuits sort the all input sequence $(a_n)_{n=1}^{N_0}$. When we need only k -th element of the sorted sequence i.e. $a_{\pi(k)}$ (for example when rank statistics is computed), we can significantly simplify sorting systolic array, cutting surplus parts of the circuits, using the principle shown in Fig. 10. For example when the median is computed, we need only 1/4 part of the area of the full sorting circuit.

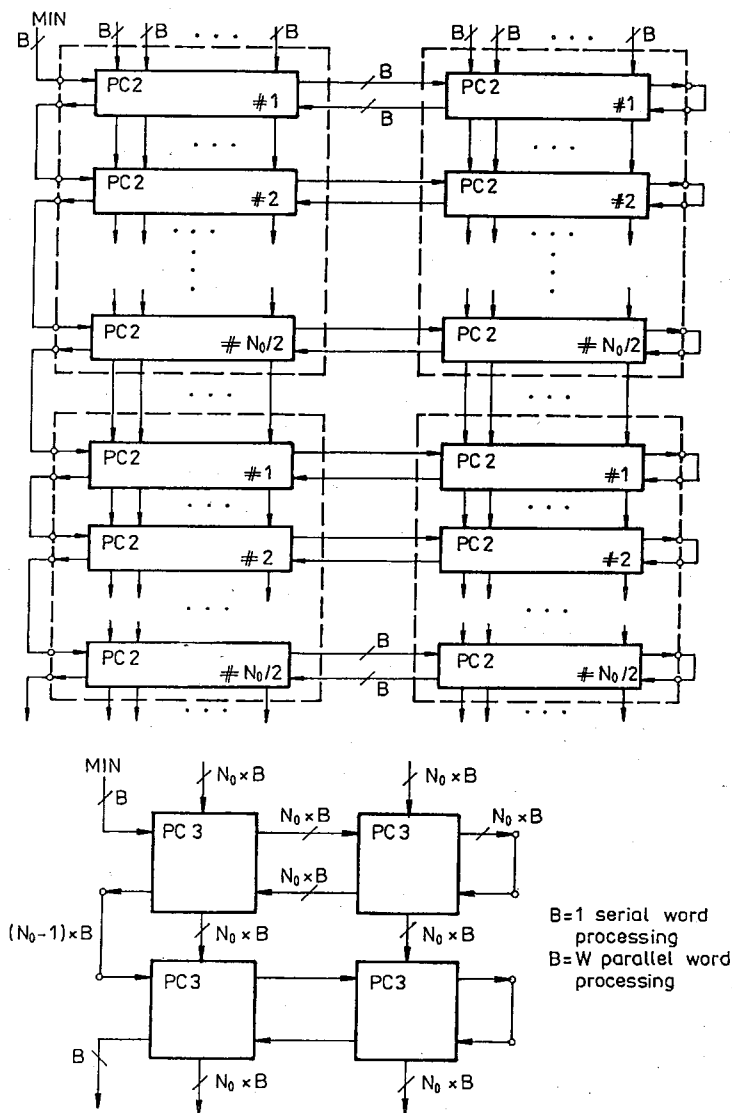


Fig. 8. The connection of our PC3 cells extending the input two times

4. APPLICATION TO SAMPLING TIME JITTER CORRECTORS

Applications of circuits for median and quantile computation, to image processing (so called median filters) are well known. In the sequel we will concentrate on applications to sampling time jitter correctors.

Assume, that an analog signal f is sampled with the frequency $1/\Delta t$. Sampling time points $k\Delta t$; $k \in \{1, 2, \dots, M\}$ are not exactly real sampling points. Because of noise phenomena in sampling circuit real sampling points are randomly shifted on the time axis. This shifting is called sampling time jitter.

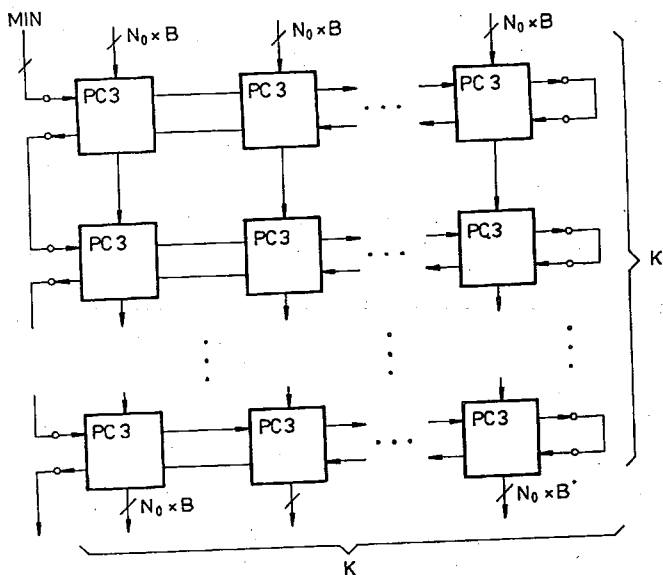


Fig. 9. The connection of K^2 PC3 cells extending the input K times

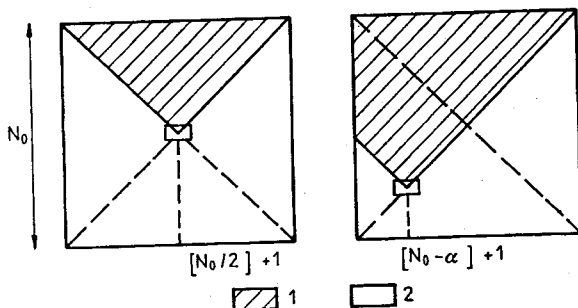


Fig. 10. Cutting sorting systolic array for a) median computation (choosing the element number $[N_0/2] + 1$ from $(a_n)_{n=1}^{N_0}$) b) quantile of the order α computation (choosing the element number $[N_0\alpha] + 1$).

Not hatched part of the circuit is surplus

In relatively slow systems these shifts are not critical and errors caused by jitter can be neglected. In fast systems, like broad band equivalent-time waveform recorders (ETWR) or sampling oscilloscopes random shifting is in the range 10–200 ps and causes very troublesome distortion of the measured signal. The aim of the sampling time jitter corrector is to reduce distortion introduced by jitter. In general, sampling in ETWR can be described for every $k=1,2,\dots,M$ by the finite random sequence $(j(k\Delta t + T_n(k\Delta t)))_{n=1}^{N_0}$ where N_0 is a number of samples taken for the point $k\Delta t$ and $T_n(k\Delta t)$ denotes the random shifting of the sampling time $k\Delta t$. It is reasonable to assume that for each $k=1,2,\dots,M$ random variables $T_1(k\Delta t), \dots, T_2(k\Delta t), \dots, T_{N_0}(k\Delta t)$ are independent with the same probability distribution. Statistics taken for every $k=1,2,\dots,M$ allow to correct distortion

caused by jitter but restoration needs deconvolution algorithm (see [3]) or in case when jitter statistics depends on k , solution of the integral Fredholm equation of the first kind. Then correction leads to the ill posed problem and needs complicated numerical regularization techniques (see [3],[4]). In this context, simple methods, like mean or quantiles computations, working quite correctly only under certain assumptions becomes attractive.

The simplest method is the mean value computation. In this method for each $k=1,2,...,M$ the number

$$MN(k\Delta t, N_0) = (f(k\Delta t + T_1(k\Delta t)) + f(k\Delta t) T_2(k\Delta t) + \dots + f(k\Delta t + T_{N_0}(k\Delta t))) / N_0$$

is computed. The method is quite correct for linear signal $f(t) = At + B$, $A, B \in R$ i.e. $MN(k\Delta t, N_0) \rightarrow f(k\Delta t)$ with probability 1 but in general (for example for pulse signals) it introduces errors which can be easily assessed in the following way.

Let $h(k\Delta t) \triangleq E(f(k\Delta t + T_n(k\Delta t)))$. If the sampling time jitter is limited (i.e. exist such a number $\varepsilon > 0$, that for every $k=1,2,...,M$

$|T_n(k\Delta t)| < \varepsilon$, $E(T_n(k\Delta t)) = 0$ and sampled signal $f \in C^{(2)}(R)$, $f, f', f'' \in L^1(R)$, then

$$\sup_{k \in \{1, \dots, M\}} |h(k\Delta t) - f(k\Delta t)| \leq \varepsilon^2 \sup_{t \in R} |f''(t)| / 2 \leq \varepsilon^2 \int_R \omega^2 |F(\omega)|_1 (d\omega) / (2\pi),$$

where $F(\omega)$ is the Fourier transform of the function f and l_1 Lebesgue measure on R .

Proposed below method of sampling jitter correction is based on quantiles computation and is a consequence of the following simple fact:

Fact 1. If for fixed $k \in \{1, 2, \dots, M\}$ zero is a quantile of order p , (where $0 < p < 1$) of the random variable $T_n(k\Delta t)$, there exists a real number $\varepsilon > 0$ such that $|T_n(k\Delta t)| \leq \varepsilon$ (i.e. jitter is limited) and function $f: [k\Delta t - \varepsilon, k\Delta t + \varepsilon] \rightarrow R$ is measurable and satisfy one of the following conditions (*) or (**)

(*) for every $t_1, t_2 \in R$; $(k\Delta t - \varepsilon \leq t_1 \leq k\Delta t \leq t_2 \leq k\Delta t + \varepsilon) \Rightarrow (f(t_1) \leq f(k\Delta t) \leq f(t_2))$

(**) for every $t_1, t_2 \in R$; $(k\Delta t - \varepsilon \leq t_1 \leq k\Delta t \leq t_2 \leq k\Delta t + \varepsilon) \Rightarrow (f(t_2) \geq f(k\Delta t) \geq f(t_1))$,

then 1) $(f(k\Delta t))$ is a quantile of the order p of the random variable $(f(k\Delta t + T_n(k\Delta t)))$,

2) If additionally distribution function of the random variable $(f(k\Delta t + T_n(k\Delta t)))$ is continuous and strict monotone increasing in a certain neighbourhood of $k\Delta t$ and f is continuous in a certain neighbourhood of $k\Delta t$ then $f(k\Delta t)$ is a unique quantile of order p of the random variable $f(k\Delta t + T_n(k\Delta t))$.

Let $p_1, p_2, \dots, p_M$ be numbers satisfying for every $k=1, 2, \dots, M$ conditions $P(T_n(k\Delta t) \leq 0) \geq p_k$, $P(T_n(k\Delta t) \geq 0) \geq 1 - p_k$ then we can say that for every $k \in \{1, 2, \dots, M\}$ zero is a quantile of order p_k of the random variable $T_n(k\Delta t)$. Let for every $k \in \{1, 2, \dots, M\}$ $\xi_j^{(N_0)}(k)$ denotes the rank statistics of order j for the random vector

$$(f(k\Delta t + T_1(k\Delta t)), f(k\Delta t + T_2(k\Delta t)), \dots, f(k\Delta t + T_{N_0}(k\Delta t))).$$

The algorithm of sampling jitter correction is the following. We compute statistics $\xi_{j(k)}^{(N_0)}(k)$ for $j(k) = [N_0 p_k] + 1$, for every $k=1, 2, \dots, M$ and admit that $f(t_k) \simeq \xi_{j(k)}^{(N_0)}(k)$. From above mentioned fact 1, we obtain that under additional (not specially restrictive) assumption $\xi_{j(k)}^{(N_0)}(k) \rightarrow f(k\Delta t)$ in probability P (all random variables are considered as

defined on the same probabilistic space $(\Omega, \mathcal{M}, P)$ then it is reasonable to admit $f(k\Delta t) \simeq \xi_{jk}^{(N_0)}(k)$ for sufficiently large N_0 . Mostly, random variables $T_n(k\Delta t)$ for $n \in N$ and $k=1,2,\dots,M$ are symmetrical then we have $p_1, p_2, \dots, p_M = 1/2$.

Software implementation of the discribed sampling jitter correction algorithm is very time consuming then using specialized systolic arrays for rank statistics computation seems to be very useful.

CONCLUSIONS

In the paper a variety of systolic arrays for sorting procedure were proposed. The high degree of modularity significantly facilitates their on chip implementation. The designer can easily exchange complexity of the system and its troughput. Presented circuits can be also used to quantiles computation. This feature give posibility of design sampling time jitter correctors based on quantiles computation and working in real time.

List of Symbols

$m|n$ — m is a divisor of n

Z — set of integers,

N — set of natural numbers,

R — set of real numbers,

REFERENCES

1. A. Aho, J. Hopcroft, J. Ullman: *The Design and Analysis of Computer Algorithms*; Addison-Weslay 1974
2. D. E. Knuth: *The Art of Computer Programing*; vol. 3. *Sorting and Searching*; Addison-Weslay 1973
3. W. Gans: *The Measurements and Deconvolution of Time Jitter in Equivalent-Time Waveform Samplers*; IEEE Proceedings; March 1983
4. C. Groetsch: *The Theory of Tikhonov Regularization for Fredholm Equations of the First Kind*; Pitman, 1984

T. ADAMSKI

UKŁADY SYSTOLICZNE DO OBLICZANIA MEDIANY I ICH ZASTOSOWANIE W KOREKTORACH JITTERU MOMENTU PRÓBKOWANIA

Streszczenie

W pracy zaproponowano kilka układów typu systolic array do sortowania skończonego ciągu liczb rzeczywistych lub w ujęciu statystycznym do obliczania statystyk porządkowych a w szczególności mediany i kwantyli z próby. Pokazano, że przedstawione układy mogą służyć w systemach pomiarowych z próbkowaniem do korekcji zniekształceń mierzonego sygnału spowodowanych tzw. jitterem momentu próbkowania. Zniekształcenia tego typu występują w rejestratorach szybkich impulsów i szerokopasmowych oscyloskopach próbkujących.

Eliminacja interferencji międzykanałowej

JERZY KISILEWICZ

Instytut Sterowania i Techniki Systemów, Politechnika Wrocławska

Otrzymano 1990.06.20

Autoryzowano do druku 1990.07.06

Opisano interferencję międzysymbolową występującą w kanale transmisji danych i pomiędzy kanałami. Rozpatrzono przypadek, gdy interferencja opisuje się macierzową transmitancją Z , której elementami są wielomiany. Udowodniono 3 twierdzenia dotyczące całkowitej eliminacji interferencji międzykanałowej. Dwa twierdzenia dotyczą przypadku, gdy elementy macierzowej transmitancji Z opisującej interferencję są dowolnymi wielomianami ograniczonego stopnia. W trzecim twierdzeniu elementy transmitancji są funkcjami dowolnymi, ale jednakowymi dla wszystkich elementów kolumny. Elementy jednej kolumny różnią się od siebie jedynie stałymi współczynnikami.

1. WSTĘP

W transmisji danych często zakłada się przesyłanie ciągu impulsów reprezentujących dane binarne. W kanale impulsy te są zniekształcane tak, że czas ich trwania (czas trwania odpowiedzi kanału na impuls) jest większy niż czas trwania impulsu wejściowego. Spróbkowany sygnał odbiorczy jest więc obciążony interferencją międzysymbolową, która może być eliminowana przy użyciu korektorów transwersalnych. Interferencja, która uwidacznia się próbkami odpowiedzi kanału na wejściowy impuls, może być opisana transformacją Z spróbkowanej odpowiedzi kanału $y(z) = y_0 + y_1 z^{-1} + y_2 z^{-2} + \dots$. Korekcja interferencji polega na skonstruowaniu korektora o transmitancji dyskretniej $c(z)$ takiej, że $c(z) \cdot y(z) \approx z^{-h}$ [1]. Dla liniowego prostego korektora transwersalnego $c(z) = c_0 + c_1 z^{-1} + \dots + c_{m-1} z^{-(m-1)}$ gdzie c_j jest wzmocnieniem na j -tym odczepie a m — liczbą odczepów. Parametry c_j korektora dobiera się tak, aby zminimalizować: maksymalną interferencję, średniokwadratową interferencję, lub średniokwadratowy błąd przy założonym poziomie zakłóceń [1]. W ogólnym przypadku na wejściu kanału cyfrowego nie musi pojawiać się ciąg bitów, lecz np.: ciąg dibitów (par bitów), tribitów czy znaków 5, 7 lub 8-bitowych. Na taki kanał można spojrzeć jak na kanał transmitujący równocześnie pewną liczbę ciągów binarnych. Poszczególne elementy binarne mogą interferować nie tylko z elementami własnego ciągu ale też z elementami innych ciągów. Mówimy wtedy o interferencji

wielowymiarowej. Z taką interferencją będziemy mieli też do czynienia w przypadku transmisji kilku strumieni danych w różnych pasmach częstotliwości, gdy składowe harmoniczne sygnałów przenikają pomiędzy kanałami [2]. Dwuwymiarowe korektory stosuje się np.: do oddzielnej korekcji składowej rzeczywistej i urojonej sygnału [3]. W niniejszej pracy rozważa się problem całkowitej eliminacji interferencji międzykanałowej.

2. INTERFERENCJA WIELOWYMIAROWA

Niech interferencję n -wymiarową opisuje macierz dyskretnych transmitancji $y_{ij}(z)$

$$Y(z) = \begin{bmatrix} y_{11}(z) & y_{12}(z) & \dots & y_{1n}(z) \\ y_{21}(z) & y_{22}(z) & \dots & y_{2n}(z) \\ \dots & \dots & \dots & \dots \\ y_{n1}(z) & y_{n2}(z) & \dots & y_{nn}(z) \end{bmatrix}.$$

Funkcje $y_{ij}(z)$ opisują wpływ elementów strumienia j na elementy strumienia i . Są to transformacje Z spróbkowanej odpowiedzi kanału i na wejściowy impuls w kanale j . Zadanie korekcji polega tu na skonstruowaniu takiego realizowalnego korektora

$$C(z) = \begin{bmatrix} c_{11}(z) & c_{12}(z) & \dots & c_{1n}(z) \\ c_{21}(z) & c_{22}(z) & \dots & c_{2n}(z) \\ \dots & \dots & \dots & \dots \\ c_{n1}(z) & c_{n2}(z) & \dots & c_{nn}(z) \end{bmatrix}$$

gdzie $c_{ij}(z)$ opisują korekcję interferencji ze strumienia j do strumienia i , aby iloczyn $E(z) = C(z)Y(z)$ opisujący wynikową interferencję jak najlepiej (w sensie wybranego kryterium) przybliżał $z^{-b} \cdot I$, gdzie I jest macierzą jednostkową.

Przykład 1.

Niech

$$Y(z) = \begin{bmatrix} 1 - wz^{-1} & a_{12}(1 + vz^{-1}) \\ a_{21}(1 - vz^{-1}) & 1 + wz^{-1} \end{bmatrix},$$

gdzie $|a_{12}| < 1$, $|a_{21}| < 1$, $|w| < 1$ i $|v| < 1$.

Ponieważ wielomiany $y_{ij}(z)$ posiadają tu maksymalne współczynniki przy wyrazie wolnym, przyjmujemy $h=0$. Aby $C(z) \cdot Y(z) \simeq I$, $C(z)$ powinno przybliżać macierz odwrotną do $Y(z)$, to jest

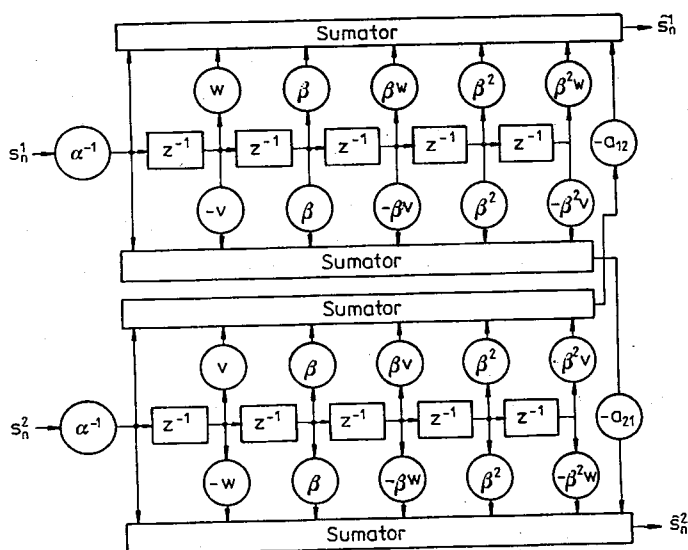
$$Y^{-1}(z) = \frac{1}{\alpha(1 - \beta z^{-2})} \begin{bmatrix} 1 + wz^{-1} & -a_{12}(1 + vz^{-1}) \\ -a_{21}(1 - vz^{-1}) & -1 + wz^{-1} \end{bmatrix}.$$

gdzie $\alpha = 1 - a_{12}a_{21}$ i $\alpha\beta = w^2 - v^2 a_{12}a_{21}$.

Zakładając, że $C(z)$ opisuje -2 wymiarowy prosty korektor transversalny o $m=6$ odczepach, poszukuje się wielomianów $c_{ij}(z)$ piątego stopnia przybliżających odpowiadające im elementy macierzy $Y^{-1}(z)$, które są tu funkcjami wymiernymi. Jedną z metod znajdowania $c_{ij}(z)$ polega na rozwinięciu elementów macierzy $Y^{-1}(z)$ w szereg dzieląc wielomian licznika przez wielomian mianownika $\alpha(1 - \beta z^{-2})$ i ograniczając wynik do m pierwszych wyrazów [1]. W ten sposób otrzymuje się

$$\begin{aligned}c_{11}(z) &= \alpha^{-1}(1 + wz^{-1} + \beta z^{-2} + w\beta z^{-3} + \beta^2 z^{-4} + w\beta^2 z^{-5}), \\c_{12}(z) &= -\alpha^{-1}(1 + vz^{-1} + \beta z^{-2} + v\beta z^{-3} + \beta^2 z^{-4} + v\beta^2 z^{-5})\alpha_{12}, \\c_{21}(z) &= -\alpha^{-1}(1 + vz^{-1} + \beta z^{-2} - v\beta z^{-3} + \beta^2 z^{-4} - v\beta^2 z^{-5})\alpha_{21}, \\c_{22}(z) &= \alpha^{-1}(1 - wz^{-1} + \beta z^{-2} - w\beta z^{-3} + \beta^2 z^{-4} - w\beta^2 z^{-5}).\end{aligned}$$

Schemat korektora pokazano na rys. 1.



Rys. 1. Schemat korektora 2 – wymiarowego do przykładu 1

Wypadkowa interferencja obliczona z iloczynu $C(z)Y(z)$ wynosi

$$E(z) = \begin{bmatrix} 1 - \beta^3 z^{-6} & 0 \\ 0 & 1 - \beta^3 z^{-6} \end{bmatrix} = (1 - \beta^3 z^{-6}) I.$$

Przyjmując np.: $a_{12} = a_{21} = 1/2$, $w = 1/2$, i $v = 1/4$ otrzymuje się $\alpha = 3/4$, $\beta = 5/16$ oraz $E(z) \cong (1 - 0.03z^{-6}) I$.

W powyższym przykładzie została całkowicie skorygowana interferencja pomiędzy strumieniami s_n^1 i s_n^2 , a interferencja pomiędzy symbolami tych samych strumieni została zredukowana z wartości w na pierwszym symbolu do wartości β^3 na szóstym symbolu. Zachodzi pytanie kiedy jest możliwe całkowite skorygowanie interferencji między strumieniami, czyli sprowadzenie interferencji wielowymiarowej do wielu interferencji jednowymiarowych.

Twierdzenie 1.

Jeżeli interferencja 2-wymiarowa $Y(z)$ opisuje się wielomianami $y_{ij}(z)$ stopnia mniejszego od m oraz istnieje macierz odwrotna $Y^{-1}(z)$, to możliwe jest całkowite skorygowanie interferencji pomiędzy strumieniami przy użyciu m odczepowego 2-wymiarowego prostego korektora transwersalnego (o strukturze pokazanej na rys. 1).

Dowód:

$$\text{Z założenia mamy: } Y(z) = \begin{bmatrix} y_{11}(z) & y_{12}(z) \\ y_{21}(z) & y_{22}(z) \end{bmatrix}$$

$$\text{oraz } Y^{-1}(z) = \frac{1}{y_{11}(z)y_{22}(z) - y_{12}(z)y_{21}(z)} \begin{bmatrix} y_{22}(z) & -y_{12}(z) \\ -y_{21}(z) & y_{11}(z) \end{bmatrix}.$$

Aproksymując ułamek przy $Y^{-1}(z)$ wielomianem $C_0(z)$ otrzymuje się

$$C(z) = \begin{bmatrix} C_0(z)y_{22}(z) & -C_0(z)y_{12}(z) \\ -C_0(z)y_{21}(z) & C_0(z)y_{11}(z) \end{bmatrix}$$

oraz $E(z) = C(z)Y(z) = C_0(z)[y_{11}(z)y_{22}(z) - y_{12}(z)y_{21}(z)]I$.

Jak widać macierz $E(z)$ jest diagonalna, co oznacza brak interferencji pomiędzy strumieniami danych. Zakładając najprostszy przypadek, że $C_0(z)$ jest wielomianem zerowego stopnia oraz $C_0(z) = c_0$ otrzymuje się elementy macierzy $C(z)$ w postaci wielomianów stopnia mniejszego od m , które są realizowane za pomocą korektorów o co najwyżej m odczepach, c.b.d.o.

Powyższe twierdzenie jest szczególnym przypadkiem twierdzenia dotyczącego interferencji wielowymiarowej.

Twierdzenie 2.

Jeżeli interferencja n wymiarowa $Y(z)$ opisuje się wielomianami $y_{ij}(z)$ stopnia mniejszego od m oraz istnieje macierz odwrotna $Y^{-1}(z)$, to możliwe jest całkowite skorygowanie interferencji pomiędzy strumieniami (skorygowanie interferencji n wymiarowej do n interferencji jednowymiarowych) przy użyciu $(n-1)(m-1)+1$ odczepowego n wymiarowego prostego korektora transwersalnego.

Dowód:

$$\text{Z założenia mamy [4]: } Y(z) = \begin{bmatrix} y_{11}(z) & \dots & y_{1n}(z) \\ \dots & \dots & \dots \\ y_{n1}(z) & \dots & y_{nn}(z) \end{bmatrix}$$

$$\text{oraz } Y^{-1}(z) = \frac{1}{\det Y(z)} \begin{bmatrix} Y_{11}(z) & \dots & Y_{1n}(z) \\ \dots & \dots & \dots \\ Y_{n1}(z) & \dots & Y_{nn}(z) \end{bmatrix},$$

gdzie $Y_{ij}(z)$ są algebraicznymi dopełnieniami elementów y_{ij} oraz [4]

$$Y_{ij}(z) = (-1)^{i+j} \sum_{< a_k} (-1)^{I_g} \prod_{k=i} y_k, \alpha_k(z),$$

$\alpha_1 \dots \alpha_n$ – permutacje liczb $1 \dots n$, w których $\alpha_1 = j$.

I_g – liczba całkowita zależna od permutacji $\alpha_1 \dots \alpha_n$.

Zatem $Y_{ij}(z)$ są sumami iloczynów $n-1$ wielomianów stopnia nie większego od $m-1$, czyli $Y_{ij}(z)$ są wielomianami stopnia nie większego od $(n-1)(m-1)$. Aproksymując ułamek $1/\det Y(z)$ wielomianem $C_0(z)$ otrzymuje się

$$C(z) = C_0(z) \begin{bmatrix} Y_{11}(z) & \dots & Y_{1n}(z) \\ \dots & \dots & \dots \\ Y_{n1}(z) & \dots & Y_{nn}(z) \end{bmatrix}$$

oraz $E(z) = C_0(z) \cdot \det Y(z) \cdot T$

Jak widać, macierz $E(z)$ opisująca wynikową interferencję jest diagonalna, a zakładając, że wielomian $C_0(z) = C_0$ jest zerowego stopnia, otrzymuje się elementy macierzy $C(z)$ w postaci wielomianów stopnia nie większego od $(n-1)(m-1)$, które opisują n wymiarowy prosty korektor transwersalny o $(n-1)(m-1)+1$ odczepach. c.b.d.o.

Twierdzenie 3.

Jeżeli n wymiarowa interferencja opisana macierzą $Y(z)$ jest taka, że dla każdego i oraz j zachodzi $y_{ij}(z) = a_{ij} y_{jj}(z)$ gdzie a_{ij} są liczbami oraz jeżeli istnieje macierz odwrotna $[a_{ij}]^{-1}$, to możliwe jest całkowite skorygowanie interferencji pomiędzy strumieniami za pomocą n wymiarowych korektorów transwersalnych o dowolnej liczbie odczepów. W szczególności za pomocą korektora 1 odczepowego.

Dowód:

Z założenia mamy

$$Y(z) = \begin{bmatrix} y_{11}(z) & a_{12}y_{22}(z) & \dots & a_{1n}y_{nn}(z) \\ a_{21}y_{11}(z) & y_{22}(z) & \dots & a_{2n}y_{nn}(z) \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1}y_{11}(z) & a_{n2}y_{22}(z) & \dots & y_{nn}(z) \end{bmatrix} =$$

$$= \begin{bmatrix} 1 & a_{12} & \dots & a_{1n} \\ a_{21} & 1 & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & 1 \end{bmatrix} \begin{bmatrix} y_{11}(z) & 0 & \dots & 0 \\ 0 & y_{22}(z) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & y_{nn}(z) \end{bmatrix}$$

$$\text{oraz } \begin{bmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \dots & \alpha_{nn} \end{bmatrix} = \begin{bmatrix} 1 & a_{12} & \dots & a_{1n} \\ a_{21} & 1 & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & 1 \end{bmatrix}^{-1}.$$

Łatwo wykazać, że

$$Y^{-1}(z) = \begin{bmatrix} y_{11}^{-1}(z) & 0 & \dots & 0 \\ 0 & y_{22}^{-1}(z) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & y_{nn}^{-1}(z) \end{bmatrix} \begin{bmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \dots & \alpha_{nn} \end{bmatrix}.$$

Niech $C_i(z)$ będzie realizowalną aproksymacją funkcji $y_{ii}^{-1}(z)$. Wówczas

$$C(z) = \begin{bmatrix} C_1(z) & 0 & \dots & 0 \\ 0 & C_2(z) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & C_n(z) \end{bmatrix} \begin{bmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \dots & \alpha_{nn} \end{bmatrix}$$

oraz

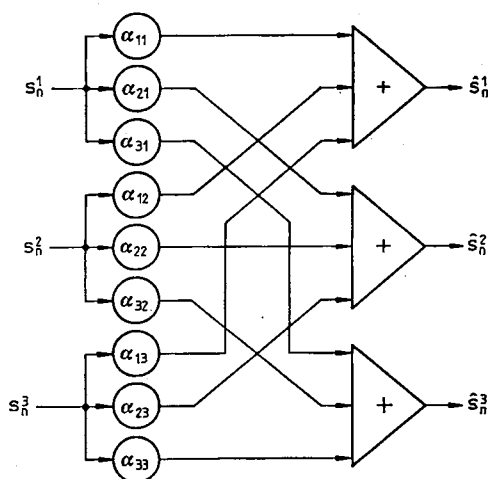
$$C(z)Y(z) = \begin{bmatrix} C_1(z)y_{11}(z) & 0 & \dots & 0 \\ 0 & C_2(z)y_{22}(z) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & C_n(z)y_{nn}(z) \end{bmatrix}.$$

W szczególności gdy $C_i(z) = c_i$ są stałymi, to $C(z)$ jest realizowane przy użyciu jednodziewowego korektora. c.b.d.o.

Schemat korektora jednodrzewowego korygującego 3 – wymiarową interferencję do 3 interferencji 1 – wymiarowych

$$C(z)Y(z) = \begin{bmatrix} y_{11}(z) & 0 & 0 \\ 0 & y_{22}(z) & 0 \\ 0 & 0 & y_{33}(z) \end{bmatrix}.$$

dla $C_1(z) = C_2(z) = C_3(z) = 1$ pokazano na rys. 2.



Rys. 2. Jednodrzewowy korektor 3 wymiarowy

ZAKOŃCZENIE

W pracy wykazano, że n wymiarową interferencję można skorygować do n interferencji jednowymiarowych jeśli tylko n wymiarowy korektor posiada odpowiednio dużą liczbę odczepów. W szczególnym przypadku gdy interferencja spełnia założenia twierdzenia 3, liczba odczepów nie jest istotna. W przykładzie 1 już dwuodrzewowy dwuwymiarowy korektor całkowicie kompensuje interferencję pomiędzy strumieniami. I tak dla

$$C(z) = \alpha^{-1} \begin{bmatrix} 1 + wz^{-1} & -a_{12}(1 + vz^{-1}) \\ -a_{21}(1 - vz^{-1}) & 1 - wz^{-1} \end{bmatrix}$$

otrzymuje się $C(z)Y(z) = (1 - \beta z^{-2})I$. Redukcja interferencji własnej w strumieniach jest tu jednak mniejsza niż w przykładzie 1.

Zastosowanie korektora jednodrzepowego

$$C(z) = \alpha^{-1} \begin{bmatrix} 1 & -a_{12} \\ -a_{21} & 1 \end{bmatrix}$$

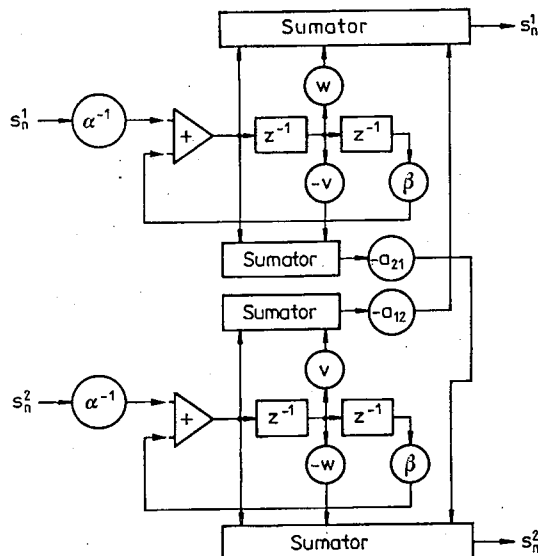
nie kompensuje całkowicie interferencji pomiędzy strumieniami i daje

$$C(z)Y(z) = \begin{bmatrix} 1 - (w - a_{12}a_{21}v)z^{-1} & -a_{12}(w - v)\alpha^{-1}z^{-1} \\ a_{21}(w - v)\alpha^{-1}z^{-1} & 1 + (w - a_{12}a_{21}v)\alpha^{-1}z^{-1} \end{bmatrix}.$$

co dla przykładowych danych $a_{12}=a_{21}=1/2$, $w=1/2$, i $v=1/4$ wynosi

$$C(z)Y(z) = \begin{bmatrix} 1 - 7/12z^{-1} & -1/6z^{-1} \\ 1/6z^{-1} & 1 + 7/12z^{-1} \end{bmatrix}.$$

Jakkolwiek założono w pracy użycie prostego korektora transversalnego, to udowodnione twierdzenia zachowują swoją moc dla innych realizacji korektorów niż za pomocą prostych filtrów transversalnych. Inne (niż liczba odczepów) będą warunki dokładnej realizacji transmitancji wielomianowych. Różnie też można realizować transmitancję wymierną $1/\det Y(z)$. Gdy np.: w tw. 2 założy się, że wielomian $\det Y(z)$ posiada wszystkie zera na płaszczyźnie zespolonej wewnątrz koła jednostkowego, to transmitancję $1/\det Y(z)$ można dokładnie zrealizować za pomocą korektora ze sprzężeniem zwrotnym o liczbie odczepów równej stopniowi wielomianu $\det Y(z)$ plus 1 [1]. Odpowiedni korektor $C(z)=Y^{-1}(z)$ z przykładu 1 miałby strukturę pokazaną na rys. 3. Korektor ten jest stabilny wtedy i tylko wtedy gdy $|\beta| < 1$ [1].



Rys. 3. Schemat 2-wymiarowego korektora ze sprzężeniem zwrotnym — do przykładu 1

BIBLIOGRAFIA

1. A. P. Clark: *Advanced Data-Transmission Systems*. London Press Limited, 1977
2. I. Korn: *Simple Expression for Interchannel and Intersymbol Interference Degradation in MSK Systems with Application to Systems with Gaussian Filters*. IEEE Trans. on Com., August 1982, vol. COM-30, No. 8, p. 1968-1972
3. D. P. Taylor, M. Shafi: *Decision Feedback Equalization for Multipath Induced Interference in Digital Microwave LOS Links*. IEEE Trans. on Com., March 1984, vol. COM-32, No. 3, p. 267-279
4. A. Mostowski, M. Stark: *Elementy algebry wyższej*. PWN Warszawa 1965

J. KISILEWICZ

INTERCHANNEL INTERFERENCE ELIMINATION

Summary

Intersymbol and interchannel interference is presented in this paper. The case when the interference is described by Z-transform matrix which contains only polynomials is considered. Three theorems concerned with interchannel interference elimination have been proofed. In two of them a greatest degree of polynomials are preassumed and limited. In the third theorem common functions for each column of Z-transform matrix are preassumed — different elements of each columns are equal to different constants multiplied by common and preassumed function.

Organizacja zadania wizualizacji w komputerowych systemach automatyki

ANNA CZEMPLIK

Instytut Cybernetyki Technicznej, Politechnika Wrocławska

Otrzymano 1990.05.25

Autoryzowano do druku 1990.07.10

Praktyka wizualizacji procesów w komputerowych systemach automatyki przemysłowej doprowadziła do wyłonienia się pewnych standardów. W artykule opisano te standardy w zakresie organizowania wizualizacji i dostępu do informacji. Przedstawiono również konkretne rozwiązanie problemu w systemie IP – 2000, opracowanym w Instytucie Cybernetyki Technicznej Politechniki Wrocławskiej.

1. WSTĘP

Wizualizacja w systemach automatyki przemysłowej polega na przejrzystym zobrazowaniu stanu automatyzowanego procesu technologicznego i stanu samego systemu automatyki. W systemach analogowych zadanie to było realizowane poprzez rozbudowane tablice synoptyczne i pulpity sterownicze. W układach komputerowych odbywa się to przede wszystkim poprzez konsole operatorskie, zawierające monitory i klawiatury funkcyjne.

Na początku lat 70-tych amerykańska firma Honeywell przeprowadziła intensywny program badań, w trakcie którego obserwowano pracę operatorów procesu technologicznego. Badania polegały między innymi na ilościowym określeniu czynności jakie są wykonywane w dyspozytorni. Zwracano uwagę na technikę wywoływania obrazów używaną przez różnych operatorów, liczono częstotliwość wykonywania poszczególnych operacji, analizowano, które z dostępnych informacji są wykorzystywane. Tymi sposobami starano się stworzyć jak najpełniejszy obraz pracy operatora w przyszłej dyspozytorni. Najważniejsze spostrzeżenia i wnioski jakie wypłynęły z tych badań są następujące:

- operator nie jest w stanie przyjąć zbyt dużej ilości informacji wyświetlanych jednocześnie,
- operacje jakie wykonuje człowiek w odpowiedzi na żądania systemu są przeprowadzane sekwencyjnie nawet jeśli informacja o nich jest wyświetlana równolegle,

- pożądane jest przedstawienie w zwartej formie informacji o całości procesu,
- korzystne jest wykorzystywanie obrazów symbolicznych i graficznych form przedstawiania wartości pozwalających na szybkie orientowanie się w sytuacji na obiekcie,
- wskazania cyfrowe są wykorzystywane tylko wtedy gdy istotne są informacje szczegółowe,
- dane dostępne dla operatora powinny być przystosowane do wyświetlania w grupach tworzących małe podsystemy automatyzowanego procesu,
- pomocne w kierowaniu procesem są wykresy czasowe (krzywe trendów i historii),
- używanie wielkości pochodnych może być efektywniejsze niż korzystanie bezpośrednio z wartości mierzonych,
- wszystkie funkcje na konsoli operatorskiej powinny być dostępne w zasięgu ręki operatora.

Wyniki badań firma uwzględniła przy konstrukcji własnego systemu automatyki TDC 2000. Od tego czasu powstało wiele różnych systemów i ich modyfikacji, opracowanych przez różne inne firmy. Pomimo tej różnorodności istnieje pewien standard obrazowania informacji dla zastosowań w automatyce przemysłowej. Jest on wynikiem zebranych w tym czasie doświadczeń oraz nawiązaniem do sprawdzonych rozwiązań stosowanych jeszcze w systemach analogowych. Poniżej zostaną przedstawione te właśnie wspólne cechy różnych systemów. Są one związane z rodzajami informacji występującymi w systemie automatyki, sposobami dostępu do niej i z formami przedstawiania tej informacji na obrazie.

2. RODZAJE INFORMACJI W SYSTEMIE AUTMATYKI

Istotną cechą systemów komputerowych stosowanych w automatyce jest możliwość przedstawienia człowiekowi (operatorowi procesu) stanu automatyzowanego procesu lub obiektu technologicznego za pomocą wartości już przetworzonych i uporządkowanych. Ponieważ dąży się do tego aby opracowane systemy mogły być wykorzystywane dla automatyzacji różnych procesów, istnieje pewien uniwersalny sposób uporządkowania danych dostępnych w systemie. Sposób ten pozwala wyróżnić i korzystać z następujących typów informacji:

- opis stanu punktu automatyki,
- zdarzenie.
- opis stanu elementu systemu.

Opis stanu punktu automatyki zawiera pewien zestaw parametrów ściśle związanych ze sobą. Mogą to być wartości liczbowe, wartości dwustanowe i teksty. Punkt automatyki może stanowić na przykład urządzenie (zawór z siłownikiem, który go otwiera i zamyka, silnik, pompa), obwód regulacji, element automatyzowanego procesu, zestaw obliczanych wskaźników, itp. Punkt stanowi podstawowy element wizualizacji w systemie. Wyróżnia się kilka typów punktów w zależności od zestawu informacji, który jest konieczny do pełnego opisu stanu punktu. Jednym z podstawowych typów jest punkt regulacji, opisany przez następujące parametry:

wartość zadana, wartość bieżąca, wartość sterująca, zakres zmienności, proggi ograniczeń, jednostka pomiarowa i nastawy regulatora.

Innym typem jest punkt binarny służący do zobrazowania stanu zaworów, napędów dwukierunkowych, itp. Stan takiego punktu opisują dwie wartości dwustanowe odpowiadające stanom stabilnym (np. zamknięty, otwarty) i dwie wartości dwustanowe określające kierunek zmian (np. zamyka się, otwiera się).

Wszystkie punkty automatyki dostępne w systemie są podzielone na grupy, np. po osiem punktów. Grupa najczęściej składa się z punktów powiązanych ze sobą technologicznie i stanowi mały podsystem automatyzowanego procesu. Każdy punkt może być przyporządkowany do kilku grup jednocześnie. Jeśli w systemie występuje wiele grup mogą zostać one podzielone na bloki (najczęściej również z uwzględnieniem powiązań technologicznych).

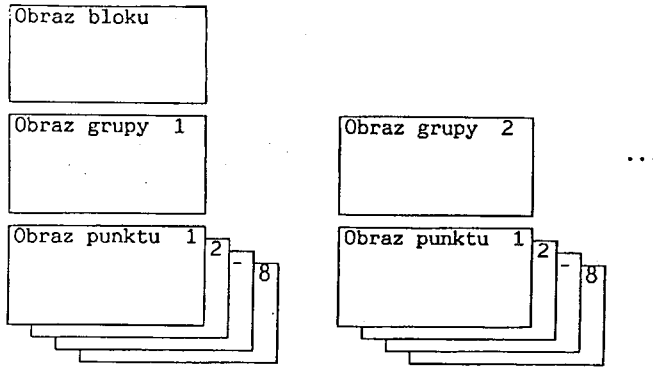
Drugim wyróżnionym typem informacji są zdarzenia. Są one wynikiem ciągłej, automatycznej kontroli stanu procesu, prowadzonej przez system komputerowy i sygnalizują zaistnienie określonych sytuacji, np. *„przekroczenie górnego ograniczenia wartości w określonym punkcie”*, *„zmiana wartości zadanej w danym obwodzie regulacji”*, *„zamknięcie zaworu”*, *„zmiana nastaw regulatora”*, itp. Informacje tego typu zawierają czas wystąpienia i słowny opis zdarzenia, jakie miało miejsce. Są one wykorzystywane do tworzenia chronologicznych raportów, opisujących historię przebiegu procesu. Wystąpienie zdarzenia może być związane z koniecznością zasygnalizowania stanu alarmowego. Za pomocą zdarzeń opisuje się nie tylko zaistnienie sytuacji alarmowych ale również wprowadzenie czy pojawienie się w systemie istotnych zmian wartości lub stanów określonych punktów automatyki.

Trzeci typ informacji dostępny w komputerowych systemach automatyki jest związany z opisem stanu samego systemu. Informacje tego typu są przeznaczone dla obsługi systemu do celów serwisowych. Określają one konfigurację sieci systemu, położenie poszczególnych elementów systemu na obiekcie i stan sygnałów pomocniczych opisujących elementy systemu (uszkodzenie modułu systemu, przerwanie połączenia, itp.).

3. ORGANIZACJA DOSTĘPU DO INFORMACJI

Do standardu wizualizacji można zaliczyć kolejną wspólną cechę różnych systemów automatyki – hierarchiczną organizację dostępu do informacji. Ponieważ operator i tak nie jest w stanie śledzić dużej ilości informacji jednocześnie, ze względów oszczędnościowych stosuje się ograniczoną liczbę monitorów a informacje przedstawia się w formie wygodnych zestawień. Powszechnie stosowana jest zasada przeglądu obrazów od ogółu do szczegółu. Im obraz dotyczy większego zakresu automatyzowanego procesu tym zawiera informacje ogólniejsze. Szczegółowe wartości są dostępne na odpowiednio niższych poziomach. Taka organizacja dostępu do informacji w systemie pozwala uwolnić operatora od konieczności interpretowania wielu wskaźników charakteryzujących obserwowany proces. Zadanie to może być realizowane przez komputer a operator może ograniczyć się do obserwacji gotowych

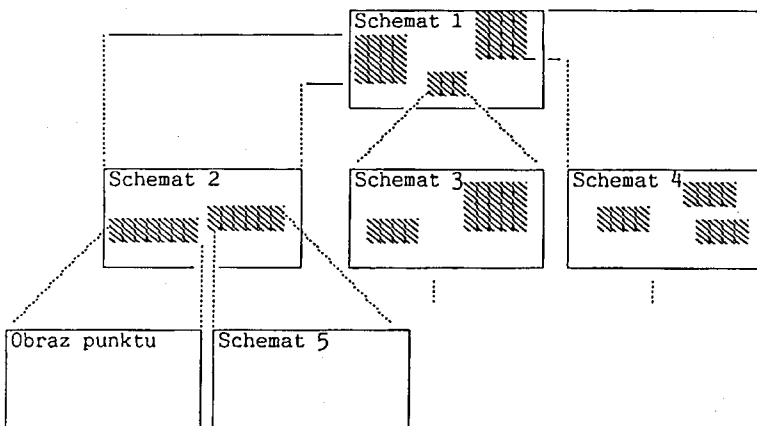
rezultatów i dopiero konieczność jego interwencji wymaga przejrzenia dokładniejszych informacji.



Rys. 1. Typy obrazów dla różnych poziomów obserwacji

Uniwersalna hierarchia dostępu do informacji w systemie opiera się na podziale wszystkich parametrów na punkty automatyki, grupy i bloki. Dla każdego takiego poziomu obserwacji przewidziany jest oddzielny typ obrazu (rys. 1). Obraz bloku, jako najbardziej ogólny, zawiera tylko informacje jakościowe — jest to poglądowe przedstawienie aktualnego stanu wszystkich punktów danego bloku. Na tym obrazie można zobaczyć np., że w określonym punkcie automatyki wystąpiło przekroczenie dolnego progu ograniczenia wartości mierzonej ale nie jest podane ile to przekroczenie wynosi. Więcej informacji można uzyskać na kolejnych obrazach. Ważniejsze wartości parametrów są podawane na obrazie grupy a na obrazie punktu można zobaczyć wszystkie informacje jakie są dostępne w systemie o określonym punkcie.

Oprócz uniwersalnych obrazów bloku, grupy i punktu, w komputerowych systemach automatyki, stosowane są również obrazy synoptyczne (schematy techno-



Rys. 2. Przykładowa struktura wizualizacji z wykorzystaniem obrazów synoptycznych

logiczne), ściśle związane z automatyzowanym procesem. Obrazy tego typu zawierają schematyczny rysunek procesu i pola w których wyświetlane są wartości pomiarów lub stany punktów binarnych. Obrazy synoptyczne można również wykorzystać do tworzenia hierarchicznych struktur dostępu do informacji. Można tak zaprojektować schematy aby stanowiły one coraz dokładniejsze przedstawienia fragmentów na kolejnych poziomach obserwacji i uzupełnić obrazami grup lub pojedynczych punktów. Przykładowa struktura tego typu została przedstawiona na rysunku 2.

Sprzętowo, wywoływanie żadanego zestawu informacji jest realizowane poprzez wykorzystanie klawiatury funkcyjnej, która umożliwia wprowadzanie określonych rozkazów do systemu. Podstawowy sposób zmiany wyświetlanego obrazu, polega na wybraniu klawisza funkcyjnego, który określa typ obrazu i podaniu odpowiednich parametrów (np. numer grupy, numer schematu).

4. FORMY GRAFICZNE

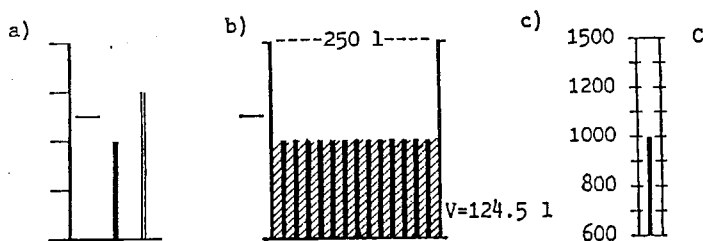
Opisane powyżej typy informacji, sposoby jej uporządkowania i dostępu do niej można określić jako podstawowe elementy stosowanego powszechnie standardu wizualizacji w komputerowych systemach automatyki. Standaryzacja jednak została posunięta dalej aż do zastosowania określonych form graficznych w konstrukcji poszczególnych typów obrazów.

Podstawowym elementem obrazów jest wartość wyświetlana w postaci liczby. Dla przekazania dodatkowych informacji wykorzystuje się możliwość kolorystycznego zróżnicowania wyświetlanych wartości zgodnie z zasadą ważności informacji lub oznaczeń technologicznych. Przy stosowaniu zasady ważności informacji, kolorami rozróżnia się stany punktu automatyki, które są określane na podstawie zmierzonej wartości aktualnej i założonych granic zmienności tej wartości. Wszystkie wartości, które aktualnie mieszczą się w założonych granicach są wyświetlane w tym samym kolorze. Przekroczenie tego zakresu jest sygnalizowane zmianą barwy. Najczęściej przyjmuje się wyświetlanie wartości poprawnych w kolorze zielonym, wartości przekraczających zadane progi w kolorze żółtym i wartości, które wywołują alarm w kolorze czerwonym. Gdy stosowana jest zasada oznaczeń technologicznych poszczególne wartości wyświetlane są w kolorze określonym dla materiału technologicznego, jaki występuje w danym punkcie automatyki, np. wartości dotyczące wody oznacza się kolorem niebieskim, powietrza — białym, itd. Dla wielu materiałów odpowiednie kolory są określone przez normy państwowe.

Liczbowa postać wartości jest to jednak forma w pełni wykorzystywana tylko wtedy, gdy operator potrzebuje szczegółowej informacji o danej wielkości. W trakcie normalnej pracy informacja musi być tak przekazana aby operator miał możliwość szybkiego zorientowania się w stanie prowadzonego procesu. Do tego celu wykorzystuje się graficzne formy przedstawiania danych. Pozwalają one właśnie na szybkie oszacowanie wartości, porównanie kilku parametrów, określenie tendencji i szybkości zmian wartości. Powszechnie stosowane są do tego celu wykresy czasowe jednej lub kilku zmiennych i wykresy słupkowe.

Przebieg czasowy pojedynczej wielkości jest wygodną formą przedstawiania informacji o wartości np. podczas regulacji ręcznej, gdy operator zmieniając parametr w obwodzie regulacji może zobaczyć w odpowiedzi reakcję obiektu na wykresie czasowym. Wykresy czasowe kilku zmiennych są często wykorzystywane na obrazach grupy, ponieważ zawiera ona zwykle obraz punktów powiązanych ze sobą technologicznie i korzystna jest możliwość porównania wartości tych punktów w czasie. Wybór czasu obserwacji i częstotliwości pobierania wartości jest ograniczony i zależy od wymagań technologicznych oraz od możliwości samego systemu automatyki. Najczęściej czas zbierania wartości wynosi 20 s, 1 min, 20 min, 1 h, 2 h, 8 h lub 24 h. W bardziej zaawansowanych systemach tworzy się jednocześnie historię krótko-, średnio- i długoterminową wartości wybranych punktów, np. 20 s, 20 min i 24 h. Maksymalna częstotliwość pobierania wartości do tworzenia historii jest uzależniona od częstotliwości uaktualniania danych w systemie. Ograniczeniem jest również ilość pamięci jaką można w danym systemie przeznaczyć na przechowywanie historii. Im więcej punktów, których wartości należy zbierać, im dłuższy czas obserwacji i większa częstotliwość zbierania wartości tym oczywiście potrzeba więcej pamięci. Z tego względu, jeśli nie ma specjalnych wymagań technologicznych, nie przechowuje się długoterminowej, dokładnej historii wartości punktów a wykorzystuje się raczej możliwość jednoczesnego zbierania krótkoterminowej historii z dużą częstotliwością uaktualniania i mniej dokładnej historii z dłuższych okresów czasu. Dla historii 24-godzinnej przy częstotliwości zbierania danych 20 s trzeba pamiętać 4320 wartości. Natomiast jeśli będziemy przechowywać tylko dokładną historię z 20 minut i historię z 24 godzin uaktualnianą co 20 minut to wystarczą 132 wartości. Przedstawienie zmian wartości w wizualizacji z taką dokładnością jest najczęściej wystarczające.

Wykresy słupkowe służą do wygodnego zobrazowania bieżącego stanu punktu analogowego, np. punktu regulacji (rys. 3a). Na jednym wykresie może być przed-

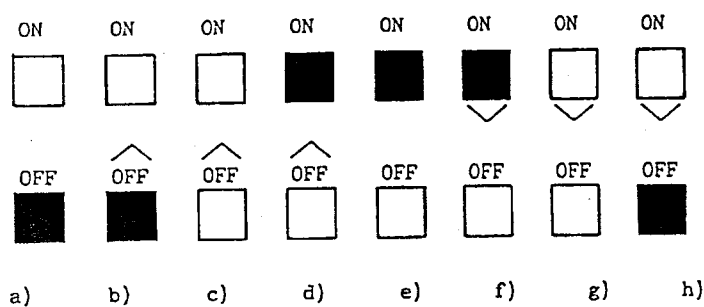


Rys. 3. Wykresy słupkowe do zobrazowania bieżących stanów punktów analogowych: a) wartości, b) stopnia wypełnienia, c) temperatury lub ciśnienia

stawionych kilka parametrów dotyczących jednego punktu. Dla punktu regulacji umieszcza się tam, jak to pokazano na rysunku, wartość zadaną (oznaczoną poziomą kreską), wartość bieżącą (lewy słupek) i wartość sterującą na wyjściu regulatora (prawy słupek). Parametry mogą być uzupełnione przedstawieniem progów ograniczeń. Ten sposób prezentacji pozwala operatorowi zorientować się we wzajemnych relacjach pomiędzy różnymi wielkościami bez konieczności porównywania wartości liczbowych. Wartości są przedstawiane w przedziale 0–100% zakresu zmienności.

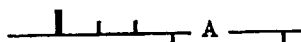
Wykresy słupkowe są stosowane zarówno na uniwersalnych obrazach bloku, grupy i punktu jak i na schematach technologicznych. Szczególnym przypadkiem tych wykresów jest, stosowany na obrazach synoptycznych, sposób oznaczenia stopnia wypełnienia zbiorników (rys. 3b), wskazania temperatury lub ciśnienia (rys. 3c), itp.

Jako odpowiednik wykresów słupkowych dla punktów binarnych stosuje się formę graficzną przedstawioną na rysunku 4a...h. Jest ona przeznaczona dla obrazowania stanów punktów typu napęd dwukierunkowy i pozwala pokazać stany stabilne urządzenia i kierunek zmian. W formie uproszczonej jest wykorzystywana do przedstawienia punktów dla których określone są tylko stany stabilne (rys. 4a,e). Na obrazach synoptycznych stosuje się również symbole technologiczne odpowiednich urządzeń, np. zawór, a ich stan określa się poprzez zróżnicowanie użytych kolorów i słowny opis stanu.



Rys. 4. Forma graficzna przedstawienia stanu punktu binarnego

Wśród powszechnie stosowanych form graficznych można wyróżnić jeszcze jedną. W systemach, w których występuje obraz bloku stosuje się symboliczne przedstawienie stanu grupy (rys. 5). Cała grupa (8 punktów) przedstawiona jest w postaci jednego odcinka. Składa się on z pojedynczych fragmentów odpowiadających



Rys. 5. Symboliczne przedstawienie stanu grupy punktów

kolejnym punktom. Jeśli wszystkie punkty w grupie są poprawne to przedstawienie tej grupy będzie miało postać zielonej linii. Wystąpienie stanu alarmowego lub innego sygnalizowanego stanu powoduje zastąpienie odpowiedniego fragmentu odcinka innym symbolem lub innym kolorem. Na przykład do oznaczenia stanów przekroczenia zadanych progów ograniczeń stosuje się różnej wielkości słupki pojawiające się na linii. W każdym systemie jest określony zestaw informacji i odpowiadających im symboli, które są wykorzystywane do określenia stanu punktu. Tego typu przedstawienie pozwala na jednym ekranie monitora pokazać jednocześnie stan kilkuset punktów. Przy odpowiednim wyborze tych punktów można w ten sposób przedstawić ogólny stan dużego fragmentu automatyzowanego procesu.

Bardzo ważnym elementem wizualizacji jest sygnalizowanie wystąpienia stanów alarmowych. Funkcja ta jest realizowana w systemach na standardowych obrazach

punktu, grupy i bloku oraz na obrazach synoptycznych, poprzez wykorzystanie umownych symboli i kolorów. Standardowo do oznaczenia alarmu wykorzystuje się kolor czerwony. W związku z tym praktycznie nie używa się tego koloru do innych celów, np. rysowania elementów schematu synoptycznego. Wystąpienie zdarzenia, które jest związane z włączeniem alarmu powoduje włączenie sygnału dźwiękowego i pojawienie się w odpowiednim miejscu na obrazie czerwonego, migającego znacznika. Potwierdzenie przyjęcia informacji o alarmie przez operatora kasuje sygnał dźwiękowy i wyłącza miganie elementu na ekranie. Znacznik w stałym, czerwonym kolorze pozostaje na obrazie do chwili zaniku stanu alarmowego.

5. WIZUALIZACJA W SYSTEMIE IP – 2000

Jako przykład rozwiązania wizualizacji w zastosowaniach przemysłowych zostanie przedstawiona realizacja tego zadania w komputerowym systemie automatyki przemysłowej IP – 2000. stanowi ona opracowanie własne Instytutu Cybernetyki Technicznej Politechniki Wrocławskiej. Rozwiązanie to jest zgodne z przedstawionym standardem a ponadto oferuje dodatkowe możliwości, które zostaną przedstawione w dalszej kolejności.

Podstawowymi zadaniami programu wizualizacji są:

- wywołanie obrazu wybranego przez operatora,
- okresowe uaktualnianie wyświetlanych obrazów.

Program ten jest przystosowany do pracy w środowisku wielozadaniowym. Dla małych obiektów cały system automatyki, łącznie z wizualizacją może być realizowany w konfiguracji z jednym mikroprocesorem. Przy większych obiektach program wizualizacji realizują wydzielone stacje graficzne, które są połączone z systemem poprzez łącze transmisji szeregowej.

Z punktu widzenia operatora interesujące są dwie wielkości, określające czas generowania obrazu i częstotliwość uaktualniania wyświetlanego obrazu. Określenie minimalnych wartości tych czasów w znacznym stopniu zależy od konkretnej konfiguracji systemu, tzn. od ilości punktów jakie system obsługuje i ilości zadań, które są wykonywane przez jeden mikroprocesor. Program wizualizacji ma stosunkowo niski priorytet wśród zadań wykonywanych w systemie. Ważniejsze są programy realizujące transmisję i przekazywanie wartości do bazy danych, procedury kontroli stanu procesu i generowania zdarzeń lub programy archiwizacji wartości. Wykonywanie programu wizualizacji może być przerywane przez każde zadanie o większym priorytecie, dlatego czas generowania obrazu jest wielkością losową. Zależy ona od ilości zmiennych elementów jakie należy wyświetlić na obrazie i czasu, zajętego przez ważniejsze zadania, które przerywały generowanie obrazu. Przyjmuje się że czas tworzenia obrazu jest mniejszy niż 1 s. Jedynie tworzenie obrazów w graficznym trybie pracy monitora, przy dużej ilości danych, takich jak obraz historii grupy pomiarowej (8 wykresów czasowych z dużej ilości pomiarów), jest wykonywane w czasie 1 – 2 s.

Częstotliwość uaktualniania obrazu jest związana z wielkością odpowiadającą stałym czasowym najszybciej zmieniających się wartości, dla których chcemy

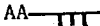
obserwować zachodzące zmiany na ekranie. Minimalny czas jest określany w trakcie konfigurowania i strojenia systemu, tzn. dobierania wielkości priorytetów i częstotliwości wywoływania poszczególnych zadań w systemie. Przyjmuje się że częstotliwość odtwarzania obrazu jest minimalnie rzędu co 1 s. Okres zbierania wartości w bazie danych czy w archiwacji jest od niego niezależny i może być krótszy.

BLOK : 1

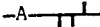
90.08.21 11:01



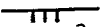
grupa 1
Zbiorniki wody



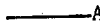
AA-grupa 6
Masa podstawowa



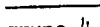
grupa 2
Domieszki suche 1



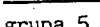
grupa 3
Domieszki suche 2



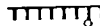
A-grupa 7
Masa zwrotna



grupa 4
Rezerwa 1



grupa 5
Rezerwa 2



grupa 8
Zawory P1



grupa 9
Zawory P2

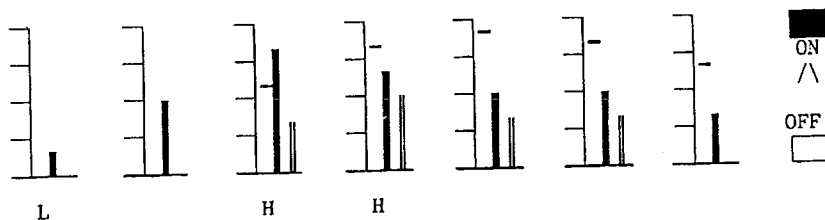


grupa 10
Zawory GP

Rys. 6. Przykładowa forma obrazu bloku

BLOK : 1
GRUPA: 1

90.08.21 11:01



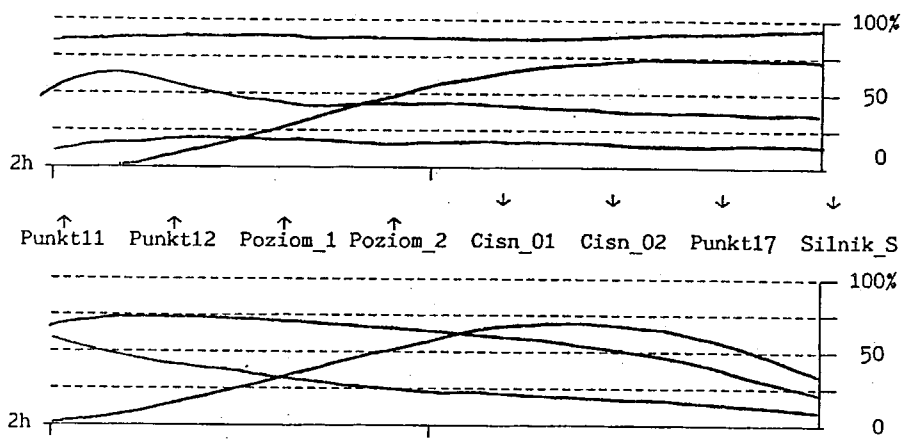
Punkt11 Punkt12 Poziom_1 Poziom_2 Cisl_01 Cisl_02 Punkt17 Silnik_S

SP			60.05	79.00	92.98	80.01	65.00	
PV	15.01	50.00	88.00	70.05	50.10	50.00	31.01	ON
OUT			30.00	48.90	31.00	30.92		

Rys. 7. Przykładowa forma obrazu grupy punktów

BLOK : 1
GRUPA: 1

90.08.21 11:01



Rys. 8. Przykładowa forma obrazu punktu

BLOK : 1
GRUPA: 1
PUNKT: 3

90.08.21 11:01



H

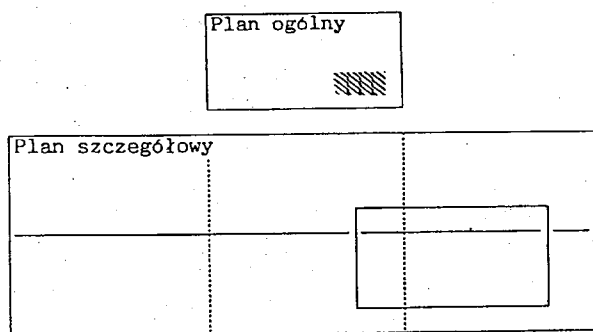
Poziom_1 Poziom wody zimnej w zbiorniku dolnym

SP	60.00 %	12.00 [m]	RL	0.000	RH	20.00 [m]
PV	88.00 %	17.60 [m]	L1	3.000	H1	17.00 [m]
OUT	30.00 %		L2	1.000	H2	18.50 [m]

Rys. 9. Przykładowa forma obrazu punktu

Podstawowy, uniwersalny zestaw obrazów systemu stanowią przedstawienia punktu automatyki, grupy punktów i bloku. Na rysunkach 6÷9 zostały pokazane przykładowe formy tych obrazów. Maksymalna zawartość grupy i bloku jest ograniczona przez rozmiar ekranu monitora. Przy założeniu że monitor pracuje w trybie 80*25 znaków (640*200 punktów graficznych), grupa składa się maksymal-

nie z 8 punktów a blok może zawierać do 45 grup. Zestaw obrazów uniwersalnych może być modyfikowany oraz uzupełniony przez schematy technologiczne konkretnych procesów. Sposoby dostępu do informacji są zgodne ze standardem opisanym powyżej, czyli uniwersalna hierarchia dostępu na bazie obrazów punktu, grupy i bloku (rys. 1) oraz poprzez wykorzystanie obrazów synoptycznych, tworzących ciąg kolejnych przybliżeń przedstawianego procesu technologicznego (rys. 2). Tworzenie i poprawianie wszystkich typów obrazów jest wykonywane przy pomocy specjalizowanego programu edytora. W systemie IP – 2000 istnieje możliwość tworzenia i wykorzystania dużych obrazów synoptycznych, tzn. większych niż powierzchnia jednego ekranu. Wybór fragmentu, który ma być widoczny na ekranie jest dokonywany poprzez płynne przesunięcie lub bezpośrednie określenie współrzędnych. Możliwość ta pozwala zorganizować wizualizację procesu na bazie ogólnego i szczegółowego planu (rys. 10). Wszystkie istotne informacje o procesie są zawarte na dużym schemacie technologicznym. Stanowi on szczegółowe rozwinięcie ogólnego planu procesu, który jest pomocny przy wywołaniu na ekran monitora określonego fragmentu całości. Maksymalny wymiar dużego schematu może odpowiadać powierzchni 16 ekranów monitora.



Rys. 10. Schemat wizualizacji procesu tworzenia obrazów synoptycznych

Rozszerzenie funkcji systemu umożliwia, oprócz standardowego sposobu zmiany wyświetlanego obrazu przy pomocy klawiszy funkcyjnych, wykorzystanie klawiszy kierunkowych kursora. Na każdym z obrazów dostępnych w systemie można zdefiniować pola aktywne dla kursora i przyporządkować im dowolny obraz, który zostanie wyświetlony po potwierdzeniu położenia kursora w obrębie tego pola. Daje to możliwość uproszczonego wywoływania kolejnych obrazów w hierarchii dostępu do informacji według określonej ścieżki.

Program realizujący wizualizację procesu technologicznego w systemie IP – 2000 pozwala, poprzez system okien programowych, wywoływać na jednym ekranie kilka różnych obrazów jednocześnie. Wszystkie wyświetlane w oknach obrazy są na bieżąco uaktualniane. Ponieważ program wizualizacji jest przystosowany do działania w środowisku wielozadaniowym, system wyświetlania informacji na ekranie uwzględnia występowanie tzw. okien systemowych. Są one przeznaczone dla innych programów systemu, które działają w tym samym czasie, umożliwiając im również niezależne wyświetlanie własnych informacji na ekranie – na przykład raport

zdarzeń, czyli chronologiczne zestawienie pojawiających się w procesie stanów awaryjnych i wprowadzanych zmian.

Rozwiązanie wizualizacji w systemie IP-2000 opiera się na stosowaniu uniwersalnego programu i pewnej liczby zbiorów konfiguracyjnych, które definiują ostateczną postać obrazowania informacji dla konkretnego automatyzowanego procesu.

PODSUMOWANIE

Obserwując rozwój komputerowych systemów automatyki przemysłowej w zakresie wizualizacji procesu można powiedzieć że zmiany, które się pojawiają mają charakter jedynie niewielkich modyfikacji, przystosowania do aktualnej mody (system okien i menu) i wykorzystania specyficznych możliwości danego systemu. W zastosowaniach przemysłowych nie wykorzystuje się wyrafinowanych programów graficznych, zwłaszcza dla niewielkich aplikacji, ze względu na koszt specjalnych monitorów i programów graficznych. Głębsze zmiany w kolejnych modyfikacjach systemów wiążą się z wewnętrzną organizacją zbierania, przetwarzania i ilości danych w systemie.

BIBLIOGRAFIA

1. TDC 2000: Total Distributed Control With Centralized Operations. Honeywell 1978
2. Total Distributed Control TDC 2000. An evolution any look at centralized operations (2. Honeywell 1978
3. C o n t r o n i c: Digitales Automatisierungssystem fur die Prozesstechnik. Hartmann & Braun 1980
4. TDC 3000: The Totally Integrated Information and Control System. Honeywell 1983
5. Retrofit von Kraftwerkanlagen mit dem flexiblen Leitsystem BBC-PROCONTROL P. BBC Brown Boveri 1987
6. COROS 2000 Einheitliches Bedienen und Beobachten fur die Automatisierungs-Systeme. Siemens 1988
7. P. P r o t r o n i c: Das digitale Kompaktregel-system. Hartmann & Braun 1989
8. Central Operation for Industrial Controller KS 4580. Philips 1989
9. The Process Control System „audatec”. VEM 1989
10. Prozesleit-system auf PC-Basis (SUCOS-Proze leit-System PLS-PC). Klockner-Moeller 1989

A. CZEMPLIK

VISUALIZATION TASK IN INDUSTRIAL COMPUTER CONTROL SYSTEMS

S u m m a r y

The practice of process visualization in industrial computer control systems in emergence of some standards. The paper describes these standards with regard to display configuration and access organization. As an example, the visualization concept used in IP-2000 system, which was developed by the Institute of Engineering Cybernetics of the Technical University of Wrocław, is presented.

Moc bierna w obwodzie z odbiornikiem nieliniowym

MIROSLAW WCIŚLIK

Instytut Automatyki, Politechnika Świętokrzyska

Otrzymano 1990.01.26

Autoryzowano do druku 1990.08.18

W pracy przedstawiono analizę obwodu prądu przemiennego zawierającego rezystancję, indukcyjność oraz odbiornik nieliniowy. Analizę charakterystyk quasistatycznych przeprowadzono uwzględniając wszystkie wyższe harmoniczne. Określono schematy zastępcze elementu nieliniowego. W pracy stwierdzono, że dla danego obwodu moc bierna k -tej harmonicznej, określona wyrażeniem $kU_k I_k \sin \varphi_k$, posiada interpretację fizyczną i umożliwia bilansowanie mocy biernej w obwodzie. Dla rozważanego obwodu wyznaczono także moc deformacji. Rozważany jest także rozkład mocy czynnej i biernej w badanym obwodzie.

1. WPROWADZENIE

Przegląd podstawowych definicji mocy biernej opublikowano w [1] w 1980 roku. Dyskusja na ten temat trwa od 1927 r., kiedy to Budeanu wychodząc z zapisu mocy czynnej dla przebiegów odkształconych

$$P = U_0 I_0 + \sum_{k=1}^{\infty} U_k I_k \cos \varphi_k \quad (1)$$

przyjął arbitralnie określenie mocy biernej

$$Q_B = \sum_{k=1}^{\infty} U_k I_k \sin \varphi_k \quad (2)$$

i na podstawie definicji mocy pozornej

$$S = UI \quad (3)$$

wprowadził moc deformacji określoną relacją

$$D = \sqrt{S^2 - Q_B^2 - P^2}. \quad (4)$$

Z kolei w propozycji S. Fryzego przedstawionej w 1931 roku moc bierna definiowana była na podstawie mocy pozornej i mocy czynnej

$$Q_F = \sqrt{S^2 - P^2} \quad (5)$$

W tym układzie definicji moc bierna nie posiadała określonego znaku, a ponadto nie wprowadzono pojęcia mocy deformacji związanej ze zniekształceniami przebiegów prądu i napięcia.

Należy podkreślić, że w obu przypadkach nie znaleziono interpretacji fizycznej definiowanej mocy biernej i jej użyteczności do rozliczeń ekonomicznych. Najczęściej pojęcia te odnoszono do dwójników [1,2] lub do obwodów zawierających elementy liniowe, zasilanych przebiegami odkształconymi [3]. Zakładano więc, że wyższe harmoniczne generowane są w otoczeniu rozważanego obwodu.

Moc bierna powinna być tak zdefiniowana, by jej wartość posiadała interpretację fizyczną i była użyteczna w wyznaczaniu kosztów dostarczania i zużywania energii elektrycznej z uwzględnieniem wpływu odbiorcy na jakość energii w systemie.

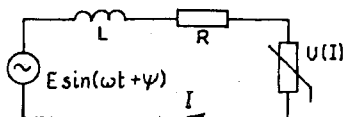
Dla odzwierciedlenia sytuacji podobnej do najczęściej występującej w systemie energetycznym należy rozważyć przypadek obwodu zasilanego napięciem sinusoidalnym, zawierającego elementy R , L i odbiornik nieliniowy. Niżej przedstawiono analizę takiego obwodu. Analityczne określenie własności tego obwodu ułatwiło przedstawienie zjawisk występujących w tym obwodzie.

2. ANALIZA OBWODU Z ODBIORKIEM NIELINIOWYM

W pracy rozważany jest obwód zawierający szeregowo połączone indukcyjność, rezystancję i odbiornik nieliniowy o charakterystyce opisanej następująco

$$U(t) = U_0 \text{sign}[I(t)] = \begin{cases} U_0 & I(t) > 0 \\ 0 & I(t) = 0; U_0 > 0 \\ -U_0 & I(t) < 0 \end{cases} \quad (6)$$

Dla takiej charakterystyki napięcie $U(t)$ jest symetryczną falą prostokątną o amplitudzie U_0 , a polaryzacji takiej jak prąd. Charakterystyka ta może być użyta jako opis własności łuku elektrycznego. Schemat rozważanego obwodu przedstawiono na rys. 1. Stosując zmienne bezwymiarowe, równanie tego obwodu można przedstawić w postaci



Rys. 1. Schemat rozważanego obwodu

$$\frac{di(\tau)}{d\tau} + ri(\tau) + u(\tau) = \sin(\tau + \psi), \quad (7)$$

gdzie:

$$i(\tau) = \frac{I(\omega t)}{E/\omega L}; u(\tau) = \frac{U(\omega t)}{E}; u_0 = \frac{U_0}{E}; r = \frac{R}{\omega L}; \tau = \omega t \quad (7a)$$

ψ – kąt przesunięcia fazowego napięcia zasilającego względem harmonicznej podstawowej napięcia łuku.

Dalej rozważany jest przypadek przeliczalnej ilości „prześć przez zero” prądu $i(\tau)$. Dla tego przypadku w stanie ustalonym napięcie $u(\tau)$ jest symetryczną falą prostokątną o amplitudzie u_0 i pulsacji harmonicznej podstawowej równej 1. Napięcie to można przedstawić w postaci szeregu Fouriera:

$$u(\tau) = u_1 \sum_{n=1}^{\infty} \frac{1}{2n-1} \sin[2n-1)\tau]; \quad u_1 = \frac{4u_0}{\Pi}, \quad (8)$$

gdzie u_1 oznacza amplitudę pierwszej harmonicznej napięcia. Podobnie można przedstawić prąd

$$i(\tau) = \sum_{n=1}^{\infty} i_{(2n-1)} \sin[2n-1)\tau + \varphi_{(2n-1)}]. \quad (9)$$

Z równań bilansu harmoniczych dla $n \geq 2$ wynikają zależności

$$i_{(2n-1)} = \frac{u_1}{(2n-1) \sqrt{(2n-1)^2 + r^2}}, \quad (10)$$

$$\varphi_{(2n-1)} = \Pi - \arctg \frac{2n-1}{r}. \quad (11)$$

Dla $\tau = k\Pi$, $k=0,1,2,\dots$ z (9) wynika zależność

$$\sum_{n=1}^{\infty} i_{(2n-1)} \sin \varphi_{(2n-1)} = 0. \quad (12)$$

Wyznaczając z (11) $\sin \varphi_{(2n-1)}$ i podstawiając go wraz z (10) do (12) otrzymuje się zależność

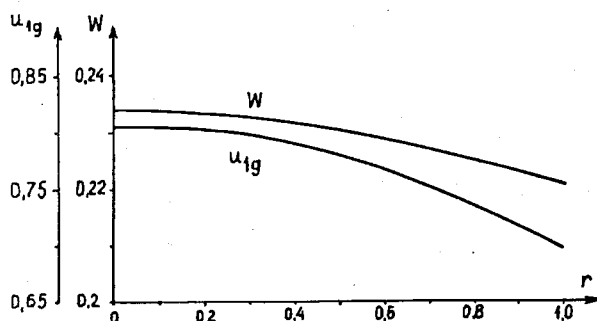
$$i_1 \sin \varphi_1 = -u_1 W, \quad (13)$$

gdzie:

$$W = \sum_{n=2}^{\infty} \frac{1}{(2n-1)^2 + r^2} = \frac{\Pi}{4r} \operatorname{tgh} \left(\frac{\Pi r}{2} \right) - \frac{1}{1+r^2}. \quad (14)$$

Na rys. 2 przedstawiono wykres wielkości W w funkcji bezwymiarowej rezystancji r . Jak wynika z tego wykresu dla $r > 0.3$ W jest stałe z odchyleniem $\leq 0.5\%$. Równania (13) i (14) wyznaczono dla r rzeczywistych, stosując następujący wzór [4]

$$\sum_{n=1}^{\infty} \frac{1}{r^2 + (2n)^2} = \frac{1}{4r} \operatorname{cth} \frac{r}{2} - \frac{1}{2r^2}. \quad (15)$$

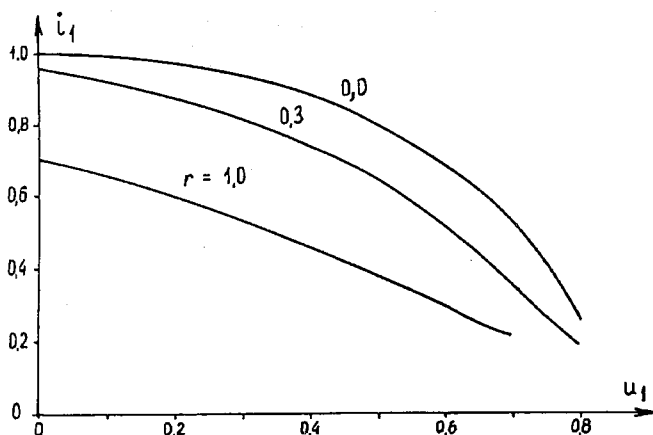


Rys. 2. Wykresy wielkości W i napięcia granicznego u_{1g} w funkcji rezystancji r

Na podstawie równania bilansu pierwszej harmonicznej oraz (13) określono jej amplitudę

$$i_1 = \sqrt{\left[\sqrt{\frac{1}{r^2 + 1}} \left[1 - u_1^2 \frac{1}{r^2 + 1} [1 + W(r^2 + 1)]^2 - u_1 \frac{r}{r^2 + 1} \right]^2 + (u_1 W)^2 \right]} \quad (16)$$

Wykres amplitudy pierwszej harmonicznej prądu w funkcji u_1 dla różnych wartości parametru r przedstawiony jest na rys. 3.



Rys. 3. Amplituda pierwszej harmonicznej prądu w funkcji napięcia u_1 i rezystancji r

Wykres ten ma charakter zbliżony do krzywych wyznaczonych dla przypadku obwodu z obciążeniem liniowym. Występuje tylko inny współczynnik skali względnej amplitudy pierwszej harmonicznej napięcia obciążenia.

Należy podkreślić, że powyższe zależności obowiązują tylko przy spełnieniu warunku przeliczalności zbioru punktów „przejścia przez zero” bezwymiarowego

prądu $i(\tau)$. Warunek ten oznacza, że prawostronna granica pochodnej prądu (9) względem czasu w punkcie $\tau = k\Pi$ jest większa od zera. Dzieje się tak, gdy

$$u_1 > u_{1g} = \frac{1}{1 + W(1 + r^2)}, \quad (17)$$

Wykres amplitudy pierwszej harmonicznej granicznego napięcia łuku u_{1g} w funkcji bezwymiarowej rezystancji r przedstawiono na rys. 2. Dla $r < 0.3$ napięcie u_1 jest stałe z odchyleniem $\leq 1\%$.

Na podstawie (10) i (14) wyznaczono relację

$$\sum_{n=2}^{\infty} i_{(2n-1)}^2 = \frac{u_1^2}{r^2} \left(\frac{\Pi^2}{8} - 1 - W \right). \quad (18)$$

Jak wynika z powyższego wzoru wartość skuteczna sumy wyższych harmonicznych prądu jest proporcjonalna do amplitudy napięcia obciążenia.

Ze względu na złożoną postać równania (16), w dalszych rozważaniach traktuje się amplitudę pierwszej harmonicznej jako daną i nie dokonuje się podstawień za i_1 .

Otrzymane zależności (16) i (18) pozwalają na proste określenie wartości skutecznej prądu w obwodzie. Nieco bardziej złożone jest wyznaczenie wartości średniej prądu (wyprostowanego dwupołkowo). W tym celu wykorzystuje się informację wynikającą z (6). Zależność ta oznacza, że polaryzacja prądu i polaryzacja napięcia na odbiorniku są tego samego znaku. Stąd

$$i_{sr} = \frac{2}{\Pi} i_1 \cos \varphi_1 + \frac{2}{\Pi} \sum_{n=2}^{\infty} \frac{1}{2n-1} i_{(2n-1)} \cos \varphi_{(2n-1)} = \frac{2}{\Pi} i_1 \cos \varphi_1 + \frac{u_1}{4\Pi} \left(\frac{\Pi^2}{8} - 1 - W \right) \quad (19a)$$

Wykorzystując (10), (11) i (14) można określić pierwszy ze składników powyższego równania. Pó podstawieniu do (19a) otrzymuje się zależność

$$i_{sr} = \frac{2}{\Pi} \left[\sqrt{\frac{1}{r^2 + 1} \left[1 - u_1^2 \frac{1}{r^2 + 1} (1 + W(1 + r^2))^2 \right]} - u_1 \frac{r}{r^2 + 1} - u_1 \frac{1}{r} \left(\frac{\Pi^2}{8} - 1 - W \right) \right] \quad (19b)$$

3. SCHEMATY ZASTĘPCZE ODBIORNIKA NIELINIOWEGO

Z równania (13) wynika, że istnieje ujemne przesunięcie fazowe pierwszej harmonicznej prądu w stosunku do pierwszej harmonicznej napięcia ($W > 0$). Stąd też, z punktu widzenia źródła sinusoidalnego napięcia zasilającego, nieliniowe obciążenie można zastąpić obciążeniem liniowym, które składa się z równoległe lub

szeregowo połączonych indukcyjności i rezystancji. Dla połączenia równoległego wartości elementów wynoszą

$$L_{zr} = L \frac{u_1}{i_1 \sin(-\varphi_1)} = \frac{1}{W} L \quad (20)$$

$$R_{zr} = \omega L \frac{u_1}{i_1 \cos \varphi_1} = \frac{\omega L}{\sqrt{(i_1/u_1)^2 - W^2}} \quad (21)$$

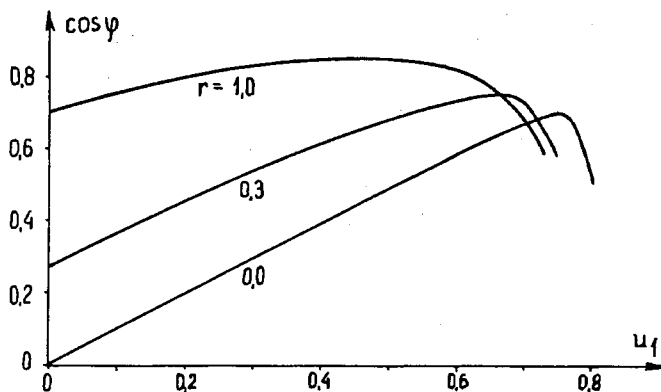
Jak wynika z (20), indukcyjność L_{zr} nie zależy od napięcia odbiornika i jest proporcjonalna do indukcyjności L , przy czym współczynnik proporcjonalności jest funkcją r .

Do dalszych rozważań bardzo użyteczny jest szeregowy schemat zastępczy, którego elementy są następujące

$$L_{zs} = \frac{u_1}{i_1} \sin(-\varphi_1) = W L \left[\frac{u_1}{i_1} \right]^2 \quad (22)$$

$$R_{zs} = \omega L \frac{u_1}{i_1} \cos \varphi_1 = \omega L \frac{u_1}{i_1} \sqrt{1 - \left[W \frac{u_1}{i_1} \right]^2} \quad (23)$$

Z powyższych zależności wynika, że zastępcza rezystancja szeregową dla małych wartości u_1/i_1 jest prawie proporcjonalna do tego ilorazu. Dla jego większych wartości rezystancja może się zmniejszać. Natomiast indukcyjność zastępcza rośnie proporcjonalnie do kwadratu tego ilorazu, przy czym należy pamiętać, że ta indukcyjność jest widoczna tylko z zacisków źródła zasilania i wynika z istnienia nieliniowości w rozważanym obwodzie. Interpretacja tego zjawiska przedstawiona jest w następnym rozdziale.



Rys. 4. Współczynnik mocy w funkcji napięcia u_1 i rezystancji r

Na podstawie szeregowego schematu zastępczego układu nieliniowego łatwo jest wyznaczyć kąt przesunięcia fazowego pierwszej harmonicznej prądu względem napięcia zasilającego

$$-\varphi = -\psi + \varphi_1 = \arctg \frac{\omega(L + L_{zs})}{R + R_{zs}}. \quad (24)$$

Wykresy cosinusa tego kąta w funkcji u_1 i parametru r przedstawiono na rys. 4. Cechą charakterystyczną tych wykresów jest występowanie maksimów w funkcji u_1 . I tak dla $r=0$ maksymalna wartość współczynnika mocy wynosi ok. 0,7.

Istnienie tych ekstremów łatwo uzasadnić na podstawie równoległego schematu zastępczego. Dla małych wartości u_1 indukcyjność równoległą można pominąć i dlatego obserwuje się wzrost $\cos\varphi$ wraz ze wzrostem u_1 . Z drugiej strony dla dużych wartości u_1 , o prądzie obwodu decyduje zastępcza równoległa indukcyjność odbiornika. Jej znaczenie maleje, gdy zmniejsza się zastępcza rezystancja obciążenia. Dlatego dla dużych wartości obciążenia, wraz ze zmniejszaniem się napięcia u_1 wzrasta $\cos\varphi$. A więc pośrednich wartości tego napięcia istnieje maksimum wartości współczynnika mocy.

4. MOC BIERNA I MOC CZYNNA W OBWODZIE

Mnożąc (12) przez $u_1 E^2 / (2\omega L)$ i uwzględniając, że amplituda $(2n-1)$ harmonicznej napięcia wynosi

$$u_{(2n-1)} = \frac{1}{2n-1} u_1 \quad (25)$$

otrzymuje się równanie

$$\frac{E^2}{2\omega L_{n=1}} \sum_{n=1}^{\infty} (2n-1) i_{(2n-1)} u_{(2n-1)} \sin\varphi_{(2n-1)} = 0. \quad (26)$$

Jeżeli przyjąć definicję mocy biernej dla k -tej harmonicznej w postaci

$$Q_k = \frac{E^2}{2\omega L} k i_k u_k \sin\varphi_k \quad (27)$$

oraz całkowitą moc bierną określić jako sumę mocy biernych poszczególnych harmonicznych to z (26) wynika, że całkowita moc bierna odbiornika jest równa zero. Oznacza to, że moc bierna pierwszej harmonicznej pobierana ze źródła zasilania i wydzielająca się z zastępczej indukcyjności odbiornika,

$$Q_{o1} = \frac{E^2}{2\omega L} u_1 i_1 \sin\varphi_1 = \frac{E^2}{2\omega L} u_1^2 W \quad (28)$$

jest w całości przekazywana do indukcyjności L w postaci mocy biernej wyższych harmonicznych. Stąd całkowita moc bierna obwodu, wydzielana w indukcyjności L wynosi

$$Q_1 = \frac{E^2}{2\omega L} i_1^2 + \frac{E^2}{2\omega L} u_1^2 W = \frac{E^2}{2\omega L} \left(1 + \frac{L_{zs}}{L}\right) i_1^2 = \frac{E^2}{2\omega L} i_1 \sin\varphi \quad (29)$$

A więc odbiornik nieliniowy w rozważanym obwodzie nie jest odbiornikiem mocy biernej, pomimo że jego schemat zastępczy dla pierwszej harmonicznej zawiera indukcyjność. Indukcyjność ta pojawia się w związku z transformacją mocy biernej doprowadzonej pierwszą harmoniczną do obciążenia nieliniowego na moc bierną odprowadzaną z nieliniowości do części liniowej odvodu wyższe harmoniczne prądu.

Dla przyjętej charakterystyki (6) łatwo sprawdzić, że moc czynna wydzielana w nieliniowości jest iloczynem wartości średnich prądu i napięcia obciążenia, wyprostowanych dwupółkowo

$$P_0 = \frac{E^2}{2\omega L} i_{sr} u_{sr} = \frac{E^2}{2\omega L} i_{sr} u_0. \quad (30)$$

Na podstawie (19a) i (8) otrzymuje się zależności

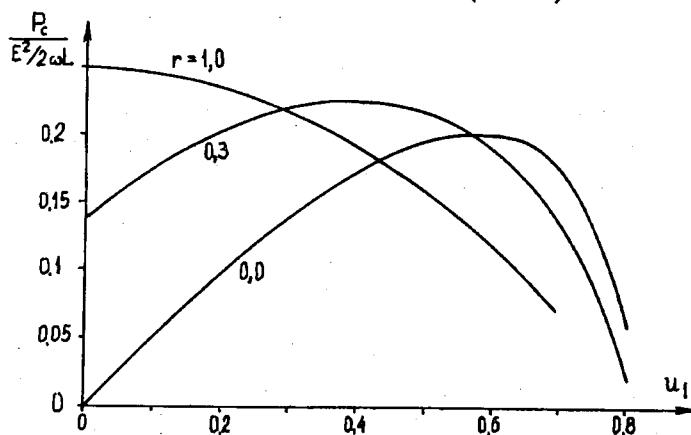
$$P_0 = \frac{E^2}{2\omega L} i_1 u_1 \cos \varphi_1 - \frac{E^2}{2\omega L} \frac{u_1^2}{r} \left(\frac{\Pi^2}{8} - 1 - W \right). \quad (31)$$

Jak wynika z powyższego równania moc czynna wydzielana w nieliniowości jest mniejsza od mocy doprowadzanej pierwszą harmoniczną. Moc czynna wydzielana na rezystancji r wynosi

$$P_{st} = \frac{E^2}{2\omega L} r \sum_{n=1}^{\infty} i_{(2n^2-1)}^2 = \frac{E^2}{2\omega L} i_1^2 r + \frac{E^2}{2\omega L} \frac{u_1^2}{r} \left(\frac{\Pi^2}{8} - 1 - W \right) \quad (32)$$

i składa się z mocy czynnej dostarczanej ze źródła zasilania oraz mocy czynnej generowanej w elemencie nieliniowym i przenoszonej wyższymi harmonicznymi. Całkowita moc obwodu wynosi

$$P_1 = P_0 + P_{st} = \frac{E^2}{2\omega L} (i_1 u_1 \cos \varphi_1 + i_1^2 r) = \frac{E^2}{2\omega L} \left(\frac{R_{zs}}{\omega L} + r \right) i_1^2 = \frac{E^2}{2\omega L} i_1 \cos \varphi. \quad (33)$$



Rys. 5. Całkowita moc czynna w funkcji napięcia u_1 i rezystancji r

Wykresy całkowitej mocy czynnej w funkcji u_1, r obwodu przedstawiono na rys. 5. Jak wynika z tego rysunku maksymalna wartość mocy (w funkcji u) zależy od r . Dla

$r=0$ maksymalna moc P_1 wydzielana jest w obwodzie dla $\cos\varphi=0,55$ i wartość tego maksimum jest o 20% mniejsza niż w przypadku obciążenia liniowego, dla którego maksymalna moc czynna wydzielana w obwodzie jest stała, przy $0 < r < 1$. Zjawisko to związane jest ze „wzrostem” reaktancji przy nieliniowym obciążeniu obwodu. Nieliniowość ta generując wyższe harmoniczne zmniejsza sprawność przesyłania energii do obciążenia oraz zwiększa energię magnetyczną indukcyjności L .

Zjawisko wzrostu mocy czynnej strat i mocy biernej obwodu należy także rozpatrywać z punktu widzenia rozliczeń energetycznych. Przy stosowaniu przedstawionych w pracy wyrażeń na moc czynną i bierną odbiorca energii o nieliniowym, odbiorniku nie jest rozliczany ze zwiększenia strat energii i wzrostu mocy biernej obwodu. Aby to uwzględnić należałoby rozliczać energię pierwszej harmonicznej, co można stosunkowo łatwo zrealizować. W ten sposób zwiększone straty w rezystancji obwodu oraz wzrost mocy biernej obciążałaby sprawcę tych niekorzystnych zjawisk [5].

Opierając się na określeniu mocy czynnej i biernej, dla pierwszej harmonicznej w [6] zaproponowano „*ekonomiczny układ pojęć mocy*”, w którym w oparciu o wartości skuteczne zdefiniowano

$$\text{— moc pozorną} \quad S = UI \quad (34)$$

$$\text{— moc czynną podstawową} \quad P_1 = U_1 I_1 \cos\varphi_1 \quad (35)$$

$$\text{— moc bierną podstawową} \quad Q_1 = U_1 I_1 \sin\varphi_1 \quad (36)$$

$$\text{— moc deformacji podstawową} \quad D_1 = \sqrt{S^2 - P_1^2 - Q_1^2} = \sqrt{(U \cdot I)^2 - (U_1 \cdot I_1)^2} \quad (37)$$

Dla rozpatrywanego wyżej obwodu łatwo stwierdzić, że moc pozorna wynosi

$$S = \frac{E^2}{2\omega L} \sqrt{i_1^2 + \frac{u_1}{r} \left(\frac{\Pi^2}{8} - 1 - W \right)} \quad (38)$$

Stąd moc deformacji podstawowej wynosi

$$D_1 = \frac{E^2}{2\omega L} \frac{u_1}{r} \sqrt{\frac{\Pi^2}{8} - 1 - W}. \quad (39)$$

Moc ta jest proporcjonalna do amplitudy napięcia elementu nieliniowego i może być przyjęta jako miara nieliniowości obwodu.

5. UWAGI KOŃCOWE

Przyjęta w pracy postać definicji mocy biernej jest różna zarówno od definicji Budeanu jak i definicji Fryzego. Związana jest ona z indukcyjnością występującą w obwodzie. Dla przebiegów okresowych można ją określić na podstawie wyrażenia

$$\frac{E^2}{2\omega L} \frac{1}{T} \int_0^T u \frac{di}{dt} dt, \quad (40)$$

które jest łatwo zrealizować fizycznie. Ale powyższej definicji nie wolno uogólniać. Można przypuszczać, że moc bierna opisana będzie ogólnie zależnością

$$Q = \frac{E^2}{2\omega L} \sum_{k=1}^{\infty} f(k) i_k u_k \sin \varphi_k \quad (41)$$

przy czym liniowy schemat zastępczy odbiornika nieliniowego może zawierać jednocześnie indukcyjność i pojemność [7].

Wpływ nieliniowości i zjawiska występujące w obwodzie z odbiornikiem nieliniowym widoczne są, gdy ten obwód porównuje się z obwodem zawierającym odbiornik rezystancyjny o mocy czynnej równej mocy odbiornika nieliniowego. Z porównania tego wynika, że nieliniowość, nawet o jednoznacznej charakterystyce może powodować:

- wzrost reaktancji obwodu,
- wzrost mocy biernej obwodu,
- wzrost strat mocy czynnej w obwodzie,
- zmniejszenie wartości współczynnika mocy.

W rozważanym obwodzie, schemat zastępczy odbiornika zawiera rezystancję i indukcyjność, ale całkowita moc bierna odbiornika jest równa zero.

Przy obecnie przyjętych definicjach mocy, odbiorca energii o odbiorniku nieliniowym nie jest rozliczany ze zwiększenia strat mocy czynnej i biernej obwodu, które obciążają ekonomicznie dostawcę energii lub innych odbiorców. Dla uniknięcia tego zjawiska, w rozliczeniach ekonomicznych należy stosować moc czynną i bierną pierwszej harmonicznej [5]. Aspekt ten oraz problem jakości zasilania obejmuje „*ekonomiczny układ pojęć mocy*” [6].

Należy dodać że, nieliniowość w obwodzie prądu przemiennego powoduje obniżenie wartości maksymalnej mocy czynnej obwodu, zależnie od rezystancji obwodu, oraz obniżenie (w stosunku do obwodu liniowego) wartości współczynnika mocy, przy której występuje to maksimum.

BIBLIOGRAFIA

1. T. Cholewicki: *Rodzaje mocy przy przebiegach odkształconych — stan obecny badań*. Przegląd Elektrotechniczny R. LVI. z. 3, 1980
2. R. Porada: *Kilka uwag o metodzie częstotliwościowej w teorii mocy obwodów z przebiegami odkształconymi*. Materiały z Wiosennego Seminarium Elektrotechniki Prądów Niesinusoidalnych, Drzonków, maj 1989
3. L. S. Czarniecki: *Moce w obwodach z przebiegami niesinusoidalnymi, ich interpretacja i pomiar*. Materiały z Wiosennego Seminarium Elektrotechniki Prądów Niesinusoidalnych, Drzonków, maj 1989
4. K. K n o p p: *Szeregi nieskończone*. Warszawa: PWN, 1956
5. M. Wciślík: *Analiza stanu quasistatycznego obwodu prądu przemiennego z łukiem elektrycznym*. Materiały z XV Sympozjum „Zjawiska Nieliniowe”, Białeżyńsko, 1987
6. J. S a w i c k i, M. G a l e w s k i: *Ekonomiczna definicja mocy deformacji*. Materiały z Wiosennego Seminarium Elektrotechniki Prądów Niesinusoidalnych, Drzonków, maj 1989
7. N. L. K u s t e r s, W. J. M. M o o r e: *On the definition of reactive power under nonsinusoidal conditions*. IEEE Trans. Power Appl. Syst. vol. PAS—99, pp. 1845—1854, Sept./Okt. 1980

M. WCIŚLIK

REACTIVE POWER IN A CIRCUIT WITH A NONLINEAR LOAD

S u m m a r y

The paper deals with an analysis of an a.c. circuit which consists of a resistance, an inductance and a nonlinear load. Analysis of the quasistatic characteristic of the circuit was performed, taking into consideration all higher harmonics. Equivalent linear circuits of the nonlinear element in the circuit are developed. Distribution of active and reactive power in the circuit is discussed. There was confirmed that in the circuit, reactive power, described for k -th harmonic as follows $kU_k I_k \sin \phi_k$, has a physical interpretation and may be balanced in the circuit. The deformation power is also described for that circuit.

Zabezpieczenie procesów elektromedycznych przed zagrożeniami elektrycznymi

JAN ŁYSKANOWSKI

Politechnika Warszawska

Otrzymano 1990.07.11

Autoryzowano do druku 1990.11.12

W pracy przedstawiono charakterystykę zagrożeń elektrycznych występujących w procesach elektromedycznych, oraz przeprowadzono ich analizę na podstawie schematów zastępczych aparatów elektromedycznych i sieci elektrycznych. Podano szereg przykładów zagrożeń elektrycznych dla różnych konfiguracji sieci elektrycznych.

W wyniku przeprowadzonej analizy uznano, że bardziej skutecznym sposobem eliminacji zagrożeń elektrycznych występujących w procesach elektromedycznych jest wprowadzenie niezależnych od już istniejących zabezpieczeń zapobiegawczych, również zabezpieczeń eliminujących zagrożenia porażeniowe — dla których określono wstępne założenia techniczne.

1. WSTĘP

W technice medycznej pojęcie bezpieczeństwa aparatury elektromedycznej jest używane w celu określenia wymagań zapobiegających przekształceniu się elektrycznych urządzeń medycznych w źródła zagrożeń zdrowia lub życia organizmu pacjenta, lekarza, obsługi. [3,4,5]

W pracy [9] wykazano, że o bezpieczeństwie pacjenta w procesie elektromedycznym decydują trzy zasadnicze czynniki: aparat elektromedyczny, jego funkcja działania i obsługa. Dlatego uznano, że zbiór wszystkich narażeń powodujących zagrożenia pacjenta zawiera trzy podzbiory tworząc triadę narażeń.

Pierwszy podzbiór triady zawiera zakłócenia wewnętrzne i zewnętrzne oddziałujące na konstrukcję aparatu, w wyniku których aparat elektromedyczny może powodować zagrożenia procesu elektromedycznego. Takie zagrożenia nazwano zagrożeniami aparatowymi.

Drugi podzbiór triady tworzą zakłócenia zewnętrzne i wewnętrzne oddziałujące na funkcję działania aparatu elektromedycznego. Takie zagrożenia nazwano zagrożeniami funkcjonalnymi.

Trzeci składnik triady tworzą zakłócenia zewnętrzne i wewnętrzne oddziałujące na obsługę aparatu elektromedycznego. Takie zagrożenia nazwano zagrożeniami obsługowymi.

W artykule omówiono część pierwszego podzbioru triady — zagrożenia aparatuowe elektryczne i zabezpieczenia procesów elektromedycznych.

2. CHARAKTERYSTYKA ZAGROŻEŃ APARATOWYCH

W procesach użytkowania aparatury elektromedycznej zagrożenia występują w przypadku, gdy w wyniku działalności ludzkiej lub zdarzenia losowego ciało człowieka jest elementem pola elektromagnetycznego lub obwodu elektrycznego, w którym jest prąd elektryczny o wartości przekraczającej dopuszczalny poziom lub staje się elementem reakcji chemicznych lub też znajduje się pod działaniem energii mechanicznej. [14,15]

Narażenia elektryczne powodujące zagrożenia porażeniowe powstają głównie w wyniku niedoskonałej izolacji elektrycznej. [6,8,16] Zagrożenia te mogą się ujawniać zarówno w stanach normalnych jak i w stanach zakłóceń (zwarciovych, przepięciowych) sieci elektrycznych.

Zagrożenia aparatuowe powstają również w wyniku oddziaływania środowiska otaczającego aparaturę i pacjenta, które umownie podzielono na trzy strefy: [13]

I — otoczenie najbliższe aparatury, którego wielkość można umownie ograniczyć do wielkości koła o promieniu równym np. długości sznura sieciowego.

II — otoczenie bliskie obejmujące obiekt medyczny zawierające pomieszczenia lekarzy, obsługi pacjentów, źródeł zasilania oraz rury, powierzchnie metalowe, metale przewodzące itp.

III — otoczenie dalekie obejmujące przestrzeń poza obiektowe, będące źródłem promieniowania ultrafioletowego, radarowego (UKF), elektromagnetycznego wiązki, ładunków elektrostatycznych itp.

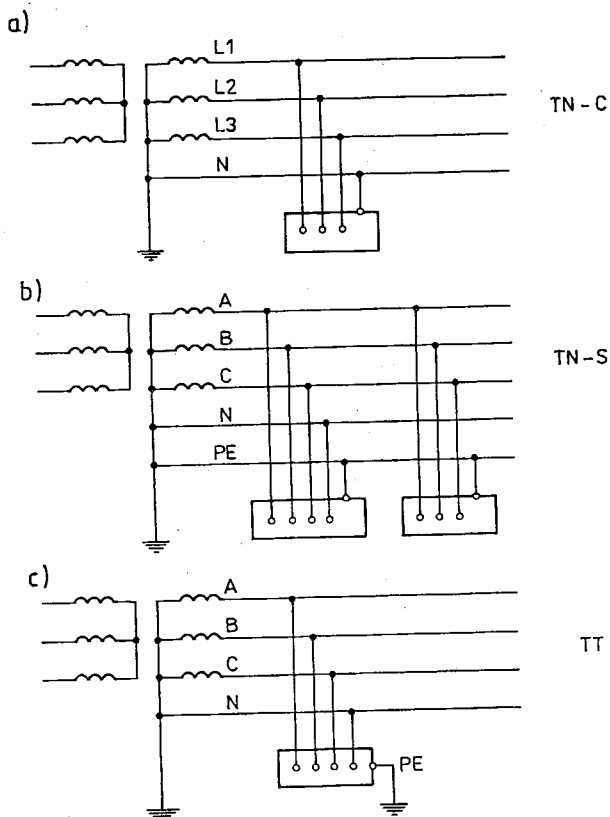
Na działanie aparatury elektromedycznej ma wpływ nie tylko środowisko strefy I. W strefie II może wystąpić wzajemne oddziaływanie nie tylko sieci i aparatury, wilgotnych powierzchni, rur ale również pacjentów, obsługi i aparatury. Środowisko strefy III reprezentuje elementy zewnętrzne, które oddziałują na środowisko II i I nie tylko bezpośrednio, lecz także za pomocą pola magnetycznego (interferencje, światło, podczerwień, promieniowanie w.cz. itp. Środowisko III może wpływać na bezpieczeństwo i charakterystyki aparatury elektromedycznej i dlatego badania bezpieczeństwa użytkowania powinny uwzględniać te wpływy.

W analizie zagrożeń elektrycznych zwraca się szczególną uwagę na prąd upływu aparatu elektromedycznego, którego wartości wyznaczają zakresy zastosowań tej aparatury. Zagrożenia porażeniowe są związane z systemem zasilania elektrycznego i wynikają przede wszystkim z faktu zasilania napięciem elektrycznym, które występuje między przewodem a ziemią lub między przewodami. Mogą także występować w wyniku przenoszenia przez system zasilania sygnałów zakłócających pochodzących z obcych źródeł nie należących do systemu zasilania elektrycznego (np. wyładowania atmosferyczne lub wyładowania elektrostatyczne). Zagrożenia zarówno pacjenta, obsługi jak i aparatu powstają w warunkach, gdy są oni uziemieni lub połączeni do ziemi przez małą rezystancję (w wyniku np. uszkodzenia izolacji

aparatu) i następuje połączenie ich z przewodem zasilania, napięciem stałym lub przemiennym m.c.z. [7,14,16]

Do analizy zagrożeń zostaną wykorzystane schematy zastępcze aparatów elektromedycznych oraz zostanie przyjęty podział sieci obowiązujący w technice zabezpieczeń elektroenergetycznych, a mianowicie na:

1. Sieci, których transformatory mają punkty zerowe bezpośrednio połączone z ziemią (charakteryzują się dużymi prądami upływu lub zwarcia z ziemią). Sieci takie są określane symbolami TN-C, TN-S, TT (rys. 1). [2]

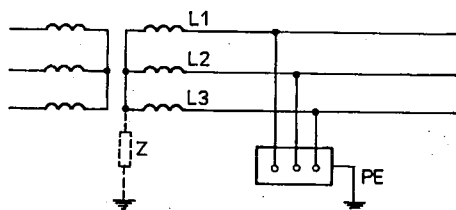


Rys. 1. Sieci elektryczne. a) typ TN-C, b) typ TN-S, c) typ TT

2. Sieci, w których punkty zerowe transformatorów nie są bezpośrednio połączone z ziemią (charakteryzują się małymi prądami upływu lub zwarcia z ziemią). Sieci takie są określane symbolem IT (rys. 2). [2]

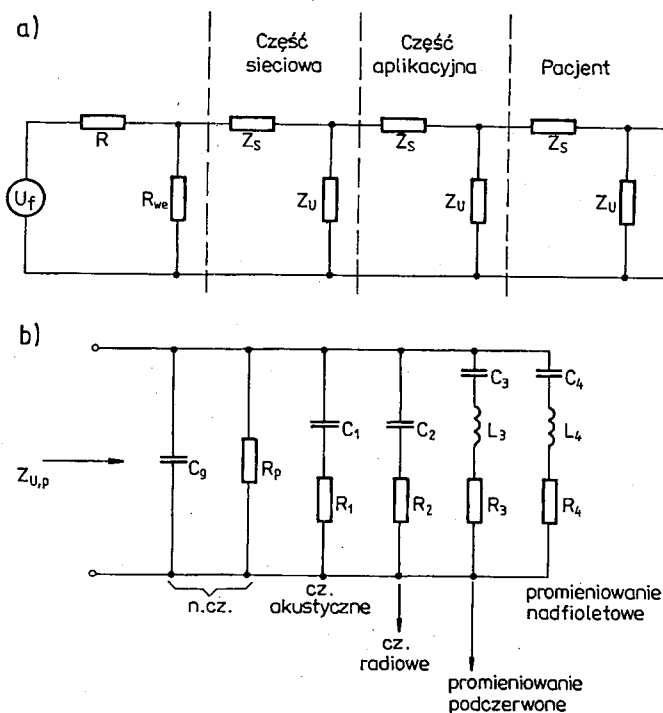
3. SCHEMATY ZASTĘPCZE SYSTEMU ELEKTROMEDYCZNEGO

Aparatura elektromedyczna jest zasilana w naszym kraju najczęściej z dwóch rodzajów sieci elektrycznej: czteroprzewodowej z uziemionym punktem zerowym



Rys. 2. Sieć elektryczna typu IT

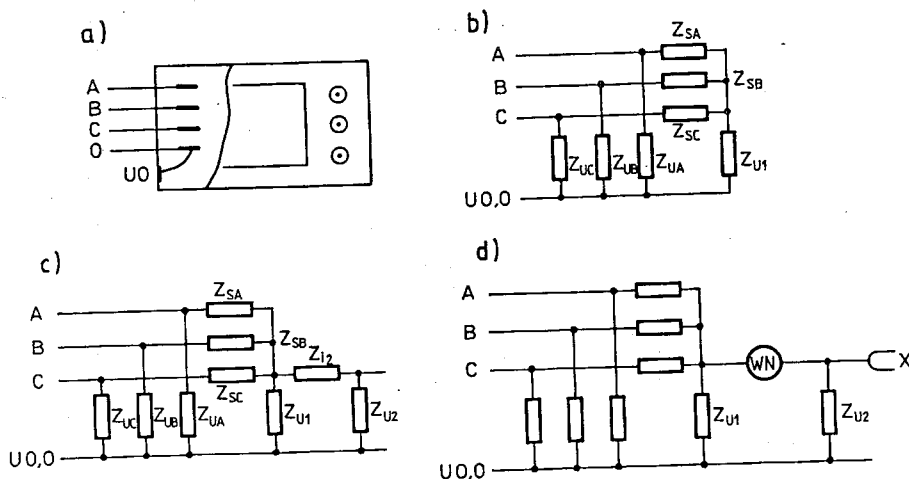
i czteroprzewodowej z izolowanym punktem zerowym. Dla aparatu zasilanego jednofazowo schemat zastępczy systemu elektromedycznego obowiązujący dla napięć sinusoidalnie zmiennych przedstawiono na rys. 3a. System podzielono na część sieciową (zasilania sieciowego), część sterowniczą i aplikacyjną aparatu elektromedycznego oraz na obwód pacjenta. Dla niskich częstotliwości można przyjąć, że pokazane na rys. 3a impedancje zawierają równoległe połączone rezystancję i reak-



Rys. 3. Schemat zastępczy systemu elektromedycznego. a) dla niskich częstotliwości, b) dla wysokich częstotliwości

Dla impulsów elektrycznych o małym czasie narastania i oscylacji w.c.z. schemat zastępczy ma postać podobną do schematu przedstawionego na rys. 3a z tym, że w schematach zastępczych impedancje izolacji Z_u i Z_s należy uwzględnić schemat zastępczy dielektryka odwzorowującego przebiegi składowych zespolonej przenikal-

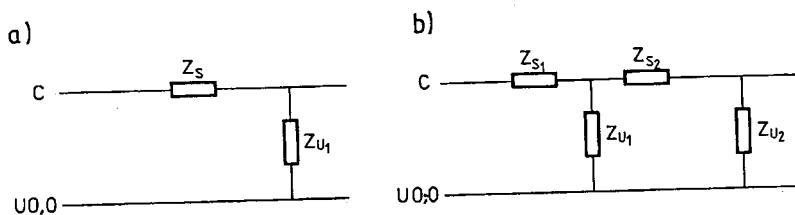
ności elektrycznej (rys. 3b). W schemacie zastępczym aparatu impedancje szeregowe obrazują impedancję izolacji między obwodami (wykonanej w postaci transformatora, transoptora itp.) zaś impedancje równoległe obrazują impedancje izolacji danego obwodu do ziemi lub obudowy aparatu.



Rys. 4. Schematy zastępcze aparatu I klasy zasilanego napięciem trójfazowym. a) schemat poglądowy, b) schemat aparatu I klasy, c) z uwzględnieniem barier izolacyjnych, d) aparatu rtg ze źródłem WN

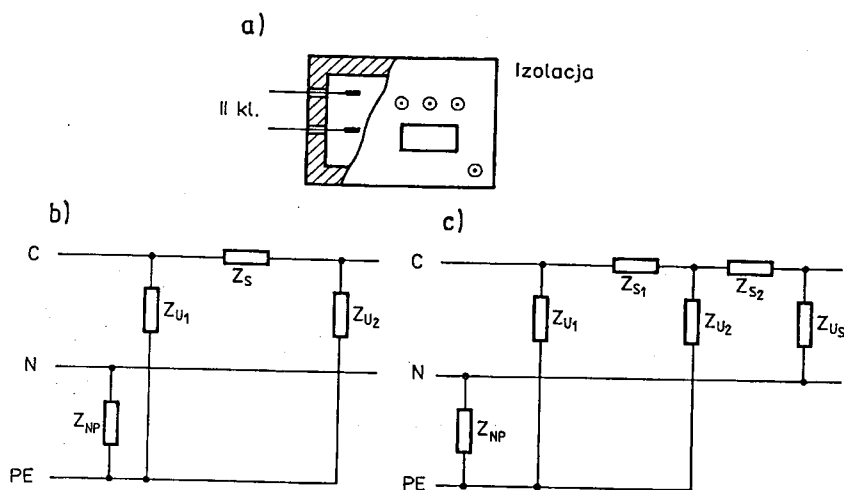
Na rys. 4 przedstawiono schematy zastępcze obwodów sieciowych aparatów elektromedycznych pokazanych na rys. 3. Tak jak w przypadku sieci, impedancje te stanowią równoległe połączenie rezystancji i reaktancji. Schemat zastępczy aparatów elektromedycznych I klasy ochronności bez barier izolacyjnych z sieci trójfazowej, pokazano na rys. 4ab, natomiast z uwzględnieniem barier, na rys. 4c. Schematy te nie dotyczą aparatów rtg, które wymagają uwzględnienia wewnętrznych źródeł WN przeznaczonych do zasilania lamp rentgenowskich. Przedstawiono to przykładowo na rys. 4d.

Na rys. 5ab uwidoczniono odpowiednio schemat zastępczy aparatu I klasy ochronności bez barier izolacyjnych i z uwzględnieniem tych barier zasilanych z sieci jednofazowej.



Rys. 5. Schemat zastępczy aparatu I klasy zasilanego napięciem jednofazowym

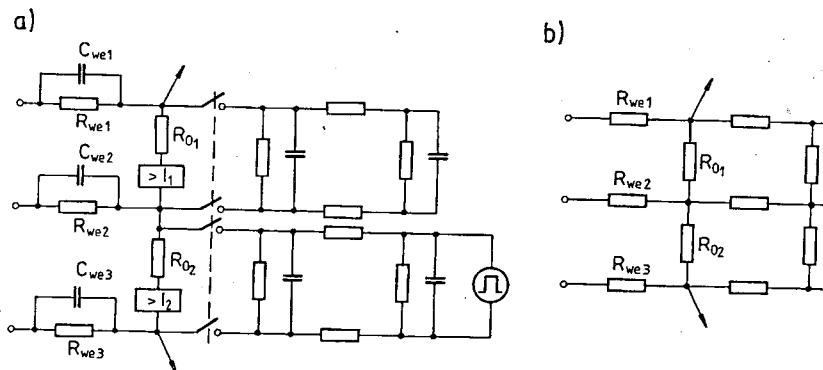
Na rys. 6 przedstawiono analogicznie jak na rys. 5 schematy zastępcze aparatów II klasy ochronności.



Rys. 6. Schematy zastępcze aparatu II klasy zasilanego napięciem jednofazowym

Z przedstawionych schematów widać, że każde ogniwo izolacyjne (transformator, bariera) stanowi czwórnik typu 1/2 T. Każdy schemat zastępczy organizmu podawany w literaturze stanowi niedoskonałe przybliżenie organizmu. Wychodząc ze schematu zastępczego komórki [15], schemat zastępczy organizmu powinien uwzględniać właściwości progowe, straty energii, zależności czasowe i częstotliwościowe.

Schemat zastępczy dla małych prądów i małych częstotliwości pokazano na rys. 7a. Schemat zawiera elementy R_{wei} i C_{wei} przy czym szeregowo z R_{wei} , R_{wej} (gdzie: $ij = 1, 2, \dots, n$) włączono elementy I_{ij} obrazujące człony nadprądowe (progowe) komórek reagujące na przekroczenie określonej wartości prądu elektrycznego. Z pośród tych elementów można wyróżnić dwa: jeden I_{1i} przedstawiający element progowy nerwu i drugi I_{2i} element mięśnia sercowego. Wartość rozruchowa nerwu wynosi kilka miliamperów, zaś mięśnia sercowego rzędu kilku mikroamperów.



Rys. 7. Schematy zastępcze organizmu. (Opis w tekście)

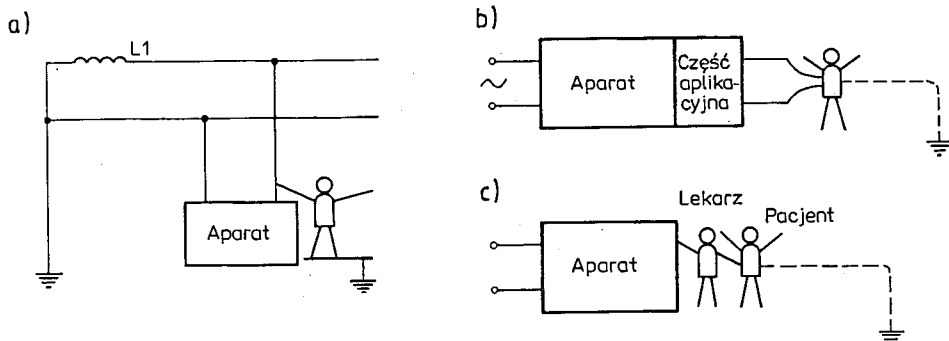
W praktyce można korzystać z uproszczonego schematu zastępczego pokazanego na rys. 7a po uwzględnieniu różnych właściwości prądowych dróg przepływu prądu przez organizm. W zakresie wyższych częstotliwości element nadprądowy, z uwagi na bezwładność przestaje reagować na zmiany prądu i dlatego schemat zastępczy przybiera postać jak na rys. 7b.

4. ZAGROŻENIA W STANIE NORMALNYM SIECI

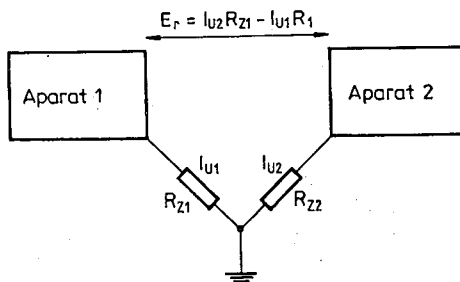
4.1. ZAGROŻENIA W SIECIACH TN-C

Na rys. 8 pokazano trzy przypadki możliwego zagrożenia pacjenta: Pierwszy (rys. 8a) dotyczy przypadku bezpośredniego dotknięcia przewodu zasilającego aparat, drugi (rys. 8b) – w wyniku skończonej wartości izolacji, trzeci (rys. 8c) – w wyniku bezpośredniego dotknięcia przez np. lekarza izolacji o skończonej wartości rezystancji między nim a drugą osobą, która jest uziemiona.

Na rys. 9 przedstawiono przypadek powstania różnicy potencjałów występującej między obudowami aparatów przyłączonych do wspólnego punktu zerowego. Wartość tej różnicy można obliczyć ze wzoru:



Rys. 8. Typowe przypadki zagrożeń porażeniowych. a) bezpośrednie dotknięcie przewodów sieci elektrycznej, b), c) pośrednie przez uszkodzoną izolację części aplikacyjnej (b) i przez uszkodzoną obudowę aparatu (c)



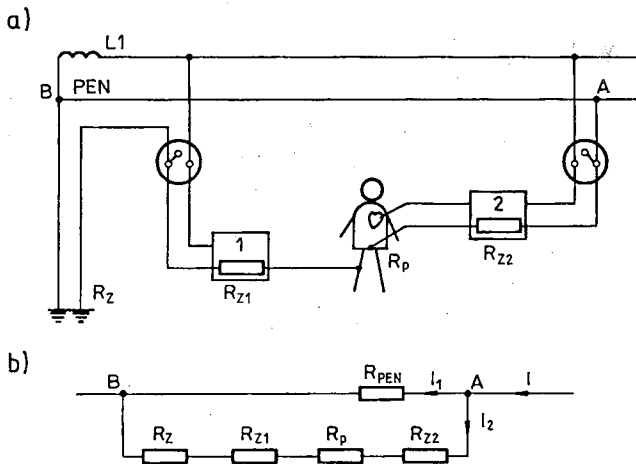
Rys. 9. Różnica potencjałów między obudowami aparatów

$$E_r = I_{u1} R_{z1} - I_{u2} R_{z2}, \quad (1)$$

w którym: I_{u1} , I_{u2} — prądy upływu aparatów

R_{z1} , R_{z2} — rezystancje obwodów ochronnych aparatów.

Dla dopuszczalnych prądów upływu i rezystancji obwodów ochronnych: $I_u = 5 \cdot 10^{-3} \text{ A}$, $R_z = 0,2 \, \Omega$ wartość tej różnicy napięć może wynosić $E_r = 1 \text{ mV}$. Zagrożenie pacjenta może wystąpić już w zakresie prądów obciążeniowych sieci. Przyjmijmy, że pacjent jest uziemiony z jednej strony przez aparat 1, zaś z drugiej strony przez aparat 2 z doprowadzoną do serca sondą (rys. 10).



Rys. 10. Przykład zagrożenia w czasie normalnej pracy sieci elektrycznej. a) schemat ideowy, b) schemat zastępczy

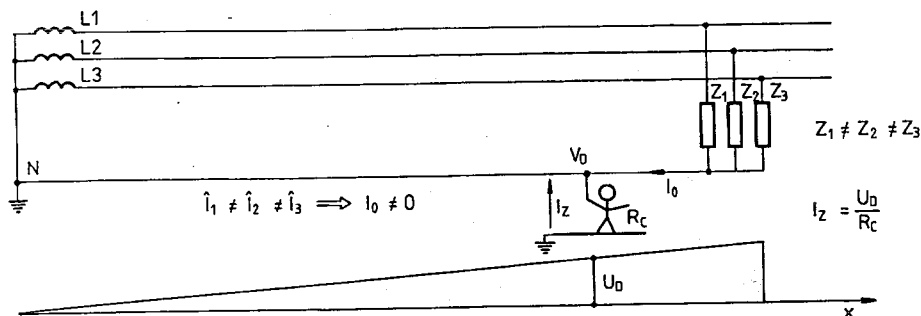
Prąd I_2 płynący przez pacjenta można obliczyć na podstawie schematu zastępczego (rys. 10), ze wzoru:

$$I_2 = I_1 \frac{R_{PEN}}{R_z + R_{z2} + R_p + R_z}, \quad (2)$$

w którym: I_1 — prąd w przewodzie PEN, R_z — rezystancja uziemienia, R_{z1} , R_{z2} — rezystancje obwodów ochronnych aparatów 1 i 2, R_p — rezystancja pacjenta, R_{PEN} — rezystancja obwodu ochronnego.

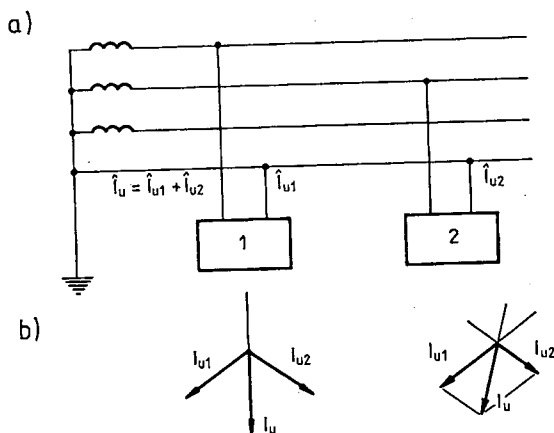
Przyjmując, że $R_{PEN} = 0,2 \, \Omega$, $R_z = 0,1 \, \Omega$, $R_z = R_{z2} = 0,15 \, \Omega$, $R_p = 1000 \, \Omega$, oraz $I_1 = 1 \text{ A}$, otrzymuje się $I_2 = 0,2 \text{ mA}$. Jest to wartość przekraczająca 20-krotnie wartość prądu przyjętą za dopuszczalną ($I_{dop} = 10 \, \mu\text{A}$). [16]

Zagrożenie może wystąpić przy zetknięciu się człowieka z przewodem zerowym uważanym za przewód ochronny lub z obudową aparatu, przyłączoną do tego przewodu. Zagrożenie to wyjaśnia rys. 11, na którym narysowano rozwiniętą sieć czteroprzewodową i wykres chwilowego spadku napięcia wzdłuż przewodu zerowego — w wybranym dostatecznie małym przedziale czasu, w którym można uznać za stałe wartości napięć i prądów płynących przez linię. W przewodzie zerowym prąd o wartości wynikającej z różnicy geometrycznej prądów fazowych, które różnią się



Rys. 11. Spadek napięcia wzdłuż przewodu zerowego

od siebie w wyniku niesymetrycznego obciążenia linii trójfazowej. Wzdłuż przewodu występuje więc spadek napięcia proporcjonalny do iloczynu różnicy prądów I_0 i rezystancji przewodu zerowego R_c , przy czym wartość bezwzględna prądu I_0 nie przewyższa wartości jego składowych, potwierdzeniem tego jest przykład pokazany na rys. 12.



Rys. 12. Przykład zagrożenia przy zasilaniu aparatów z sieci z różnych faz, a) schemat ideowy, b) wykres wskazowy

W przykładzie tym (rys. 12) aparaty są przyłączone do różnych faz sieci elektrycznej. Można wykazać, że wartość wypadkowa natężenia prądu wpływu I_u nie przewyższa wartości prądu składowego o najwyższej wartości:

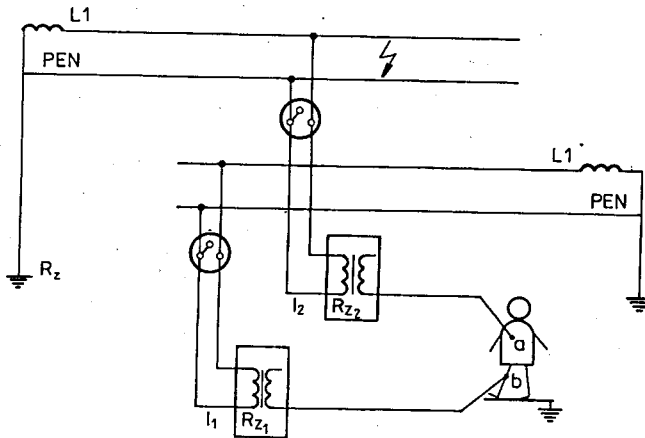
$$I_u \leq \max(I_{u1}, I_{u2}). \quad (3)$$

Kolejny przykład możliwości wystąpienia różnicy potencjałów między obudowami aparatów pokazano na rys. 13. Aparat 1 jest zasilany z sieci, natomiast aparat 2 z sieci rezerwowej lub z innej fazy sieci. Aparaty te wykonane w klasie I typu B, są przyłączone do pacjenta przez części aplikacyjne, jeżeli części te są połączone

z przewodami PEN, to między punktami a i b , w których są przyłączone do elektrody pacjenta, może powstać różnica potencjałów, określona wzorem:

$$\Delta U_{ab} = I_2 R_{z2} - I_1 R_{z1}, \quad (4)$$

w którym I_1, I_2 — prądy płynące w częściach sieciowych aparatów,
 R_{z1}, R_{z2} — rezystancje obwodów ochronnych tych aparatów.



Rys. 13. Przykład zagrożenia przy zasilaniu aparatów z sieci wewnętrznej i rezerwowej

4.2. ZAGROŻENIA W SIECIACH TN-S

W sieciach TN-S funkcje przewodu zerowego i ochronnego spełniają oddzielne przewody, odpowiednio przewód N i PN (rys. 1b). Rozdzielenie tych funkcji eliminuje niektóre z zagrożeń występujących w sieciach TN-C. W stanie normalnym sieci TN-S występują zagrożenia związane z bezpośrednim lub pośrednim dotknięciu (przez niewielką rezystancję) przewodu sieciowego, oraz w wyniku powstania różnicy potencjałów między obudowami, szczegółowo opisano w poprzednim rozdziale. Zredukowane są natomiast te zagrożenia, które są związane z pojawieniem się napięcia wzdłuż przewodu zerowego. Wynika to z tej przyczyny, że przez przewód ochronny nie płyną prądy obciążeniowe, a jedynie prądy upływu aparatów elektro-medycznych. Napięcie występujące wzdłuż przewodu jest małe i może wynosić $\Delta U = \sum I_{un} R_{PE} = 1,2 \text{ mV}$ dla $n=1$, oraz $\Delta U = 6 \text{ mV}$ dla $n=5$. Napięcie to może stanowić zagrożenie dla pacjenta jedynie w przypadku przyłączenia aparatu bezpośrednio do serca oraz gdy w przewodzie ochronnym płynie prąd $I_u = 50 \text{ mA}$.

4.3. ZAGROŻENIA W SIECIACH TT

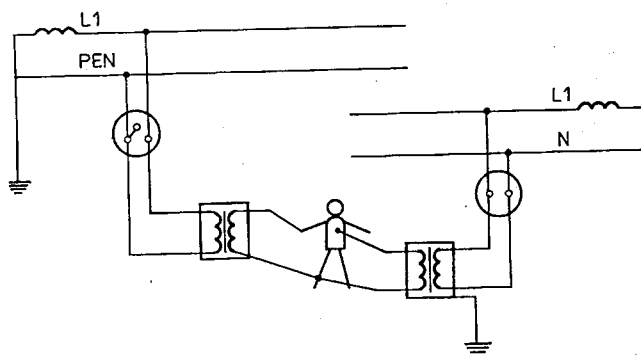
Sieć elektryczna TT odróżnia się od sieci TN-C indywidualnymi połączeniami obudowy każdego aparatu do ziemi. Sieć ta wymaga dużej liczby uziemień, co nie jest łatwe do uzyskania w praktyce. W stanie normalnym sieci zagrożenie występuje

w przypadku bezpośredniego przewodu fazowego przez pacjenta lub obsługę aparatu. Przyłączenie obudowy aparatu bezpośrednio do ziemi eliminuje powstanie groźnych napięć istniejących wzdłuż przewodu zerowego w sieciach TN-C i TN-S.

4.4. ZAGROŻENIA W SIECIACH IT Z TRANSFORMATOREM Z NIEUZIEMIANYM PUNKTEM ZEROWYM

W sieciach z transformatorem o nieziemionym bezpośrednio punktem zerowym, w przypadku zwarcia z ziemią przewodu jednej fazy nie popłynie duży prąd zwarciowy, ponieważ nie istnieje dla niego obwód zamknięty galwanicznie. Prąd może płynąć w obwodzie zawierającym jedynie pojemności przewodów fazowych w stosunku do ziemi. W obliczeniach praktycznych sieci energetycznych korzysta się z wykresów przedstawiających wartość prądu ziemnozwarciowego przypadającego na 1 km długości linii. Dla sieci medycznych brak jest takich danych.

Bezpośrednie dotknięcie przewodu fazowego nie stanowi zagrożenia dla życia pacjenta, ponieważ prąd rażeniowy zależy przede wszystkim od wartości impedancji izolacji sieci i przez to jest dużo mniejszy od prądu, który płynie w analogicznych przypadkach w sieci z transformatorem z uziemionym punktem zerowym. Odmianną sieć IT (rys. 2) jest sieć z przyłączonym punktem zerowym transformatora do ziemi przez impedancję Z .



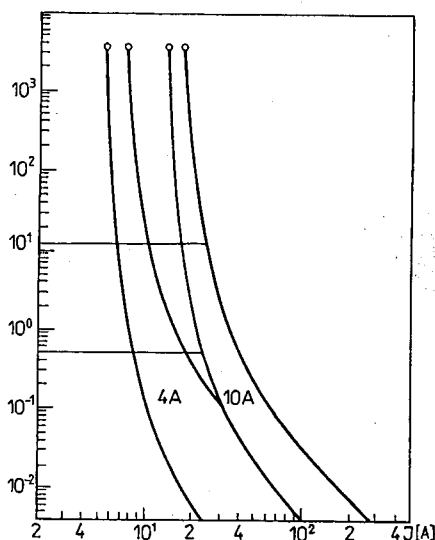
Rys. 14. Przykład zagrożenia przy zasilaniu aparatów z sieci TN-C i IT

Jeżeli dla sieci IT przyjmujemy, że prąd rażeniowy wynosi I_{rz} to dla sieci z izolowanym punktem zerowym prąd rażeniowy jest mniejszy w stosunku $R_c/(R_c + R_i)$, przy czym R_c – rezystancja pacjenta, R_i – rezystancja izolacji sieci lub impedancja Z . Zagrożenie porażeniowe może powstać w wyniku błędów projektowych, gdy w otoczeniu pacjenta oprócz sieci IT jest zainstalowana sieć TN do zasilania urządzeń niemedyceńskich np. lamp oświetleniowych. Dotknięcie przez pacjenta, który np. ma doprowadzoną sondę dosercową (rys. 14), obudowy urządzenia może spowodować migotanie jego komór serca w wyniku przepływu prądu o natężeniu większym niż 10 μA .

4.5. ZAGROŻENIA W STANACH ZWARCIOWYCH SIECI

W stanach zwarciovych sieci występują głównie zagrożenia porażeniowe. Stan zwarciovych sieci występuje w przypadkach, gdy powstaje zetknięcie (zwarcie) dwóch przewodów, pomiędzy którymi istnieje różnica potencjałów powodująca przepływ prądu przekraczającego wartości dopuszczalne. W sieciach medycznych mamy do czynienia przede wszystkim ze zwarciami jednofazowymi.

W stanach zwarciovych przewod zerowy w sieci TN-C przestaje pełnić funkcję przewodu ochronnego, ponieważ może wystąpić na nim napięcie równe napięciu fazowemu sieci. Napięcie to, przez przewód zerowy, może przenieść się na aparat elektromedyczny (rys. 11) w wyniku czego może on się znaleźć pod niebezpiecznym napięciem względem ziemi. Stopień zagrożenia dla pacjenta zależy od czasu zwarcia w sieci elektrycznej. Czas ten jest zmienny i zależy przede wszystkim od rodzaju zabezpieczeń przeciwzwarciowych, a ściślej od charakterystyki czasowo-prądowej. Zgodnie z obowiązującymi przepisami budowy urządzeń elektrycznych (PBUE) w większości pomieszczeń medycznych stosuje się bezpieczniki topikowe instalacyjne. Charakterystyki czasowo-prądowe pasmowe tych bezpieczników określa polska norma PN-87/E-93100/05. [11]



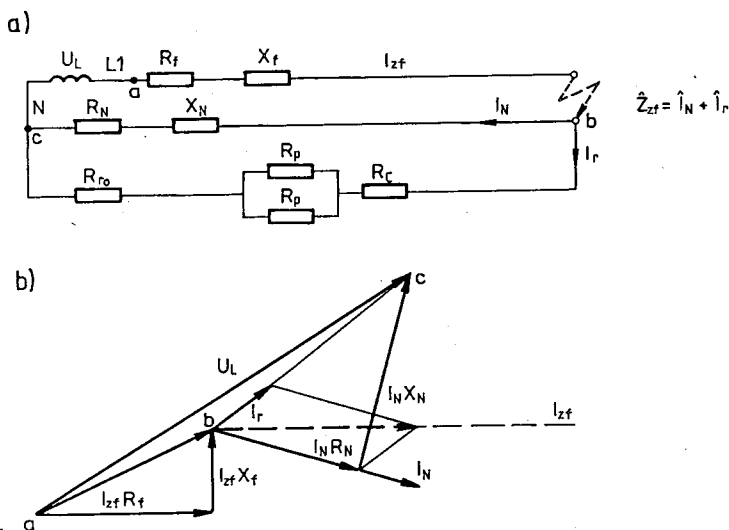
Rys. 15. Charakterystyki czasowo-prądowe, pasmowe wkładek topikowych według PN-87/E-93100/05-6

Rozpatrzmy przypadek przedstawiony na rys. 14 obrazujący pomieszczenie, w którym jedno gniazdo jest zasilane z sieci a drugie z sieci rezerwowej. Pomieszczenie jest zabezpieczone bezpiecznikiem o prądzie nominalnym $I_n = 10$ A. Pacjent jest przyłączony do aparatów klasy I typu B. Części aplikacyjne są połączone z przewodami PEN przez aparaty. Przyjmijmy, że zwarcie powstało poza pomieszczeniem, a prąd zwarciovych wynosi $I_z = 25$ A. Zgodnie z charakterystyką bezpiecznika (rys. 15)

czas zwarcia może trwać w przedziale 0,5–12 s, a natężenie prądu płynącego przez pacjenta może wynosić:

$$I_p = \frac{\Delta U_z - I_1 R_{z1}}{R_p} = \frac{25(0,2 + 0,15) - 20,15 \Omega}{1000} = 9 \text{ mA}. \quad (5)$$

W przypadku zwarcia wewnątrz pomieszczenia natężenie prądu płynącego przez pacjenta może być większy niż 9 mA w wyniku większego natężenia prądu zwarciovego (rys. 13). Schemat zastępczy pętli zwarcioviej jest przedstawiony na rys. 16a, a na rys. 16b wykres wskazowy napięć występujących w czasie zwarcia.



Rys. 16. Pętla zwarciovia. a) schemat zastępczy, b) wykres wskazowy. U_1 – napięcie fazowe, R_f – rezystancja przewodu fazowego, X_f reaktancja przewodu fazowego, R_N – rezystancja przewodu zerowego, R_{r0} – zesystancja uziemienia punktu zerowego transformatora i aparatu, R_p – rezystancja pacjenta w miejscu przyłożenia elektrody, R_c – rezystancja przejścia pacjent-ziemia

Prąd rażeniowy pacjenta jest odwrotnie proporcjonalny do sumy rezystancji: R_{r0} – uziemienia roboczego, R_a – aparatu, R_p – przejścia stopa-ziemia, R_{cz} – przejścia ręka-stopy:

$$I_r = \frac{Z_0}{Z_p R_{r0} + R_a + R_p + R_{cz}}, \quad (6)$$

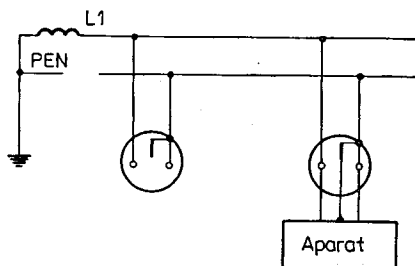
w którym Z_0 – impedancja przewodu zerowego między punktami AC,
 Z_p – impedancja pętli zwarcioviej,
 U_f – napięcie fazowe sieci.

Przewód zerowy spełnia swe zadanie, jeżeli podczas zwarcia faza-zero, przy przepływie prądu zwarciovego $I_{zf} = U_f / Z_p$, prąd rażeniowy pacjenta jest mniejszy od wartości prądu przyjętej za dopuszczalną dla danej klasy aparatu elektromedycznego.

W sieci TN-S zagrożenie porażeniowe występuje w przypadku zwarcia przewodu fazowego z przewodem ochronnym. Może to wystąpić np. w chwili przebicia izolacji obwodu sieciowego aparatu I klasy. W przewodzie ochronnym płynie prąd zwarcia, powodując pojawienie się na obudowie aparatu napięcia elektrycznego o wartości większej od dopuszczalnej. Zagrożenie to, jeżeli prąd zwarcia jest większy od prądów rozruchowych zabezpieczeń, znika po czasie ich działania. Groźna sytuacja występuje, gdy prądy zwarcia są mniejsze od prądów rozruchowych zabezpieczeń. Niebezpieczne napięcie na obudowie aparatu może występować trwale. W sieciach typu IT, w stanach zwarcia zagrożenia może wystąpić jedynie w przypadku zwarcia przewodu fazowego bezpośrednio do obudowy, jednak stopień tego zagrożenia zależy od rezystancji obudowy aparatu i uziemienia.

4.6. PRZERWANIE OBWODU

Przerwanie obwodu w sieciach trójfazowych np. jednego przewodu fazowego, powoduje dużą asymetrię prądów. Spowodować to może niedopuszczalne przegrzanie urządzeń odbiorczych trójfazowych np. aparatury rentgenowskiej oraz zmniejszenie ich wydajności a także, co jest groźniejsze dla organizmu, pojawienie się spadku napięcia na przewodzie PEN w sieci TN-C.



Rys. 17. Sieć TN-C z przerwą w obwodzie zerowym

Daleko groźniejsza w skutkach może być przerwa w przewodzie ochronnym PEN (rys. 17). Na bolcach ochronnych gniazd wtyczkowych może się pojawić pełne napięcie sieci. Z tych względów układ sieci TN-C powinien być całkowicie wyeliminowany z pomieszczeń medycznych.

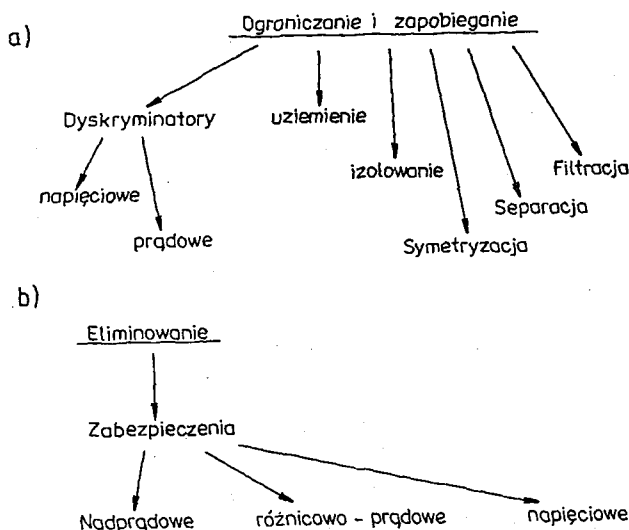
4.7. PRZECIĄŻENIA

Przeciążenia mogą wystąpić tylko w niektórych elementach aparatów elektromedycznych pracujących równolegle, gdy np. zostanie wyłączony z pracy jeden element przenoszący energię elektryczną. Wtedy pozostały w pracy element (np. tyrystor, tranzystor mocy, transformator, silnik) musi przenieść energię elektryczną, którą poprzednio przenosiły dwa elementy. W przeciążonym elemencie może wystąpić niedopuszczalny przyrost temperatury, powodując zagrożenia cieplne.

5. ZABEZPIECZENIA ELEKTROMEDYCZNE

5.1. WSTĘP

W technice medycznej do podstawowych zabezpieczeń przed skutkami pogorszenia lub zwarcia izolacji zalicza się jedynie zabezpieczenia zapobiegawcze i ograniczające (rys. 18a). Stosowanie w praktyce tego rodzaju zabezpieczeń zaleca się w sposób obligatoryjny. Oprócz zabezpieczeń zapobiegawczych powinno się stosować zabezpieczenia eliminujące zagrożenia porażeniowe (rys. 18b). Zastosowanie obu rodzajów zabezpieczeń pozwoli na podział zagrożeń na dwie grupy.



Rys. 18. Podział środków ochronnych przed zagrożeniami

Pierwsza grupa obejmuje zakłócenia, które powinny być natychmiast usunięte, a druga obejmuje te zakłócenia, które mogą być usunięte po upływie dłuższego czasu. Do pierwszej grupy zaliczamy pogorszenia izolacji powodujące przekroczenie wartości prądu upływu uznanej za wartość niebezpieczną, oraz znaczne obniżenie napięcia. Do drugiej grupy zaliczymy pogorszenie izolacji z ziemią w sieci z izolowanym punktem zerowym oraz zwiększanie się impedancji obwodów ochronnych. Zakłócenia należące do pierwszej grupy powinny być wyłączone przez zabezpieczenia, działające samoczynnie i powodujące otwarcie wyłączników i w ten sposób wyłączenie elementów będących źródłami zakłóceń. Działanie takich zabezpieczeń jest o wiele szybsze od reakcji człowieka. Zakłócenia należące do drugiej grupy nie muszą być natychmiast likwidowane i wystarczy, że personel obsługi ruchu elektrycznego jest informowany przez układy sygnalizacyjne o wystąpieniu odpowiednich rodzajów zakłóceń.

O stopniu zagrożenia decyduje wartość prądu upływu, przy czym w systemach elektromedycznych wartość tego prądu powodująca zagrożenie stanowi ułamek

wartości prądu przyjętej w elektryce za wartość zagrożeniową. Wartości te wynoszą odpowiednio — $10 \mu\text{A}$ (dla części aplikacyjnej tegoż aparatu) oraz około 20–30 mA (dla urządzeń elektrycznych).

5.2. UKŁADY ZABEZPIECZEŃ SYSTEMU ELEKTROMEDYCZNEGO

Od prawidłowej pracy układów zabezpieczeniowych stosowanych w sieci elektrycznej zależy niezawodność dostawy energii elektrycznej oraz sprawna praca poszczególnych instalacji i aparatury. Od prawidłowej pracy układów zabezpieczeniowych aparatury zależy zaś bezpieczeństwo życia pacjenta.

Pierwszy problem został rozwiązany w pełni w energetyce zawodowej, w mniejszym stopniu w elektromedycynie. Drugi problem jest na świecie dostrzegany, w Polsce zaś do rozwiązania ponieważ podstawowym zabezpieczeniem aparatu elektromedycznego jest tylko izolacja elektryczna.

Z dotychczasowych rozważań wynika, że zwarcie jest niewątpliwie najcięższym i najgroźniejszym uszkodzeniem w układzie elektrycznym. Wystąpienie zwarcia w systemie elektromedycznym powoduje nie tylko przerwę w zasilaniu energią elektryczną, ale również w wielu przypadkach zagrożenie życia pacjenta i personelu lekarskiego, a także zniszczenia w obwodach lub aparaturze. Jeżeli w systemie zostanie przyłączony np. aparat o uszkodzonej izolacji Z_z to prąd I_z , który popłynie w obwodzie tego aparatu można wyrazić wzorem:

$$I_z \approx \frac{U_c}{Z_t + Z_{pc} + Z_{pc} + (Z_z \parallel Z_{vc})} = \frac{U_c}{Z_t + Z_{pc} + Z_{pN} + Z_z}, \quad (7)$$

przy czym Z_t — impedancja uzwojeń transformatora zasilającego,

Z_{pc} — impedancja przewodu fazowego,

Z_{pN} — impedancja przewodu zerowego i zerującego,

Z_z — impedancja uszkodzonej izolacji.

Dla typowej sieci w obiektach medycznych o napięciu $U=220 \text{ V}$ można przyjąć, że $Z_t=0,02 \Omega$, $R_p=0,24 \Omega$, $Z_z=1 \Omega$ (przy pełnym zwarciu). W przypadku sieci z uziemionym punktem zerowym $Z_{pN1}=0,25 \Omega$, prąd zwarcia osiąga wartość $I_{z1}=145,7 \text{ A}$. W przypadku zaś sieci z izolowanym punktem zerowym, dla której $Z_{pN1}=5000 \Omega$, prąd zwarcia osiąga wartość $I_{z2}=0,044 \text{ A}$.

Jeżeli w systemie elektromedycznym nastąpi uszkodzenie izolacji szeregowej Z_{zs} np. w części aplikacyjnej aparatu, to popłynie prąd w obwodzie, którego elementem jest m.in. pacjent. Wartość tego prądu można obliczyć ze wzoru:

$$I_{zp} = \frac{U_c}{Z_t + Z_{pc} + Z_{pN} + Z_{zs} + Z_c}. \quad (8)$$

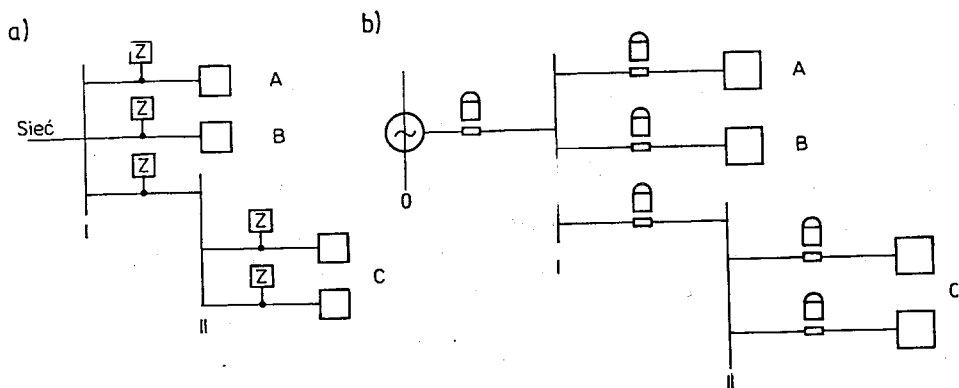
Dla sieci o takich samych parametrach przyjętych w równaniu (7) oraz dla $Z_{zs}=1 \Omega$ i $Z_c=1000 \Omega$, otrzymuje się $I_{zp}=0,229 \text{ A}$. Wartość tego prądu znacznie przekracza wartość $I=0,005 \text{ A}$ uznaną w przepisach międzynarodowych za dopuszczalną [4]. Celowe jest więc opracowanie zabezpieczeń, sygnalizujących, a następnie

eliminujących to zagrożenie. Poprawiłoby to stopień bezpieczeństwa użytkowania aparatury elektromedycznej. Zdarzające się zwarcia w aparaturze elektromedycznej są wyłączane przez zabezpieczenie nadprądowe sieciowe, które wyłącza napięcie z całego zabezpieczanego odcinka, bez względu na ważność działającej aparatury. Np. w przypadku sal operacyjnych lub intensywnego nadzoru korzystne byłoby w wielu przypadkach selektywne odłączenie uszkodzonego aparatu i pozostawienie pozostałych aparatów pod napięciem.

Oprócz zwarć w aparaturze elektromedycznej rozróżniamy dwa rodzaje przeciążeń: obwodu zasilającego i obwodu izolacyjnego. Przeciążenie obwodu zasilającego występuje wtedy, gdy wartość prądu zasilającego przekracza wartość znamionową, przeciążenie zaś obwodu izolacyjnego, jeżeli prąd upływu przekracza wartość przypisaną do konkretnej części układu aparatu elektromedycznego.

Zabezpieczenie przed przeciążeniami jest związane z trwałością aparatury chronionej, ponieważ pracując w warunkach znamionowych aparatura powinna pracować co najmniej przez okres wyznaczony przez jej trwałość. Tę z kolei wyznacza izolacja aparatu. Zabezpieczenia elektromedyczne powinny mieć cztery podstawowe cechy, w tym charakterystyczną dla zabezpieczeń elektromedycznych: wybiórczość sterowaną oraz cechy podobne do cech zabezpieczeń elektroenergetycznych: szybkość, czułość i niezawodność. Wybiórczość sterowana jest cechą, która pozwala na odłączenie zasilania przez obsługę lub automatycznie zasilania jedynie z aparatów uszkodzonych i na pozostawieniu zasilania aparatów nieuszkodzonych.

Na rys. 19 przedstawiono schemat zasilania pomieszczeń szpitalnych np. sali operacyjnej lub intensywnego nadzoru nasyconych aparaturą elektromedyczną. Wybiórczość zastosowanych zabezpieczeń można uzyskać grupując aparaturę i zabezpieczenia pod względem ważności — dla pacjenta i zasilając z dwóch grup I i II gniazd (rys. 19a)



Rys. 19. Schemat zasilania pomieszczeń szpitalnych

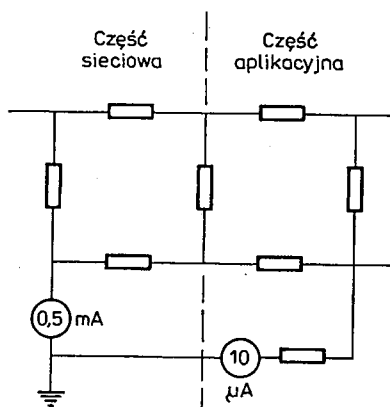
Jeżeli uszkodzenie powstanie w grupie A aparatów, wówczas popłynie duży prąd przez odcinek AI i IO. Zabezpieczenie zainstalowane w punktach A i I pobudzą się, gdyż jednocześnie stwierdzą przepływ dużego prądu. Ponieważ uszkodzeniu uległ

odcinek AI, wyłączenie powinno dotyczyć tylko tej grupy lub aparatu, natomiast grupa B i C powinna być zasilana z sieci.

Szybkość. Podczas uszkodzenia powstaje zagrożenie nie tylko dla pacjenta lub personelu medycznego, ale również dla urządzeń elektromedycznych — np. wzrost prądu upływu może rozwinąć się w zwarcie i spowodować termiczne ich zniszczenie. Należy więc dążyć do ograniczenia czasu trwania uszkodzenia, w ten sposób, aby zabezpieczenia działały tak szybko jak jest to możliwe. [1]

Czułość. Czułością zabezpieczenia nazywa się zdolność jego reagowania na zakłócenia występujące w obwodach podlegających zabezpieczeniu. [11,17]. Zabezpieczenie elektromedyczne powinno być tak czułe, aby jego prąd rozruchu był mniejszy od najmniejszego prądu uszkodzeniowego jaki może się pojawić w aparatach zabezpieczanych, przy czym należy odróżnić dwie klasy zabezpieczeń.

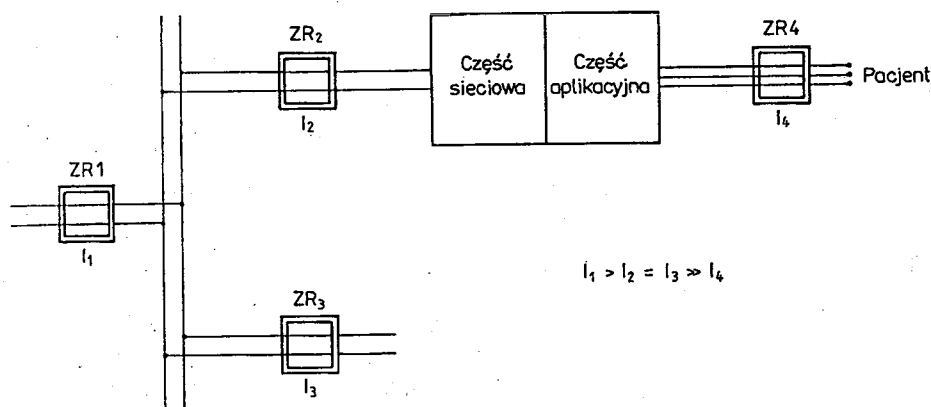
Klasę I stanowi zabezpieczenie części sieciowej (rys. 20) charakteryzującej się dużym prądem dopuszczalnym upływu ($0,5...5,0$ mA) a klasę II — zabezpieczenie części aplikacyjnej aparatu charakteryzującej się małym dopuszczalnym prądem upływu ($10\text{ }\mu\text{A}$).



Rys. 20. Schematy zastępcze części sieciowej i aplikacyjnej aparatu

Zabezpieczeniem głównym systemu elektromedycznego powinno stać się, stosowane już w niektórych krajach i zalecane przez przepisy IEC, przeciwporażeniowe zabezpieczenie różnicowo prądowe wyposażone np. w człon nadprądowy czasowy. Ideą pracy zabezpieczenia byłoby reagowanie szybkie bezwzględne w przypadku pojawienia się zbyt dużego prądu upływu i wybiórczego w przypadku zwarcia. Zabezpieczenie to powinno charakteryzować się dużo większą czułością od czułości obecnie stosowanych zabezpieczeń. To ostatnie wymaganie może spowodować zdecydowaną zmianę rozwiązania konstrukcyjnego stosowanych rozwiązań, a nawet zmianę zasady działania zabezpieczenia.

Schemat ideowy zabezpieczenia systemu elektromedycznego pokazano na rys. 21. System, w tym i aparat elektromedyczny podzielono na odcinki które wymagają ochrony. Ciąg instalacji i aparatury jest zasilany z jednego punktu. Zgodnie



Rys. 21. Schemat zabezpieczeń systemu elektromedycznego

z przepisami Budowy Urządzeń Elektrycznych zabezpieczenia powinny być zainstalowane na każdym elemencie układu w punkcie najbardziej zbliżonym do źródła.

Pierwszym zabezpieczeniem systemu elektromedycznego powinien być przełącznik różnicowo prądowy zainstalowany przy gnieździe ściennym pomieszczenia medycznego. Drugim zabezpieczeniem powinien być przełącznik zaistalowany w miejscu zasilania (np. w listwie gniazd) grupy aparatów elektromedycznych I strefy wyłączania itd.) Aparat elektromedyczny, oprócz zabezpieczenia części sieciowej, powinien być wyposażony w zabezpieczenie członów aplikacyjnych oraz układy automatycznego testowania aparatu przed jego użyciem.

ZAKOŃCZENIE

W celu przeprowadzenia analizy zagrożeń elektrycznych w systemach elektromedycznych wprowadzono pojęcie schematu zastępczego bezpieczeństwa aparatu elektromedycznego I i II klasy. W schematach zastępczych uwzględniono przede wszystkim impedancje równoległe i szeregowe decydujące o bezpieczeństwie działania aparatu. Przeprowadzono także analizę stanów pracy sieci elektrycznych i aparatury w przypadkach, w których może powstać szczególne zagrożenie dla pacjentów i personelu. Ponadto przedstawiono propozycję wprowadzenia do praktyki oprócz zabezpieczeń zapobiegawczych również zabezpieczenia eliminujące zagrożenia porażeniowe.

Z podanych przykładów wynika konieczność opracowania szeregu układów zabezpieczających uwzględniających specyfikę procesów elektromedycznych. Stosowanie znanych w elektroenergetyce zabezpieczeń może okazać się w praktyce medycznej mało przydatne.

Uzasadnione jest więc przystąpienie do indywidualnego, bardziej kompleksowego czy nawet odmiennego niż w warunkach nie medycznych opracowania zagadnień bezpieczeństwa elektrycznego aparatury elektromedycznej. Artykuł niniejszy stanowi wstęp do tych opracowań.

BIBLIOGRAFIA

1. E. D. Cooper: *AC Power for Electronic Equipment*. Med. Electron, 1982, Iss. 75
2. IEC 62A (Secretariat) 35: *Electrical Installations in Hospitals*
3. IEC 513 (1976): *Basic aspects of the safety philosophy of electrical equipment used in medical practice*
4. IEC 601-1 (1988): *Medical Electrical Equipment. Part 1: General Requirements for Safety*
5. IEC 479-1 (1984): *Effects of Current passing through the Human Body. Part 1: General aspects*
6. K. Kobyliński, I. Słowikowski: *Materiałoznawstwo elektrotechniczne*, Warszawa: WNT, 1978
7. Latos, A. Szczepański: *Problemy systemu sieci elektrycznej w środowisku pacjenta*. Probl. Techn. Med., 1979, nr. 10
8. J. Łyskanowski: *Kompatybilność pomiarowa aparatury elektromedycznej*. Sympozjum: Nowe metody pomiarowe stosowane w diagnostyce medycznej. Warszawa, 1988, OBRTM ORMED
9. J. Łyskanowski: *Współczesne problemy bezpieczeństwa aparatury elektromedycznej*. Rozpr. Elektrotechn., (w druku)
10. E. Matula, M. Sych: *Zapobieganie porażeniom elektrycznym w przemyśle*. Warszawa: WNT, 1980
11. E. Peeron: *Electrical safety*. Med. Electron., 1978, June
12. PN-66/E-05009: *Urządzenia elektroenergetyczne. Ochrona przeciwporażeniowa w urządzeniach na napięcie znamionowe do 1000 V*
13. D. Rovetti: *Testing Medical Equipment for Electrical Safety*. Med. Electron., 1977, Iss. 45
14. Z. Teresiak: *Ochrona przeciwporażeniowa w niskonapięciowych urządzeniach elektroenergetycznych*. Niektóre zagadnienia podstawowe. Rozpr. hab., Polit. Wr., Wrocław, 1966
15. Z. Teresiak, J. Kupfer: *Wartości graniczne wielkości warunkujących fizjologiczne zagrożenie porażeniom elektrycznym*. Prz. Elektrotechn., 1986, z. 10-11-12
16. C. W. Walter: *Safe Electric Environment in hospitals*. Proc. of the First National Conference on Electronics in Medicine. Mc Grow-Hill Inc, New York
17. J. Zydanowicz: *Zabezpieczenia elektroenergetyczne*

J. ŁYSKANOWSKI

PROTECTIONS OF ELECTROMEDICAL PROCESSES AGAINST ELECTRICAL HAZARDS

Summary

In the paper the characteristics and analysis of electrical hazards occurring in electromedical processes are presented. The equivalent circuits of apparatus and mains are also discussed. Several principles of electrical hazards in different kind of mains are presented.

An emphasis is on the most important investigations which should out in order to increase safety of medical electrical equipment. Such investigations should result in elaboration and introduction into practice eliminating protections apart from protections preventing.

Rotatory cyfrowe — teoria i zastosowanie

ROMUALD DRAHOKAUPIL, JAN STANISŁAWSKI, CZESŁAW MICHALIK

Instytut Telekomunikacji i Akustyki, Politechnika Wrocławska

Otrzymano 1990.10.26

Autoryzowano do druku 1990.11.1

W artykule opisano własności, zasadę działania oraz budowę rotatorów cyfrowych działających wg algorytmu Voldera. Rotatory te realizują procedurę przeliczania współrzędnych wektorów obracanych na płaszczyźnie za pomocą trzech równań różnicowych I. rzędu. Podano przykład realizacji tego rotatora przy użyciu typowych układów techniki cyfrowej. Opisano zmodyfikowaną przez Wathera postać tych równań, umożliwiającą realizację rotacji kołowej, liniowej i hiperbolicznej. Wskazano na dużą uniwersalność procesorów rotatorowych działających wg tego algorytmu. Podano zastosowania rotatora do wykonywania podstawowych operacji i funkcji matematycznych. W zakończeniu omówiono przykłady zastosowań przy przeliczaniu współrzędnych geograficznych, do obliczania szybkiej transformaty Fouriera (FFT) i cyfrowej filtracji ortogonalnej.

1. WSTĘP

Rotatory są procesorami, za pomocą których można wykonać obrót wektora o zadany kąt, tzn. wyliczyć nowe współrzędne tego wektora po obrocie na płaszczyźnie lub w przestrzeni.

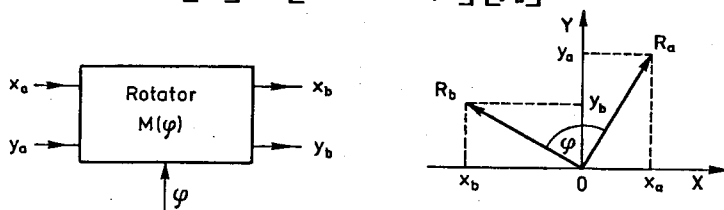
Zainteresowanie rotatorami datuje się od roku 1959, w którym Volder opublikował pracę [1], omawiającą iteracyjną postać równań obrotu wektorów, dostosowaną do specyfiki pracy układów cyfrowych. Od tego czasu opublikowano wiele prac, w których omawiano nie tylko rozwinięcia teorii rotatora (np. [12,13]), ale również możliwości zastosowań rotatora do rozwiązywania problemów matematycznych [5,7,9] i technicznych [6,9,13]. Dotychczasowe wyniki wskazują na dużą uniwersalność algorytmu Voldera, nazywanego w literaturze algorytmem CORDIC (*Coordinate Rotation Digital Computer*) i potrzebę dalszych badań w tej dziedzinie.

W pracy podano zasadę działania rotatora cyfrowego, budowę oraz wybrane przykłady zastosowań rotatora.

2. OPIS ROTATORA I RÓWNANIA OBROTU VOLDERA

Rotator, zwany także rotatorem Givensa, jest zdefiniowany za pomocą równania [1]:

$$\begin{bmatrix} x_b \\ y_b \end{bmatrix} = \begin{bmatrix} \cos\varphi & -\sin\varphi \\ \sin\varphi & \cos\varphi \end{bmatrix} \begin{bmatrix} x_a \\ y_a \end{bmatrix}. \quad (1)$$



Rys. 1. Rotator dwuwymiarowy

Jeżeli x_a i y_a są współrzędnymi wektora R_a leżącego na płaszczyźnie, to x_b i y_b są współrzędnymi nowego wektora R_b , obróconego o kąt φ . Tak więc rotator wykonuje obrót wektora $r_a = [x_a, y_a]^T$ o kąt φ (rys. 1) bez zmiany jego długości i oblicza współrzędne wektora $R_b = [x_b, y_b]^T$ z równania

$$R_b = M(\varphi) R_a,$$

gdzie macierz rotacji

$$M(\varphi) = \begin{bmatrix} \cos\varphi & -\sin\varphi \\ \sin\varphi & \cos\varphi \end{bmatrix}. \quad (2)$$

Można łatwo wykazać, że macierz $M(\varphi)$ ma następujące właściwości:

$$M^T(\varphi) M(\varphi) = 1_2, \quad (3a)$$

$$M(\varphi_1 + \varphi_2) = M(\varphi_1) M(\varphi_2), \quad (3b)$$

gdzie 1_2 jest macierzą jednostkową stopnia 2. Właściwość (3b) została wykorzystana przez Voldera do sformułowania iteracyjnej postaci równań rotacji. W tym celu kąt rotacji φ został podzielony na n odpowiednio dobranych kątów elementarnych, zgodnie ze wzorem:

$$\varphi = \sum_{i=0}^{n-1} \mu_i \alpha_i \quad (4)$$

gdzie $\mu_i = \pm 1$, α_i – kąt rotacji cząstkowej.

Równanie (1) można wtedy zapisać w postaci:

$$\begin{bmatrix} x_b \\ y_b \end{bmatrix} = \prod_{i=0}^{n-1} \cos\alpha_i \begin{bmatrix} 1 & -\mu_i \tan\alpha_i \\ \mu_i \tan\alpha_i & 1 \end{bmatrix} \begin{bmatrix} x_a \\ y_a \end{bmatrix}. \quad (5)$$

Idea metody Voldera polega na przyjęciu takiego zbioru kątów $\{\alpha_i\}, i=0, \dots, n-1$, aby $\operatorname{tg} \alpha_i = 2^{-i}$. Wtedy $\cos \alpha_i = (1 + 2^{-2i})^{-1/2}$ i równanie (5) przekształca się do postaci:

$$\begin{bmatrix} x_b \\ y_b \end{bmatrix} = K_n \prod_{i=0}^{n-1} \begin{bmatrix} 1 & -\mu_{i2}^{-i} \\ \mu_{i2}^{-i} & 1 \end{bmatrix} \begin{bmatrix} x_a \\ y_a \end{bmatrix} \quad (6)$$

gdzie K_n jest stałą, zależną od liczby obrotów elementarnych n , wyznaczaną ze wzoru:

$$K_n = \prod_{i=0}^{n-1} (1 + 2^{-2i})^{-1/2}. \quad (7)$$

Ostatecznie, równania rotacji przyjmują następującą postać iteracyjną:

$$\begin{aligned} x_{i+1} &= x_i - \mu_i 2^{-i} y_i, \\ y_{i+1} &= y_i + \mu_i 2^{-i} x_i, \\ z_{i+1} &= z_i - \mu_i \alpha_i. \end{aligned} \quad (8)$$

gdzie $x_0 = K_n x_a$, $y_0 = K_n y_a$, $z_0 = \varphi$, $\mu_i = \operatorname{sign} z_1$, $i=0, 1, \dots, n-1$.

Równanie (8) pozwala wyznaczyć w ostatnim $(n-1)$ kroku nowe współrzędne: $x_b = x_n$, $y_b = y_n$, przy czym maksymalny kąt rotacji jest dla $n \geq 4$ równy w przybliżeniu $\pi/2$ rad. Aby wykonać rotację w pełnym zakresie $(\pi, -\pi)$ rad należy, przed wykonaniem rotacji na podstawie równań (8), dokonać wstępnego obrotu wektora o kąt $\pm \Pi/2$ zgodnie z równaniami

$$\begin{aligned} x_0 &= -\mu_{-1} K_n y_0 \\ y_0 &= \mu_{-1} K_n x_0 \\ z_0 &= \varphi - \mu_{-1} \Pi/2, \end{aligned} \quad (9)$$

gdzie $\mu_{-1} = \operatorname{sign} \varphi$.

Stała K_n występuje we wzorach na x_0 , y_0 w związku z wydłużeniem długości wektora w trakcie iteracyjnego wykonywania obrotu zgodnie z równaniami (8). Korekta wydłużenia może być przeprowadzona zarówno na początku (jak wyżej), w trakcie, lub na końcu obliczeń. Ze względu na możliwość wystąpienia przepelnienia w sumatorze, lepiej przeprowadzić ją na samym wstępie. Realizacja tej operacji w trakcie obliczeń komplikuje układ rotatora. Sama korekta, zwana także skalowaniem wektorów, może również być wykonywana iteracyjnie. W tym celu może posłużyć się algorytmem podanym przez Udo, Deprettere i Devilde w [5]. Opiera się on na rozkładzie liczby K_n z przedziału $(0,5, 1)$ na iloczyn czynników $(1 - \gamma_i 2^{-i})$ zgodnie ze wzorem:

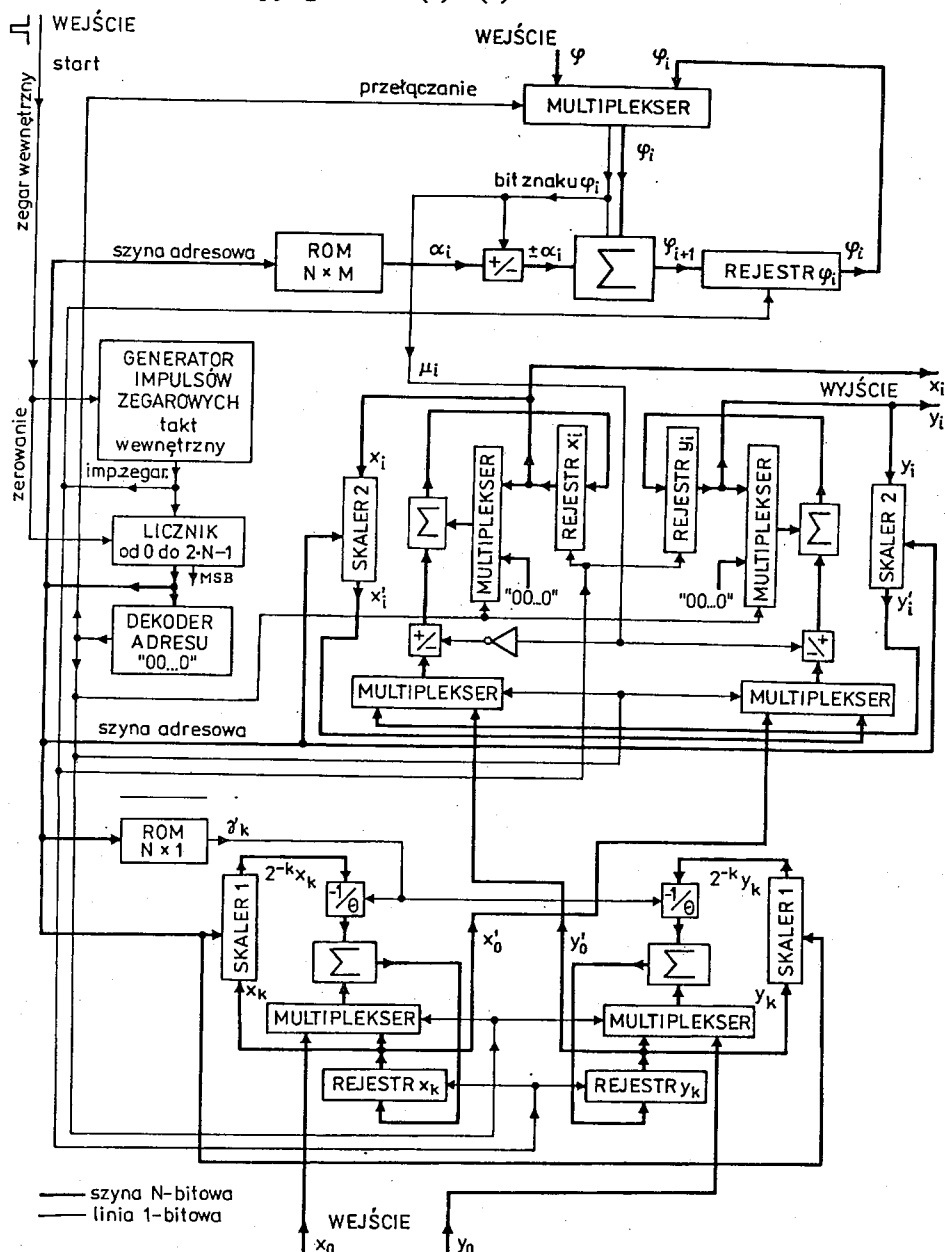
$$K_n \simeq \prod_{i=1}^n (1 - \gamma_i 2^{-i}). \quad (10)$$

gdzie $\gamma_i \in \{0,1\}$. I tak, np. dla $n=8$, $K_n=0,607259112$, otrzymuje się: $\gamma_1=\gamma_5=\gamma_6=0$, $\gamma_2=\gamma_3=\gamma_4=\gamma_7=\gamma_8=1$.

Pomnożenie współrzędnych wektora $\mathbf{R}_a$ przez K_n wymaga maksymalnie n pomnożeń przez czynniki $(1 - 2^{-i})$, co podobnie jak równania (8), wykonuje się prosto przy użyciu sumatorów i rejestrów przesuwnych.

Reasumując powyższe, należy stwierdzić, że układ działający wg algorytmu Voldera powinien wykonywać następujące funkcje:

1. wyznaczać współczynniki μ_i ,
2. dokonywać korekcji wydłużenia,
3. wykonywać rotację zgodnie z (8) i (9).



Rys. Rotator N-bitowy oparty na algorytmie CORDIC

Schemat blokowy rotatora realizującego wszystkie wymienione tu funkcje pod nadzorem układu sterującego pokazano na rys. 2 [4]. W skład układu sterującego wchodzi tu generator impulsów zegarowych, licznik impulsów i dekodery adresów „00...0”. Impuls startowy przychodzący z zewnątrz powoduje wyzerowanie i przedstawienie licznika impulsów w stan gotowości, a także krótkotrwale zatrzymanie (tylko na czas trwania impulsu startowego) generatora zegarowego. Zerowy stan licznika zostaje wykryty przez wspomniany dekodery adresowy, a ten przełącza multiplexery w sposób umożliwiający wprowadzenie danych wejściowych: kąta rotacji φ , oraz współrzędnych x_a , y_a odpowiednio do układu wyznaczającego współczynniki μ_i i do układu korekcji. Wpisanie tych wielkości do rejestrów następuje w pierwszym takcie po odblokowaniu zegara. Równocześnie multiplexery powracają do pracy w pętliach. Układ przechodzi do realizacji równań (8) i (9), wykorzystując wartości α_i i γ_i przechowywane w pamięciach ROM. Operacje przesuwania o i pozycji w prawo wykonują układy skalerów (SKALER 1, SKALER 2). Po n -tym impulsie dekodery znów wykrywa stan „00...0” (do dekodera nie jest dołączona najstarsza pozycja licznika) i cały cykl powtarza się. W przedziale czasu od zablokowania licznika do momentu podania kolejnego impulsu startowego możliwy jest odczyt wielkości wyjściowych. Układ rotatora w taktach od 0 do $n-1$ realizuje korekcję długości składowych wektora, a od taktu n do $2n-1$ obrót o kąt $\pm \Pi/2$ oraz właściwą rotację. Praca układu skalującego w taktach n do $2n$, oraz układu rotującego w czasie wykonywania korekcji wydłużenia nie wpływa na poprawne działanie procesora.

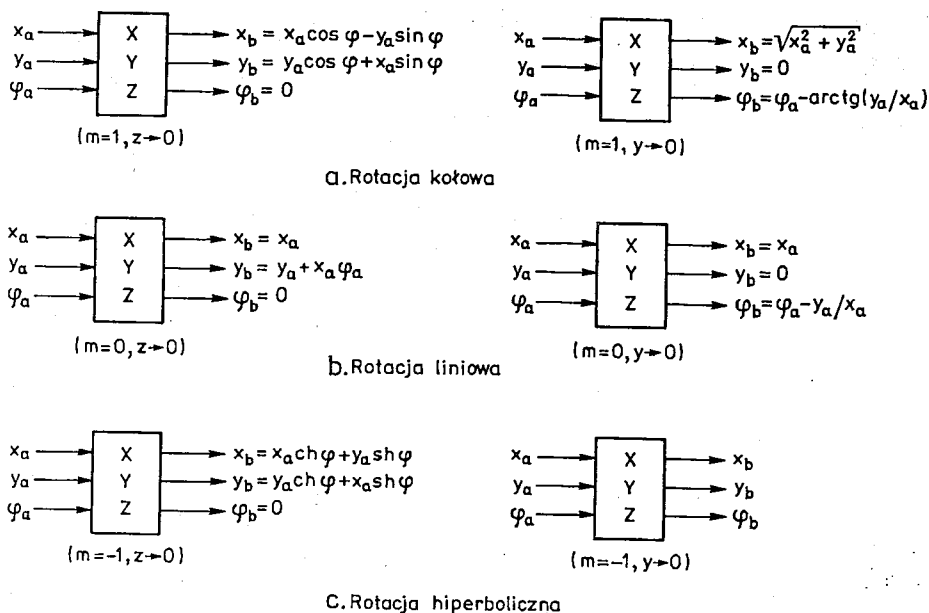
Opisany wyżej rotator jest specjalizowanym układem, którego wewnętrzna architektura jest dostosowana do pełnej funkcji. Oznacza to, że wszystkie działania wykonywane są tylko pod kontrolą układu sterującego (sprzętowo), a nie programu umieszczonego w pamięci typu ROM. Dzięki temu możliwe jest uzyskanie znacznie krótszych czasów rotacji niż przy programowych realizacjach algorytmu CORDIC.

3. UOGÓLNIONY ALGORYTM VOLDERA

W 1971 r. Walther [11] zaproponował nieznaczny modyfikację algorytmu Voldera, dzięki której stał się on bardziej uniwersalny. Polegała ona na wprowadzeniu parametru $m \in \{-1, 0, 1\}$ do równania (8):

$$\begin{aligned} x_{i+1} &= x_i - m \mu_i 2^{-i} y_i, \\ y_{i+1} &= y_i + \mu_i 2^{-i} x_i, \\ z_{i+1} &= z_i - \mu_i \alpha_i, \quad i = 0, 1, \dots, n-1. \end{aligned} \quad (11)$$

oraz takiej możliwości wyboru współczynników μ_i , aby w końcowej fazie obliczeń w rejestrze z lub y uzyskać wartość możliwie bliską zero. Ta zmiana umożliwia realizację rotacji kołowej, liniowej i hiperbolicznej. Odpowiednie operacje dla różnych parametrów m pokazano na rys. 3. Jak wynika z tego rysunku za pomocą „uogólnionego CORDICA” można realizować różne operacje matematyczne, takie jak: dodawanie, odejmowanie, mnożenie i dzielenie, obliczanie wartości wielu funkcji



Rys. 3. Funkcje uogólnionego algorytmu CORDIC

standardowych. Sposób wykonywania obliczeń polega na wyznaczaniu dokładnych wartości kolejnych pozycji binarnych wyniku (metoda „*digit by digit*”). Na takiej zasadzie działa m.in. procesor arytmetyczny który wykonuje operacje na liczbach 24-bitowych w czasie 40 μ s [6], a także kalkulator HP-35 firmy Hewlett-Packard [8]. Do podstawowych wad algorytmu CORDIC należy zaliczyć: potrzebę stosowania szybkich skalerów [10], konieczność wykonywania korekcji wydłużenia, oraz tzw. „*podwójne cykle*”, niezbędne przy rotacji hiperbolicznej [9]. Wpływa to niewątpliwie na stopień komplikacji budowy wewnętrznej procesora. Problemy te znikają przy wykorzystaniu rotatora z użyciem pamięci typu ROM, oraz przy wprowadzeniu dodatkowej pętli działającej iteracyjnie [14]. Sam zmodyfikowany algorytm, realizowany programowo, ułatwia tworzenie spójnego zbioru procedur dla operacji standardowych wykonywanych przez kalkulator lub komputer.

4. ZASTOSOWANIA

Procesor arytmetyczny oparty na uogólnionym algorytmie CORDIC umożliwia rozwiązanie wielu problemów naukowych i technicznych. Oto krótkie zestawienie niektórych z nich [9]:

1. Obliczanie funkcji standardowych: trygonometrycznych i hiperbolicznych oraz ich odwrotności, logarytmicznych i wykładniczych, zamiana współrzędnych prostokątnych na biegunowe i odwrotnie, wykonywanie typowych operacji arytmetycznych.

2. Operacje macierzowe, w tym rozwiązanie układu równań liniowych, wyznaczanie wartości i wektorów własnych macierzy, FFT (szybka transformata Fouriera), operacje na tensorach, wyznaczanie pierwiastków wielomianu, itp.

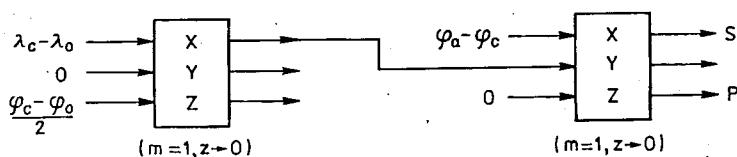
3. Przetwarzanie i rozpoznawanie obrazów, radiolokacja, biomedycyna, nawigacja i kierowanie ruchem, cyfrowa obróbka sygnałów oraz ich filtracja.

Poniżej omówimy kilka konkretnych przykładów zastosowań. Weźmy jako pierwszy problem wyznaczania odległości S_{cr} i namiaru P obiektu przy znanych współrzędnych geograficznych miejsca obserwacji (φ_c, λ_c) i obserwowanego obiektu (φ_0, λ_0) . W przypadku niewielkich odległości poszukiwane wielkości można znaleźć z podanych niżej zależności [9]:

$$S_{cr} = \sqrt{(\varphi_0 - \varphi_c)^2 + (\lambda_0 - \lambda_c)^2 \cos^2((\varphi_0 + \varphi_c)/2)}.$$

$$P = \arctg \frac{(\lambda_c - \lambda_0) \cos((\varphi_0 + \varphi_c)/2)}{\varphi_0 - \varphi_c}.$$

Na rys. 4 pokazano schemat układu zawierającego 2 rotatory, pozwalającego na szybkie wyznaczenie poszukiwanych wielkości.



Rys. 4. Rotatorowy układ obliczający odległość S_{cr} i namiar P obiektu przy znanych współrzędnych geograficznych miejsca obserwacji (φ_c, λ_c) i miejsca obserwowanego (φ_0, λ_0)

Innym przykładem uniwersalności zastosowań rotatorów jest możliwość ich wykorzystania do wyznaczenia szybkiej transformaty Fouriera [7,9]. W trakcie obliczeń wielokrotnie wykonuje się operacje na liczbach zespolonych wg podanego niżej schematu:

$$A_{l+1}(j) = A_l(j) + A_l(p) e^{-2\pi i p r / N}, \quad (12a)$$

$$A_{l+1}(p) = A_l(j) - A_l(p) e^{-2\pi i p r / N}, \quad (12b)$$

gdzie $l = 0, 1, \dots, \log_2 N - 1$, $r = 0, 1, \dots, N - 1$, N – liczba próbek, $i = \sqrt{-1}$. Uwzględniając, że

$$\begin{aligned} \operatorname{Re}[A_1(p)] e^{-2\pi i p r / N} &= \operatorname{Re}[A_l(p)] \cos(2\pi p r / N) \\ &\quad - \operatorname{Im}[A_l(p)] \sin(2\pi p r / N), \\ \operatorname{Im}[A_1(p)] e^{-2\pi i p r / N} &= \operatorname{Im}[A_l(p)] \cos(2\pi p r / N) \\ &\quad + \operatorname{Re}[A_l(p)] \sin(2\pi p r / N). \end{aligned}$$

otrzymuje się:

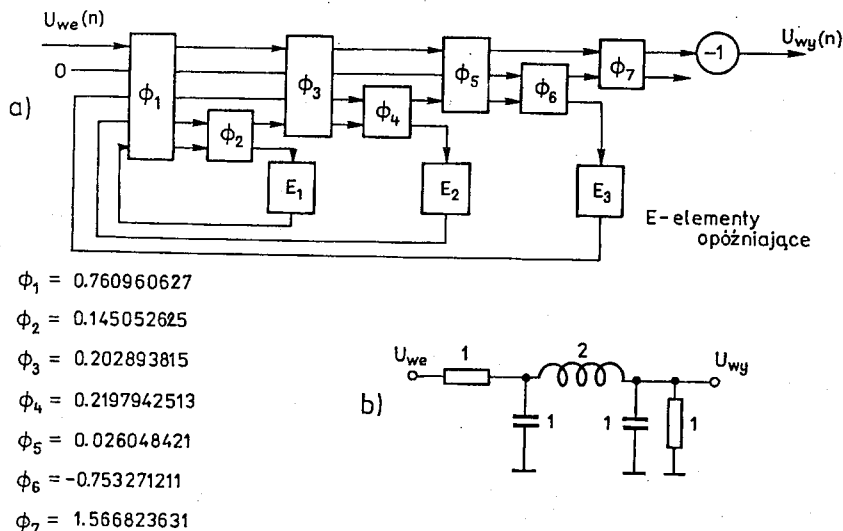
$$\begin{bmatrix} \operatorname{Re}[A_{l+1}(j)] \\ \operatorname{Im}[A_{l+1}(j)] \end{bmatrix} = \begin{bmatrix} \operatorname{Re}[A_l(j)] \\ \operatorname{Im}[A_l(j)] \end{bmatrix} + \begin{bmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{bmatrix} \begin{bmatrix} \operatorname{Re}[A_l(p)] \\ \operatorname{Im}[A_l(p)] \end{bmatrix}. \quad (13a)$$

$$\begin{bmatrix} \operatorname{Re}[A_{l+1}(p)] \\ \operatorname{Im}[A_{l+1}(p)] \end{bmatrix} = \begin{bmatrix} \operatorname{Re}[A_l(j)] \\ \operatorname{Im}[A_l(j)] \end{bmatrix} - \begin{bmatrix} \cos\varphi & -\sin\varphi \\ \sin\varphi & \cos\varphi \end{bmatrix} \begin{bmatrix} \operatorname{Re}[A_l(p)] \\ \operatorname{Im}[A_l(p)] \end{bmatrix}, \quad (13b)$$

gdzie $\varphi = 2\pi pr/N$.

Równania (13) zawierają, jako składowe rotowanych wektorów, części rzeczywiste i urojone j -tego i p -tego współczynnika wyznaczanego w $l+1$ etapie przetwarzania danych wejściowych (wartości próbek). Dla ich wyznaczania wystarczy wykonywać operacje rotacji i sumowania arytmetycznego. Nie trzeba mnożyć ani obliczać funkcji trygonometrycznych. Po wykonaniu całości obliczeń zwykle wykonuje się przejście ze współrzędnych prostokątnych na biegunowe dla otrzymania widma amplitudowego i fazowego analizowanego przebiegu. Odpowiednia operacja, zwana wektoryzacją kołową, jest również wykonywana przez rotator przy $m=1$ i $y \rightarrow 0$ (rys. 3a). Takie same operacje wykorzystywane są przy sporządzaniu wykresów Bodego i Nyquista.

Interującym przykładem zastosowań rotatorów jest ich użycie do konstrukcji filtrów, a w szczególności tzw. ortogonalnych filtrów cyfrowych, które są odpowiednikiem analogowych filtrów bezstratnych. Filtry te cechuje mała wrażliwość funkcji transmitancji na zmiany parametrów charakteryzujących te filtry, niski poziom szumów kwantyzacji i zniekształceń nadmiarowych, przy jednoczesnym małym prawdopodobieństwie wzbudzenia się oscylacji pasożytniczych [2,3,5].



Rys. 5. Ortogonalny filtr cyfrowy 3-go stopnia (a) i jego odpowiednik analogowy

Ortogonalny filtr cyfrowy składa się z pewnej liczby rotatorów o ustalonych kątach rotacji, z elementów opóźniających, którymi są rejestry przechowujące dane, oraz mnożników przez -1 . Istnieje kilka opracowanych procedur syntezy tych filtrów, opisanych m.in. w pracach [2,3,5]. Jako przykład zastosowań rotatorów na

rys. 4 pokazano filtr będący cyfrowym odpowiednikiem analogowego dolno-przepustowego filtra Butterwortha 3. rzędu. Z rysunku wynika, że filtr ten wymaga użycia 9. rotatorów, 3. elementów opóźniających (rejestrów) oraz 1. mnożnika przez -1 . Pracą tego dość złożonego układu zarządza, nie pokazany na rysunku, układ sterujący. Kosztem zmniejszenia szybkości przetwarzania sygnałów można uprościć budowę filtra. Osiąga się to dzięki multipleksowej pracy rotatora oraz wykorzystanie jego rejestrów jako elementów opóźniających. W sumie, mimo tych uproszczeń, jest to nadal bardzo rozbudowany układ, który, aby mógł poprawnie pracować, powinien być wykonany jako jeden układ scalony VLSI.

PODSUMOWANIE

Rotatory cyfrowe, jak to pokazano w artykule, są specjalizowanymi procesorami, wykonującymi różnorodne operacje matematyczne przy zastąpieniu mnożenia rotacją wektorów. W odróżnieniu od układów mikroprocesorowych wykonują złożone operacje nie programowo, ale pod nadzorem układu sterującego. Dzięki temu cechują się dużą szybkością przetwarzania danych. Od strony funkcjonalnej są to układy wykonujące w sposób iteracyjny 3 proste równania liniowe z uwzględnieniem pozycyjnego (dwójkowego) systemu zapisu danych liczbowych. Wykorzystują przy tym najprostsze operacje jakie mogą wykonywać układy cyfrowe: przesyłanie i przesuwanie danych, rozpoznawanie bitów znakowych i sumowanie algebraiczne. Rotatory mogą być konstruowane z podstawowych układów cyfrowych takich jak rejestry, multipleksery, sumatory itp. Mogą wchodzić w skład procesorów macierzowych i struktur systolitycznych [12], a także, jak w ortogonalnych filtrach cyfrowych, być jednym z elementów składowych realizujących odpowiednie procedury obliczeniowe. Duża różnorodność wykonywanych operacji, szeroki zakres zastosowań i nieskomplikowana budowa czynią z rotatora jeden z bardziej interesujących układów cyfrowych.

BIBLIOGRAFIA

1. J. V. Volder: *The CORDIC trigonometric computing technique*. IRE Trans., Vol. E. C-8, p. 330-334, Sept. 1959
2. M. S. Piekar ski: *Stanowa metoda syntezy macierzy transmitancji cyfrowych filtrów ortogonalnych*. Materiały X Krajowej Konferencji TOiUE, t. I, cz. I, str. 145-152
3. M. S. Deprettere, P. Dewilde: *Orthogonal Cascade Realization of Real Multiport Digital Filters*. Circuit Theory Appl., 8; 245-272, 1980
4. R. Draho ka up il, Cz. Micha li k, J. Sta ni s ł a w s k i: *Programowa i sprzętowa realizacja algorytmu Voldera*. Raport ITA PWr. I-28/S-051/88, 1988
5. R. Udo, E. F. Deprettere, P. Dewilde: *On the implementation of orthogonal and orthogonalizing algorithm using pipelined CORDIC architectures*. Signal Processing, Ed. H. W. Schussler, North-Holland, (EURASIP), pp. 847-850, 1983

6. G. L. Haviland, A. Tuszyński: *A CORDIC Arithmetic Processor Chip*. IEEE Trans. on Computers, 1980, vol. C-29, no. 2, pp. 68-78
7. A. M. Despain: *Fourier transforms computers using CORDIC iterations*. IEEE Trans. on Computers, 1974, vol. C-23, no. 10, pp. 993-10001
8. W. E. Egbert: *Personal calculator algorithm II: Trigonometric functions*. Hewlett Packard J., 1977, vol. 28, no. 10, pp. 17-20
9. W. D. Bajkov, W. B. Smolov: *Specjalizowanyje processory - iteracionnyje algoritmy i struktury*. Moskva, Radio i Sviaz. 1985
10. R. Drahokaupil, C. Michalik, J. Stanisławski: *Skaler binarny*. Materiały XI Krajowej Konferencji TOiU, str. 218-223, 1988
11. J. S. Walther: *A Unified Algorithm for Elementary Functions*. AFIPS Conf., vol. 38, 1971, SJCC, pp. 379-385
12. S. Y. Kung, H. J. Whitehouse, T. Kailath (editors): *VLSI and Modern Signal Processing*. Prentice-Hall, Inc., Englewood Cliffs, New Jersey 07632 lub tłumaczenie rosyjskie
13. M. S. Piekarski: *Koncepcja i zastosowanie rotatora zespolonego do syntezy filtrów cyfrowych*. Materiały XIII Konferencji TOiUE; Bielsko-Biała, 1990, s. 295-300
14. R. Drahokaupil, C. Michalik, J. Stanisławski: *Rotatorowy moduł arytmetyczny*. Raport ITA PWr. I-28/P-046/89

R. DRAHOKAUPIL, CZ. MICHALIK, J. STANISŁAWSKI

DIGITAL ROTATORE - THEORY AND APPLICATION

Summary

In this paper we present properties, work principle and structures of digital rotors working according to the Volder's algorithm. The rotors realize the procedure of calculation of vector coordinates rotated on the plane by means of three differential equations of first degree.

We gave an example of this rotor realization, using typical digital circuits. We also discussed the nature of these equations, modified by Wather, enabling the realization of circular, linear and hyperbolic rotation. Universality of the rotation processors working according to this algorithm was pointed out. We also gave the typical application of rotors for making basic mathematical operations and functions. In the end, we discussed some examples of rotor application for calculation of geographical coordinates, FFT end orthogonal digital filtering.

SPIS TREŚCI

J. B a l i c k i: Polioptymalizacja sprzętowej struktury sieci komputerowej	3
M. K o w a l s k i, S. B e k: Analiza stabilności liniowych obwodów aktywnych zawierających oporniki o okresowo przełączanych rezystancjach	13
H. B u d z i s z: Reprezentacja wiedzy o układach elektronicznych. Część I: Reprezentacja w języku klauzul	23
Z. W r ó b e l: Funkcje strukturalne grafów układów dwójników i czwórników składających się z dwóch rodzajów elementów	35
Z. W r ó b e l: Homeomorfizm grafów układów dwójników i czwórników składających się z dwóch rodzajów elementów	53
Z. W r ó b e l, Synteza struktur kanonicznych bi-grafów dwójników składających się z dwóch rodzajów elementów. Część 1.: Metoda przeglądu bezpośredniego	67
Z. W r ó b e l: Synteza struktur kanonicznych bi-grafów dwójników składających się z dwóch rodzajów elementów. Część 2: Metoda przeglądu hierarchicznego	85
T. W y s o c k i, L. J. W e i s s, F. W y s o c k a - S c h i l l a k: Aproksymacja statycznej charakterystyki wyjściowej tranzystora bipolarnego za pomocą funkcji wykładniczej	99
A. J a w o r e k: Wpływ parametrów wzmacniacza pomiarowego na pomiar składowych immitancji dwójników za pomocą detektorów fazoczułych	117
Z. L i s i k: Modelowanie generacji ciepła w przyrządach półprzewodnikowych. Część I: Bilans energetyczny struktury półprzewodnikowej	141
Z. L i s i k: Modelowanie generacji ciepła w przyrządach półprzewodnikowych. Część II: Funkcje rozpraszania ciepła	167
Z. L i s i k: Modelowanie generacji ciepła w przyrządach półprzewodnikowych: Część III: Wkład efektów termoelektrycznych do bilansu energetycznego przyrządu półprzewodnikowego	189
A. M a t e r k a, W. G ó r n i s i e w i c z, M. K u k u ł a: Skalowanie i korekcja zniekształceń geometrycznych w systemie cyfrowego przetwarzania obrazów z kamerą widikonową	211
N. K i m S a c h: Cyfrowe odtwarzanie i uwydatnianie krawędzi dwuwymiarowych obrazów zaszumionych	231
K. A. K a s s e m, P. D y m a r s k i: Kwantyzer adaptacyjny w modulacji DPCM	239
A. C y s e w s k a - S o b u s i a k: Model czujnika fotoelektrycznego jako fotodetektora w układzie otwartego łącza optycznego	253
N. K i m S a c h: Możliwości poprawy jakości cyfrowych obrazów telewizyjnych w czasie rzeczywistym	267
T. A d a m s k i: Układy systoliczne do obliczania mediany i ich zastosowanie w korektach jitteru momentu próbkowania	273
J. K i s i l e w i c z: Eliminacja interferencji międzykanałowej	283
A. C z e m p l i k: Organizacja zadania wizualizacji w komputerowych systemach automatyki ..	293
M. W c i ś l i k: Moc bierna w obwodzie z odbiornikiem nieliniowym	305
J. Ł y s k a n o w s k i: Zabezpieczenie procesów elektromedycznych przed zagrożeniami elektrycznymi	317
R. D r a h o k a u p i l, J. S t a n i s ł a w s k i, C z. M i c h a l i k: Rotatory cyfrowe — teoria i zastosowanie	337

CONTENTS

J. B a l i c k i: The vector optimization of the hardware structure of computer networks.....	3
M. K o w a l s k i, S. B e k: Stability analysis of linear active circuits containing periodically switched resistors.....	13
H. B u d z i s z: Knowledge representation in the domain of electronic circuits. Part I: Representation in the clause language.....	23
Z. W r ó b e l: Structural functions of graphs of one-port and two-port circuits which consist of two kinds of elements.....	35
Z. W r ó b e l: Homeomorfizm of graphs of one-port and two-port circuits consisting of two kinds elements.....	53
Z. W r ó b e l: Synthesis of canonical structures of bigraphs one-port circuits consisting of two-kinds of elements. Part 1. Direct reviewing method.....	67
Z. W r ó b e l: Synthesis of canonical structures of bigraphs one-port circuits consisting of two-kinds of elements. Part 2. Hierarchical reviewing method.....	85
T. W y s o c k i, L. J. W e i s s, F. W y s o c k a - S c h i l l a k: The exponential approximation of bipolar transistor static output characteristics.....	99
A. J a w o r e k: Operational amplifier parameters influence on results of measurement of a two-terminal immittance components by means of phase-sensitive detectors.....	117
Z. L i s i k: Modelling of heat generation in semiconductor devices. Part I: Energy ballance of semiconductor structure.....	141
Z. L i s i k: Modelling of heat generation in semiconductor devices. Part II: Functions of heat generation.....	167
Z. L i s i k: Modelling of heat generation in semiconductor devices. Part III: The contribution of thermoelectric effects to the energy ballance of semiconductor device.....	189
A. M a t e r k a, W. G ó r n i s i e w i c z, M. K u k u ł a: Calibration and geometric distortion correction in a digital image processing system with a vidicon camera.....	211
N. K i m S a c h: Digital 2—D picture restoration and enhancement.....	231
K. A. K a s s e m, P. D y m a r s k i: Adaptive quantizers in DPCM codes.....	239
A. C y s e w s k a - S o b u s i a k: Model of the photoelectric sensor as the photodetector in the system of the open optical coupler.....	253
N. K i m S a c h: Possibilities of digital TV picture processing in real time.....	267
T. A d a m s k i: Systolic circuits for fast median computation with application to sampling time jitter correctors.....	273
J. K i s i l e w i c z: Interchannel interference elimination.....	283
A. C z e m p l i k: Visualization task in industrial computer control system.....	293
M. W c i ś l i k: Reactive power in a circuit with a nonlinear load.....	305
J. Ł y s k a n o w s k i: Protection of electromedical process against electrical hazards.....	317
R. D r a h o k a u p i l, J. S t a n i s ł a w s k i, C z. M i c h a l i k: Digital rotators — theory and application.....	337

Informujemy, że z przyczyn technicznych numer 1-2
z 1991 r ukaże się przed numerem 3-4 z 1990 r.

Kwartalnik Elektroniki i Telekomunikacji

Wpłaty na prenumeratę przyjmowane są na okresy kwartalne:

- na terenie kraju — jednostki kolportażowe „Ruch” i urzędy pocztowe właściwe dla miejsca zamieszkania lub siedziby prenumeratora,
- na zagranicę — Zakład Kolportażu Prasy i Wydawnictw, 00-958 Warszawa, konto PBK XIII Oddział Warszawa 370044-1195-139-11.

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

Dostawa zamówionej prasy następuje:

- przez jednostki kolportażowe „Ruch” — w sposób uzgodniony z zamawiającym,
- przez urzędy pocztowe — pocztą zwykłą na wskazany adres, w ramach opłaconej prenumeraty z wyjątkiem zlecenia dostawy za granicę pocztą lotniczą do odbiorcy zagranicznego, której koszt w pełni pokrywa prenumerator.

Terminy przyjmowania wpłat na prenumeratę:

- krajową i zagraniczną — 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.

Bieżące i archiwalne numery można nabyć w Księgarni Wydawnictwa Naukowego PWN Sp. z o.o. ul. Miodowa 10. 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.

Subscription orders for all the magazines published in Poland available through the local press distributors or directly through the

Foreign Trade Enterprise

ARS POLONA

00-068 Warszawa, Krakowskie Przedmieście 7, Poland

Our bankers:

BANK HANDLOWY S.A. 201061-710-13100

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

PL ISSN 0035-9386

Nr Indeksu 363189

KWARTALNIK ELEKTRONIKI I TELEKOMUNIKACJI

ELECTRONICS AND TELECOMMUNICATIONS QUARTERLY

dawniej

ROZPRAWY ELEKTROTECHNICZNE

TOM XXXVII — ZESZYT 3-4

WYDAWNICTWO NAUKOWE PWN
WARSZAWA 1991

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

**KWARTALNIK
ELEKTRONIKI I TELEKOMUNIKACJI**

**ELECTRONICS AND
TELECOMMUNICATIONS
QUARTERLY**

TOM XXXVII — ZESZYT 3-4

WYDAWNICTWO NAUKOWE PWN

WARSZAWA 1991

RADA REDAKCYJNA

Przewodniczący
prof. dr ADAM SMOLIŃSKI
członek rzeczywisty PAN

Członkowie
prof. dr hab. DANIEL JÓZEF BEM, prof. dr hab. 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 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
Redakcja naczelny
prof. dr hab. inż. WIESŁAW WOLIŃSKI

Zastępca Redaktora Naczelnego
doc. dr inż. KRYSTYN PLEWKO

Sekretarz Odpowiedzialny
mgr inż. KRYSTYNA LELAKOWSKA

ADRES REDAKCJI
00-665 Warszawa, ul. Nowowiejska 15/19 Politechnika,
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
Sekretarz Odpowiedzialnego: 25 29 18

WYDAWNICTWO NAUKOWE PWN
Warszawa, ul. Miodowa 10

Ark. wyd. 28,25 Ark. druk. 22,75	Podpisano do druku w maju 1992 r.
Papier offsetowy kl. III 70 g. B-1	Druk ukończono w czerwcu 1992 r.

Skład „GP-BIT” — W-wa ul. Marymoncka 34. Druk i oprawa: Drukarnia — W-wa, ul. Hajoty 50.

On the constitutive distributed parameter modelling of the electromagnetic field

Part I. The constitutive distributed parameter mathematical model of the electromagnetic field

WACŁAW NIEMIEC

Politechnika Śląska, Gliwice

Received 1990.06.20

Authorized 1991.02.05

An original constitutive distributed parameter mathematical model of the electromagnetic field has been suggested in the article. The system of the partial differential constitutive state equations in the forms of the following derivatives for: — charge and electric induction, — electromagnetic energy, — momentum of electromagnetic medium, has been deduced. This system contains all literature motivated physical phenomena and parameters modelling real space-time distributed problems of the electromagnetic field for the physical control influences of this field.

1. INTRODUCTION

The literature concerning the theory and applications of the electromagnetic field does not contain its complete description by the use of: — charge, — energy and — momentum of this field [2–11]. A new proposition of the description of the electromagnetic field suggested in this article to my previous work of the constitutive distributed parameter modelling of the electric field is related [1]. This approach has its basis on [2–11]:

- analysis of the physical phenomena of the yield and quality aspects of the electromagnetic field,
 - the source of the emission or adsorption of the electromagnetic field,
 - the environment conditions of the existence of the electromagnetic field,
- and the mathematical tools of the constitutive distributed parameter modelling:
- reciprocity principle for the isotropic and anisotropic nonhomogeneous media with the space and time memories [15],

– the definition of the following derivative and the definition of the partial differential equation of the continuity for every state variable of the electromagnetic field [12], [1],

– the constitutive invariance for the continuous media with the space and time memories [13], [1],

– the elements of the mathematical field theory [14],

– the properties of the charge, energy and momentum system of the partial differential constitutive state equations of the electric field [1].

In the considerations of this article we chose the state variables of the electromagnetic field using the method of the mathematical modelling for the distributed parameter control theory [16], [1].

2. SOME LOCALLY DISTRIBUTED PARAMETER ASPECTS OF THE ELECTROMAGNETIC FIELD

For the theory of apparatuses for the emission and adsorption of the electromagnetic field the underlying significance for the profile of the electromagnetic field have the locally distributed parameters surrounding its source [2–11]. Making use of the literature of electromagnetic field studies we can suggest the following approach to the constitutive distributed parameter modelling of this field:

– the emission or adsorption equivalent source point $Z(x, y, z, t)$,

– the cuboidal locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ (phenomenally closed – every phenomenon can be treated in the separate way) of the physical conditions of the emission or adsorption intensity and quality questions,

– the so called „working” space volume-time Ω_{Rt} defined as limited or unlimited of the distribution of the electromagnetic field.

As a consequence of the utilization of these ideas the following relation can be stated [1]:

$$Z(x, y, z, t) \in \bar{\Omega}_{(a,b,c)t} \in \Omega_{Rt} \quad (2.1)$$

Let Ω_{Rt} be determined by the so called „variable point” $Q^R(\xi^R, \eta^R, \beta^R, \tau)$ and is fulfilled

$$Q^R(\xi^R, \eta^R, \beta^R, \tau) \in \Omega_{Rt} \quad (2.2)$$

The following reciprocal constant image is delivered [15], [12]

$$Z(x, y, z, t) \overset{\text{in } \Omega_{Rt}}{\leftrightarrow} Q^R(\xi^R, \eta^R, \beta^R, \tau) \quad (2.3)$$

The reciprocal idea of the locally distributed parameters is presented in the figure Fig. 1.

The variable point $Q^R(\xi^R, \eta^R, \beta^R, \tau)$ assures its presentation for the potential fields $Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q(\xi, \eta, \beta, \tau)$ and for the rotational fields $Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q'(\xi', \eta', \beta', \tau)$ pertinent to the physical phenomena of the electromagnetic field. The circulation of the normal external surface orientation vector $n^{(z)}$ of the surface $F_{(a,b,c)t}$ of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ in relation

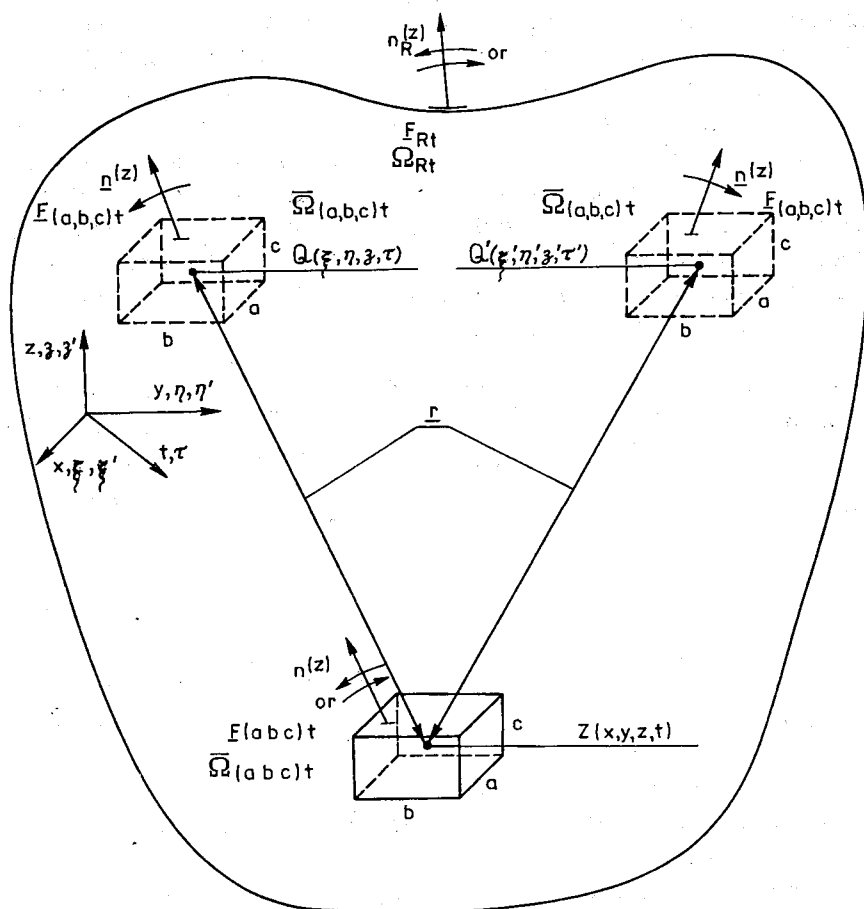


Fig. 1. A geometric interpretation of the locally distributed parameter problems of the modelling of the electromagnetic field

$Z(x, y, z, t)$ — point of the activity of the function transformation of the electromagnetic field,
 $\bar{\Omega}_{(x, y, z)t}$ — the locally selected volume-time (phenomenally closed — every phenomenon can be treated in the separate way) around point $Z(x, y, z, t)$,

$F_{(a, b, c)}$ — the external oriented surface of $\bar{\Omega}_{(a, b, c)t}$,

$\mathbf{n}^{(z)}$ — the normal external surface orientation vector of $F_{(a, b, c)t}$,

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

Ω_{Rt} — the working space volume-time for the electromagnetic field,

F_{Rt} — the external oriented surface of Ω_{Rt} ,

$\mathbf{n}_R^{(z)}$ — the normal external surface orientation vector of F_{Rt} ,

$Q^R(\xi^R, \eta^R, \beta^R, \tau)$ — the picture of point of the transformation of the electromagnetic field from the point view of its space and time memories,

$Q(\xi, \eta, \beta, \tau) \equiv Q^R(\xi^R, \eta^R, \beta^R, \tau)$ — for the physical phenomena of the potential fields,

$Q'(\xi', \eta', \beta', \tau) \equiv Q^R(\xi^R, \eta^R, \beta^R, \tau)$ — for the physical phenomena of the rotational fields,

r — the radius of the activity of the forces of the electromagnetic field,

The circulation of the vector $\mathbf{n}_R^{(z)}$ in relation to the circulation of the vector $\mathbf{n}^{(z)}$ depends to the boundary control problems on the surface F_{Rt} and the signs of the summations of the physical phenomena of the potential fields and rotational fields in adequate charge, energy and momentum balances of the electromagnetic field.

to the circulation of the normal external surface orientation vector $\mathbf{n}_R^{(z)}$ of the surface F_{Rt} of the working space volume-time Ω_{Rt} determines the signs of the balance summations of the effects of the physical phenomena of the electromagnetic field:

- for the potential fields „+” in charge, energy and momentum balances,
- for the rotational fields „–” in the charge, energy and momentum balances.

On the basis of the literature studies of the electromagnetic field problems of the many quoted authors it is necessary to determine the important state variables of this field giving the ability to perform further considerations of this article [2–11], [13]. The reciprocal interpretation of the state vectors of the potential fields and the rotational fields can be presented as follows:

$$r < \begin{bmatrix} \rho(\xi, \eta, \beta, \tau) \\ D(\xi, \eta, \beta, \tau) \\ B(\xi, \eta, \beta, \tau) \\ v(\xi, \eta, \beta, \tau) \end{bmatrix} \leftrightarrow r < \begin{bmatrix} \rho(x, y, z, t) \\ D(x, y, z, t) \\ B(x, y, z, t) \\ v(x, y, z, t) \end{bmatrix} \leftrightarrow r < \begin{bmatrix} \rho(\xi', \eta', \beta', \tau) \\ D(\xi', \eta', \beta', \tau) \\ B(\xi', \eta', \beta', \tau) \\ v(\xi', \eta', \beta', \tau) \end{bmatrix} \quad \text{in } \Omega_{Rt} \quad (2.4)$$

of points

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

potential fields

rotational fields

r – related double formed state variables.

The state variables D , B and v have been chosen for the constitutive description of the electromagnetic field, because [1–12]:

- the vector of the induction of the electric field „ D ” is stright connected to the electric charge „ ρ ”

$$\operatorname{div} D = \rho \quad \left[\frac{C}{m^3} \right] \quad (2.5)$$

and

$$D = \varepsilon_0 \varepsilon E \quad \left[\frac{C}{m^2} \right] \quad (2.6)$$

- the scalar product „ DB ” is related to the module form of the Poynting vector

$$EH = |E \times H| \quad \left[\frac{W}{m^2} \right] \quad (2.7)$$

with

$$B = \mu_0 \mu H \quad [T] \quad (2.8)$$

the field vector velocity „ v ” with the mass of the electric charge is a momentum relation

$$f\rho v \quad \left[\frac{\text{kg}}{\text{m}^2\text{s}} \right] \quad (2.9)$$

The state variables D , B and v are invariant in respect to the properties of the medium in which the electromagnetic field exists, [2—11].

3. THE PHYSICAL CONDITIONS OF THE ELECTROMAGNETIC FIELD

In the range of the macroscopic physical phenomena there exists inseparable alliance between the electric field and the magnetic field determining both together the electromagnetic field. This is expressed that by the change of the electric field in time appears magnetic field, and on the contrary, this kind of the change of the magnetic field is accompanied by the electric field. To formulate these relations between electric and magnetic fields a necessary task is to consider possible all physical phenomena of the electric field and magnetic field separately and their space-time distributed mutual aspects. The proposed approach to the constitutive distributed parameter modelling of the electromagnetic field has its basis which can be divided to:

1. — the tasks from the domain of the analysis of the physical phenomena being decisive for the yield and quality aspects of the electromagnetic field,
2. — the tasks of the scope of the synthesis of the above physical phenomenal information to the constitutive distributed parameter model of the electromagnetic field by the use of adequate balance rules for its charge, energy and momentum balances.

The grounds of our considerations on the following assumptions are established [2—11], [1], [12]; [13]:

- ASSUMPTION 1. The environment is anisotropic and nonhomogeneous.
- ASSUMPTION 2. The static and dynamic aspects of the electric and magnetic problems remain valid.
- ASSUMPTION 3. The propagation of the electromagnetic wave is small in relation to the other energy forms of the electromagnetic field.
- ASSUMPTION 4. The thermic energies of the transformations of the electric field to the magnetic field and the magnetic field to the electric field can be neglected.
- ASSUMPTION 5. All literature motived physical phenomena being decisive for the yield and quality aspects of the electromagnetic field are used into our considerations.
- ASSUMPTION 6. The continuity partial differential formulae for the electric field and magnetic field

$$\text{div } D = \rho \quad (3.1)$$

and

$$\text{div } B = 0 \quad (3.2)$$

for the content of this work remain valid.

ASSUMPTION 7. The state variables D, B and v characteristic for the electromagnetic field are all functions of four constant space time coordinates (x, y, z, t) of Ω_{R^4} .

ASSUMPTION 8. For the constant coordinates system (x, y, z, t) we introduce the variable coordinates system $(\xi^R, \eta^R, \beta^R, \tau)$ moving with the velocity „ v ” and „ v ” $\ll c$ (velocity of light).

An arbitrary chosen vector state variable A has its presentations:
 – in constant coordinates system

$$A = A(x, y, z, t) \quad (3.3)$$

and

– in variable coordinates system

$$A^R = A^R(\xi^R, \eta^R, \beta^R, \tau) \quad (3.4)$$

with the velocity vector

$$v = \frac{dx}{dt}i + \frac{dy}{dt}j + \frac{dz}{dt}k, \quad (3.5)$$

where (i, j, k) – versors.

The derivative of the vector state variable A is:

– for the constant coordinates system because $v=0$ we have

$$\frac{DA}{Dt} = \frac{\partial A}{\partial t} \quad (3.6)$$

and consequently

– for the variable coordinates system

$$\frac{DA}{Dt} = \frac{\partial A}{\partial t} + (v \text{grad})A \quad (3.7)$$

with

$$(v \text{grad})A = v \text{div}A - \text{rot}(v \times A) \quad (3.8)$$

REMARK 1. Formula (3.7) is a ground form of the definition of the following derivative for the vector case.

ASSUMPTION 9. The space and time memories of the transformations of the electric field and magnetic field for the appearance and existence of the electromagnetic field exist.

According to the cited literature of the electromagnetic field [2–11], the physical elements of the charge, energy and momentum balances of the electromagnetic field in Table 1 are placed.

For the construction of the partial differential constitutive state equations of the electromagnetic field the below presented balance formula is useful [1], [12]:

The mathematical formulae of the charge, energy and momentum balances of the electromagnetic field

Balances Phenomena		Charge	Energy	Momentum
Diffusion	electric	$j_D^e = -D_e \text{ grad } \rho \left[\frac{\text{A}}{\text{m}^2} \right]$	$I_D^e = W_{-D} \iiint_{\Omega(x,y,z)t} \text{div}(D_e \text{ grad}) D \cdot d\Omega_{(x,y,z)} \left[\frac{\text{Ws}}{\text{m}} \right]$	$N_D^e = \iiint_{\Omega(x,y,z)t} (N_D \text{ grad}) D dF_{(x,y,z)} [\text{N}]$
	magnetic		$I_D^m = \iiint_{\Omega(x,y,z)t} \text{div}(U_B \text{ grad}) B \cdot d\Omega_{(x,y,z)} \left[\frac{\text{Ws}}{\text{m}} \right]$	$N_D^m = \iiint_{\Omega(x,y,z)t} (M_B \text{ grad}) B dF_{(x,y,z)} [\text{N}]$
Conductivity		$j_D^e = \gamma E \left[\frac{\text{A}}{\text{m}^2} \right]$	$i_p^e = V_D \gamma E \left[\frac{\text{Ws}}{\text{m}^4} \right]$	
Convection	electric	$j_c^e = \rho v \left[\frac{\text{A}}{\text{m}^2} \right]$	$i_c^{e,m} = D B v \left[\frac{\text{Ws}}{\text{m}^4} \right]$	$N_k^e = -f \iiint_{\Omega(x,y,z)t} [v \text{div}(\rho v) + (\rho v \text{grad}) v] d\Omega_{(x,y,z)t} [\text{N}]$
	magnetic			
Material properties				$N_k^{e,m} = \iiint_{\Omega(x,y,z)t} [a \nabla^2 v + b \text{grad div} v] d\Omega_{(x,y,z)} [\text{N}]$
Source	electric	$J_\Omega^e = \iiint_{\Omega(x,y,z)t} J(t) d\Omega_{(x,y,z)} [\text{A}]$	$I_\Omega^{e,m} = \iiint_{\Omega(x,y,z)t} S(t) d\Omega_{(x,y,z)} \left[\frac{\text{Ws}}{\text{m}} \right]$	$N_\Omega^{e,m} = \iiint_{\Omega(x,y,z)t} T(t) d\Omega_{(x,y,z)t} [\text{N}]$
	magnetic			
Increase	electric	$\frac{\partial}{\partial t} \iiint_{\Omega(x,y,z)t} \rho d\Omega_{(x,y,z)} [\text{A}]$	$\frac{\partial}{\partial t} \iiint_{\Omega(x,y,z)t} D B d\Omega_{(x,y,z)} \left[\frac{\text{Ws}}{\text{m}} \right]$	$\frac{\partial}{\partial t} \iiint_{\Omega(x,y,z)t} \rho v d\Omega_{(x,y,z)} [\text{N}]$
	magnetic			

$$\frac{\partial}{\partial t} \iiint_{\Omega_{(x,y,z)}} \left(\begin{array}{c} \text{Increase of elements} \\ \text{of balance.} \end{array} \right) d\Omega_{(x,y,z)} = \iint_{F_{(x,y,z)}} \left(\begin{array}{c} \text{Fluxes of physical phenomena} \\ \text{pertinent to balance.} \end{array} \right)$$

$$\begin{aligned} dF_{(x,y,z)} + \iiint_{\Omega_{(x,y,z)}} \left(\begin{array}{c} \text{Values of physical phenomena} \\ \text{pertinent to balance.} \end{array} \right) d\Omega_{(x,y,z)} + \\ + \iiint_{\Omega_{(x,y,z)}} \left(\begin{array}{c} \text{Source element} \\ \text{pertinent to balance.} \end{array} \right) d\Omega_{(x,y,z)} \end{aligned} \quad (3.9)$$

By the moving coordinates system we use for the charge, energy and momentum balances the definition of the following derivative [12], [1], [13]

$$\frac{D}{Dt} = \frac{\partial}{\partial t} - \mathbf{v} \text{grad} + (\text{sum of adequate gradient operations}) \quad (3.10)$$

and the partial differential equations of the continuity as the constitutive invariance forms fulfil for every balance general condition [12], [1], [13]

$$\frac{D}{Dt} = 0 \quad (3.11)$$

which precisely can be written as [12], [1]

$$\frac{\partial}{\partial t} = \mathbf{v} \text{grad} - (\text{sum of adequate gradient operations}) \quad (3.12)$$

by the application of the condition (3.11) into the eq. (3.10).

4. THE CONSTRUCTION OF CHARGE, ENERGY AND MOMENTUM PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATIONS OF THE ELECTROMAGNETIC FIELD

4.1. THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATION OF THE CHARGE BALANCE OF THE ELECTROMAGNETIC FIELD

From the content of the charge coordinates of the Table 1 the outset formula of the charge balance is:

$$\frac{\partial}{\partial t} \iiint_{\Omega_{(x,y,z)}} \rho d\Omega_{(x,y,z)} = \iint_{F_{(x,y,z)}} [D_e \text{grad } \rho - \rho \mathbf{v}] dF_{(x,y,z)} + \iint_{F_{(x,y,z)}} (\gamma \mathbf{E}) dF_{(x,y,z)} +$$

$$+ \iiint_{\Omega(x,y,z)} J(t) d\Omega_{(x,y,z)} \quad (4.1.1)$$

Making use of the Gauss–Ostrogradski law [14], [12], we have:

$$\frac{\partial \rho}{\partial t} = \operatorname{div}(D_e \operatorname{grad} \rho) - \operatorname{div}(\rho \mathbf{v}) + \operatorname{div}(\gamma \mathbf{E}) + J(t) \quad (4.1.2)$$

and after simple transformations, the last following derivative form of the above equation can be written as:

$$\frac{D\rho}{Dt} = D_e \nabla^2 \rho - \rho \operatorname{div} \mathbf{v} + \gamma \operatorname{div} \mathbf{E} + J(t). \quad (4.1.3)$$

The partial differential equations of the charge continuity of the electromagnetic field from the condition $\frac{D\rho}{Dt} = 0$ [12], [1] are:

– for the variable coefficients

$$D_e = D_e(x, y, z, t) \quad \text{and} \quad \gamma = \gamma(x, y, z, t) \\ \frac{\partial \rho}{\partial t} = \operatorname{grad} D_e \operatorname{grad} \rho - \mathbf{v} \operatorname{grad} \rho + \mathbf{E} \operatorname{grad} \gamma, \quad (4.1.4)$$

or

– for the constant coefficients

$$D_e \neq D_e(x, y, z, t) \quad \text{and} \quad \gamma = \gamma(x, y, z, t) \\ \frac{\partial \rho}{\partial t} = \mathbf{v} \operatorname{grad} \rho. \quad (4.1.5)$$

In accordance with the above partial differential equations of the continuity the eq. (4.1.3) is invariant to its coefficients D_e and γ and *these coefficients are treated as being constant for the future considerations.*

Introducing now the definition [2–11]

$$\rho = \operatorname{div} \mathbf{D} \quad (4.1.6)$$

into the eq. (4.1.3) its resultant form is:

$$\frac{D \operatorname{div} \mathbf{D}}{Dt} = D_e \nabla^2 \operatorname{div} \mathbf{D} - \operatorname{div} \mathbf{D} \operatorname{div} \mathbf{v} + \frac{\gamma}{\varepsilon_0 \varepsilon} \operatorname{div} \mathbf{D} + J(t). \quad (4.1.7)$$

From the elimination of the operation ∇D at last we have:

$$\frac{D \mathbf{D}}{Dt} = D_e \nabla^2 \mathbf{D} \operatorname{div} \mathbf{v} + \frac{\gamma}{\varepsilon_0 \varepsilon} \mathbf{D} + J(t). \quad (4.1.8)$$

4.2. THE PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATION OF THE ENERGY BALANCE FOR THE ELECTROMAGNETIC FIELD

The „energy” coordinates of the Table 1 allow on the formulation of the balance equation:

$$\begin{aligned} \frac{\partial \rho}{\partial t} \iiint_{\Omega_{(x,y,z)}} (DB) d\Omega_{(x,y,z)} &= \iiint_{F_{(x,y,z)}} [(U_B \text{grad}) B - DB \mathbf{v}] dF_{(x,y,z)} + \\ &+ \iiint_{F_{(x,y,z)}} \mathbf{v} \frac{\gamma}{D \varepsilon_0 \varepsilon} D dF_{(x,y,z)} + \iiint_{\Omega_{(x,y,z)}} S(t) d\Omega_{(x,y,z)} + \\ &+ W_D \iiint_{\Omega_{(x,y,z)}} \text{div} (D_e \text{grad}) D d\Omega_{(x,y,z)}. \end{aligned} \quad (4.2.1)$$

As a consequence of the utilization of the Gauss – Ostrogradski law [14], the eq. (4.2.1) obtains its new form:

$$\begin{aligned} B \frac{\partial D}{\partial t} + D \frac{\partial B}{\partial t} &= \text{div} U_B \text{div} B + U_B \nabla^2 B - DB \text{div} \mathbf{v} - D(\mathbf{v} \text{grad}) B - \\ &- B(\mathbf{v} \text{grad}) D + D \text{grad} V \frac{\gamma}{D \varepsilon_0 \varepsilon} + V \frac{\gamma}{D \varepsilon_0 \varepsilon} \text{div} D \pm S(t) + W_D \text{grad} D_e \text{div} D + \\ &+ W_D D_e \nabla^2 D. \end{aligned} \quad (4.2.2)$$

Introducing the partial differential equation of the continuity multiplied both handsides by the vector „B”

$$B \frac{\partial D}{\partial t} = -B(\mathbf{v} \text{grad}) D \quad (4.2.3)$$

into the eq. (4.2.2) the result is:

$$\begin{aligned} D \frac{\partial B}{\partial t} &= \text{div} U_B \text{div} B + U_B \nabla^2 B - DB \text{div} \mathbf{v} - D(\mathbf{v} \text{grad}) B + D \text{grad} V_D \frac{\gamma}{\varepsilon_0 \varepsilon} + \\ &+ V_D \frac{\gamma}{\varepsilon_0 \varepsilon} \text{div} D \pm S(t) + W_D \text{grad} D_e \text{div} D + W_D D_e \nabla^2 D. \end{aligned} \quad (4.2.4)$$

Now, let be defined a function

$$DB = \varphi(D, B) \quad (4.2.5)$$

from which we can written:

– the time derivative

$$D \frac{\partial B}{\partial t} = D \frac{\partial B}{\partial t} + W_D \frac{\partial D}{\partial t} \quad \text{and} \quad W_D = \frac{\partial (DB)}{\partial D} \Big|_B, \quad (4.2.6)$$

and consequently

– the space derivative

$$D(\text{grad})B = D(\text{grad})B + W_D \text{div} D. \quad (4.2.7)$$

The eq. (4.2.7) is multiplied both handsides by the field vector velocity „ $\mathbf{v}$ ” to the form

$$D(\mathbf{v} \text{grad})B = D(\mathbf{v} \text{grad})B + \mathbf{v} W_D \text{div} D. \quad (4.2.8)$$

Making use of the eq. (4.2.6) and the eq. (4.2.8) into the eq. (4.2.4) the resulting its form can be written as:

$$\begin{aligned} D \frac{\partial B}{\partial t} + W_D \frac{\partial D}{\partial t} = & \text{div} U_B \text{div} B + U_B \nabla^2 B - DB \text{div} \mathbf{v} - D(\mathbf{v} \text{grad})B - \mathbf{v} W_D \text{div} D + \\ & + D \text{grad} V_D \frac{\gamma}{\varepsilon_0 \varepsilon} + V_D \frac{\gamma}{\varepsilon_0 \varepsilon} \text{div} D \pm S(t) + W_D \text{grad} D_e \text{div} D + \\ & + W_D D_e \nabla^2 D. \end{aligned} \quad (4.2.9)$$

Now, we recall the part form of the partial differential equation of the continuity (4.1.4) multiplied both handsides by „ W_D ”

$$W_D \frac{\partial D}{\partial t} = W_D \text{grad} D_e \text{div} D - \mathbf{v} W_D \text{div} D \quad (4.2.10)$$

The subtraction of the eq. (4.2.10) from the eq. (4.2.9) leads to the formula

$$\begin{aligned} B \frac{\partial D}{\partial t} = & \text{div} U_B \text{div} B + U_B \nabla^2 B - DB \text{div} \mathbf{v} - D(\mathbf{v} \text{grad})B + D \text{grad} V_D \frac{\gamma}{\varepsilon_0 \varepsilon} + \\ & + V_D \frac{\gamma}{\varepsilon_0 \varepsilon} \text{div} D \pm S(t) + W_D D_e \nabla^2 D \end{aligned} \quad (4.2.11)$$

For the simplification we introduce additionally

$$G_D = V_D \frac{\gamma}{\varepsilon_0 \varepsilon} \quad (4.2.12)$$

and the eq. (4.2.11) can be rewritten in the form

$$\begin{aligned} D \frac{\partial B}{\partial t} = & (\text{grad}) U_B (\text{grad}) B + U_B \nabla^2 B - DB \text{div} \mathbf{v} - D(\mathbf{v} \text{grad})B + D \text{grad} G_D + \\ & + G_D \text{div} D \pm S(t) + W_D D_e \nabla^2 D. \end{aligned} \quad (4.2.13)$$

As a consequence of the utilization of the definition of the following derivative [12], [1], the eq. (4.2.13) assumes its final form:

$$D \frac{DB}{Dt} = U_B \nabla^2 B - DB \text{div} \mathbf{v} \pm S(t) + G_D \text{div} D + W_D D_e \nabla^2 D. \quad (4.2.14)$$

For the eq. (4.2.14) we have the partial differential equations of the continuity, obtained from the condition $D \frac{DB}{Dt} = 0$ [12], [1], pertinent to the cases:

— for the variable coefficients

$$\begin{aligned}
 U_B &= U_B(x, y, z, t), V_D = V_D(x, y, z, t), \gamma = \gamma(x, y, z, t), \\
 \varepsilon &= \varepsilon(x, y, z, t), G_D = G_D(x, y, z, t) \text{ and } D_e = D_e(x, y, z, t), \\
 D \frac{\partial B}{\partial t} &= (\text{grad})U_B(\text{grad})B - D(\text{vgrad})B + D\text{grad}G_D[+ W_D\text{grad}D_e(\text{grad})D]^*
 \end{aligned}
 \tag{4.2.15}$$

[]* — the included element only for the complete presentation of the electromagnetic field energy continuity from the diffusion phenomenon.

REMARK 2. The energy continuity partial differential equation for the case of the variable coefficients contains its permanent existing elements — without brackets and included only element — inside brackets []*.

The permanent existing elements of the energy continuity partial differential equation are related to the definition of the following derivative but the included only element is eliminated from this definition by the use of the charge partial differential equation element of the specific energy of the phenomenon.

— for the constant coefficients

$$\begin{aligned}
 U_B &\neq U_B(x, y, z, t), V_D \neq V_D(x, y, z, t), \gamma = \gamma(x, y, z, t), \varepsilon \neq \varepsilon(x, y, z, t), \\
 G_D &\neq G_D(x, y, z, t) \text{ and } D_e \neq D_e(x, y, z, t),
 \end{aligned}$$

$$\frac{\partial B}{\partial t} = -(\text{vgrad})B.
 \tag{4.2.20}$$

4.3. THE MOMENTUM BALANCE OF THE ELECTROMAGNETIC FIELD

The „momentum” elements of the Table 1 enable us on the formulation of the balance equation

$$\begin{aligned}
 \frac{\partial}{\partial t} \iiint_{\Omega_{(x,y,z)}} (\rho \mathbf{v}) d\Omega_{(x,y,z)} &= \iiint_{\Omega_{(x,y,z)}} [\hat{a} \nabla^2 \mathbf{v} + \hat{b} \text{grad div} \mathbf{v}] d\Omega_{(x,y,z)} - \\
 - f \iiint_{\Omega_{(x,y,z)}} [\mathbf{v} \text{div}(\rho \mathbf{v}) + (\rho \mathbf{v} \text{grad}) \mathbf{v}] d\Omega_{(x,y,z)} &\pm \iiint_{\Omega_{(x,y,z)}} T(t) d\Omega_{(x,y,z)} - \\
 - \iint_{F_{(x,y,z)}} (M_B \text{grad}) B dF_{(x,y,z)} &+ \iint_{F_{(x,y,z)}} (N_D \text{grad}) D dF_{(x,y,z)}.
 \end{aligned}
 \tag{4.3.1}$$

Making use of the Gauss—Ostrogradski law [14], after simple transformations; we have

$$\begin{aligned} \frac{\partial}{\partial t} (f\rho\mathbf{v}) = & \hat{a}\nabla^2\mathbf{v} + \hat{b}\text{grad div}\mathbf{v} - f\mathbf{v}\text{div}(\rho\mathbf{v}) - f(\rho\mathbf{v}\text{grad})\mathbf{v} \pm \\ & \pm T(t) - \text{grad}M_B(\text{grad})\mathbf{B} - M_B\nabla^2\mathbf{B} + \text{grad}N_D(\text{grad})\mathbf{D} + \\ & + N_D\nabla^2\mathbf{D}. \end{aligned} \quad (4.3.2)$$

Continuing, the eq. (4.3.2) can be presented in the form

$$\begin{aligned} f\mathbf{v}\frac{\partial\rho}{\partial t} + f\rho\frac{\partial\mathbf{v}}{\partial t} = & \hat{a}\nabla^2\mathbf{v} + \hat{b}\text{grad div}\mathbf{v} - f\mathbf{v}\rho\text{div}\mathbf{v} - f\mathbf{v}(\mathbf{v}\text{grad}\rho) - \\ & - f(\rho\mathbf{v}\text{grad})\mathbf{v} \pm T(t) - \text{grad}M_B(\text{grad})\mathbf{B} - M_B\nabla^2\mathbf{B} + \\ & + \text{grad}N_D(\text{grad})\mathbf{D} + N_D\nabla^2\mathbf{D}. \end{aligned} \quad (4.3.3)$$

Now, we introduce the partial differential equation of the continuity (4.1.5) multiplied both handsides by the element „ $f\mathbf{v}$ ” to the form

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

Substraction of the eq. (4.3.4) from the eq. (4.3.3) leads to the formula

$$\begin{aligned} f\rho\frac{\partial\mathbf{v}}{\partial t} = & \hat{a}\nabla^2\mathbf{v} + \hat{b}\text{grad div}\mathbf{v} - f\rho\mathbf{v}\text{div}\mathbf{v} - f\rho(\mathbf{v}\text{grad})\mathbf{v} \pm T(t) - \\ & - \text{grad}M_B(\text{grad})\mathbf{B} - M_B\nabla^2\mathbf{B} + \text{grad}N_D(\text{grad})\mathbf{D} + N_D\nabla^2\mathbf{D}. \end{aligned} \quad (4.3.5)$$

As a consequence of the utilization of the definition of the following derivative [12], to the eq. (4.3.5) its last form is:

$$f\rho\frac{D\mathbf{v}}{Dt} = \hat{a}\nabla^2\mathbf{v} + \hat{b}\text{grad div}\mathbf{v} - f\rho\mathbf{v}\text{div}\mathbf{v} \pm T(t) - M_B\nabla^2\mathbf{B} + N_D\nabla^2\mathbf{D}. \quad (4.3.6)$$

Applying the definition [2–11]

$$\text{div}\mathbf{D} = \rho \quad (4.3.7)$$

to the eq. (4.3.6) this equation is modified to

$$\begin{aligned} f\text{div}\mathbf{D}\frac{D\mathbf{v}}{Dt} = & \hat{a}\nabla^2\mathbf{v} - f\mathbf{v}\text{div}\mathbf{D}\text{div}\mathbf{v} + \hat{b}\text{grad div}\mathbf{v} \pm T(t) - M_B\nabla^2\mathbf{B} + \\ & + N_D\nabla^2\mathbf{D}. \end{aligned} \quad (4.3.8)$$

From the condition for the following derivative $f\text{div}\mathbf{D} = \frac{D\mathbf{v}}{Dt} = 0$ [12], [1], the partial differential equation of the field vector velocity fulfils two below cases:

– for the variable coefficients

$$M_B = M_B(x, y, z, t) \quad \text{and} \quad N_D = N_D(x, y, z, t)$$

$$f\text{div}\mathbf{D}\frac{\partial\mathbf{v}}{\partial t} = -f\text{div}\mathbf{D}(\mathbf{v}\text{grad})\mathbf{v} - \text{grad}M_B(\text{grad})\mathbf{B} + \text{grad}N_D(\text{grad})\mathbf{D}, \quad (4.3.9)$$

— for the constant coefficients

$$M_B = M_B(x, y, z, t) \quad \text{and} \quad N_D = N_D(x, y, z, t)$$

$$\frac{\partial \mathbf{v}}{\partial t} = -(\mathbf{v} \text{grad}) \mathbf{v}. \quad (4.3.10)$$

5. THE COMPLETE CONSTITUTIVE DESCRIPTION OF THE ELECTROMAGNETIC FIELD

Making use of the content of the section 4 the complete constitutive description of the electromagnetic field can be written in the system form:

$$\frac{D\mathbf{D}}{Dt} = D_e \nabla^2 \mathbf{D} - \mathbf{D} \text{div} \mathbf{v} + \frac{\gamma}{\varepsilon_0 \varepsilon} \mathbf{D} \pm \mathbf{J}(t) \quad (1)$$

$$(I) \quad \mathbf{D} \frac{D\mathbf{B}}{Dt} = U_B \nabla^2 \mathbf{B} - \mathbf{D} \mathbf{B} \text{div} \mathbf{v} \pm \mathbf{S}(t) + G_D \text{div} \mathbf{D} + W_D D_e \nabla^2 \mathbf{D} \quad (2)$$

$$f \text{div} \mathbf{D} \frac{D\mathbf{v}}{Dt} = \hat{a} \nabla^2 \mathbf{v} - f \mathbf{v} \text{div} \mathbf{D} \text{div} \mathbf{v} + \hat{b} \text{grad} \text{div} \mathbf{v} \pm \mathbf{T}(t) -$$

$$- M_B \nabla^2 \mathbf{B} + N_D \nabla^2 \mathbf{D}, \quad (3)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} + \mathbf{v} \text{grad} - (\text{sum of adequate gradient operations}) [12], [1].$$

From the condition for the following derivative $\frac{D}{Dt} = 0$; [12], [1], for every state variable of the electromagnetic field two continuity system cases can be analyzed:

— for the variable coefficients

$$D_e = D_e(x, y, z, t), U_B = U_B(x, y, z, t), V_D = V_D(x, y, z, t), \gamma = \gamma(x, y, z, t),$$

$$\varepsilon = \varepsilon(x, y, z, t), G_D = G_D(x, y, z, t), M_B = M_B(x, y, z, t) \text{ and } N_D = N_D(x, y, z, t)$$

$$\frac{\partial \mathbf{D}}{\partial t} = \text{grad} D_e (\text{grad}) \mathbf{D} - (\mathbf{v} \text{grad}) \mathbf{D} + \mathbf{D} \text{grad} \frac{\gamma}{\varepsilon_0 \varepsilon}, \quad (1')$$

$$(II) \quad \mathbf{D} \frac{\partial \mathbf{B}}{\partial t} = (\text{grad}) U_B (\text{grad}) \mathbf{B} - \mathbf{D} (\mathbf{v} \text{grad}) \mathbf{B} + \mathbf{D} \text{grad} G_D +$$

$$[+ W_D \text{grad} D_e (\text{grad}) \mathbf{D}]^*, \quad (2')$$

$$f \text{div} \mathbf{D} \frac{\partial \mathbf{v}}{\partial t} = -f \mathbf{v} \text{div} \mathbf{D} (\mathbf{v} \text{grad}) \mathbf{v} - \text{grad} M_B (\text{grad}) \mathbf{B} + \text{grad} N_D (\text{grad}) \mathbf{D} \quad (3')$$

[]* — the included element only for the complete presentation of the electromagnetic field energy continuity from the diffusion phenomenon.

— for the constant coefficients

$$D_e \neq D_e(x, y, z, t), U_B \neq U_B(x, y, z, t), V_D \neq V_D(x, y, z, t), \gamma \neq \gamma(x, y, z, t),$$

$$\varepsilon \neq \varepsilon(x, y, z, t), G_D \neq G_D(x, y, z, t), M_B \neq M_B(x, y, z, t) \text{ and } N_D \neq N_D(x, y, z, t)$$

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

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

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

The system of the partial differential constitutive state equations of the electromagnetic field (I) is invariant with respect to its physical coefficients according to its continuity systems of the partial differential equations (II) — variable physical coefficients or (III) — constant physical coefficients. The systems (II) or (III) represent the continuous groups of the transformations of the charge, energy and momentum coordinates of the electromagnetic field [13], [12], [1].

While the system of the partial differential constitutive state equations (I) possesses all its properties independently to its variable or constant coefficients, these physical coefficients are treated as being constant for the analytical solution of the system (I) in the Part II of this article.

6. CONCLUSIONS

Deduced in the article constitutive distributed parameter mathematical model of the electromagnetic field is based on all known physical phenomena of this field and the general properties of the environment. This approach is the first in the present world literature of the electromagnetic field. The application of the moving coordinates system and the problems of the following derivative makes possible to conceive the locally distributed parameter problems of both electric and magnetic fields. The chosen state variables „ $(\rho, D), B, \mathbf{v}$ ” characterize in the complete and synonymous form the electromagnetic field. Suggested in the article the constitutive invariance of the system (I) assures constancy of the physical coefficients of this system, independent to its physical phenomena taking part in the continuity systems of the partial differential equations (II) or (III).

APPENDIX 1

Pertinent to the properties of the electric conductor, in my work [1], the coefficients „ $\hat{a}, \hat{b}$ ” are deduced as follows:

— for the conductor in fluid phase

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

and η_w — coefficient of dynamic viscosity $\frac{\text{Ns}}{\text{m}^2}$,

— for the conductor in crystal phase

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

with h — crystal press module $\frac{N}{m^2}$, $\bar{k}$ — Poisson constant of crystal, p — unitary time transformations $\frac{l}{s}$, for the crystal having cuboidal shape.

NOTATION

ρ — density of the electric charge	$\frac{C}{m^3}$
D — vector of the induction of the electric field	$\frac{C}{m^2}$
B — vector of the induction of the magnetic field	T
v — field vector velocity of the electromagnetic medium	$\frac{m}{s}$
D_e — coefficient of the electric charge diffusion	$\frac{m^2}{s}$
γ — coefficient of the electric conductivity	$\frac{A}{Vm}$
ε_0 — coefficient of the absolute electric permeability	$\frac{C}{Vm}$
ε — coefficient of the relative electric permability	A
U_B — vector coefficient of the energy distribution of the magnetic induction	Tm
V_D — coefficient of the energy of the conductivity and permability for the electric induction	T
W_D — vector coefficient of the energy distribution of the charge diffusion	$\frac{kg}{C}$
f — coefficient of the mass the electric charge	$\frac{Ns}{m^2}$
$\hat{a}, \hat{b}$ — coefficient of the material properties of the electromagnetic medium	$\frac{Nm^2}{T}$
M_B — coefficient of the force of the magnetic induction	$\frac{Nm^4}{C}$
N_D — coefficient of the force of the electric induction	$\frac{Ns}{m^2}$
η_w — coefficient of the dynamic viscosity	$\frac{N}{m^2}$
h — the crystal press module	$\frac{N}{m^2}$
$\bar{k}$ — the Poisson constant of the crystal	$\frac{N}{m^2}$

p — coefficient of the time transformation of the electric field	$\frac{1}{s}$
$J(t)$ — the outside source of the electric charge on the chosen direction	$\frac{A}{m^2}$
$J(t)$ — the outside source of the electric charge	$\frac{A}{m^3}$
$S(t)$ — the outside source of the electromagnetic energy	$\frac{J}{m^4}$
$T(t)$ — the outside source of the momentum of the electromagnetic medium	$\frac{N}{m^3}$
$\bar{\Omega}_{(x,y,z)t}$ — the locally selected volume-time element phenomenally closed (every phenomenon can be treated in the separate way) around point $Z(x, y, z, t)$ of the electromagnetic field	$m^3 - s$
Ω_{Rt} — the so called working space volume-time of the apparatuses of the electromagnetic field	$m^3 - s$
$(a,b,c)t$ — the dimensions of the locally selected volume-time element $\bar{\Omega}_{(x,y,z)t}$ and: $x=a, y=b, z=c$	m, m, m, s
(x,y,z,t) — the constant space-time coordinates of Ω_{Rt}	m, m, m, s
$F_{(a,b,c)t}$ — the external oriented surface the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$	$m^2 - s$
$n^{(z)}$ — the normal external surface orientation vector of the surface $F_{(a,b,c)t}$	,
F_{Rt} — the external oriented surface of Ω_{Rt}	$m^2 - s$
$n_R^{(z)}$ — the normal external surface orientation vector of the surface F_{Rt}	,
t — time	s

REFERENCES

1. W. Niemiec: *On the constitutive distributed parameter modelling of the electric field. Part I. The construction of partial differential constitutive state equations of the electric field.* *Rozprawy Elektroniczne*, 1986, 32, z. 4. ss. 1095—1108
2. R. Sikora: *Teoria pola elektromagnetycznego.* WNT, Warszawa, 1985
3. R. Matysia k: *Teoria pola elektromagnetycznego.* WNT, Warszawa, 1976
4. J. D. Jackson: *Elektrodynamika klasyczna.* PWN, Warszawa, 1982
5. E. M. Purcell: *Elektryczność i magnetyzm.* PWN, Warszawa, 1974
6. R. P. Feynman, R. B. Leighton, M. Sands: *Feynmana wykłady z fizyki. Tom II — Część I.* PWN, Warszawa, 1974
7. A. N. Tichonow, A. A. Samarski: *Równania fizyki matematycznej.* PWN, Warszawa, 1963
8. L. D. Landau, E. M. Lifszyc: *Krótki kurs fizyki teoretycznej. Tom I. Mechanika. Elektrodynamika.* PWN, Warszawa, 1986
9. N. A. Krall, A. W. Trivelpiece: *Fizyka plazmy.* PWN, Warszawa, 1979
10. A. W. Czerniński: *Wstęp do fizyki plazmy.* PWN, Warszawa, 1971
11. C. A. Coulson, A. Jeffrey: *Modele matematyczne.* WNT, 1982
12. B. Średniawa: *Hydrodynamika i teoria sprężystości.* PWN, Warszawa, 1977

13. A. J. A. M o r g a n: *On the construction of constitutive equations for continuous media*. Archiwum Mechaniki Stosowanej, vol. 17, 1965, No. 1, pp. 145–174
14. T. T r a j d o s: *Matematyka dla inżynierów*. Tom 1 i 2. WNT, Warszawa, 1987
15. W. N o w a c k i: *Dynamiczne zagadnienia termosprężystości*. PWN, Warszawa, 1986
16. S. W ę g r z y n: *Podstawy automatyki*. PWN, Warszawa 1974

W. NIEMIEC

O KONSTYTUTYWNYM ROZŁOŻONYM PARAMETRYCZNIE MODELOWANIU
POLA ELEKTROMAGNETYCZNEGO
CZĘŚĆ I. KONSTYTUTYWNY ROZŁOŻONY PARAMETRYCZNIE MATEMATYCZNY
MODEL POLA ELEKTROMAGNETYCZNEGO

Streszczenie

W artykule zaproponowano oryginalny, konstruktywny, rozłożony parametrycznie matematyczny model pola elektromagnetycznego. Wyprowadzono układ równań różniczkowych cząstkowych konstytutywnych stanu w postaciach pochodnych śledczych dla: — ładunku i indukcji elektrycznej, — energii elektromagnetycznej, — pędu ośrodka elektromagnetycznego.

Ten układ zawiera wszystkie literaturowo umotywowane zjawiska fizyczne i parametry modelujące rzeczywiste przestrzenno-czasowe zagadnienia pola elektromagnetycznego dla fizycznych wpływów sterujących tym polem.

On the constitutive distributed parameter modelling of the electromagnetic field

Part II. The solution of the model partial differential constitutive state equations of the electromagnetic field

WACŁAW NIEMIEC

Politechnika Śląska, Gliwice

Received 1990.06.20

Authorized 1991.02.05

The constitutive distributed parameter mathematical model of the electromagnetic field deduced in Part I of the article has been solved in an original way, analytically. Two cases of this analytical solution have been considered:

A — existence of the initial conditions,

B — existence of the initial boundary conditions,

for the physical phenomena of the electromagnetic field.

Some control influences by the application of the physical phenomena on the space-time distributed important parameters of the electromagnetic field have been presented.

INTRODUCTION

We consider the constitutive distributed parameter mathematical model of the electromagnetic field, deduced in Part I of this article [1] from the point of view of its analytical solution. In this part of the article our focus is concentrated on the solution of the system of partial differential constitutive state equations having from Part I [1], the below form:

$$\frac{DD}{Dt} = D_e \nabla^2 D \operatorname{divv} + \frac{\gamma}{\varepsilon_0 \varepsilon} D + J(t). \quad (1)$$

$$(I) \quad D \frac{DB}{Dt} = U_B \nabla^2 B - DB \operatorname{divv} \pm S(t) + G_D \operatorname{div} D + W_D D_2 \nabla^2 D. \quad (2)$$

$$f \operatorname{div} \mathbf{D} \frac{D\mathbf{v}}{Dt} = \hat{a} \nabla^2 \mathbf{v} - f \mathbf{v} \operatorname{div} \mathbf{D} \operatorname{div} \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} \pm T(t) - M_B \nabla^2 \mathbf{B} + \\ + N_D \nabla^2 \mathbf{D}. \quad (3)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} = + \mathbf{v} \operatorname{grad} - (\text{sum of adequate gradient operations}) [3], [2], [1].$$

For the system (I) the constitutive invariance [4] is given by the system of partial differential equations of the continuity for all state variables of the electromagnetic field for two different cases:

– for the variable coefficients

$$D_e = D_e(x, y, z, t), U_B = U_B(x, y, z, t), \gamma = \gamma(x, y, z, t), \\ \varepsilon = \varepsilon(x, y, z, t), G_D = G_D(x, y, z, t), M_B = M_B(x, y, z, t) \text{ and } N_D = N_D(x, y, z, t)$$

$$\frac{\partial \mathbf{D}}{\partial t} = \operatorname{grad} D_e (\operatorname{grad}) \mathbf{D} - (\mathbf{v} \operatorname{grad}) \mathbf{D} + \mathbf{D} \operatorname{grad} \frac{\gamma}{\varepsilon_0 \varepsilon}, \quad (1')$$

$$(II) \quad \mathbf{D} \frac{\partial \mathbf{B}}{\partial t} = (\operatorname{grad}) U_B - (\operatorname{grad}) \mathbf{B} - \mathbf{D} (\mathbf{v} \operatorname{grad}) \mathbf{B} + \mathbf{D} \operatorname{grad} G_D + \\ [+ W_D \operatorname{grad} D_e (\operatorname{grad}) \mathbf{D}]^*, \quad (2')$$

$$f \operatorname{div} \mathbf{D} \frac{\partial \mathbf{v}}{\partial t} = - f \mathbf{v} \operatorname{div} \mathbf{D} (\mathbf{v} \operatorname{grad}) \mathbf{v} - \operatorname{grad} M_B (\operatorname{grad}) \mathbf{B} + \operatorname{grad} N_D (\operatorname{grad}) \mathbf{D}, \quad (3)$$

[*] – the included element only for the complete presentation of the electromagnetic field energy continuity from the diffusion phenomenon.

– for the constant coefficients

$$D_e \neq D_e(x, y, z, t), U_B \neq U_B(x, y, z, t), \gamma \neq \gamma(x, y, z, t), \\ \varepsilon \neq \varepsilon(x, y, z, t), G_D \neq G_D(x, y, z, t), M_B \neq M_B(x, y, z, t) \text{ and } N_D \neq N_D(x, y, z, t)$$

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

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

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

In accordance with the fact that the system (I) is invariant to its systems of continuity partial differential equations (II) or (III) the physical coefficients of the system (I) are assumed as being constant for its analytical solution.

THE DISCUSSION OF THE SYSTEM OF PARTIAL DIFFERENTIAL CONSTITUTIVE STATE EQUATIONS OF THE ELECTROMAGNETIC FIELD (I)

Let us recall the complete form of the system of partial differential constitutive state equations describing the electromagnetic field:

$$\frac{DD}{Dt} = D_e \nabla^2 D \operatorname{div} + \frac{\gamma}{\varepsilon_0 \varepsilon} D + J(t). \quad (1)$$

$$(I) \quad D \frac{DB}{Dt} = U_B \nabla^2 B - DB \operatorname{div} \pm S(t) + G_D \operatorname{div} D + W_D D_2 \nabla^2 D. \quad (2)$$

$$f \operatorname{div} D \frac{Dv}{Dt} = \hat{a} \nabla^2 v - f v \operatorname{div} D \operatorname{div} + \hat{b} \operatorname{grad} \operatorname{div} \pm T(t) - M_B \nabla^2 B + N_D \nabla^2 D. \quad (3)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} = v + \operatorname{grad} - (\text{sum of adequate gradient operations}) [3], [2], [1].$$

1. THE SOURCE FUNCTIONS FOR THE ELECTROMAGNETIC FIELD

We need to consider the source functions acting inside the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ and transformed to the point $Z(x, y, z, t)$ and $Z(x, y, z, t) \in \bar{\Omega}_{(a,b,c)t}$. It is evident that for the constant coordinates values $x=a, y=b, z=c$, the below state vector of electromagnetic field in constant

$$\begin{bmatrix} \rho(a, b, c) \\ D(a, b, c) \\ B(a, b, c) \\ v(a, b, c) \end{bmatrix} = \text{constant}$$

and the space partial differential operations are equal to zero

$$\nabla \begin{bmatrix} \rho(a, b, c) \\ D(a, b, c) \\ B(a, b, c) \\ v(a, b, c) \end{bmatrix} = 0$$

and with respect to this approach the system (I) is modified to:

(I^s)

The source functions:

$$\frac{dD}{dt} = \frac{\gamma}{\varepsilon_0 \varepsilon} D \pm J(t) \int_0^t (1) \quad D^z(t), B^z(t), v^z(t), \rho^z(t)$$

$$D \frac{dB}{dt} = \pm S(t) \quad (2)$$

$$f \rho \frac{dv}{dt} = \pm T(t) \quad (3)$$

$$\frac{d\rho}{dt} = \pm J(t) \quad (4)$$

The initial conditions: D_0, B_0, v_0, ρ_0 .

REMARK. For the integration of the eq. (I^s.2) sign modulus of the eq. (I^s.1) is used, and consequently for the eq. (I^s.3) is applied the modulus of the eq. (I^s.4) — both cases with respect to the possibility of inverse transformation of the electromagnetic field.

Making use of the above approach, the solution of the system (I^s) leads to the below source functions:

— for the electric induction vector (see Appendix 1)

$$D^z(t) = e^{-\frac{\gamma}{\epsilon_0 \epsilon} t} \left[\pm \int J(t) e^{\frac{\gamma}{\epsilon_0 \epsilon} t} dt + D_C \right] \quad (S.1)$$

D_C — the integration constant vector,

— for the magnetic induction vector

$$B^z(t) = B_0 \pm \int_0^t \frac{S(t)}{|D^z(t)|_s} dt \quad (S.2)$$

where: $|_s$ — sign modulus only,

— for the field vector velocity

$$v^z(t) = v_0 \pm \int_0^t \frac{T(t)}{f|\rho^z(t)|} dt \quad (S.3)$$

and

$$\rho^z(t) = \rho_0 + \int_0^t J(t) dt. \quad (S.4)$$

2. THE ANALYSIS OF THE HOMOGENEOUS PART OF THE SYSTEM (I)

We present the homogeneous part of the system (I) as having the below form:

$$\frac{DD}{Dt} = D_e \nabla^2 D - D \operatorname{div} v \quad (1)$$

$$(I^h) \quad D \frac{DB}{Dt} = U_B \nabla^2 B - DB \operatorname{div} v \pm S(t) + G_D \operatorname{div} D + W_D D_2 \nabla^2 D. \quad (2)$$

$$f \operatorname{div} D \frac{Dv}{Dt} = \hat{a} \nabla^2 v - f v \operatorname{div} D \operatorname{div} v + \hat{b} \operatorname{grad} \operatorname{div} v - M_B \nabla^2 B + N_D \nabla^2 D. \quad (3)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} = + \mathbf{v} \text{grad} - (\text{sum of adequate gradient operations}) [3], [2], [1].$$

In concord with the mathematical field theory [5], [3], a necessary problem is to lead into our considerations: A – vector potential of the flow field and $\mathbf{v} = \text{rot } A$, with $\text{div rot } A \equiv 0$, so that the system (I^h) may be modified to:

– the partial differential equations of the potential fields (variable point $Q(\xi, \eta, \beta, \tau)$)

$$\frac{\partial D}{\partial t} = D_e \tau \nabla^2 D \quad (1)$$

$$(I^{hp1}) \quad D \frac{\partial B}{\partial t} = U_B \nabla^2 B + G_D \text{div} D + W_D D_2 \nabla^2 D. \quad (2)$$

$$f \text{div} D \frac{\partial \mathbf{v}}{\partial t} = - M_B \nabla^2 B + N_D \nabla^2 D. \quad (3)$$

For the analysis of the system (I^{hp1}) the following approaches remain possible [2], [5]:

1. idealization of the medium – only one physical coefficient or one state variable is different than zero but other are assumed as being zero,

2. separation of the single physical phenomena – the identical space operations on the same state variables are valid.

Making use of the method 1 or 2 to the system (I^{hp1}) – we have:

– the partial differential equations type diffusion of the electric induction vector

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

$$(I^{hp2}) \quad D \frac{\partial B}{\partial t} = W_D D_e \nabla^2 D \quad (2)$$

$$f \text{div} D \frac{\partial \mathbf{v}}{\partial t} = N_D \nabla^2 D, \quad (3)$$

– the partial differential equation of the electric induction vector of the electromagnetic energy continuity

$$(I^{hp3}) \quad \dots\dots\dots (1)$$

$$D \frac{\partial B}{\partial t} = G_D \text{div} D, \quad (2)$$

$$\dots\dots\dots (3)$$

– the partial differential equations type diffusion of the magnetic induction vector

$$(I^{hp4}) \quad \dots\dots\dots (1)$$

$$D \frac{\partial B}{\partial t} = U_B \nabla^2 B \quad (2)$$

$$f \text{div} D \frac{\partial \mathbf{v}}{\partial t} = - M_B \nabla^2 B, \quad (3)$$

— the partial differential equations in the forms of the following derivatives of the rotation field (variable point $Q'(\xi', \eta', \beta', \tau)$).

From the system (I^h) it is possible to obtain:

$$(I^{hr}) \quad \frac{DD}{Dt} = -D \operatorname{div} v \quad (1)$$

$$D \frac{DB}{Dt} = -DB \operatorname{div} v \quad (2)$$

$$f \operatorname{div} D \frac{Dv}{Dt} = \hat{a} \nabla^2 v - f v \operatorname{div} D \operatorname{div} v + \hat{b} \operatorname{grad} \operatorname{div} v \quad (3)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} + v \operatorname{grad} \cdot [3], [2], [1].$$

THE ANALYTICAL SOLUTION OF THE CONSTITUTIVE DISTRIBUTED PARAMETER MODEL OF ELECTROMAGNETIC FIELD

A. EXISTENCE OF THE INITIAL CONDITIONS

A. 1. THE PARTIAL DIFFERENTIAL EQUATIONS OF THE POTENTIAL FIELDS (VARIABLE POINT $Q(\xi, \eta, \beta, \tau)$)

A. 1. 1. The partial differential equations type diffusion of the electric induction vector — system (I^{hp2}) and its sources

A. 1. 1. 1. The diffusion partial differential equation of the electric induction vector

Let us recall the partial differential equation ($I^{hp2.1}$) rewritten in the form:

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

The electric induction vector „ D ” can be presented as

$$D(x, y, z, t) = C(x, y, z) g(t) \quad (2)$$

and from the eq. (1) one obtains

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

and

$$\frac{dg}{dt} = -\lambda D_e g. \quad (4)$$

The application of theory of the superposition of potential fields to rotational fields leads to the zero conditions on the brim of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ [6], [5], [3], [2]

$$\begin{aligned}
 C(0, y, 0) &= 0 & C(a, y, 0) &= 0 \\
 C(x, 0, 0) &= 0 & \text{and} & C(x, b, 0) = 0 \\
 C(0, 0, z) &= 0 & C(0, y, c) &= 0.
 \end{aligned} \tag{5}$$

We observe that the introduction of function

$$C(x, y, z) = X(x) Y(y) Z(z) \tag{6}$$

into the eq. (3), makes possible to modified this equation to its system form:

$$\begin{aligned}
 X'' + \theta X &= 0 & X(0) &= 0 & X(a) &= 0 \\
 Y'' + \mu Y &= 0 & \text{with} & Y(0) &= 0 & \text{and} & Y(b) &= 0 \\
 Z'' + \kappa Z &= 0 & Z(0) &= 0 & Z(c) &= 0.
 \end{aligned} \tag{7}$$

The system (7) has the solution placed below, by x^0, y^0, z^0 — versors

$$\begin{aligned}
 X(x) &= x^0 \sin \frac{n\pi x}{a}; & \theta &= \left(\frac{n\pi}{a}\right)^2 \\
 Y(y) &= y^0 \sin \frac{m\pi y}{b}; & \mu &= \left(\frac{m\pi}{b}\right)^2 \\
 Z(z) &= z^0 \sin \frac{k\pi z}{c}; & \kappa &= \left(\frac{k\pi}{c}\right)^2.
 \end{aligned} \tag{8}$$

Hence, we have the eigenvalues written as

$$\lambda_{n,m,k} = \left(\frac{n\pi}{a}\right)^2 + \left(\frac{m\pi}{b}\right)^2 + \left(\frac{k\pi}{c}\right)^2 \tag{9}$$

for which the eigenfuctions become a suitable form

$$C_{n,m,k}(x, y, z) = A_{n,m,k} x^0 \sin \frac{n\pi x}{a} y^0 \sin \frac{m\pi y}{b} z^0 \sin \frac{k\pi z}{c}. \tag{10}$$

The amplitude of eigenfunction „ $A_{n,m,k}$ ” can be determined from the point of view of dimensions % a, b, c ” — of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$, thus the norm N_e of eigenfunctions $C_{n,m,k}(x, y, z)$ with the weight equal to 1 is equal to 1 [6], [3], [2],

$$\begin{aligned}
 N_e &= \iiint_{0 \leq x \leq a, 0 \leq y \leq b, 0 \leq z \leq c} C_{n,m,k}^2(x, y, z) dx dy dz = A_{n,m,k}^2 \int_0^a x^0 \sin^2 \frac{n\pi x}{a} dx \int_0^b y^0 \sin^2 \frac{m\pi y}{b} dy \cdot \\
 &\quad \cdot \int_0^c z^0 \sin^2 \frac{k\pi z}{c} dz = 1.
 \end{aligned} \tag{11}$$

It is evident that from the integration of the term (11) the amplitude „ $A_{n,m,k}$ ” is:

$$A_{n,m,k} = \sqrt{\frac{8}{abc}}. \quad (12)$$

We note that the eigenfunction for the constant coordinates point $Z(x, y, z, t)$ are:

$$C_{n,m,k}(x, y, z) = \sqrt{\frac{8}{abc}} x^0 \sin \frac{n\pi x}{a} y^0 \sin \frac{m\pi y}{b} z^0 \sin \frac{k\pi z}{c}. \quad (13)$$

Consequently for the variable coordinates point $Q(\xi, \eta, \beta, \tau)$ we have the eigenfunction

$$C_{n,m,k}(\xi, \eta, \beta) = \sqrt{\frac{8}{abc}} x^0 \sin \frac{n\pi \xi}{a} y^0 \sin \frac{m\pi \eta}{b} z^0 \sin \frac{k\pi \beta}{c}. \quad (14)$$

The integration of the eq. (4) allows to obtain

$$- \lambda_{n,m,k} D_e t \quad (15)$$

$g(t) = D_{n,m,k} e$.

The integration constant „ $D_{n,m,k}$ ” was been found as follows:

$$D_{n,m,k} = \iiint_{\Omega_R} D \sqrt{\frac{8}{abc}} x^0 \sin \frac{n\pi \xi}{a} y^0 \sin \frac{m\pi \eta}{b} z^0 \sin \frac{k\pi \beta}{c} d\xi d\eta d\beta. \quad (16)$$

We recall that from the eq. (S.1) the source function for this case is (Appendix 1):

$$D^z(t) = E(t). \quad (17)$$

In concord with the above considerations the resultant solution of the topic problem may be written as:

$$D(Z, t) = \int_0^t \iiint_{\Omega_R} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k} e^{-\lambda_{n,m,k} D_e (t-\tau)} \frac{8}{abc} \left(x^0 \sin \frac{n\pi x}{a} \cdot x^0 \sin \frac{n\pi \xi}{a} \right) \cdot \left(y^0 \sin \frac{m\pi y}{b} \cdot y^0 \sin \frac{m\pi \eta}{b} \right) \left(z^0 \sin \frac{k\pi z}{c} \cdot z^0 \sin \frac{k\pi \beta}{c} \right) \right\} \cdot D^z(\tau) d\xi d\eta d\beta d\tau. \quad (18)$$

A. 1. 1. 2. The diffusion partial differential equation of the electric induction vector for the electromagnetic field energy

The outset point of our considerations is the formula ($I^{hp2.2}$) having the form

$$D \frac{\partial B}{\partial t} = W_D D_e \nabla^2 D \quad (1)$$

and the eq. ($I^{hp2.1}$) rewritten as

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

Combining the eq. (1) and the eq. (2) one obtains

$$\mathbf{D} \frac{\partial \mathbf{B}}{\partial t} = \mathbf{W}_D \frac{\partial \mathbf{D}}{\partial t} \quad (3)$$

The solution of the magnetic diffusion topic problem presented by the formula (3) is:

$$\mathbf{B}(\mathbf{Z}, t) = \mathbf{B}_0(\mathbf{Z}, t) + \mathbf{W}_D \left[\ln \left| \int_0^t \int_{\Omega_n} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} D_{n,m,k} e^{-\lambda_{n,m,k} D_0(t-\tau)} \frac{8}{abc} \left(x^0 \sin \frac{n\pi x}{a} \cdot x^0 \sin \frac{n\pi \xi}{a} \right) \right. \right. \right. \\ \left. \left. \left(y^0 \sin \frac{m\pi y}{b} \cdot y^0 \sin \frac{m\pi \eta}{b} \right) \left(z^0 \sin \frac{k\pi z}{c} \cdot z^0 \sin \frac{k\pi \beta}{c} \right) \right\} \cdot D^z(\tau) d\xi d\eta d\beta d\tau \right| - \ln \left| D_0(\mathbf{Z}, t) \right| \right] \quad (4)$$

In the formula (4) the initial conditions problems are: $\mathbf{B}_0(\mathbf{Z}, t)$ – initial distributed magnetic induction vector, $\mathbf{D}_0(\mathbf{Z}, t)$ – initial distributed electric induction vector, $\mathbf{D}_0$ – initial source electric induction vector. For the integration constant $D_{n,m,k}$ and eigenvalues $\lambda_{n,m,k}$ see chapter A.1.1.1.

A. 1. 1. 3. The diffusions partial differential equation of the electric induction vector for the electromagnetic medium field vector velocity

We introduce the partial differential equation ($I^{hp2}.3$) possessing the form

$$f \operatorname{div} \mathbf{D} \frac{\partial \mathbf{v}}{\partial t} = N_D \nabla^2 \mathbf{D} \quad (1)$$

which subsequently can be modified to

$$f \frac{\partial \mathbf{v}}{\partial t} = N_D \operatorname{grad}(\mathbf{D}_n(\mathbf{z})) \quad (2)$$

A. 1. 2. The partial differential equation of the electric induction vector of the electromagnetic field energy continuity

The following term of the topic problem remains valid

$$\mathbf{D} \frac{\partial \mathbf{B}}{\partial t} = G_D \operatorname{div} \mathbf{D} \quad (1)$$

A. 1. 3. 1. The diffusion partial differential equation of the magnetic induction vector for the electromagnetic energy

The outset partial differential equation is:

$$\mathbf{D} \frac{\partial \mathbf{B}}{\partial t} = U_B \nabla^2 \mathbf{B} = 0 \quad (1)$$

The magnetic induction vector „ $\mathbf{B}$ ” can be presented as:

$$\mathbf{B}(\mathbf{x}, \mathbf{y}, \mathbf{z}, t) = \mathbf{L}(\mathbf{x}, \mathbf{y}, \mathbf{z}) \mathbf{h}(t) \quad (2)$$

and the eq. (1) gets its system form

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

and

$$\frac{dh}{dt} = -\lambda U_B h \frac{1}{D^z(t)} \quad (4)$$

The zero brim conditions on the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ from the superposition of potential fields to rotational fields are [6], [5], [3], [2]

$$\begin{aligned} L(0, y, 0) &= 0 & L(a, y, 0) &= 0 \\ L(x, 0, 0) &= 0 & L(x, b, 0) &= 0 \\ L(0, 0, z) &= 0 & L(0, y, c) &= 0 \end{aligned} \quad (5)$$

Consequently to the section A. 1. 1. 1. the eigenvalues are written in the form

$$\lambda_{n,m,k} = \left(\frac{n\pi}{a}\right)^2 + \left(\frac{m\pi}{b}\right)^2 + \left(\frac{k\pi}{c}\right)^2 \quad (6)$$

and proceeding in the above manner the eigenfunctions may be written in the below form:

– for the constant coordinates point $Z(x, y, z, t)$

$$L_{n,m,k}(x, y, z) = \sqrt{\frac{8}{abc}} x^0 \sin \frac{n\pi x}{a} y^0 \sin \frac{m\pi y}{b} z^0 \sin \frac{k\pi z}{c}, \quad (7)$$

– for the variable coordinates point $Q(\xi, \eta, \beta, \tau)$

$$L_{n,m,k}(\xi, \eta, \beta) = \sqrt{\frac{8}{abc}} x^0 \sin \frac{n\pi \xi}{a} y^0 \sin \frac{m\pi \eta}{b} z^0 \sin \frac{k\pi \beta}{c}. \quad (8)$$

From the integration of the eq. (4) one gets

$$h(t) = F_{n,m,k} e^{-\lambda_{n,m,k} U_B [K(t)]}, \quad (9)$$

where:

$$[K(t)] = \int \frac{1}{D^z(t)} dt$$

and „ $D^z(t)$ ” is the source function from the formula (S.1). The integration constant for the eq. (9) is:

$$F_{n,m,k} = \iiint_{\Omega_n} B_0 \sqrt{\frac{8}{abc}} x^0 \sin \frac{n\pi \xi}{a} y^0 \sin \frac{m\pi \eta}{b} z^0 \sin \frac{k\pi \rho}{c} d\xi d\eta d\beta. \quad (10)$$

Subsequently we recall the source function from the eq. (S.2) as follows

$$B^z(t) = B_0 \pm \int_0^t \frac{S(t)}{|D^z(t)|_s} dt, \quad (11)$$

where: $|_s$ — sign modulus only and the resultant solution of the topic problem with respect to above considerations may be written in the form:

$$B(Z, t) = \int_0^t \iiint_{\Omega_x} \left\{ \sum_{\substack{n=1 \\ m=1 \\ k=1}}^{\infty} F_{n,m,k} e^{-\lambda_{n,m,k} U_B / K(t-\tau)} \frac{8}{abc} \left(x^0 \sin \frac{n\pi x}{a} \cdot x^0 \sin \frac{n\pi \xi}{a} \right) \cdot \left(y^0 \sin \frac{m\pi y}{b} \cdot y^0 \sin \frac{m\pi \eta}{b} \right) \left(z^0 \sin \frac{k\pi z}{c} \cdot z^0 \sin \frac{k\pi \beta}{c} \right) \right\} \cdot \left[B_0 \pm \int_0^{t-\tau} \frac{S(\tau)}{|D^z(\tau)|_s} d\tau \right] d\xi d\eta d\beta d\tau. \quad (12)$$

where: $|_s$ — sign modulus only.

A. 1. 3. 2. The diffusion partial differential equation of the magnetic induction vector for the electromagnetic medium field vector velocity

The outset of the considerations of this section is the formula ($I^{hp4.3}$) having the form

$$f \operatorname{div} D \frac{\partial v}{\partial t} = -M_B \nabla^2 B \quad (1)$$

to which is related the formula ($I^{hp4.2}$)

$$D \frac{\partial B}{\partial t} = -U_B \nabla^2 B. \quad (2)$$

Combining the eq. (2) and the eq. (1) the result possesses the form

$$f \operatorname{div} D \frac{\partial v}{\partial t} = \frac{M_B}{U_B} D \frac{\partial B}{\partial t}. \quad (3)$$

The result given by the eq. (3) is not solvable analytically so we can use the general formula of the electromagnetic energy continuity ($I^{hp3.2}$) in the form:

$$D \frac{\partial B}{\partial t} = G_D \operatorname{div} D \quad (4)$$

tho modified the eq. (3) to

$$f \operatorname{div} D \frac{\partial v}{\partial t} = -\frac{M_B}{U_B} G_D \operatorname{div} D. \quad (5)$$

In accordance with the above approach the last form of the eq. (5) is:

$$f \frac{\partial \mathbf{v}}{\partial t} = -\frac{M_B}{U_B} G_D. \quad (6)$$

The solution of the eq. (6) may be written as:

$$\mathbf{v}(Z, t) = \mathbf{v}_0(Z, t) - \frac{M_B}{f U_B} G_D t \quad (7)$$

A. 2. THE PARTIAL DIFFERENTIAL EQUATIONS OF THW ROTATIONAL FIELD (VARIABLE POINT $Q(\xi', \eta', \beta', \tau)$)

A. 2. 1. The partial differential equations of the field vector velocity for the medium of the electromagnetic field

We need to recall to our considerations the system (I^h) in its complete form:

$$\frac{D\mathbf{D}}{Dt} = -D \operatorname{div} \mathbf{v} \quad (1)$$

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

$$f \operatorname{div} D \frac{D\mathbf{v}}{Dt} = \hat{a} \nabla^2 \mathbf{v} - f \mathbf{v} \operatorname{div} D \operatorname{div} \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v}, \quad (3)$$

where:

$$\frac{D}{Dt} = \frac{\partial}{\partial t} + \mathbf{v} \operatorname{grad} \quad [3], [1-2]$$

It is possible to apply the eq. (1) into the eq. (3) obtaining after modification the result:

$$f \operatorname{div} D \frac{D\mathbf{v}}{Dt} - f \mathbf{v} \operatorname{div} D \frac{D\mathbf{D}}{Dt} = \hat{a} \nabla^2 \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v}. \quad (4)$$

For the statical reciprocal constant image of the element $\bar{\Omega}_{(a,b,c)t}$ the eq. (4) has its system form [7], [3], [6], [2]

$$\hat{a} \nabla^2 \mathbf{v} + \hat{b} \operatorname{grad} \operatorname{div} \mathbf{v} = 0 \quad (5)$$

and

$$f \operatorname{div} D \frac{D\mathbf{v}}{Dt} - f \mathbf{v} \operatorname{div} D \frac{D\mathbf{D}}{Dt} = 0. \quad (6)$$

Applying for the eq. (5) the mathematical rule [5-6], [3], [2]

$$\operatorname{rot} \operatorname{rot} \mathbf{v} = \operatorname{grad} \operatorname{div} \mathbf{v} - \nabla^2 \mathbf{v} \quad (7)$$

and the eq. (5) may be rewritten as

$$(\hat{a} + \hat{b}) \operatorname{grad} \operatorname{div} \mathbf{v} - \hat{a} \operatorname{rot} \operatorname{rot} \mathbf{v} = 0. \quad (8)$$

The eq. (8) has its another form

$$\nabla^2 \mathbf{v} = 0 = \hat{\alpha} \operatorname{grad} \operatorname{div} \mathbf{v} - \hat{\beta} \operatorname{rot} \operatorname{rot} \mathbf{v}. \quad (9)$$

From the reciprocal constant image and assuming that the statical state forces are transferred as in the elastic media, the statical isotropic Green tensor for $\bar{\Omega}_{(a,b,c)t}$ is written in the below form [3], [2], [9]

$$\Gamma_G(Z, Q') = \frac{1}{8\pi} \left[\frac{1}{\hat{\alpha}} \Gamma_{\hat{\alpha}}(Z, Q') + \frac{1}{\hat{\beta}} \Gamma_{\hat{\beta}}(Z, Q') \right] \quad (10)$$

for the points $Z(x, y, z)$ and $Q'(\xi', \eta', \beta')$ (Appendix 2).

In my previous work [2] was deduced an anisotropic Green tensor for the different deformations of the cuboidal locally selected volume element $\bar{\Omega}_{(a,b,c)t}$ in the series form:

$$\Gamma_{GA}(Z, Q') = \Gamma_{ln}(Z, Q') = \Gamma_{lp}(Z, Q') * \hat{e}_{pq}(Z, Q') \Gamma_{qn}(Z, Q') \quad (11)$$

$k \qquad \qquad \qquad 0 \qquad \qquad \qquad l \qquad \qquad \qquad k-l$

for $k=l$, and

$\Gamma_{lp}(Z, Q')$ – the statical Green tensor for the cuboidal locally selected volume element $\bar{\Omega}_{(a,b,c)t}$

$\hat{e}_{pq}(Z, Q')$ – the modified tensor of the statical deformations of the cuboidal locally selected volume element $\bar{\Omega}_{(a,b,c)t}$ according to the first order small parameter,

$\Gamma_{qn}(Z, Q')$ – the statical Green tensor according to the $(k-1)$ deformations pertinent to the small parameter $(k-1)$ order,

w – number of power of the small parameter.

For the detailed forms of the elements of the statical Green tensor (11) see Appendix 3.

Assuming that

$$\Gamma_G(Z, Q') \equiv \Gamma_{GA}(Z, Q') \equiv \Gamma_G^w(Z, Q') \quad (12)$$

the considerations of this article have the universal form.

For the statical state in point $Z(x, y, z, 0)$ we have the velocity vector

$$\mathbf{v}_0(Z, 0) = \mathbf{v}_0 \quad (13)$$

and consequently for the statical distribution of the field vector velocity „ $\mathbf{v}$ ” one gets

$$\mathbf{v}(Z, 0) = \iiint_{\Omega_x} \Gamma_G^w(Z, Q') \mathbf{v}(Q') d\xi' d\eta' d\beta'. \quad (14)$$

In order to continue our considerations let us take to the analysis the dynamical aspects of the topic problem, described by the formula (6) and rewritten as:

$$\frac{D\mathbf{v}}{Dt} - \frac{\mathbf{v}}{D} \frac{DD}{Dt} = 0 \quad (15)$$

and subsequently modified to:

$$\frac{D\mathbf{v}}{\mathbf{v}} = \frac{DD}{D} \quad (16)$$

The integration of the eq. (16) leads to the formula

$$\ln |N| + \ln |\mathbf{v}| = \ln |D| \quad (17)$$

and after simple transformations

$$\mathbf{v}(t) = \frac{D^z(t)}{N_t}, \quad (18)$$

where:

N_t — constant, taken from eq. (14), as $N_t = |\mathbf{v}(Z, 0)|$.

$D^z(t)$ — the source function of the electric induction vector given by the eq. (S.1)

$\mathbf{v}(t)$ — the dynamic part of the field vector velocity for the medium of the electromagnetic field.

In accordance with the above considerations the resultant solution of the field vector velocity of the medium of the electromagnetic field is:

$$\mathbf{v}(Z, t) = \int_0^t \iiint_{\Omega_R} \left\{ \frac{1}{N_t} \left[D^z(t - \tau) \right] \Gamma_G^W(Z, Q) \right\} \left[\mathbf{v}_0 \pm \int_0^{t-\tau} \frac{T(\tau)}{f|\rho^z(\tau)|} d\xi' d\eta' d\beta' d\tau \right] \quad (19)$$

A. 3. THE COMPLETE INITIAL CONDITIONS SOLUTION OF THE CONSTITUTIVE DISTRIBUTED PARAMETER MATHEMATICAL MODEL OF THE ELECTROMAGNETIC FIELD

We concentrate our focus on the analytical solutions of the source system of the time ordinary differential equations (I^p) and their responses by the partial differential equations of the physical phenomena of the potential fields; systems (I^{hp2}), (I^{hp3}) and (I^{hp4}) and the rotational field; system (I^{hr}). The complete analytical solution is based on summation of the effects of the influences of the physical phenomena of potential fields and rotational field pertinent to the coordinates of the state vector of the electromagnetic field. In accordance with above, we have:

The complete form
of the solution

The transformation of the electromagnetic
field in point $Z(x, y, z, t)$

$$\begin{array}{c}
 \left[\begin{array}{l}
 \text{Density of electric charge} \\
 \text{Induction vector of electric field} \\
 \text{Induction vector of magnetic field} \\
 \text{Field vector velocity of medium of electro-magnetic field}
 \end{array} \right] = \left[\begin{array}{l}
 \text{Density of electric charge changes (formula S.4)} \\
 \text{Induction vector of electric field changes (formula S.1)} \\
 \text{Induction vector of magnetic field changes (formula S.2)} \\
 \text{Field vector velocity of medium of electromagnetic field changes (formula S.3)}
 \end{array} \right] + \\
 \text{Electric diffusion} \quad \text{Electromagnetic field energy} \\
 Q \rightarrow Z \quad Q \rightarrow Z \\
 + \left[\begin{array}{l}
 \Delta \\
 \text{Induction vector of electric field increase (formula A. 1. 1. 18)} \\
 X \\
 \text{Field vector velocity of medium of electromagnetic field by electric induction continuity (formula A. 1. 1. 3. 2)}
 \end{array} \right] + \left[\begin{array}{l}
 \Delta \\
 X \\
 \text{Induction vectors of electric and magnetic continuity + increase (formula A}^0\text{.1.2.1)} \\
 X \text{ (formula A}^0\text{.1.1.2.4)}
 \end{array} \right] + \\
 \text{Magnetic diffusion} \quad \text{Field vector velocity of medium of electromagnetic field} \\
 Q \rightarrow Z \quad Q' \rightarrow Z \\
 + \left[\begin{array}{l}
 \Delta \\
 X \\
 \text{Induction vector of magnetic field increase (formula A.1.3.1.12)} \\
 \text{Field vector velocity of medium of electromagnetic field by magnetic induction increase (formula A.1.3.2.7)}
 \end{array} \right] + \left[\begin{array}{l}
 \Delta \\
 \text{Induction vector electric field continuity (formula A.2.1.1)} \\
 \text{Induction vector of magnetic field continuity (formula A.2.1.2)} \\
 \text{Field vector velocity of medium of electromagnetic field increase (formula A.2.1.19)}
 \end{array} \right]
 \end{array}$$

Δ — The physical phenomena not taking parts into considerations,

X — the places without relations between physical phenomena.

B. EXISTENCE OF THE INITIAL BOUNDARY CONDITIONS

B.1. THE PHENOMENAL SOLUTIONS

For the solutions of the phenomenal partial differential equations of the potential fields and rotation field one general formula governing the solution problems inside so called „working” space volume-time Ω_{Rt} and the on its outside oriented brim F_{Rt} is valid. This formula may be written as follows:

$$\begin{aligned}
S_{Ph}(Z, t) = & \int_0^t \iiint_{\Omega_r} G_{Ph}(Z, t; Q^R, \tau) m^z(Q^R, \tau) d\Omega_R(\xi^R, \eta^R, \beta^R) d\tau + \\
& + \int_0^t \iint_{F_R} G_{Ph}(Z, t; Q^R, \tau) \big|_{F_{Rr}} m_{F_{Rr}}^z(Q^R, \tau) dF_R(\xi^R, \eta^R, \beta^R) d\tau, \quad (1)
\end{aligned}$$

where:

$Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q(\xi, \eta, \beta, \tau)$ – the variable coordinates point for the potential fields,

$Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q'(\xi, \eta, \beta, \tau)$ – the variable coordinates point for the rotational field,

$Z(x, y, z, t)$ – the constant coordinates point,

$G_{Ph}(Z, t; Q^R, \tau)$ – the phenomenal Green function existing inside Ω_{Rt} ,

$m^z(Q^R, \tau)$ – the source function pertinent to the physical phenomenon acting inside Ω_{Rt} ,

$G_{Ph}(Z, t; Q^R, \tau) \big|_{F_{Rr}}$ – the phenomenal Green function existing on the brim F_{Rt} ,

$m_{F_{Rr}}^z(Q^R, \tau)$ – the source function pertinent to the physical phenomenon acting on the brim F_{Rt} ,

$S_{Ph}(Z, t)$ – resultant state variable pertinent to the physical phenomenon.

B.2. THE COMPLETE INITIAL AND BOUNDARY CONDITIONS SOLUTION OF CONSTITUTIVE DISTRIBUTED PARAMETER MATHEMATICAL MODEL OF THE ELECTROMAGNETIC FIELD

Making use of the considerations of the chapter A.3 the procedure of construction of complete solution of the topic problem (A.3.1) is valid for case of existence of initial and boundary conditions, too. By the general phenomenal solutions formula (B.1.1) we can state the complete solution as follows:

The complete form
of the solution

$$\begin{bmatrix} \text{Denisity of electric charge} \\ \text{Induction vector of electric field} \\ \text{Induction vector of magnetic field} \\ \text{Field vector velocity of medium of electro-magnetic field} \end{bmatrix} = \begin{bmatrix} \text{The state vector of source functions of every state variable for the point } Z(x, y, z, t) \text{ of } \Omega_{Rt} \text{ with the brim of } F_{Rt} \end{bmatrix} +$$

$$+ \begin{bmatrix} \text{The state vector of} \\ \text{phenomenal solutions} \\ \text{of the potential fields} \\ \text{for } \Omega_{Rt} \text{ with the} \\ \text{brim } F_{Rt} \text{ and} \\ Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q(\xi, \eta, \beta, \tau) \end{bmatrix} + \begin{bmatrix} \text{The state vector of} \\ \text{phenomenal solutions} \\ \text{of the rotational field} \\ \text{for } \Omega_{Rt} \text{ with the} \\ \text{brim } F_{Rt} \text{ and} \\ Q^R(\xi^R, \eta^R, \beta^R, \tau) \equiv Q'(\xi', \eta', \beta', \tau) \end{bmatrix}.$$

The above presented complete solution contains in this way the initial and boundary aspects of all phenomenal solutions of considered electromagnetic field.

THE INPUT-OUTPUT PHENOMENAL CONTROL PROBLEMS FOR ELECTROMAGNETIC FIELD

As a consequence of the suggested in article theory, we can defined the subsequent state vectors of electromagnetic field:

– the INPUT state vector

$$[I(Z, y)] = \begin{bmatrix} \rho_o(Z, t) \\ D_o(Z, t) \\ B_o(Z, t) \\ v_o(Z, t) \end{bmatrix} \quad (1)$$

– the SOURCE state vector, from the system (I^s)

$$[M(Z, t)] = \begin{bmatrix} \rho^z(t) \\ D^z(t) \\ B^z(t) \\ v^z(t) \end{bmatrix} Z(x, y, z, t) \quad (2)$$

– the PROCESS state vector, from the distributed parameter description and subsequently solution of the processing of electromagnetic field, determined as the phenomenal Green functions being medium phenomenal responses on the space-time. Delta impulses pertinent to the physical phenomena and the general complete solution procedure (A.3.1), so that

$$[R(Z, t)] = G_R(Z, t) \quad (3)$$

$G_R(Z, t)$ – resultant for the whole electromagnetic field processing Green function,

– the OUTPUT state vector, as the complete solution state vector

$$[S(Z, t)] = \begin{bmatrix} \Delta \\ D(Z, t) \\ B(Z, t) \\ v(Z, t) \end{bmatrix} \quad (4)$$

Δ – the places of the physical phenomena which are not taking parts into the considerations.

The application of above approach makes possible on the formulation of the INPUT→SOURCE→ELECTROMAGNETIC PROCESS→OUTPUT RELATION BY ADEQUATE STATE VECTORS:

$$\{I(Z, t)\} \rightarrow \{M(Z, t)\} \rightarrow \{R(Z, t)\} \rightarrow \{S(Z, t)\} \quad (5)$$

On the basis of the above approach one can formulate the phenomenally distributed parameter control concept for the processing of electromagnetic field [10]. The idea of this control in the figure Fig. 1. has been presented.

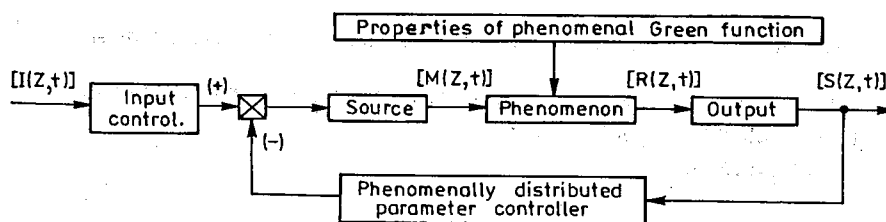


Fig. 1. The idea of the phenomenally distributed parameter control for the electromagnetic field. $[I(Z,t)]$ – the INPUT state vector, $[M(Z,t)]$ – the SOURCE state vector, $[R(Z,t)]$ – the PROCESS state vector, $[S(Z,t)]$ – the OUTPUT state vector. The relation INPUT – OUTPUT of one physical phenomenon on a few state variables is possible.

CONCLUSIONS

The results if the considerations of this article allow on the formulation of the space-time distribution problems for the state variables of electromagnetic field. The suggested in article solution method is based on the dimensions „a, b, c” of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$ in the phenomenal solutions of all state variables „(ρ, D)B, v” which characterize electromagnetic field. On the other hand the space-time distribution of characteristic state variables „(ρ, D)B, v” of electromagnetic field may have a very wide field of applications:

- for the theory of aerals of both emission and adsorbtion problems of civil and military kinds,
- the constructions of apparatuses of electromagnetic field,
- the problems of electromagnetic guns,
- the processing of electromagnetic signals,
- the plasma and magnetohydrodynamic problems.

Although our considerations to the limited Ω_{Rt} are reduced we may extend them for the unlimited case of Ω_{Rt} .

The circulation of the normal external surface orientation vector $n_R^{(z)}$ of the brim surface F_{Rt} in relation to the circulation of the normal external surface orientation vector $n^{(z)}$ of the surface $F_{(a,b,c)t}$ makes possible on the formulation of the boundary control problems on the surface F_{Rt} pertinent to the summation criterions of the effects

of physical phenomena of electromagnetic field n adequate its charge, energy and momentum balances.

The complete solution procedere based on the summations of the effects of the physical phenomena of potential fields and rotational field to the coordinates of electromagnetic field pertinent to the balances, by the existence of initial conditions

— case A and by the existence of initial and boundary conditions

— case B, from the properties of phenomenal Green functions is valid. The constitutive distributed parameter mathematical model deduced in Part I [1] of the article and solved here for cases A and B has its basis on the presented above and space-time distributed state variables of electromagnetic field which are measureable. As a consequence of this approach the idea of the phenomenally distributed parameter control of the electromagnetic field in the figure Fig. 1, has been presented.

APPENDIX 1

SOURCE FUNCTION OF THE ELECTRIC INDUCTION VECTOR IN THE POINT $Z(x, y, z, t)$

The outset formula for the source of the electric induction vector (F.1) possesses the form:

$$\frac{dD}{dt} = \frac{\kappa}{\epsilon_0 \epsilon} D \pm J(t). \quad (1)$$

The eq. (1) fulfils its general presentation

$$\frac{dD}{dt} - R(t)D = U(t), \quad (2)$$

where:

$$R(t) = \frac{\kappa}{\epsilon_0 \epsilon} \quad \text{and} \quad U(t) = \pm J(t).$$

The eq. (2) contains the following solution steps [8]:

— the integration factor

$$E_I = e^{\int R(t) dt} \quad (3)$$

— the general integral for the eq. (3) is:

$$E(t) = e^{-\int R(t) dt} \left[\int U(t) e^{\int R(t) dt} dt + D_C \right] \quad (4)$$

and the vector integration constant D_C is determined for the case $t=t_0$ and $D=D_0$. By the case that any particular case is unknown, we have:

$$E(t) = D(t). \quad (5)$$

If one particular solution is known, one gets:

$$D(t) = D_1(t) + D_C e^{-\int R(t) dt}. \quad (6)$$

Consequently, if we know two different particular solutions $D_1(t)$ and $D_2(t)$,

$$D(t) = D_1(t) + |D_C|[D_2(t) - D_1(t)]. \quad (7)$$

APPENDIX 2

STATICAL ISOTROPIC GREEN TENSOR FOR THE VOLUME ELEMENT $\bar{\Omega}_{(a,b,c)t}$

Let be defined the tensor of statical deformations for the volume element $\bar{\Omega}_{(a,b,c)t}$ [9], [3], [2]

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

For the presented above tensor of the deformations, the statical Green tensor is [9], [3], [2]:

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

The component $\Gamma_{\hat{\alpha}}(Z, Q')$ is correlated to the operation [graddivv]

$$\Gamma_{\hat{\alpha}}(Z, Q') = \frac{1}{|\mathbf{r}|} - \frac{\mathbf{r} * \mathbf{r}}{|\mathbf{r}|^3}. \quad (3)$$

The component $\Gamma_{\hat{\beta}}(Z, Q')$ is correlated to the operation [-rotrotrv]

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

for $\mathbf{r} = \mathbf{r}_Z - \mathbf{r}_{Q'}$, [m] with $\hat{\alpha} = \hat{a} + \hat{b}$ and $\hat{\beta} = \hat{\alpha}$

APPENDIX 3

STATICAL ANISOTROPIC GREEN TENSOR FOR THE VOLUME ELEMENT $\bar{\Omega}_{(a,b,c)t}$

For the electromagnetic medium coefficients

$$\hat{\alpha} = \hat{a} + \hat{b} \quad \text{and} \quad \hat{\beta} = \hat{a} \quad (1)$$

the elasticity equation of the deformations of the volume element $\bar{\Omega}_{(a,b,c)t}$ is:

$$\varepsilon_{iklm} \nabla_k \nabla_m V_1(Z) = -N_{Vi}(Z) \quad (2)$$

where: ε_{iklm} — the elasticity tensor of the electromagnetic medium, $V_1(Z)$ — the field of the deformations of the volume element $\bar{\Omega}_{(a,b,c)t}$, $N_{Vi}(Z)$ — the field of the forces of the deformations of the volume element $\bar{\Omega}_{(a,b,c)t}$.

Now let us introduce the operation

$$\varepsilon_{iklm} \nabla_k \nabla_m = \hat{\varepsilon}_{il} \quad (3)$$

which modifies the eq. (2) to the form

$$\hat{\varepsilon}_{il} V_1(Z) = -N_{Vi}(Z). \quad (4)$$

The field of the deformations $V_1(Z)$ of the locally selected volume element $\bar{\Omega}_{(a,b,c)}$ of the electromagnetic medium has the below form

$$V_1(Z) = \iiint_{\Omega_R} \Gamma_{ln}(Z - Q') N_{Vn}(Q') d\Omega_R(Q'), \quad (5)$$

where: $\Gamma_{ln}(Z - Q')$ — the statical Green tensor of the volume element $\bar{\Omega}_{(a,b,c)}$, $Q'(\xi', \eta', \beta')$ — the variable point of rotational field and $Q' \in \Omega_R$. In accordance with the above considerations the eq. (4) obtain the form

$$\hat{\varepsilon}_{il} \Gamma_{ln} = -\delta_{in} \delta_D(Z). \quad (6)$$

The splice form of the eq. (5) is:

$$V_1 = \Gamma_{ln} * N_{Vn}. \quad (7)$$

The elasticity tensor of the deformations of the volume element $\bar{\Omega}_{(a,b,c)}$ as the anisotropic medium can be presented in the series form (φ — small parameter)

$$\varepsilon_{iklm} = \sum_{r=0}^{\infty} \varphi^r \varepsilon_{iklm}^r \quad (8)$$

and

$$\varepsilon_{iklm}^0 = \hat{\alpha} \delta_{ik} \delta_{lm} + \hat{\beta} (\delta_{il} \delta_{km} + \delta_{im} \delta_{kl}). \quad (9)$$

Let the Green tensor be defined in the series form

$$\Gamma_{ln} = \sum_{r=0}^{\infty} \varphi^r \Gamma_{ln}^r. \quad (10)$$

Making use of the eq. (8) with equations (3) and (10), the eq. (6) gets its final form

$$\left[\sum_{r=0}^{\infty} \varphi^r \hat{\varepsilon}_{il}^r \right] \left[\sum_{r=0}^{\infty} \varphi^r \Gamma_{ln}^r \right] = -\delta_{in} \delta_D(Z). \quad (11)$$

After development of the series (8) and (10) in the eq. (11) and comparing its parts by the same power of the small parameter φ , the anisotropic Green tensor has the last form

$$\Gamma_{GA} = \Gamma_{ln} = \Gamma_{lp} * \hat{\varepsilon}_{lpq} \Gamma_{k-1}^{qn} \quad \text{for} \quad k = l(l)w \quad (12)$$

and w — number of power of the small parameter for $\bar{\Omega}_{(a,b,c)}$.

NOTATION

ρ — density of the electric charge	$\frac{C}{m^3},$
D — vector of the induction of the electric field	$\frac{C}{m^2},$
B — vector of the induction of the magnetic field	T,
v — field vector velocity of the electromagnetic medium	$\frac{m}{s^2},$
D_e — coefficient of the electric charge diffusion	$\frac{m^2}{s},$
κ — coefficient of the electric conductivity	$\frac{A}{Vm},$
ϵ_0 — coefficient of the absolute electric permability	$\frac{C}{Vm},$
ϵ — coefficient of the relative electric permability	,
U_B — vector coefficient of the energy distribution of the magnetic induction	A,
V_D — coefficient of the energy of the conductivity and permedability for the electric induction	Tm,
W_D — vector coefficient of the energy distribution of the charge diffusion	T,
f — coefficient of the mass of the electric charge	$\frac{kg}{C},$
$\hat{a}, \hat{b}$ — coefficients of the material properties of the electromagnetic medium	$\frac{Ns}{m^2},$
M_B — coefficient of the force pf the magnetic induction	$\frac{Nm^2}{T},$
N_D — coefficient of the force of the electric induction	$\frac{Nm^4}{C},$
$J(t)$ — the outside source of the electric charge on the chosen direction	$\frac{A}{m^2},$
$J(t)$ — the outside source of the electric charge	$\frac{A}{m^3},$
$S(t)$ — the outside source of the electromagnetic energy	$\frac{J}{m^4},$
$T(t)$ — the outside source of the momentum of the electromagnetic medium	$\frac{N}{m^3},$
$\bar{\Omega}_{(a,b,c)t}$ — the locally selected volume-time element phenomenally closed (every phenomenon can be treated in the separate way) around point $Z(x,y,z,t)$ of the transformation of electromagnetic field	$m^3-s,$
Ω_{Rt} — the so called „working” space volume-time of the apparatuses of electromagnetic field	$m^3-s,$

$(a,b,c)_t$ — the dimensions of the locally selected volume-time element	
$\bar{\Omega}_{(a,b,c)t}$ and: $x=a, y=b, z=c$	m,m,m,s
(x,y,z,t) — the constant space-time coordinates of Ω_{Rt}	m,m,m,s
$F_{(a,b,c)t}$ — the external oriented surface of the locally selected volume-time element $\bar{\Omega}_{(a,b,c)t}$	m^2-s
$n^{(z)}$ — the normal external surface orientation vector of the surface $F_{(a,b,c)t}$	,
F_{Rt} — the external oriented surface of Ω_{Rt}	m^2-s
$n_R^{(z)}$ — the normal external surface orientation vector of the surface F_{Rt}	,
t — time	s.

REFERENCES

1. W. Niemiec: Part I of this article.
2. W. Niemiec: *On the constitutive distributed parameter modelling of the electric field. Part II. The solution of partial differential constitutive state equations of the electric field.* Rozprawy Elektrotechniczne, 1986, 32, z. 4, ss. 1109–1130
3. B. Średniawa: *Hydrodynamika i teoria sprężystości.* PWN, Warszawa, 1977
4. A. J. A. Morgan: *On the construction of constitutive equations for continuous media.* Archiwum Mechaniki Stosowanej, vol. 17, 1965, No. 1, pp. 145–174
5. T. Trajdos: *Matematyka dla inżynierów.* Tom 2. WNT, Warszawa, 1987
6. A. N. Tichonow, A. A. Samarski: *Równania fizyki matematycznej.* PWN, Warszawa, 1963
7. W. Nowacki: *Dynamiczne zagadnienia termosprężystości.* PWN, Warszawa, 1966
8. J. Bronsztejn, K. Siemiendiajew: *Matematyka: poradnik encyklopedyczny.* PWN, Warszawa, 1986
9. W. Kupradze: *Wybrane zagadnienia teorii sprężystości i termosprężystości.* Wydawnictwo PAW, Wrocław—Warszawa—Kraków, 1970
10. S. Węgrzyn: *Podstawy automatyki.* PWN, Warszawa, 1974

W. NIEMIEC

O KONSTYTUTYWNYM ROZŁOŻONYM PARAMETRYCZNIE MODELOWANIU
POLA ELEKTROMAGNETYCZNEGO
CZĘŚĆ II. ROZWIĄZANIE MODELOWYCH RÓWNAŃ RÓŻNICZKOWYCH CZĄSTKOWYCH
KONSTRUKTYWNYCH STANU POLA ELEKTROMAGNETYCZNEGO

Streszczenie

Wyprowadzony w CZĘŚCI I artykułu konstytutywny rozłożony parametrycznie model matematyczny pola elektromagnetycznego został rozwiązany w oryginalny sposób — analitycznie. Zostały rozważone dwa przypadki tego analitycznego rozwiązania:

A — istnienie warunków początkowych,

B — istnienie warunków początkowych i brzegowych, dla fizycznych zjawisk pola elektromagnetycznego. Zaprezentowano pewne wpływy sterujące poprzez zastosowanie zjawisk fizycznych na przestrzennie czasowo rozłożone ważne parametry pola elektromagnetycznego.

Przetwarzanie wiedzy o układach elektronicznych

HENRYK BUDZISZ

Instytut Elektroniki, Wyższa Szkoła Inżynierska w Koszalinie

Otrzymano 1990.07.05

Autoryzowano do druku 1991.01.15

Część 1: Unifikacja i reguły przetwarzania

W pierwszej części pracy przedstawiono mechanizm dostępu do informacji zawartych w klauzulowej reprezentacji wiedzy o układach elektronicznych. Mechanizm ten jest oparty na uzgadnianiu (unifikacji) predykatów wchodzących w skład klauzul. Dwa wyrażenia uważa się za uzgodnione, jeżeli za zmienne w obu wyrażeniach można podstawić takie wartości, że wyrażenia te stają się identyczne. Przedstawiono także zasadę łączenia klauzul w oparciu o regułę rezolucji, w ciągi tworzące drzewo wnioskowania. Omówiono też pokrótce strategię przeszukiwania takiego drzewa.

1. WPROWADZENIE

Klauzula jest twierdzeniem pozwalającym na przyjęcie (odrzućenie) pewnych konkluzji przy założeniu, że spełnione są (nie są spełnione) określone założenia [3,8]. W zależności od tego, które spośród predykatów składających się na klauzulę mają określoną wartość, klauzula może być interpretowana następująco:

- czy prawdą jest, że z tych założeń wynikają takie konkluzje?
- jakie konkluzje wynikają z tych założeń?
- jakie założenia muszą być spełnione, aby otrzymać takie konkluzje?

Pierwsza interpretacja jest często określana mianem weryfikacji. Występuje na przykład w sytuacji sprawdzania zgodności projektu z założeniami. Drugie sformułowanie wyraża proces typowy dla projektowania. Trzecią interpretację można kojarzyć z analizą pewnych obiektów lub zjawisk.

Aby odpowiedzieć na te pytania, należy przeprowadzić pewne rozumowanie, w trakcie którego używa się klauzul pomocniczych. Proces ten nazywan jest

wnioskowaniem. Ważnym elementem wnioskowania jest dopasowanie między konkluzjami i założeniami w ciągu klauzul tworzących wywód twierdzeń pierwotnego.

2. PODSTAWIENIA I UNIFIKACJA

Klasyczna baza danych jest zorganizowana przy użyciu indeksów — często złożonego hierarchicznego systemu indeksów. Dostęp do informacji uzyskuje się przez podanie odpowiednich wartości indeksów.

Klauzule tworzące bazę wiedzy są osiągalne poprzez treść jaką zawierają. Powszechnie używanym mechanizmem dostępu do bazy wiedzy jest porównywanie z wzorcem. Porównywanie odbywa się na poziomie predykatów. Ponieważ predykaty mogą zawierać argumenty w postaci zmiennych, za zmienne te należy podstawić takie termy, aby porównywane predykaty były identyczne. Operacja ta nazywana jest unifikacją lub dopasowaniem, albo uzgadnianiem.

Wynikiem unifikacji jest zbiór podstawień:

$$\{x_1 = t_1, x_2 = t_2, \dots, x_m = t_m\},$$

gdzie: $x_1, x_2, \dots, x_m$ — są zmiennymi,

$t_1, t_2, \dots, t_m$ — są termami [3, 7, 8].

Zmienne będą dalej oznaczane małymi literami alfabetu łacińskiego, np. x, y, z , a stałe identyfikatorami wieloliterowymi.

Zastąpienie zmiennych termami nie zmienia wartości logicznej predykatu. Wynika to z reguły podstawiania [4, 8], w której stwierdza się: jeżeli $\forall x W(x)$ jest prawdziwe, to również $W(A)$ jest prawdziwe, gdzie A jest dowolną stałą. Wynikiem uzgodnienia predykatów:

układ (żyrator) oraz układ (x),

będzie podstawienie $\{x = \text{żyrator}\}$.

Podobnie, w wyniku uzgodnienia:

jest _ to (wzmacniacz _ operacyjny, x)

oraz

jest _ to (y , układ _ liniowy),

otrzymuje się:

$$\{x = \text{układ_liniowy}, y = \text{wzmacniacz_operacyjny}\}.$$

Po wykonaniu tych podstawień oba predykaty będą identyczne.

Nieco bardziej złożona będzie unifikacja predykatów:

MOS (nazwa (M_1), typ (n), źródło (V_{DD}), bramka (we), dren (wy))

oraz MOS (x, y , źródło (z), bramka (b), dren (d)).

Wynikiem będą podstawienia:

$$\{x = \text{nazwa}(M_1), y = \text{typ}(n), z = V_{DD}, b = we, d = wy\}$$

Następny przykład pokazuje uzgodnienia z udziałem list.

Predykat

$$\text{MOS}(\text{nazwa}(M_1), \text{typ}(n), \text{zaciski}([V_{DD}, we, wy]))$$

porównany z predykatem

$$\text{MOS}(x, y, z),$$

daje w wyniku:

$$\{x = \text{nazwa}(M_1), y = \text{typ}(n), z = \text{zaciski}([V_{DD}, we, wy])\}.$$

Ten sam predykat porównany z predykatem

$$\text{MOS}(\text{nazwa}(M_1), y, \text{zaciski}(z))$$

daje

$$\{y = \text{typ}(n), z = [V_{DD}, we, wy]\}.$$

Z kolei porównanie z predykatem

$$\text{MOS}(\text{nazwa}(x), y, \text{zaciski}([z, b, d]))$$

prowadzi do podstawień:

$$\{x = M_1, y = \text{typ}(n), z = V_{DD}, b = we, d = wy\}.$$

Jeżeli uzgodnione predykaty zawierają zmienne na tych samych pozycjach, to zmienne te są utożsamiane, na przykład po uzgodnieniu

$$\text{jest_to}(\text{układ_liniowy}, x) \text{ oraz } \text{jest_to}(y, z)$$

otrzymuje się

$$\{y = \text{układ_liniowy}, x = z\}$$

Natomiast predykaty:

$$\text{jest_to}(\text{układ_liniowy}, x)$$

oraz

$$\text{jest_to}(\text{wzmacniacz}, y)$$

nie dają się uzgodnić, ponieważ na pierwszej pozycji mają różne stałe i żadne podstawienia nie doprowadzą do ich identyczności.

Podstawienie termu za zmienną obowiązuje dla wszystkich wystąpień w danej klauzuli. Na przykład po dokonaniu uzgodnienia między predykatami

spełnia (dzielnik _ napięcia, zasada _ superpozycji)
 oraz spełnia (x , zasada _ superpozycji)

otrzyma się podstawienie $\{x = \text{dzielnik_napięcia}\}$.

Jeżeli ten drugi predykat jest elementem klauzuli

spełnia (x , zasada _ superpozycji) $\leftarrow$ układ (x), liniowy (x)

to stała *dzielnik_napięcia* zostaje również podstawiona za x w predykatkach *układ* i *liniowy*. Operacja ta nazywa się ukonkretnieniem predykatów.

Znanych jest kilka algorytmów uzgadniania wyrażeń logicznych [4, 10]. Mechanizm uzgadniania został wbudowany do języka programowania logicznego — Prolog [6].

Mechanizm unifikacji jest wykorzystywany w procesie wnioskowania do dopasowania założeń jednej klauzuli do konkluzji innych klauzul lub odwrotnie. Może być także użyty do wyszukiwania informacji w niewielkich bazach danych. W dodatku przedstawiony został przykład zastosowania mechanizmu unifikacji do obsługi bazy danych zawierającej krótki katalog oscyloskopów.

3. REGUŁA REZOLUCJI

Przetwarzanie wiedzy zapisanej w języku klauzul oparte jest na regule rezolucji [5,7,10]. Reguła ta sformułowana jest następująco: jeżeli prawdziwe są wyrażenia logiczne $A_1 \vee A_2$ i $\neg A_2 \vee A_3$, to wynika z nich prawdziwość $A_1 \vee A_3$.

Istotnie, jeżeli założymy prawdziwość A_2 , to $\neg A_2$ jest fałszywe. Ale jeżeli $\neg A_2$ jest fałszywe, to z drugiego założenia wynika, że A_3 musi być prawdziwe. Stąd wynika prawdziwość $A_1 \vee A_3$.

Jeżeli z kolei założymy, że A_2 jest fałszywe, to z pierwszego założenia wynika, że A_1 musi być prawdziwe. Jeżeli tak jest, to $A_1 \vee A_3$ jest znowu prawdziwe.

Wyrażenie $A_1 \vee A_3$ nazywane jest rezolwentą wyrażeń $A_1 \vee A_2$ i $\neg A_2 \vee A_3$. Regułę rezolucji można uogólnić na dowolną liczbę dysjunkcji w obu założeniach (w tym również na założenia jednoelementowe).

Reguła rezolucji nie może być bezpośrednio zastosowana do klauzul hornowskich [3,8]. W takim przypadku stosuje się dwie jej szczególne postacie — *modus ponens* (reguła odrywania) i *modus tollens*.

Jeżeli regułę rezolucji ograniczy się do postaci:

„z prawdziwości A_1 i $\neg A_1 \vee A_2$ wynika prawdziwość A_2 ”

to zastępując alternatywę implikacją otrzymuje się klasyczny zapis reguły *modus ponendo ponens*.

„z prawdziwości A_1 i $A_1 \rightarrow A_2$ wynika prawdziwość A_2 ”

Podobnie, upraszczając regułę rezolucji do postaci:

„z prawdziwości $\neg A_2$ i $\neg A_1 \vee A_2$ wynika prawdziwość $\neg A_1$ ”

po zastąpieniu alternatywy implikacją, otrzymuje się:

”z prawdziwości $\neg A_2$ i $A_1 \rightarrow A_2$ wynika prawdziwość $\neg A_1$ ”

co jest klasycznym zapisem reguły *modus tollendo tollens*.

Reguły te można w sposób opisowy sformułować następująco:

- jeżeli pewna teza jest implikowana przez założenia prawdziwe, czy też akceptowalne według jakiegoś kryterium, logika każe nam przyjąć ją jako konkluzję
- jeżeli założenia implikują niemożliwą do zaakceptowania lub fałszywą tezę, to chcąc pozostać w zgodzie z logiką – powinniśmy odrzucić przynajmniej jedno z założeń.

T a b l i c a 1

Definicja implikacji

	założenia A_1	teza A_2	implikacja $A_1 \rightarrow A_2$
*	fałszywa	fałszywa	prawdziwa
	fałszywa	prawdziwa	prawdziwa
	prawdziwe	fałszywa	fałszywa
*	prawdziwe	prawdziwa	prawdziwa

Obie reguły można również odczytać z tabelki definiującej implikację. W tabl. 1 oznaczono je znakiem *. Z tabelki tej można też określić dwie inne reguły wnioskowania, znane pod nazwami *modus ponendo tollens* i *modus tollendo ponens*.

Obie reguły wnioskowania uzupełnione mechanizmem unifikacji stanowią bazę, na której zbudowane są dwie podstawowe metody przetwarzania wiedzy zapisanej w języku klauzul hornowskich – wnioskowania wstępującego (ang. *bottom-up derivation* lub *forward chaining*) i zstępującego (ang. *top-down derivation* lub *backward chaining*).

4. STRATEGIA POSZUKIWAŃ

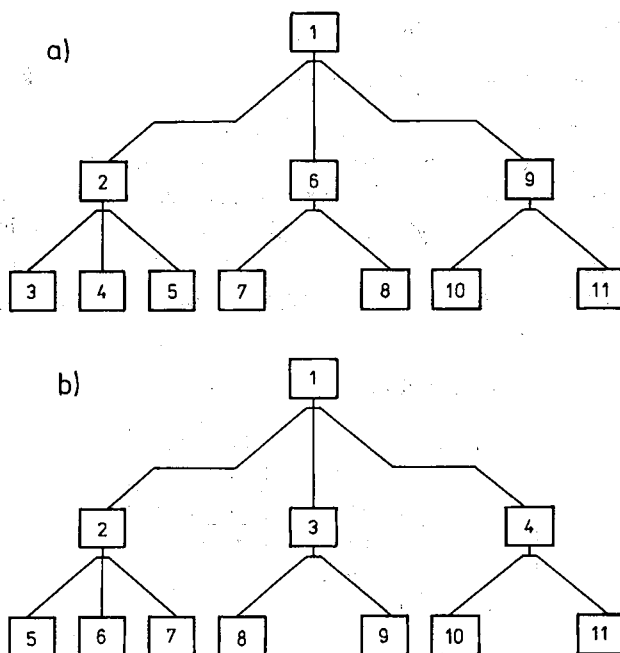
Reguły wnioskowania określają elementarne kroki procesu przetwarzania wiedzy. Wszystkie możliwe sposoby zastosowania reguł wnioskowania do początkową danego zbioru klauzul oraz klauzul z niego wyprowadzonych, tworzą dla tego zbioru klauzul przestrzeń poszukiwań (ang. *search space*). Graficznie przedstawia się ją w postaci drzewa.

W niektórych przypadkach, drzewo takie ma wiele tysięcy gałęzi, a w przypadku definicji rekursywnych, nieskończoną liczbę gałęzi. Zjawisko takie nazywane jest eksplozją kombinatoryczną.

Jeżeli rozmiar przestrzeni poszukiwań nie jest zbyt duży, to stosuje się poszukiwanie wyczerpujące, przy zastosowaniu jednej z dwóch podstawowych strategii [1,2,8]:

- poszukiwanie „w głąb” (ang. *depth-first search*)
- poszukiwanie „wszerz” (ang. *breadth-first search*).

Kolejność przechodzenia przez poszczególne węzły drzewa przy obu strategiach pokazana została na rys. 1.



Rys. 1. Zasada przeszukiwania drzewa: a) według strategii „w głąb”, b) według strategii „wszerz”

Znane są również różne modyfikacje tych strategii podstawowych [2,8]. Niezależnie od przyjętej strategii, istotną rolę w procesie przeszukiwania odgrywa też kolejność przeszukiwania gałęzi wychodzących ze wspólnego węzła, tzw. problem szeregowania zadań [8].

W warunkach eksplozji kombinatorycznej, nie można zastosować przeszukiwania wyczerpującego. Proces przeszukiwania musi być wówczas sterowany przy pomocy reguł pomocniczych, nazywanych metaregułami [1,9]. Często stosowana jest strategia odcinania gałęzi (ang. tree pruning strategies), które nie dają dobrych perspektyw na znalezienie rozwiązania [9].

PODSUMOWANIE

W pierwszej części pracy dokonano przeglądu podstawowych narzędzi umożliwiających przetwarzanie wiedzy. Są to mechanizm uzgadniania i reguła rezolucji. Zasady uzgadniania przedstawiono na przykładach. Regułę rezolucji przedstawiono w wersjach uproszczonych, umożliwiających realizację wnioskowania wstępującego i zstępującego dla klauzul hornowskich. Wspomniano też o problemach jakie powstają w procesie przeszukiwania drzewa wnioskowania.

W następnych dwóch częściach pracy przedstawione zostaną szczegółowo zasady wnioskowania wstępującego i zstępującego wraz z przykładami zastosowania do identyfikacji i weryfikacji układów elektronicznych.

BIBLIOGRAFIA

1. J. L. Alty, M. J. Coombs: *Expert System. 5 Concepts and Examples*, NCC Publ. 1984
2. A. Bonet: *Artificial Intelligence*. Prentice-Hall Inc., London, 1985
3. H. Budzisz: *Reprezentacja wiedzy o układach elektronicznych, cz. 1, Reprezentacja w języku klauzul*. Kwart. Elektr. i Telekom., 1991, 37, z. 3–4
4. E. Charniak: *Artificial Intelligence Programming*. Lawrence Erlbaum Assoc., Inc., 1980
5. E. Charniak, D. McDermott: *Introduction to artificial intelligence*. Addison-Wesley, Reading, 1985
6. W. F. Clocksin, C. S. Mellish: *Programming in Prolog*. Springer Verlag, Berlin, Heidelberg, 1984, 2-nd ed.
7. R. Kowalski: *Logic programming*. Information Processing, 1983, pp. 133–145, Invited paper
8. R. Kowalski: *Logika w rozwiązywaniu zadań*. Warszawa, WNT, 1989
9. A. Muszycka, R. Shinghal: *An Empirical Comparison of Pruning Strategies in Game Trees*. IEEE Trans. on Systems, Man, and Cybernetics, 1985, vol. SMC–15, No 3, pp. 389–399
10. N. J. Nilsson: *Principles of Artificial Intelligence*. Tioga Publishing Co., 1980

DODATEK

Poniżej przedstawiono przykład zastosowania mechanizmu uzgadniania do wyszukiwania danych katalogowych oscyloskopów firmy Philips (na podstawie katalogu z 1989 r)

domains

tak_nie = tak;nie

predicates

oscyloskop (string, integer, integer, integer,
tak_nie, tak_nie, tak_nie, tak_nie)
jest (tak_nie, tak_nie)

```

/*****
/* dane o oscyloskopie: model, szerokość pasma w MHz,
/* podstawa czasu w [ns/działkę], liczba kanałów,
/* podwójna podstawa czasu – tak/nie, znaczniki – tak/nie,
/* automatyczny wybór warunków pomiaru – t/n, pamięć – t/n
*****/

```

clauses

jest (X, X).
oscyloskop („PM 3295A/96A”, 400, 1, 2, tak, tak, tak, nie).
oscyloskop („PM 3285A/86A”, 200, 2, 2, tak, tak, tak, nie).
oscyloskop („PM 3070”, 100, 5, 2, tak, tak, tak, nie).
oscyloskop („PM 3065”, 100, 5, 2, tak, nie, tak, nie).
oscyloskop („PM 3064”, 100, 5, 2, tak, nie, nie, nie).
oscyloskop („PM 3266”, 100, 5, 2, tak, nie, nie, tak).
oscyloskop („PM 3055”, 50, 5, 2, tak, nie, tak, nie),
oscyloskop („PM 3050”, 50, 5, 2, nie, nie, tak, nie),
oscyloskop („PM 3219”, 50, 10, 2, tak, nie, nie, tak),
oscyloskop („PM 3206”, 15, 100, 2, nie, nie, nie, nie),

Sformułujemy teraz pytania dotyczące tych danych:

1. Jakie modele mają szerokość pasma ≥ 100 MHz i czy mają automatyczny wybór warunków pomiaru?

Goal: oscyloskop (Model, Szerokość_pasma, __, __, __, __, Automatyka, __),
and Szerokość_pasma > 100 .

Model=PM 3295A/96A, Szerokość_pasma=400, Automatyka=tak

Model=PM 3285A/86A, Szerokość _pasma=200, Automatyka=tak
 Model=PM 3070 , Szerokość _pasma=100, Automatyka=tak
 Model=PM 3065 , Szerokość _pasma=100, Automatyka=tak
 Model=PM 3064 , Szerokość _pasma=100, Automatyka=nie
 Model=PM 3266 , Szerokość _pasma=100, Automatyka=nie
 6 Solutions

2. Jakie modele wyposażone są w pamięć obrazu i jakie mają pasmo przenoszenia?

Goal: oscyloskop (Model, Szerokość _pasma, __, __, __, __, Pamięć),
 and jest (Pamięć, tak).

Model=PM 3266 , Szerokość _pasma=100, Pamięć=tak

Model=PM 3219 , Szerokość _pasma= 50, Pamięć=tak

2 Solutions

Prostota i zwartość realizacji jest tu wręcz zaskakująca.

H. BUDZISZ

KNOWLEDGE PROCESSING IN THE DOMAIN OF ELECTRONIC CIRCUITS

Part I: Unification and knowledge processing rules

S u m m a r y

In the first part of the paper, a retriever of knowledge from clause representation is presented. Given a request, the retriever looks through its data base of formulas until it finds one that unifies with the request. Two formulas are unified by finding values for their variables such that substituting values for variables in each formula, gives the same instance formula. The principle of chaining clauses according to the resolution rule has likewise been shown. The chained rules form a resolution tree and search strategies of the tree has also been briefly introduced.

Część 2: Wnioskowanie wstępujące

Reguły wnioskowania są ogólnymi twierdzeniami dotyczącymi związków między założeniami a konkluzjami tworzącymi klauzule. We wnioskowaniu wstępującym, przedstawionym w tej części pracy, proces przetwarzania wiedzy polega na cyklicznym stosowaniu reguły odrywania. W każdym kroku tego procesu, wyprowadzane są konkluzje, które z kolei stają się nowymi asercjami umieszczonymi w bazie wiedzy. Jako przykład ilustrujący sposób wnioskowania, przedstawiono problem identyfikacji modułów wzmacniaczy CMOS. Pokazano też sposób implementacji wnioskowania w Prologu.

1. WPROWADZENIE

W rozumowaniu wstępującym, wychodzi się od podanych założeń i wyprowadza się konkluzję, a z niej wyprowadza się dalsze konkluzje. Proces ten kończy się po osiągnięciu celu. O postępowaniu wstępującym mówi się też, że jest sterowane podanymi założeniami (ang. data-driven, data-directed) [1,6,9].