

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

KWARTALNIK
ELEKTRONIKI I TELEKOMUNIKACJI

ELECTRONICS AND
TELECOMMUNICATIONS
QUARTERLY

TOM 44 — ZESZYT 4



WYDAWNICTWO NAUKOWE PWN
WARSZAWA 1998

RADA REDAKCYJNA

Przewodniczący

prof. dr hab. inż. ALFRED ŚWIT
członek rzeczywisty PAN

Członkowie

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

REDAKCJA

Redaktor Naczelny

prof. dr hab. inż. WIESŁAW WOLIŃSKI

Zastępca Redaktora Naczelnego

doc. dr inż. KRYSTYN PLEWKO

Sekretarz Odpowiedzialny

mgr KRYSTYNA LELAKOWSKA

ADRES REDAKCJI

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

Dyżury Redakcji: środy i piątki, godz. 14–16
tel/fax (0-22) 660 77 37, 825 29 18 + aut. sekr.

Telefony domowe: Redaktora Naczelnego: 12 17 65

Zastępcy Red. Naczelnego: 826 83 41

Sekretarza Odpowiedzialnego: 825 29 18

WYDAWNICTWO NAUKOWE PWN

Warszawa, ul. Miodowa 10

Ark. wyd. 11,5 Ark. druk. 9,00

Podpisano do druku w lutym 1999 r.

Papier offsetowy kl. III 70 g. B-1

Druk ukończono w kwietniu 1999 r.

Skład: *Amel*, Warszawa

Druk i oprawa: Drukarnia Braci Grodzickich, Żabieniec, ul. Przelotowa 7

Syn
stapien
modela
lub zac
elemen

CO — czł.
HAŁAS,
hab. inż.
czł. rzecz.
OLIŃSKI
LSKI

Makromodele scalonych wzmacniaczy mocy małej częstotliwości

Józef Stanclik

Politechnika Wrocławska, Instytut Telekomunikacji i Akustyki

Otrzymano 1998.10.14.

Autoryzowano 1999.01.19

W pracy przedstawiono dwa makromodele scalonych wzmacniaczy mocy m.cz. przeznaczone do użytkowania w programie PSpice (model o większej dokładności, ale bardziej złożony i model prostszy lecz mniej dokładny). Oba makromodele cechują się stopniem złożoności wielokrotnie mniejszym, niż układy oryginalne i dokładnością modelowania podstawowych właściwości elektrycznych zadowalającą w typowych zastosowaniach. Makromodel o większej dokładności, przeznaczony do modelowania wzmacniaczy mocy o zasilaniu symetrycznym, stanowi modyfikację makromodelu zaproponowanego w [8]. Zapewnia on prawidłowe zależności prądów pobieranych ze źródeł zasilających oraz mocy: dostarczonej z zasilaczy, wyjściowej i strat, sprawności energetycznej, a po przesterowaniu również współczynnika zawartości harmonicznych od napięcia wejściowego, impedancji obciążenia i napięć zasilających. Makromodel o mniejszej złożoności jest przeznaczony do modelowania wzmacniaczy mocy o zasilaniu niesymetrycznym [10]. Przedstawiono struktury makromodeli, algorytmy identyfikacji i przykłady modelowania typowych scalonych wzmacniaczy mocy. Parametry makromodeli są wyznaczane na podstawie typowych danych katalogowych układów scalonych. Właściwości makromodeli zweryfikowano przez analizę układów próbnych i porównanie wyników z danymi katalogowymi. Ich praktyczną przydatność zilustrowano przykładami zastosowania do symulacji komputerowej elektroakustycznych wzmacniaczy mocy.

Słowa kluczowe: makromodel, wzmacniacz mocy m.cz., program PSpice.

1. WPROWADZENIE

Symulacja komputerowa układów elektronicznych jest przeprowadzana po zastąpieniu poszczególnych elementów biernych, czynnych i układów scalonych ich modelami elektrycznymi. Układy scalone można modelować na poziomie elementów lub zaciskowo (*black-box, behavioral modelling*) [2, 11]. Modelowanie na poziomie elementów jest poważnie utrudnione ze względu na niedostępność odpowiednich

danych technologicznych oraz złożoność elektrycznej sieci zastępczej. Powszechnie jest stosowane modelowanie zaciskowe gdzie układ scalony jest zastępowany makromodelem czyli siecią elektryczną, której stopień złożoności jest znacznie mniejszy niż stopień złożoności modelowanego układu a dokładność modelowania ważnych elektrycznych właściwości jest wystarczająca dla celów praktycznych. Skutkiem stosowania makromodeli jest nie tylko zmniejszenie rozmiaru sieci elektrycznej, reprezentującej analizowany układ scalony (co prowadzi do skrócenia czasu obliczeń i obniżenia kosztów symulacji), ale także zmniejszenie ewentualnych problemów ze zbieżnością procedur numerycznych symulatorów układów elektronicznych (w przypadku symulacji na poziomie elementów układu zawierającego nawet tylko jeden scalony wzmacniacz mocy m.cz. problemy takie mogą wystąpić — patrz p. 3). Stosowanie makromodeli jest zatem szczególnie atrakcyjne przy wspomaganiem komputerem projektowaniu urządzeń zawierających większą liczbę układów scalonych.

Symulacja komputerowa wzmacniaczy elektroakustycznych zawierających scalone wzmacniacze mocy m.cz. byłaby ułatwiona, gdyby w skład bibliotek modeli programów obliczeniowych wchodziły odpowiednie makromodele. Popularny i szeroko stosowany program PSpice jest wyposażony w biblioteki modeli analogowych układów scalonych. Jednak zarówno standardowe biblioteki, opracowane przez firmę MicroSim Corporation, lub uzyskiwane za pomocą programu identyfikacji modeli elementów PARTS, który wchodzi w skład pakietu MicroSim Design Center, jak i biblioteki opracowane przez poszczególnych producentów układów scalonych nie zawierają makromodeli scalonych wzmacniaczy mocy m.cz. [4].

W tej pracy problem braku modeli wzmacniaczy mocy m.cz. został rozwiązany. Przedstawiono dwa makromodele scalonych wzmacniaczy mocy m.cz. (o większej i mniejszej dokładności i złożoności). Makromodel o większej dokładności, przeznaczony do modelowania wzmacniaczy mocy o zasilaniu symetrycznym, stanowi modyfikację makromodelu prezentowanego na konferencji 95 Convention of the Audio Engineering Society [8]. W modelu tym uwzględniono wpływ napięcia sterującego, impedancji obciążenia i napięć zasilających na prądy pobierane ze źródeł zasilających oraz na moce: dostarczoną z zasilaczy, wyjściową i strat, sprawność energetyczną, a po przekroczeniu progu obcinania, również na współczynnik zawartości harmonicznych [9]. Makromodel o mniejszej złożoności jest przeznaczony do modelowania wzmacniaczy mocy o zasilaniu niesymetrycznym. Stanowi on uproszczenie makromodelu o większej dokładności i był prezentowany na konferencji OSA'96 [10].

Oba makromodele opracowano przy użyciu typowych technik stosowanych w modelowaniu: techniki upraszczania (*simplification technique*) i dobudowywania (*build-up technique*) [1, 2]. W pracy przedstawiono struktury makromodeli o większej dokładności i o mniejszej złożoności oraz procedury identyfikacji ich parametrów i elementów na podstawie typowych danych katalogowych układów scalonych. Identyfikacja makromodeli została zilustrowana przykładami modelowania wzmacniaczy TDA2030A i UL1495N. Właściwości makromodeli zweryfikowano przez

analizy
niaczy
kompu
zustr
tyczny

Str
cniacz
cnie
ciężn
w kla
równie
jściow
zasilaj
został
do kla

Str
dokład
metryc
trzy s
oraz ż
charak
niezró
wartą
malny
i opac
przyję
niezale
SR (w
w celu
odpor
modul

Sto
(elem
oraz r
wejście
modyf

analizę układów próbnych i porównanie wyników z danymi katalogowymi wzmacniaczy mocy, a w przypadku układu UL1495N również z wynikami symulacji komputerowej pełnego schematu ideowego. Praktyczną przydatność makromodeli zilustrowano przykładami zastosowania do symulacji komputerowej elektroakustycznych wzmacniaczy mocy.

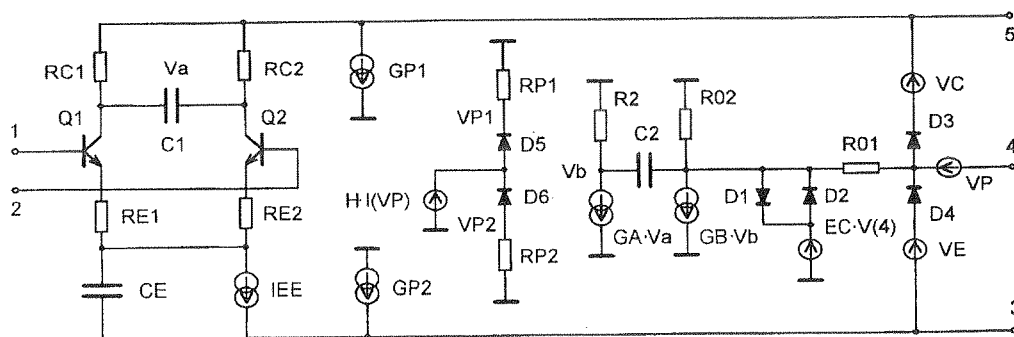
2. MAKROMODEL WZMACNIACZY MOCY M.CZ. O WIĘKSZEJ DOKŁADNOŚCI

Struktura scalonego wzmacniacza mocy m.cz. jest podobna do struktury wzmacniacza operacyjnego. W skład obu wzmacniaczy wchodzi po trzy stopnie wzmocnienia: stopień wejściowy zawierający parę różnicową, stopień sterujący z obciążeniem dynamicznym i stopień wyjściowy z wtórnikiem komplementarnym w klasie AB. Można więc przyjąć, że makromodele obu układów scalonych również mogą być podobne. Różnice dotyczą głównie wielkości prądów: wyjściowego i zasilającego oraz konieczności modelowania poboru mocy ze źródeł zasilających w przypadku wzmacniacza mocy. Na podstawie takich przesłanek został wyprowadzony makromodel scalonych wzmacniaczy mocy m.cz., podobny do klasycznego modelu Boyle'a [1].

2.1. Opis makromodelu

Struktura makromodelu scalonego wzmacniacza mocy m.cz. o większej dokładności, przeznaczonego do modelowania wzmacniaczy o zasilaniu symetrycznym, jest pokazana na rys. 1. W makromodelu tym można wyróżnić trzy stopnie. Stopień wejściowy składa się z idealnych tranzystorów Q_1 i Q_2 oraz źródła prądowego i elementów pasywnych. Stopień ten modeluje nieliniowe charakterystyki wejściowe wzmacniacza, uwzględnia wejściowe napięcie i prąd niezrównoważenia, położenie drugiego bieguna transmitancji wzmacniacza z otwartą pętlą sprzężenia zwrotnego i umożliwia uwzględnienie różnicy maksymalnych szybkości zmian napięcia wyjściowego na zboczu narastającym SR^+ i opadającym SR^- . Wzmocnienie napięciowe stopnia wejściowego dla wygody przyjęto jako jednostkowe. Rezystory R_{e1} i R_{e2} w makromodelu umożliwiają niezależny dobór pierwszego bieguna transmitancji wzmacniacza i współczynnika SR (w rzeczywistych układach wzmacniaczy rezystory takie są stosowane w celu poszerzenia zakresu liniowej pracy stopnia wejściowego i zwiększenia odporności wzmacniacza na generowanie dynamicznych zniekształceń intermodulacyjnych [3]).

Stopień pośredni (wraz ze stopniem wyjściowym) ustala wzmocnienie układu (elementy G_a , V_a i R_2), pierwszy biegun transmitancji (źródło G_a i kondensator C_2) oraz maksymalną szybkość zmian napięcia wyjściowego SR (wraz ze stopniem wejściowym). Kondensator C_2 jest dostępny na zewnątrz makromodelu i umożliwia modyfikację obwodu kompensacji charakterystyk częstotliwościowych wzmacnia-



Rys. 1. Struktura makromodelu scalonego wzmacniacza mocy m.cz. o zasilaniu symetrycznym

cza. Sposób jego włączenia do makromodelu zapewnia uzależnienie impedancji wyjściowej wzmacniacza od częstotliwości.

Stopień wyjściowy modeluje zależność impedancji wyjściowej wzmacniacza od częstotliwości. Rezystancja wyjściowa wzmacniacza przy prądzie stałym R_{out} jest sumą rezystorów R_{01} i R_{02} . W miarę wzrostu częstotliwości rezystancja R_{02} jest bocznikowana przez pojemność wyjściową stopnia pośredniego i przy większych częstotliwościach moduł impedancji wyjściowej maleje do poziomu rezystancji R_{01} . Elementy $E_C V(4)$, D_1 i D_2 zapewniają ograniczanie maksymalnego prądu wyjściowego, zaś elementy V_C i D_3 oraz V_E i D_4 — ograniczanie maksymalnego napięcia wyjściowego.

Sekcja mocy makromodelu zapewnia właściwe zależności prądów pobieranych ze źródeł zasilających (prąd spoczynkowy, prąd pobierany ze źródła napięcia dodatniego i ujemnego), oraz w konsekwencji mocy zasilania, mocy strat i sprawności energetycznej od napięcia sterującego i impedancji obciążenia. Pobór mocy ze źródeł zasilających jest modelowany za pomocą wielomianowych źródeł prądowych sterowanych napięciem G_{P1} i G_{P2} . Stałe współczynniki wielomianów źródeł odpowiadają prądowi spoczynkowemu wzmacniacza (prądy tranzystorów Q_1 i Q_2 można pominąć). Źródła G_{P1} i G_{P2} są sterowane dodatnimi i ujemnymi (odpowiednio) częściami prądu wyjściowego za pośrednictwem pomocniczego obwodu, złożonego z diod D_5 i D_6 , rezystorów R_{P1} i R_{P2} oraz ze źródła napięciowego H sterowanego prądem źródła V_P próbującego prąd wyjściowy. Diody D_5 i D_6 są idealnymi diodami półprzewodnikowymi, podobnie jak D_1 do D_4 . Transkonduktancje źródeł G_{P1} i G_{P2} , transrezystancja źródła H i rezystancje rezystorów R_{P1} i R_{P2} zostały dobrane w taki sposób, żeby spadki napięć na diodach D_5 i D_6 można było zaniedbać. Przyjęty sposób modelowania poboru mocy ze źródeł zasilających zapewnia prawidłowe wielkości prądów zasilania w zależności od prądu wyjściowego w całym zakresie napięć sterujących, napięć zasilających i impedancji obciążenia wzmacniacza.

Obni-
niacza
modelu
wią na
jściow
napięc
zasilan
(I_{os}), w
wzmoc
częstot
mitanc
charak
szybko
(SR^- ;
prądzi
nie są
ność k

Na
zestaw
rów m
[1] z
właści
uwzglę
źródle
potrze
zależn
mocy

Tr
o para
a ich v

Współ-
czynnik
równą
Jeżeli
na zbc

2.2. Identyfikacja makromodelu

Obliczenia wartości elementów i parametrów makromodelu skalonego wzmacniacza mocy m.cz. o większej dokładności wykonuje się sukcesywnie (parametry modelu są wyznaczane kolejno, niezależnie od siebie). Podstawę do obliczeń stanowią następujące parametry katalogowe układu skalonego: maksymalna moc wyjściowa wzmacniacza (P_{omax}) przy znamionowej rezystancji obciążenia (R_L) oraz napięciach zasilających: dodatnim (V_{CC}) i ujemnym (V_{EE}), spoczynkowy prąd zasilania (I_d), wejściowy prąd polaryzacji (I_b), wejściowy prąd niezrównoważenia (I_{os}), wejściowe napięcie niezrównoważenia (V_{os}), szczytowy prąd wyjściowy (I_{omax}), wzmocnienie napięciowe wzmacniacza z otwartą pętlą sprzężenia zwrotnego (A_V), częstotliwości odpowiadające pierwszemu (f_{p1}) i drugiemu (f_{p2}) biegunowi transmitancji napięciowej wzmacniacza z otwartą pętlą, wyznaczone na podstawie charakterystyk: amplitudowej $|A_V| = f(f)$ i fazowej $arg(A_V) = f(f)$, współczynniki szybkości zmian napięcia wyjściowego na zboczu narastającym (SR^+) i opadającym (SR^- ; zwykle $SR^+ = SR^- = SR$), moduł impedancji wyjściowej wzmacniacza przy prądzie zmiennym (R_{0ac}) i rezystancja przy prądzie stałym (R_{out} ; wielkości te często nie są podawane — odpowiednie parametry modelu są wtedy szacowane), pojemność kondensatora kompensacji częstotliwościowej (C_2).

Na podstawie wymienionych parametrów wzmacniacza mocy i zależności zestawionych poniżej wyznaczane są wartości poszczególnych elementów i parametrów makromodelu. Zależności te wyprowadzono w podobny sposób, jak w pracy [1] z uwzględnieniem specyfiki skalonych wzmacniaczy mocy m.cz. Pominięto właściwości wzmacniaczy dotyczące sterowania współfazowego (sumacyjnego), uwzględniono natomiast poziom wielkości prądów wyjściowych i pobieranych ze źródeł zasilających oraz wielkość rezystancji wewnętrznych, a także dostępność potrzebnych danych katalogowych. Kolejność obliczeń i forma poszczególnych zależności jest dogodna przy identyfikacji makromodeli skalonych wzmacniaczy mocy m.cz.

Tranzystory Q_1 i Q_2 (rys. 1) są modelowane za pomocą najprostszych modeli o parametrach β_1, I_{s1} oraz β_2, I_{s2} . Prądy kolektorów obu tranzystorów są jednakowe, a ich wartość wynika z wymaganej wielkości współczynnika SR^+ :

$$I_{C1} = I_{C2} = k \frac{C_2 \cdot SR^+}{2}. \quad (1)$$

Współczynnik korekcyjny k wprowadzono w celu zwiększenia dokładności współczynników SR^+ i SR^- makromodelu. Jego wartość dobrano doświadczalnie jako równą $k = 1,09$.

Jeżeli wielkość współczynnika „slew rate” na zboczu opadającym jest mniejsza, niż na zboczu narastającym, wtedy C_E jest równe:

$$C_E = \frac{2I_{C1}}{SR^-} - C_2, \quad (2)$$

natomiast gdy $SR^+ < SR^-$, to jako Q_1 i Q_2 należy użyć tranzystory $p-n-p$ i we wzorach (1) i (2) zastąpić miejscami SR^+ i SR^- , a jeśli na podstawie dostępnych danych nie jest możliwe zróżnicowanie współczynników SR , to do zależności (1) należy wstawić SR , a C_E przyjąć równe zero [1].

Wielkości współczynników SR^+ i SR^- nie zawsze są podawane. Można je wyznaczyć na podstawie pasma przenoszenia mocy f_p lub znanej amplitudy maksymalnego, nieznieskształconego napięcia wyjściowego V_{omax} wzmacniacza przy częstotliwości f_p , jako:

$$SR = 2\pi V_{omax} f_p, \quad (3)$$

gdzie: f_p — pasmo przenoszenia mocy lub górna częstotliwość graniczna wzmacniacza przy wielkim sygnale,

V_{omax} — amplituda napięcia wyjściowego przy częstotliwości f_p znana lub obliczona jako: $V_{omax} = \sqrt{2P(f_p) R_L}$.

Wejściowe napięcie i prąd niezrównoważenia są uzyskiwane w wyniku przyjęcia różnych parametrów modeli tranzystorów Q_1 i Q_2 (napięcie niezrównoważenia jest wynikiem różnych prądów nasycenia złącz, a prąd niezrównoważenia — różnych współczynników wzmocnienia prądowego):

$$\beta_1 = \frac{I_{C1}}{I_b + \frac{I_{os}}{2}} \quad \beta_2 = \frac{I_{C2}}{I_b - \frac{I_{os}}{2}} \quad (4)$$

$$I_{S1} = 8 \cdot 10^{-16} \text{ A} \quad I_{S2} = I_{S1} \cdot \exp\left(\frac{V_{os}}{V_T}\right), \quad (5)$$

gdzie $V_T = kT/q = 25,85 \text{ mV}$ ($T = 300 \text{ K}$), a wartość prądu I_{S1} przyjęto za [1].

Wartości rezystorów R_{C1} i R_{C2} są zależne od położenia pierwszego bieguna transmitancji wzmacniacza z otwartą pętlą sprzężenia zwrotnego i można je wyznaczyć jako:

$$R_{C1} = R_{C2} = \frac{1}{2\pi A_V C_2 f_{p1}}. \quad (6)$$

Jednostkowe wzmocnienie pierwszego stopnia uzyskuje się, gdy są spełnione zależności:

$$G_a = \frac{1}{R_{C1}} \quad (7)$$

$$R_{e1} = R_{e2} = \frac{\beta_1 + \beta_2}{2 + \beta_1 + \beta_2} \left(R_{C1} - \frac{V_T}{I_{C1}} \right). \quad (8)$$

Prąd źródła I_{EE} jest sumą prądów emiterowych tranzystorów Q_1 i Q_2 :

$$I_{EE} = \left(\frac{\beta_1 + 1}{\beta_1} + \frac{\beta_2 + 1}{\beta_2} \right) \cdot I_{C1}. \quad (9)$$

Kondensator C_1 wpływa na położenie drugiego bieguna transmitancji wzmacniacza z otwartą pętlą:

$$C_1 = \frac{1}{4\pi R_{C1} f_{p2}}. \quad (10)$$

Rezystory R_{01} i R_{02} łatwo jest wyznaczyć jeśli są znane rezystancje wyjściowe R_{out} i R_{0ac} ($R_{01} = R_{0ac}$, $R_{02} = R_{out} - R_{0ac}$). Często jednak brak jest tych danych. Wartość rezystora R_{02} można wtedy przyjąć równą $R_{02} = 1\Omega$ [8]. Od rezystora R_{01} zależy szczytowy prąd wyjściowy wzmacniacza. W makromodelu wzmacniacza mocy m.cz. to kryterium lepiej określa wartość rezystora R_{01} , niż inne kryteria proponowane w [1]:

$$R_{01} = \frac{V_T}{I_{0max}} \cdot \ln \left(\frac{2A_V I_{C1} R_{C1}}{I_{S1} R_{02}} \right). \quad (11)$$

Iloczyn $G_b R_2$ wpływa na wzmocnienie A_V i jedną z wielkości G_b lub R_2 można założyć. Przy wzmocnieniu modelowanego wzmacniacza z otwartą pętlą sprzężenia rzędu $10^4 - 10^5$ zalecana wartość rezystora R_2 wynosi 100 k Ω [1]. Teraz G_b można obliczyć jako:

$$G_b = A_V \cdot \frac{R_{C1}}{R_2 R_{02}}. \quad (12)$$

Prądy nasycenia diod D_1 i D_2 przyjmuje się równe I_{s1} , zaś wzmocnienie źródła sterowanego E_C jest równe jedności.

Ograniczanie napięcia wyjściowego następuje gdy dioda D_3 lub D_4 przechodzi w stan przewodzenia. Jeśli prądy nasycenia tych diod są równe $I_{sD3} = I_{sD4} = I_{s1}$ i obcinanie jest symetryczne, to napięcia źródeł V_C i V_E można obliczyć jako:

$$V_C = V_E = V_{CC} - \sqrt{2P_{0max} R_L} + V_T \cdot \ln \left(\frac{1}{I_{s1}} \sqrt{\frac{2P_{0max}}{R_L}} \right). \quad (13)$$

Elementy sekcji mocy zostały dobrane w taki sposób, aby błąd modelowania poboru prądu ze źródeł zasilających δ nie był większy od 0,02% przy prądzie wyjściowym o maksymalnej wartości 3,5 A. Aby spadki napięć U_D na przewodzących diodach D_5 i D_6 (przyjęto $U_D = 0,7$ V) mogły zostać z dokładnością do 0,02% pominięte w stosunku do napięcia źródła sterowanego H , transrezystancja źródła musi spełniać warunek:

$$\delta = \frac{U_D}{H \cdot I_{0max}} \leq 0,02\%,$$

stąd

$$H = \frac{U_D}{\delta \cdot I_{o\max}} \geq \frac{0,7 \text{ [V]}}{2 \cdot 10^{-4} \cdot 3,5 \text{ [A]}} = 1000 \Omega. \quad (14)$$

Zatem spełnienie przyjętego kryterium dokładności jest możliwe, przy wartościach pozostałych elementów sekcji mocy mieszczących się w typowych przedziałach wielkości, gdy elementy sekcji mocy są następujące:

transrezystancja źródła H :	$H = 1 \text{ k}\Omega$,
rezystancje rezystorów R_{P1} i R_{P2} :	$R_{P1} = R_{P2} = 1 \text{ M}\Omega$
diody D_5 i D_6 takie, jak D_1 do D_4 :	$I_{sD5} = I_{sD6} = I_{s1}$,
współczynniki źródła wielomianowego G_{P1} :	$P_0 = I_{d^*}$, $P_1 = 1 \text{ mS}$,
współczynniki źródła wielomianowego G_{P2} :	$P_0 = I_{d^*}$, $P_1 = -1 \text{ mS}$.

2.3. Makromodel skalonego wzmacniacza mocy TDA2030A

W celu zilustrowania przedstawionej procedury identyfikacji zostały wyznaczone elementy makromodelu układu skalonego TDA2030A. Wartości potrzebnych danych katalogowych są następujące [7]: maksymalna moc wyjściowa $P_{o\max} = 22 \text{ W}$ przy nominalnej rezystancji obciążenia $R_L = 4 \Omega$ i napięciach zasilających $V_{CC} = -V_{EE} = 18 \text{ V}$, spoczynkowy prąd zasilania $I_d = 50 \text{ mA}$, wejściowy prąd polaryzacji $I_b = 0,2 \mu\text{A}$, wejściowy prąd niezrównoważenia $I_{os} = 0,02 \mu\text{A}$, wejściowe napięcie niezrównoważenia $V_{os} = 2 \text{ mV}$, szczytowy prąd wyjściowy $I_{o\max} = 3,5 \text{ A}$, wzmocnienie napięciowe wzmacniacza z otwartą pętlą sprzężenia zwrotnego $A_V = 31800$ (90 dB), częstotliwości $f_{p1} = 300 \text{ Hz}$ i $f_{p2} = 2 \text{ MHz}$, współczynnik „slew rate” $SR^+ = SR^- = SR = 8 \text{ V}/\mu\text{s}$ (przy zamkniętej pętli sprzężenia zwrotnego i wzmocnieniu napięciowym 26 dB), R_{0ac} , R_{out} i C_2 nie są znane; zakładamy $C_2 = 30 \text{ pF}$.

Na podstawie tych danych wyznaczono następujące elementy i parametry makromodelu:

$I_{C1} = I_{C2} = 130,8 \mu\text{A}$, $C_E = 0$, $\beta_1 = 622,857$, $\beta_2 = 688,421$, $I_{s1} = 8,0 \cdot 10^{-16} \text{ A}$, $I_{s2} = 8,64353 \cdot 10^{-16} \text{ A}$ ($T = 300 \text{ K}$), $R_{C1} = R_{C2} = 556,11 \Omega$, $G_a = 1,7982 \text{ mS}$, $R_{e1} = R_{e2} = 357,94 \Omega$, $I_{EE} = 262,0 \mu\text{A}$, $C_1 = 71,55 \text{ pF}$, $R_{02} = 1,0 \Omega$, $R_{01} = 0,31907 \Omega$, $R_2 = 100,0 \text{ k}\Omega$, $G_b = 176,84 \text{ S}$, $E_C = 1$, $V_C = V_E = 5,6631 \text{ V}$, $H = 1,0 \text{ k}\Omega$, $R_{P1} = R_{P2} = 1,0 \text{ M}\Omega$, współczynniki wielomianowego źródła G_{P1} : $P_0 = 50 \text{ mA}$, $P_1 = -1,0 \text{ mS}$, współczynniki wielomianowego źródła G_{P2} : $P_0 = 50 \text{ mA}$, $P_1 = 1,0 \text{ mS}$. Obliczanie wartości elementów i parametrów makromodelu z dokładnością do sześciu lub pięciu cyfr znaczących jest konieczne aby uzyskać zadowalającą dokładność makromodelu.

Wyniki tych obliczeń w formacie dostosowanym do formatu modeli w bibliotekach podukładów liniowych programu PSpice zamieszczono w tablicy 1.

Tablica 1

Makromodel scalonego wzmacniacza mocy m.cz. TDA2030A w formacie bibliotek modeli podukładów liniowych programu PSpice

(14)

```

*
* TDA2030A audio power amplifier "macromodel" subcircuit
*
* connections:      non-inverting input
*                   |      inverting input
*                   |      |      negative power supply
*                   |      |      |      output
*                   |      |      |      |      positive power supply
*                   |      |      |      |      |
.SUBCKT TDA2030A    1    2    3    4    5
Q1    11    1    13    TDA2030AQ1
Q2    12    2    14    TDA2030AQ2
RC1    5    11    556.11
RC2    5    12    556.11
C1    11    12    71.550E-12
RE1    13    15    357.94
RE2    14    15    357.94
IEE    15    3    262.00E-6
GA    16    0    12    11    1.7982E-3
R2    16    0    100.0E3
C2    16    17    30.0E-12
GB    17    0    16    0    176.84
RO2    17    0    1.0
D1    17    18    TDA2030AD1
D2    18    17    TDA2030AD1
EC    18    0    4    0    1.0
RO1    17    21    0.31907
D3    21    19    TDA2030AD1
VC    5    19    5.6631
D4    20    21    TDA2030AD1
VE    20    3    5.6631
VP    21    4
H    22    0    VP    1.0E3
D5    22    23    TDA2030AD1
RP1    23    0    1E6
D6    24    22    TDA2030AD1
RP2    24    0    1E6
GP1    5    0    POLY(1) 23    0    .05    .001
GP2    0    3    POLY(1) 24    0    .05    -.001
.MODEL TDA2030AD1 D (IS=8.00E-16)
.MODEL TDA2030AQ1 NPN (IS=8.00000E-16 BF=622.857)
.MODEL TDA2030AQ2 NPN (IS=8.64353E-16 BF=688.421)
.ENDS

```

2.4. Weryfikacja makromodelu wzmacniacza TDA2030A

W celu zweryfikowania uzyskanego makromodelu wykonano za pomocą programu PSpice obliczenia jego parametrów elektrycznych i porównano ich wyniki z danymi katalogowymi wzmacniacza. Na rysunku 2a przedstawiono charakterystykę przejściową (liniowości) wzmacniacza z otwartą pętlą sprzężenia zwrotnego przy sterowaniu stałoprądowym. Odczytane z tej charakterystyki wzmocnienie napięciowe A_V wzmacniacza wynosi 31788, maksymalne napięcie wyjściowe, odpowiadające liniowej pracy wzmacniacza $V_{o\max}$ jest równe 13,265 V, a wejściowe napięcie niezrównoważenia — 2,008 mV. Wielkości te praktycznie nie różnią się od parametrów modelowanego wzmacniacza. Największy błąd wynosi 0,4%. Wejściowy prąd polaryzacji wynosi 200,0 nA, a wejściowy prąd niezrównoważenia — 20,0 nA. Ich wartości wyznaczono po skompensowaniu wejściowego napięcia niezrównoważenia za pomocą autonomicznego źródła napięciowego — 2,008 mV. Obie wielkości prądów są równe prądom modelowanego wzmacniacza.

Rysunek 2b przedstawia zależność modułu i argumentu transmitancji napięciowej wzmacniacza od częstotliwości. Analizę wykonano przy otwartej pętli sprzężenia zwrotnego po skompensowaniu wejściowego napięcia niezrównoważenia za pomocą źródła napięciowego o sile elektromotorycznej — 2,008 mV włączonego w szereg ze źródłem sygnału sterującego. Wzmocnienie przy najmniejszych częstotliwościach wynosi 31793 a częstotliwości dwóch najniższych leżących biegunów, odczytane z charakterystyki fazowej makromodelu, wynoszą odpowiednio 298,1 Hz (błąd — 0,7%) i 1,998 MHz (błąd — 0,1%).

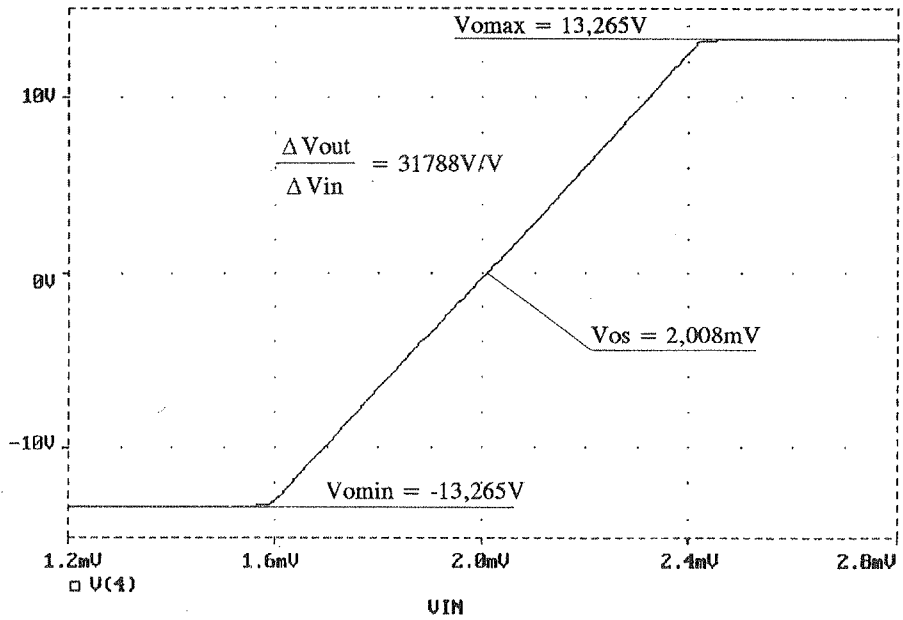
Ograniczniki napięcia i prądu wyjściowego badano przy zamkniętej pętli sprzężenia zwrotnego w układzie wzmacniacza o wzmocnieniu 26 dB. Działanie ograniczników ilustruje rys. 3. Rysunek 3a przedstawia przebiegi czasowe napięcia wyjściowego przy trzech wielkościach napięcia wejściowego odpowiadających pełnemuysterowaniu i przesterowaniu wzmacniacza. Napięcie wyjściowe jest ograniczane na poziomie $\pm 13,195$ V. Napięcie wyjściowe wyznaczone na podstawie maksymalnej mocy wyjściowej wzmacniacza powinno być równe $V_{o\max} = \sqrt{2P_{o\max}R_L} = 13,27$ V (błąd mniejszy, niż 0,6%).

Na rysunku 3b pokazano przebiegi czasowe prądu wyjściowego wzmacniacza pobudzonego sygnałem sinusoidalnym i obciążonego znamionową rezystancją obciążenia 4Ω , a następnie rezystorem 3Ω i 2Ω . Napięcie wejściowe odpowiada pełnemuysterowaniu, ale nie przesterowaniu wzmacniacza ($V_{in\max} = 0,65$ V). Ograniczanie prądu wyjściowego występuje przy wartości 3,5015 A (błąd 0,04%). Zgodność poziomów ograniczania napięcia wyjściowego i prądu wyjściowego wzmacniacza TDA2030A i jego makromodelu jest bardzo dobra.

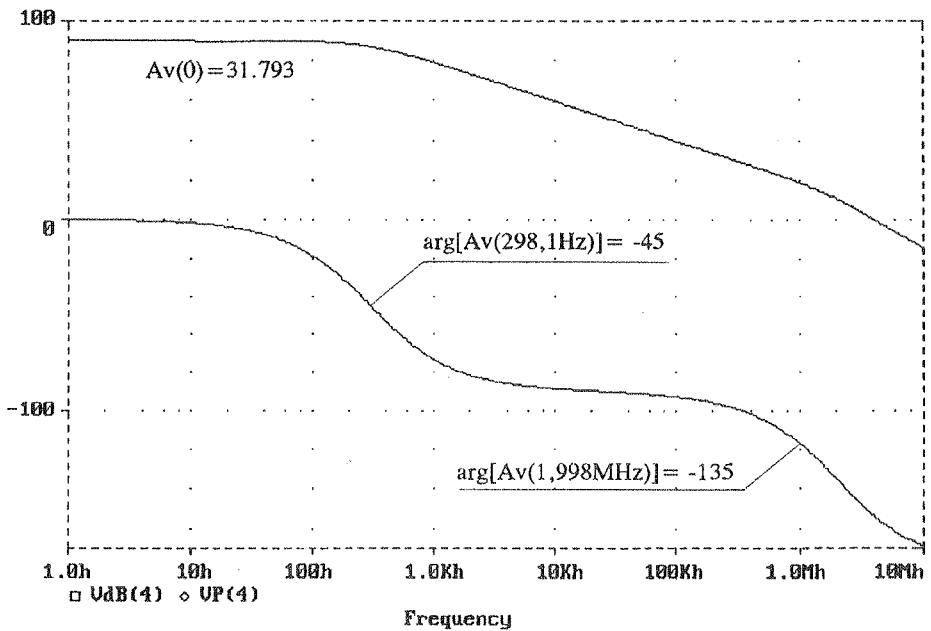
Współczynnik szybkości zmian napięcia wyjściowego badano również w układzie wzmacniacza z zamkniętą pętlą sprzężenia zwrotnego o wzmocnieniu 26dB, przy pobudzeniu falą prostokątną o amplitudzie 0,6 V odpowiadającej pełnemuysterowaniu (ale nie przesterowaniu) wzmacniacza. Na rysunku 4a są pokazane przebiegi czasowe pobudzenia i odpowiedzi wzmacniacza. Maksymalna szybkość

Rys. 2.
a) chara

a)

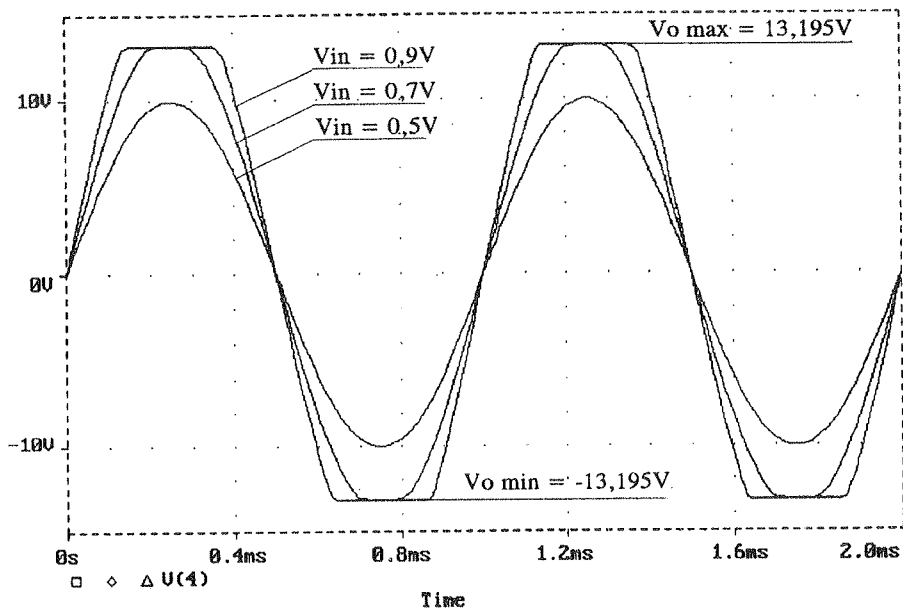


b)

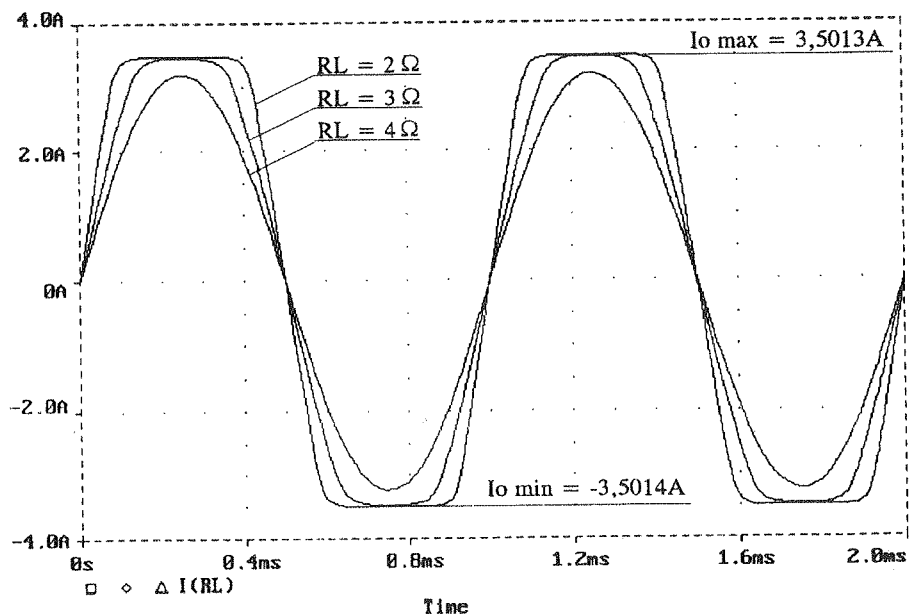


Rys. 2. Charakterystyki makromodelu wzmacniacza TDA2030A z otwartą pętlą sprzężenia:
a) charakterystyka przejściowa przy sterowaniu stałoprądowym, b) charakterystyki częstotliwościowe:
amplitudowa i fazowa

a)



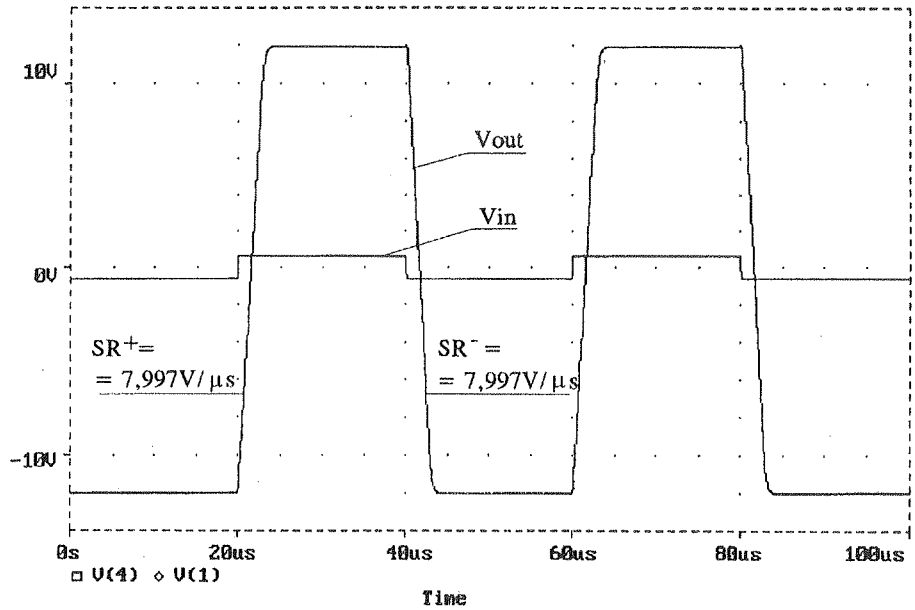
b)



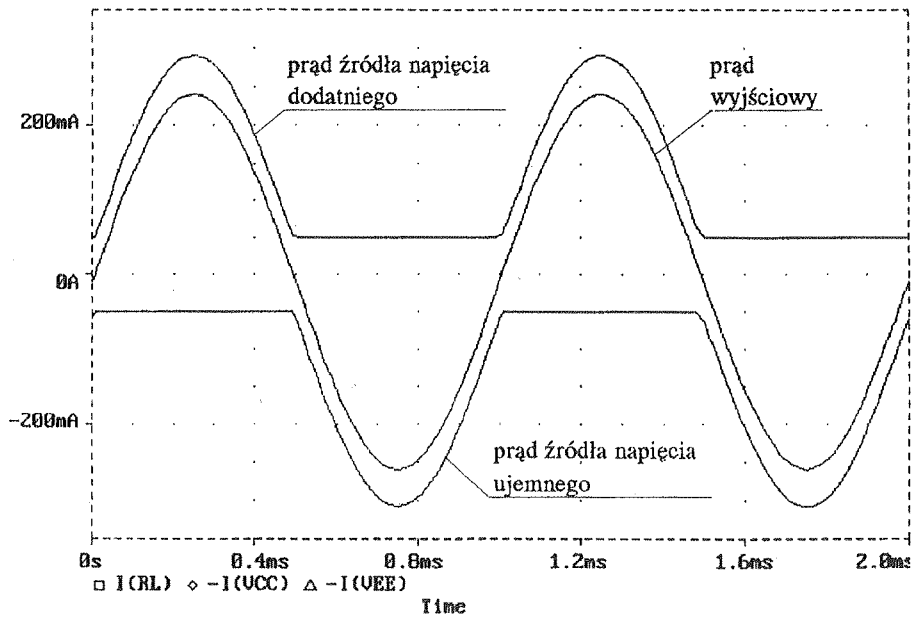
Rys. 3. Przebiegi czasowe napięcia (a) i prądu (b) wyjściowego makromodelu wzmacniacza TDA2030A (zamknięta pętla sprzężenia zwrotnego)

Rys.
wyjści

a)



b)



TDA2030A

Rys. 4. Przebiegi czasowe napięć i prądów w makromodelu wzmacniacza TDA2030A: a) napięcie wyjściowe i wejściowe — pobudzenie falą prostokątną, b) prąd wyjściowy i prądy pobierane ze źródeł zasilających

zmian napięcia wyjściowego jest jednakowa na zboczu narastającym i opadającym i wynosi $7,997 \text{ V}/\mu\text{s}$ (błąd mniejszy niż $0,04\%$).

Na rysunku 4b pokazano przebiegi czasowe prądów pobieranych ze źródeł zasilających (źródła napięcia dodatniego i ujemnego) oraz prądu wyjściowego. Ilustrują one działanie makromodelu i mogą służyć do obliczania mocy dostarczonej, mocy wyjściowej i mocy strat, oraz sprawności energetycznej i współczynnika zawartości harmonicznych w funkcji amplitudy napięcia wejściowego.

Wyniki wykonanych obliczeń potwierdzają uzyskanie bardzo dobrej zgodności właściwości makromodelu z parametrami i charakterystykami katalogowymi modelowanego układu scalonego. Różnice wielkości parametrów obliczonych dla makromodelu i parametrów katalogowych nie przekraczają jednego procenta. Można zatem uznać, że opracowany makromodel cechuje się wysoką dokładnością i nadaje się do większości zastosowań praktycznych.

2.5. Przykład zastosowania makromodelu wzmacniacza TDA2030A

Przydatność opracowanego makromodelu do komputerowej symulacji wzmacniaczy elektroakustycznych ilustruje przykład wzmacniacza o mocy 60 W , przeznaczonego do trójdrożnego aktywnego zespołu głośnikowego [7]. Schemat ideowy wzmacniacza jest pokazany na rys. 5. Każdy z torów zawiera aktywny filtr rozdzielający (zwrotnicę) i wzmacniacz mocy zbudowane na jednym układzie scalonym (tzw. „filtr aktywny mocy”). Zrealizowane są trzy filtry (dolnoprzepustowy, środkowoprzepustowy i górnoprzepustowy) drugiego rzędu o charakterystykach Butterwortha i częstotliwościach podziału 300 Hz i 3 kHz . Wzmocnienie napięciowe i moc torów średnio- i wysokotonowego są mniejsze od wzmocnienia i mocy toru niskotonowego ze względu na różnice w efektywności i impedancji głośników każdego z torów.

W torze niskotonowym wzmacniacza prąd pobierany przez obciążenie jest większy od maksymalnego prądu dostarczanego przez układ scalony. Dodatkowy prąd pochodzi z tranzystorów $Q1$ (dodatnia część) i $Q2$ (część ujemna). Tranzystory są sterowane spadkami napięć na rezystorach $1,5 \Omega$, które są wynikiem przepływu prądów zasilania układu scalonego. Ich rezystancja decyduje o podziale prądów oraz mocy strat pomiędzy układem scalonym i tranzystorami.

Analiza wzmacniacza z rys. 5 bez posłużenia się makromodelem układu scalonego jest właściwie niewykonalna, bo chociaż jest znany uproszczony schemat ideowy wzmacniacza mocy TAD2030A [7], to brak jest informacji o modelach występujących w tym schemacie tranzystorów. Użycie zaproponowanego w tej pracy makromodelu umożliwia wykonanie symulacji komputerowej wzmacniacza. Wyniki różnych analiz są przedstawione na rys. 6. Rysunek 6a przedstawia charakterystyki amplitudowe toru nisko-, średnio- i wysokotonowego wzmacniacza. Przebieg charakterystyk, częstotliwości podziału i wzmocnienia oraz moce wyjściowe torów są zgodne z [7].

wej.

Rys. 5. Scl

Przebieg
rys. 6b. I
podczas
około $0,6$
ra $Q1$ wy
z powod
prądowe
czynnika
i zamies
maksyma
napięcie

nadającym

ze źródeł
ściowego.
ia mocy
getycznej
ęcia wej-

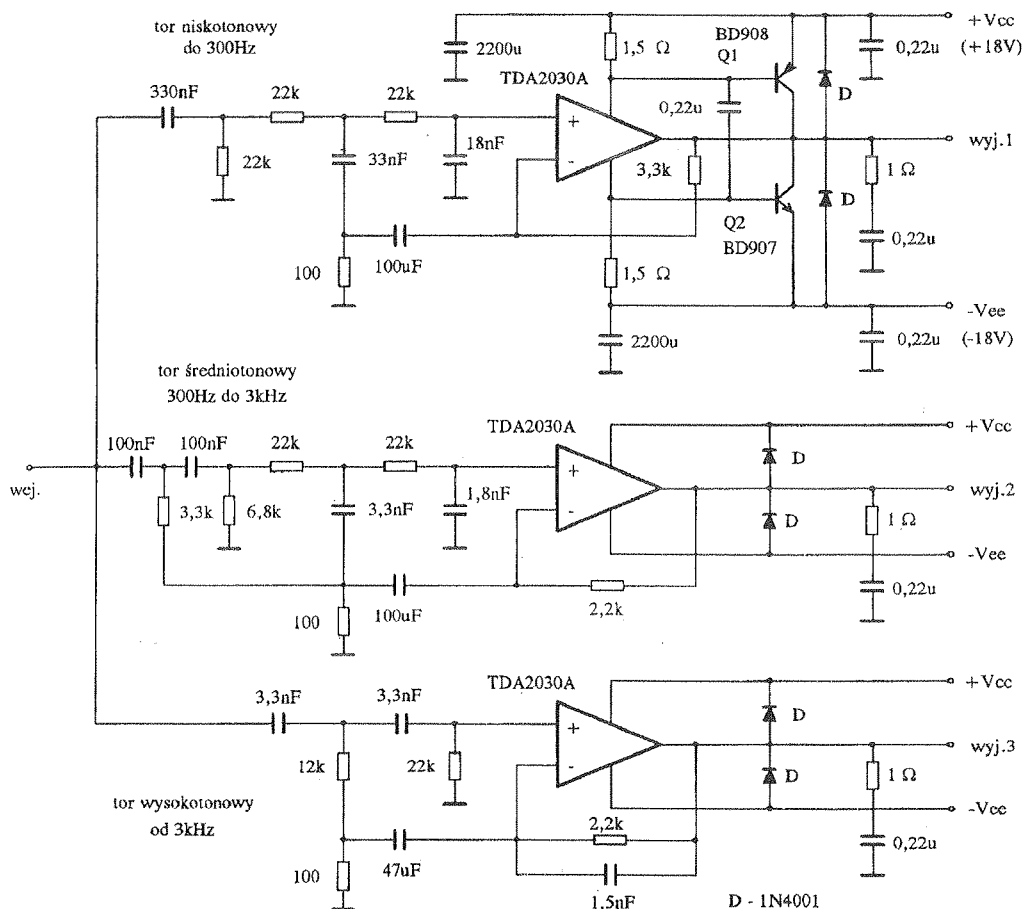
zgodności
mi mode-
la makro-
a. Można
i nadaje

A

ni wzmac-
W, prze-
at ideowy
filtr roz-
scalonym
wy, środ-
tach But-
apięciowe
mocy toru
łośników

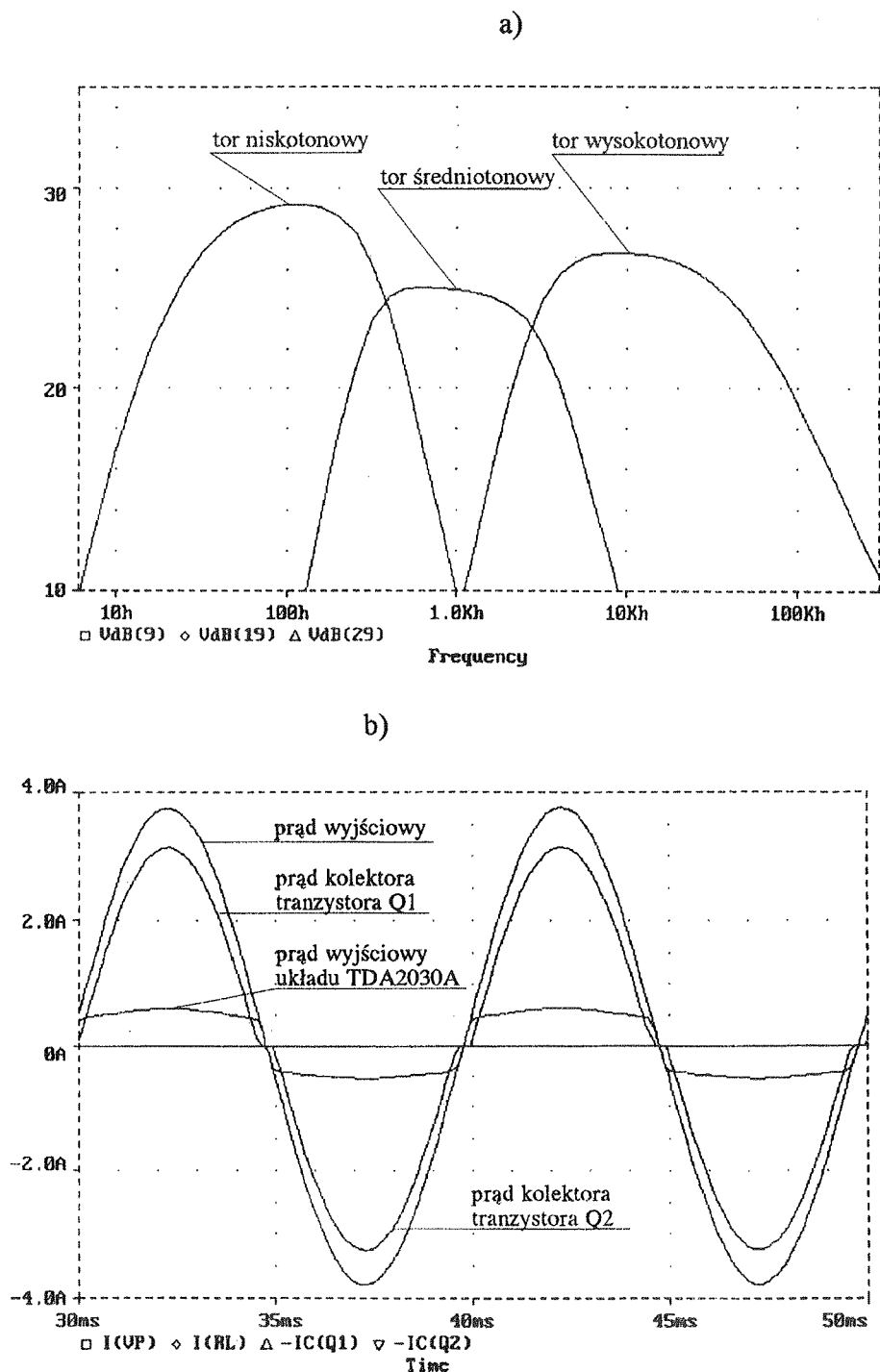
żenie jest
odatkowy
anzystory
przepływu
dów oraz

u scalone-
at ideowy
występują-
y makro-
wniki róż-
terystyki
g charak-
torów są



Rys. 5. Schemat ideowy elektroakustycznego wzmacniacza mocy do trójdrożnego aktywnego zespołu głośnikowego

Przebiegu czasowe prądów w torze niskotonowym wzmacniacza pokazano na rys. 6b. Przy napięciu wejściowym o amplitudzie 0,45 V prąd wyjściowy osiąga 3,8 A, podczas gdy maksymalna wielkość prądu wyjściowego układu scalonego wynosi około 0,6 A, a minimalna – 0,54 A. Maksymalna wielkość prądu kolektora tranzystora Q1 wynosi 3,2 A, a tranzystora Q2 – 3,25 A. Prądy tranzystorów nie są jednakowe z powodu różnic w charakterystykach wejściowych i współczynnikach wzmocnienia prądowego tranzystorów *n-p-n* i *p-n-p*. Zgodność wielkości mocy wyjściowej i współczynnika zawartości harmonicznych, uzyskanych w wyniku symulacji komputerowej i zamieszczonych w [7], uzyskuje się po zamodelowaniu układu scalonego przy maksymalnym prądzie wyjściowym 0,6 A. Od prądu wyjściowego najsilniej zależy napięcie źródeł V_C i V_E . Przy prądzie 0,6 A napięcia tych źródeł wynoszą 2,8 V.



Rys. 6. Wyniki analizy wzmacniacza do tródrożnego aktywnego zespołu głośnikowego: a) charakterystyka amplitudowa, b) przebiegi czasowe prądów

Przebiegi mocy do torach, o w funkcji sygnałów $P_{wy} = R I_{wy}^2$ jako $P_d = R I_d^2$ funkcje R pobudzenia do obciążenia nieliniowego. Sprawność Wyniki wskazują na modelu zespołu g

Makro wzmacniacza, zawierającego nieznaczny częściowy sprzężenie z zamknięciem części

Makro wzmacniacza, budowy go oraz wzmacniacza pierwiastka obwodu impedancji wzmacniacza

Współczynnik jest p technologicznych. W przypadku przeprowadzenia

Przebiegi czasowe prądów ilustrują działanie wzmacniacza i służą do obliczania mocy dostarczonej, mocy wyjściowej i mocy strat w układzie scalonym i tranzystorach, oraz sprawności energetycznej i współczynnika zawartości harmonicznych w funkcji napięcia wejściowego. Ze względu na zniekształcenia nieliniowe sygnałów przy największych mocach, moc wyjściowa była obliczana jako $P_{wy} = RMS[V(OUT)] \cdot RMS[I(RL)]$ zaś moc dostarczana ze źródeł zasilających, jako $P_d = 18 \cdot AVG[I(V_{EE}) - I(V_{CC})]$, a moc strat, jako ich różnica. Aby uśrednić funkcje RMS i AVG przebiegi napięć i prądów były całkowane w kilkunastu okresach pobudzenia. Z wykonanych obliczeń wynika, że wzmacniacz jest w stanie dostarczyć do obciążenia o rezystancji 4Ω moc 25 W przy bardzo małych zniekształceniach nieliniowych i 30 W przy współczynniku zawartości harmonicznych $h \approx 0,5\%$. Sprawność energetyczna przyjmuje wielkości typowe dla wzmacniacza mocy klasy B.

Wyniki analiz wykonanych z zastosowaniem przedstawionego makromodelu wskazują na pełną zgodność parametrów obliczonych przy użyciu tego makromodelu z parametrami elektroakustycznego wzmacniacza mocy do trójdrożnego zespołu głośnikowego zamieszczonymi w [7].

3. MAKROMODEL WZMACNIACZY MOCY M.CZ. O MNIEJSZEJ ZŁOŻONOŚCI

Makromodel o mniejszym stopniu złożoności jest przeznaczony do modelowania wzmacniaczy mocy m.cz. o zasilaniu niesymetrycznym. Wzmacniacze takie mają mniejszą moc wyjściową ($0,5\text{ W}$ do 10 W), a ich dane katalogowe zwykle nie zawierają informacji o wejściowym prądzie polaryzacji, wejściowym napięciu i prądzie niezrównoważenia, współczynniku maksymalnej szybkości zmian napięcia wyjściowego SR i położeniu biegunów transmitancji napięciowej z otwartą pętlą sprzężenia zwrotnego (jest podawana górna częstotliwość graniczna wzmacniacza z zamkniętą pętlą sprzężenia w zależności od zewnętrznego kondensatora kompensacji częstotliwościowej).

Makromodel o mniejszej złożoności stanowi uproszczenie pełnego makromodelu wzmacniacza o zasilaniu symetrycznym przedstawionego w p. 2. Pomimo prostej budowy modeluje on z zadowalającą dokładnością pobór mocy ze źródła zasilającego oraz uwzględnia następujące właściwości wzmacniacza: rezystancję wejściową, wzmocnienie napięciowe przy otwartej pętli ujemnego sprzężenia zwrotnego, położenia pierwszego bieguna transmitancji przy otwartej pętli, wpływ kondensatora obwodu kompensacji na charakterystyki częstotliwościowe wzmacniacza, zależność impedancji wyjściowej od częstotliwości oraz ograniczanie napięcia wyjściowego wzmacniacza.

Wspomagana komputerem analiza pełnego schematu ideowego układu scalonego jest praktycznie niewykonalna z powodu niedostępności odpowiednich danych technologicznych, potrzebnych do modelowania elementów półprzewodnikowych. W przypadku niektórych analogowych układów scalonych analiza taka została przeprowadzona, a wyniki opublikowane (np. analiza wzmacniacza operacyjnego

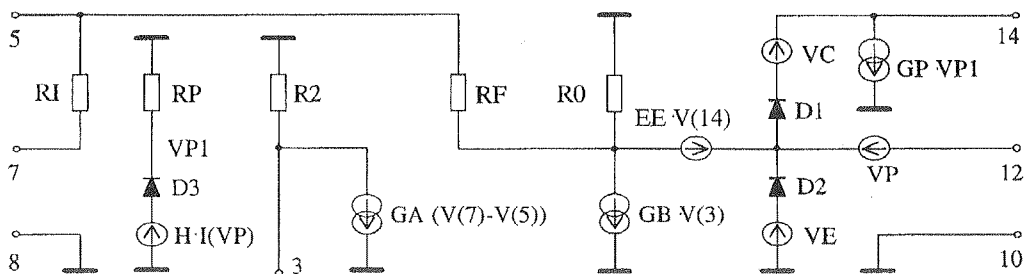
$\mu A741$). W [5] jest zamieszczony przykład analizy scalonego wzmacniacza mocy m.cz. UL1495N. Mimo pewnych niedostatków modelowania tranzystorów (pominięcie części dynamicznej modeli) wyniki analiz wzmacniacza z tym układem, podobnie jak jego dane katalogowe, mogą służyć do weryfikacji opracowanego makromodelu. Dlatego jako przykład scalonego wzmacniacza mocy m.cz. o zasilaniu niesymetrycznym, którego makromodel zidentyfikowano i zweryfikowano, został wybrany układ UL1495N.

Przyjęty w [5] punkt startowy analizy stałoprądowej wzmacniacza z układem scalonym UL1495N jest odległy od końcowego punktu pracy układu. Próby zmiany tego punktu startowego prowadziły do braku zbieżności procedur numerycznych programu PSpice. Widać stąd, że przy symulacji na poziomie elementów układu zawierającego nawet tylko jeden scalony wzmacniacz mocy m.cz. problemy ze zbieżnością procedur numerycznych symulatorów układów elektronicznych mogą wystąpić. Symulacja wzmacniacza, w którym układy scalone są zastąpione makromodelami jest wolna od takich trudności, nawet gdy w skład wzmacniacza wchodzi kilka układów scalonych sprzężonych galwanicznie (np. wzmacniacz ACE-bass [8]).

3.1. Opis makromodelu

Struktura makromodelu scalonego wzmacniacza mocy m.cz. o mniejszym stopniu złożoności jest pokazana na rys. 7. W schemacie tym stopień wejściowy został pominięty w stosunku do modelu pełnego, jedynie rezystor R_1 ustala rezystancję wejściową wzmacniacza. Stopień pośredni, składający się z elementów $G_a(V(7) - V(5))$ i R_2 , ma jednostkowe wzmocnienie napięciowe i uwzględnia położenie bieguna transmitancji napięciowej wzmacniacza z otwartą pętlą sprzężenia zwrotnego (wraz ze stopniem wyjściowym i kondensatorem C_F). Kondensator C_F dołączony na zewnątrz makromodelu (między końcówki 3 i 12), umożliwia kompensację charakterystyk częstotliwościowych wzmacniacza w taki sam sposób, jak w rzeczywistym wzmacniaczu.

Stopień wyjściowy modeluje zależność impedancji wyjściowej wzmacniacza od częstotliwości. Rezystancja wyjściowa przy bardzo małych częstotliwościach jest



Rys. 7. Struktura makromodelu wzmacniacza mocy m.cz. o zasilaniu niesymetrycznym

równa re-
częstotliw-
stopnia
 $E_E V(14)$
zacisku
niacza.

Pobó-
źródła p-
czynkow-
prądu w-
 D_3 i rezy-
(próbko-
podobnie
 H i rezy-
na przew-
poboru
spoczynk-
energety-
niacza

Wart-
wzmacn-
paramet-
moc wy-
napięciu
pięciowe
jest zna-
wzmacn-
graniczn-
densator
niacza
tłumieni-
rezystor
scaloneg

Na-
poniżej
modelu.
niaczy
danych
jest dog

acza mocy
ów (pomi-
układem,
cowanego
i. o zasil-
wano, zo-

układem
y zmiany
erycznych
w układu
blemy ze
ych mogą
astąpione
macniacza
macniacz

zym stop-
wy został
ezystancję
elementów
względnie
sprężenia
asator C_F
kompen-
osób, jak

niacza od
ciach jest

14
P VP1

12

10

ym

równa rezystancji rezystora R_o . Kondensator C_F powoduje, że w miarę wzrostu częstotliwości rezystancja R_o jest bocznikowana przez pojemność wyjściową stopnia pośredniego i impedancja wyjściowa wzmacniacza maleje. Elementy $E_E V(14)$, V_C , V_E , D_1 i D_2 zapewniają właściwą wartość napięcia stałego na zacisku wyjściowym i ograniczanie maksymalnego napięcia wyjściowego wzmacniacza.

Pobór mocy ze źródła zasilającego jest modelowany za pomocą wielomianowego źródła prądowego G_p . Stały współczynnik wielomianu odpowiada prądowi spoczynkowemu wzmacniacza. Źródło G_p jest sterowane dodatnią połówką przebiegu prądu wyjściowego za pośrednictwem pomocniczego obwodu złożonego z diody D_3 i rezystora R_p oraz ze źródła napięciowego H sterowanego prądem wyjściowym (próbkowanym przez źródło V_p). Dioda D_3 jest idealną diodą półprzewodnikową podobnie jak D_1 i D_2 . Transkonduktancja źródła G_p , transrezystancja źródła H i rezystancja rezystora R_p zostały dobrane w taki sposób, żeby spadek napięcia na przewodzącej diodzie D_3 można było zaniedbać. Przyjęty sposób modelowania poboru mocy ze źródła zasilającego gwarantuje prawidłowe wielkości prądów: spoczynkowego i zasilającego, a także mocy zasilania, mocy strat i sprawności energetycznej w całym zakresie napięć sterujących i impedancji obciążenia wzmacniacza oraz napięć zasilania.

3.2. Identyfikacja makromodelu

Wartości elementów i parametrów uproszczonego makromodelu scalonego wzmacniacza mocy m.cz. wyznaczane są kolejno, niezależnie od siebie na podstawie parametrów katalogowych wzmacniacza. Są to następujące parametry: maksymalna moc wyjściowa ($P_{o,max}$) przy znamionowej rezystancji obciążenia (R_L) oraz napięciu zasilającym (V_{CC}), spoczynkowy prąd zasilania (I_d), wzmocnienie napięciowe wzmacniacza (A_V) z otwartą pętlą sprzężenia zwrotnego (jeśli nie jest znane, można założyć wartość $1 \cdot 10^4$ do $3 \cdot 10^4$), wzmocnienie napięciowe wzmacniacza z zamkniętą pętlą sprzężenia zwrotnego (A_{VCL}), górna częstotliwość graniczna (f_g) wzmacniacza z zamkniętą pętlą sprzężenia i zewnętrznym kondensatorem kompensacji częstotliwościowej (C_F), rezystancja wyjściowa wzmacniacza (R_{oac}) przy zamkniętej pętli sprzężenia zwrotnego lub współczynnik tłumienia $K = R_L/R_{oac}$, rezystancja wejściowa wzmacniacza (R_I), rezystancja rezystora pętli sprzężenia zwrotnego (R_F) — zwykle wbudowanego do układu scalonego.

Na podstawie tych parametrów układu scalonego i zależności zestawionych poniżej wyznaczane są wartości poszczególnych elementów i parametrów makromodelu. W zależnościach tych została uwzględniona specyfika scalonych wzmacniaczy mocy m.cz. o zasilaniu niesymetrycznym (wielkości prądów, dostępność danych katalogowych, itp.). Kolejność obliczeń i forma poszczególnych zależności jest dogodna przy identyfikacji modeli takich wzmacniaczy.

wejściowej mocnienia wartości od-
f_g. Można
(15)
Rezystor
R_{0ac}:
(16)
współczyn-
owej, jako
ne:
(17)
przewodzi.
metryczne,
(18)
ówny 0,5,
cja źródła
k D₁ i D₂
S.
el scalone-
owych są
10⁴ (osza-
V_{CL} = 200
logowych
o układu
rametrów
nia rezys-
y wielkość
rezystan-

cji R₀ jest wartość 57,8 Ω uzyskana w wyniku symulacji pełnego układu wzmacniacza UL1495N. Tę wielkość przyjmujemy do dalszych obliczeń. Konduktancja źródła sterowanego G_b jest równa 840S, I_{sD1} = I_{sD2} = I_{sD3} = 8 · 10⁻¹⁶ A, V_C = V_E = 60,93 V (V_{o max} = 4,4 V), E_E = 0,5, H = 1,0 kΩ, R_p = 1,0 MΩ, współczynnik źródła G_p: P₀ = 6mA, P₁ = 2,0mS.

Wyniki tych obliczeń w formacie zgodnym z formatem modeli w bibliotekach podukładów liniowych programu PSpice zamieszczono w tablicy 2.

Tablica 2

Makromodel scalonego wzmacniacza mocy m.cz. UL1495N w formacie bibliotek modeli podukładów liniowych programu PSpice

```
*
* UL1495N audio power amplifier "macromodel" subcircuit
*
* connections: non-inverting input
*                  |      inverting input
*                  |      frequency compensation
*                  |      output ground
*                  |      positive power supply
*                  |      input ground
*                  |      output
*
.SUBCKT UL1495N      7      5      3      10      14      8      12
RI      5      7      1.4E6
GA      3      0      7      5      1.28E-3
R2      3      0      781
RF      21      5      7.7E3
GB      21      0      3      0      840
RO      21      0      57.8
EE      28      21      14      0      0.5
D1      28      22      UL1495ND
VC      14      22      0.93
D2      24      28      UL1495ND
VE      24      0      0.93
VP      28      12
H      26      0      VP      1.0E3
D3      26      27      UL1495ND
RP      27      0      1.0E6
GP      14      0      POLY(1) 27      0      .006 .002
.MODEL UL1495ND D      (IS=8.00E-16)
.ENDS
*
```

3.4. Weryfikacja i przykład zastosowania makromodelu wzmacniacza UL1495N

W celu zweryfikowania uzyskanego makromodelu wykonano obliczenia jego parametrów elektrycznych i porównano ich wyniki z danymi katalogowymi układu scalonego UL1495N i z wynikami symulacji komputerowej wzmacniacza mocy zawierającego pełny schemat ideowy układu scalonego [5]. Obliczenia wykonano za pomocą programu PSpice, a ich wyniki przedstawiono na rys. 8 i 9.

Na rysunku 8a pokazano zależność napięcia wyjściowego od wejściowego makromodelu z otwartą pętlą sprzężenia zwrotnego przy wysterowaniu stałoprądowym (zamiast kondensatora sprzęgającego z obciążeniem zastosowano źródło napięciowe 4,5 V eliminujące składową stałą). Jak wynika z rysunku wzmocnienie napięciowe wynosi 9986 V/V, a maksymalne międzyszczytowe napięcie wyjściowe — około 8,8 V (ograniczanie napięcia wyjściowego na poziomie 0,1 V i 8,9 V).

Rysunek 8b przedstawia charakterystyki amplitudowe makromodelu z otwartą i zamkniętą pętlą ujemnego sprzężenia zwrotnego. Parametrami krzywych są pojemności kondensatora kompensacji częstotliwościowej. Wzmocnienie przy małych częstotliwościach wynosi około 80 dB (pętla otwarta), lub 46 dB (pętla zamknięta), a częstotliwości graniczne przy zamkniętej pętli są zgodne z [6]. Dla porównania wykonano symulację komputerową pełnego układu wzmacniacza UL1495N. Różnice pomiędzy charakterystykami amplitudowymi makromodelu i wzmacniacza przy zamkniętej pętli sprzężenia zwrotnego nie przekraczają 1 dB.

Impedancja wejściowa makromodelu jest zgodna z przyjętą wartością i wynosi 1,4 M Ω , a moduł impedancji wejściowej obliczonej dla pełnego schematu wzmacniacza UL1495N jest zależny od częstotliwości i nieco mniejszy (1,35 M Ω —1,38 M Ω). Na rysunku 9a pokazano zależność modułu impedancji wyjściowej obliczoną dla makromodelu przy otwartej pętli sprzężenia zwrotnego z różnymi pojemnościami kondensatora kompensacyjnego. Jest widoczna silna zależność impedancji wyjściowej od pojemności tego kondensatora. Dla porównania obliczono impedancje wyjściowe pełnego układu wzmacniacza UL1495N. Różnice pomiędzy odpowiednimi modułami impedancji wyjściowych makromodelu i wzmacniacza nie przekraczają 20%.

Działanie ograniczników napięcia wyjściowego ilustruje rys. 9b. Na rysunku jest pokazana zależność napięć wyjściowych makromodelu od czasu przy trzech wartościach napięcia wejściowego w warunkach pełnego wysterowania i przesterowania wzmacniacza. Napięcie wyjściowe jest ograniczane na poziomie około 4,4 V, co odpowiada katalogowej wielkości maksymalnej mocy wyjściowej 0,65 W. Podczas symulacji pełnego schematu wzmacniacza uzyskano takie same wyniki.

Przedstawione wyniki potwierdzają uzyskanie bardzo dobrej zgodności właściwości makromodelu z charakterystykami i parametrami katalogowymi scalonego wzmacniacza mocy oraz z wynikami symulacji komputerowej pełnego schematu ideowego układu scalonego UL1495N. Dokładność zaproponowanego makromodelu można uznać jako zadowalającą w typowych zastosowaniach praktycznych.

UL1495N

zenia jego
mi układu
cza mocy
konano za

ejściowego
ałoprądo-
dło napię-
mocnienie
wyjściowe
9 V).

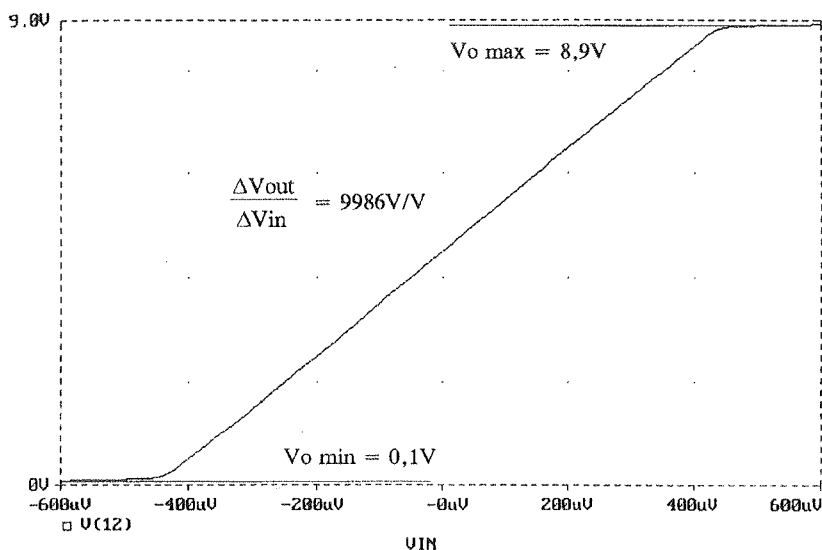
z otwartą
są pojem-
y małych
umknięta),
ównania
5N. Róż-
iacza przy

i wynosi
schematu
mniejszy
ancji wyj-
ego z róż-
zależność
nania ob-
Różnice
i wzmac-

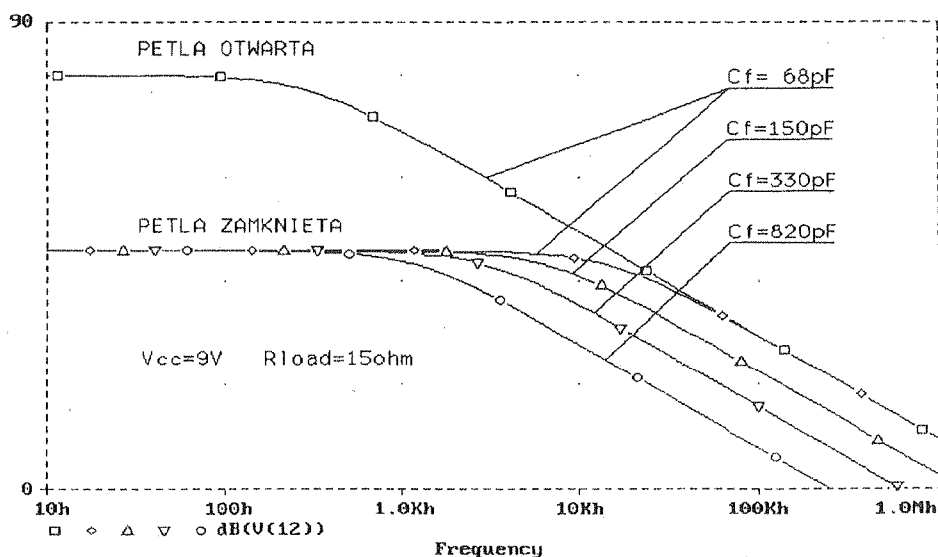
sunku jest
zech war-
terowania
4,4 V, co
Podczas

ości właś-
scalonego
schematu
makromode-
nych.

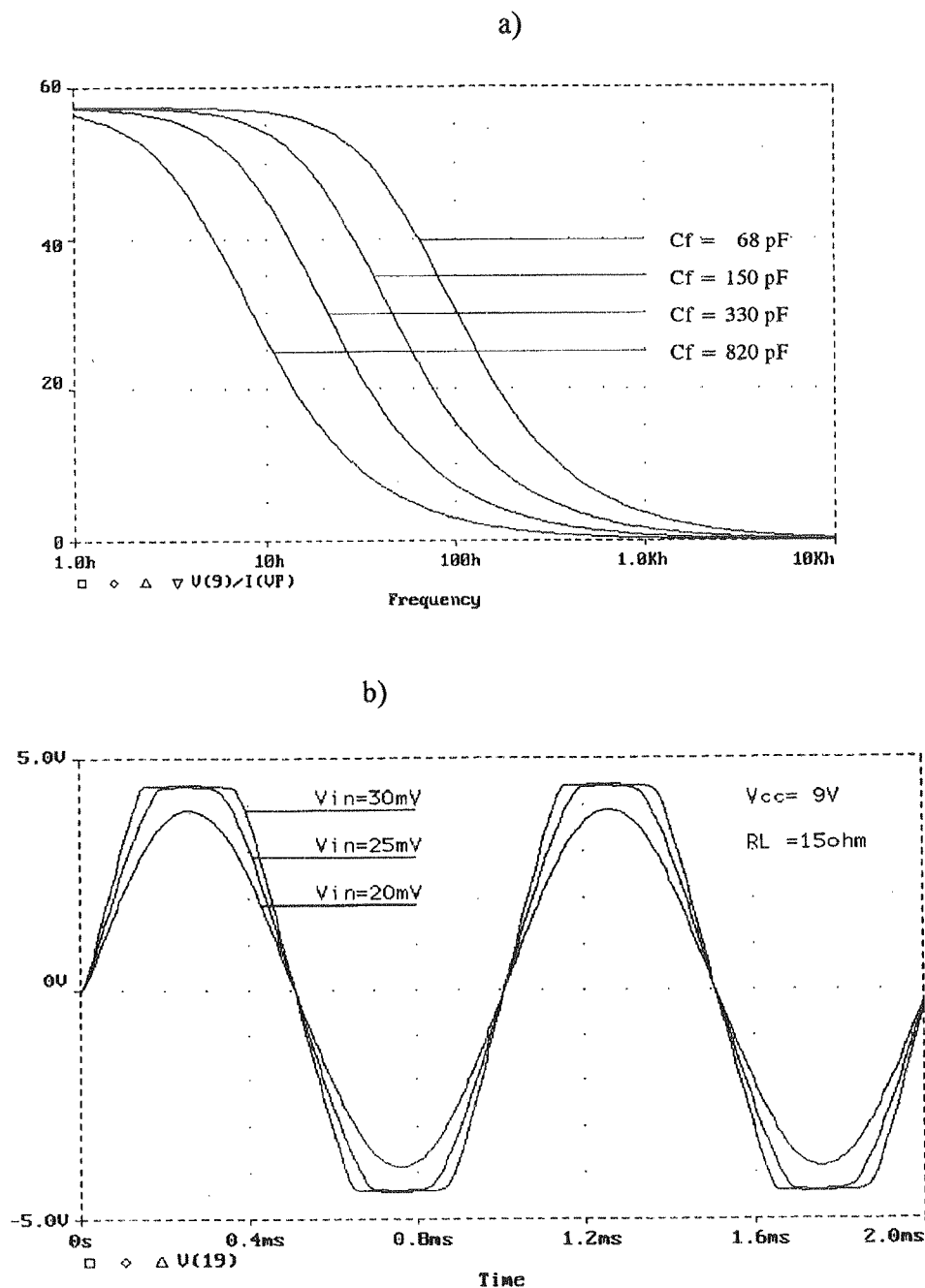
a)



b)



Rys. 8. Charakterystyki wyznaczone dla makromodelu wzmacniacza UL1495N: a) stałoprądowa charakterystyka przejściowa (otwarta pętla), b) charakterystyki amplitudowe (pętla otwarta i zamknięta)



Rys. 9. Charakterystyki obliczone dla makromodelu wzmacniacza UL1495N: a) zależność modułu impedancji wyjściowej od częstotliwości, b) przebiegi czasowe napięcia wyjściowego

Mak
znacznie
liczby w
o większ
uproszcz
symulacji
(makrom
zasadnic
symulato

Zapr
modeli u
na podsta
uzupełni
Každy p
Struktur
scalony

Mak
a dzięki
nego, ma
podukła
ściowych
mentów
niego p
m.cz. i c

Nap
zgodne
modelow
przebieg
ność mo
cji obci
komput
ści do
optymal
prostym
wzmocn
stosowa
i wzmoc
żoność

4. WNIOSKI I UWAGI KOŃCOWE

Makromodele wzmacniaczy mocy m.cz., przedstawione w tej pracy, cechują się znacznie mniejszym stopniem złożoności niż oryginalne wzmacniacze. Redukcja liczby węzłów i gałęzi jest co najmniej pięciokrotna w przypadku makromodelu o większej dokładności i co najmniej dziesięciokrotna w przypadku makromodelu uproszczonego. Dzięki temu jest możliwe uzyskanie wielokrotnie krótszego czasu symulacji i obniżenie jej kosztów. Makromodel zawiera zaledwie 10 złączy $p-n$ (makromodel o większej dokładności) lub 3 złącza (makromodel uproszczony) co zasadniczo zmniejsza problemy związane ze zbieżnością procedur numerycznych symulatorów układów elektronicznych.

Zaproponowane struktury i procedury wyznaczania parametrów i elementów modeli umożliwiają identyfikację makromodeli scalonych wzmacniaczy mocy m.cz. na podstawie ich danych katalogowych. Ewentualne niekompletne dane jest łatwo uzupełnić wynikami prostych pomiarów lub oszacowaniem niektórych parametrów. Każdy parametr wzmacniacza jest modelowany kolejno, niezależnie od pozostałych. Struktury modeli i metodyka identyfikacji ich parametrów uwzględnia specyfikę scalonych wzmacniaczy mocy m.cz.

Makromodele są zbudowane z elementów akceptowanych przez program PSpice, a dzięki zachowaniu struktury podobnej do makromodelu wzmacniacza operacyjnego, makromodele scalonych wzmacniaczy mocy mogą być dołączane do bibliotek podukładów liniowych tego programu lub bezpośrednio do plików danych wejściowych, przygotowanych w formie listy połączeń. Algorytmy wyznaczania elementów i parametrów makromodeli mogą być podstawą do napisania odpowiedniego programu komputerowego identyfikacji makromodeli wzmacniaczy mocy m.cz. i dołączenia go do segmentu PARTS pakietu MicroSim Design Center.

Napięcia i prądy stałe i zmienne końcówek prezentowanych makromodeli są zgodne z napięciami i prądami stałymi i zmiennymi odpowiednich końcówek modelowanych układów scalonych. Oba makromodele zapewniają prawidłowe przebiegi czasowe prądów pobieranych ze źródeł zasilających oraz modelują zależność mocy wyjściowej, mocy zasilania i mocy strat od napięcia sterującego, rezystancji obciążenia wzmacniacza i napięć zasilających. Przytoczone wyniki symulacji komputerowej świadczą o zadowalającej dokładności makromodeli i ich przydatności do typowych zastosowań praktycznych. Symulacja komputerowa umożliwia optymalizację makromodeli, prowadzącą do zwiększenia ich dokładności, dzięki prostym związkom między parametrami makromodeli i danymi katalogowymi wzmacniacza. Makromodele i metodyka identyfikacji ich parametrów może być stosowana równie dobrze do modelowania scalonych wzmacniaczy mocy m.cz., jak i wzmacniaczy zbudowanych z elementów dyskretnych. Rodzaj tranzystorów i złożoność struktury nie ma wpływu na przebieg i wynik modelowania.



BIBLIOGRAFIA

1. G.R. Boyle, B.M. Cohn, D.O. Pederson, J.E. Solomon: *Macromodeling of Integrated Circuit Operational Amplifiers*. IEEE Journal of Solid State Circuits, 1974, vol. SC-9, p. 353
2. A. Filipkowski (Ed.): *Computer-Aided Design and Computer-Aided Engineering in Electronic Engineering Education*. Wyd. PW, Warszawa 1996, p. 41–68
3. M. Głowacki, M. Pierzchała, J. Stanclik: *Dynamiczne zniekształcenia intermodulacyjne we wzmacniaczach mocy m.cz.* Kwartalnik Elektroniki i Telekomunikacji, 1997, Vol. 43, z. 4, s. 551
4. MicroSim PSpice A/D. Ver. 6.3. *Reference Manual*. MicroSim Corporation, April 1996
5. J. Porębski, P. Korohoda: *SPICE program analizy nieliniowej układów elektronicznych*. WNT, Warszawa 1996, Wyd. V, s. 214–217
6. C. Rudnicki: *Układy scalone w sprzęcie elektroakustycznym*. Wyd. SIGMA, Warszawa 1987, s. 69–79
7. SGS-THOMPSON Microelectronics, TDA2030A. March 1995, p. 1/15–15/15 (<http://www.st.com>)
8. J. Stanclik: *Macromodeling of Audio Power Amplifiers*. 95th Convention of the Audio Engineering Society, New York (USA), Oct. 7–10 1993, paper 3768 (B4-AM-5)
9. J. Stanclik: *Makromodel akustycznych wzmacniaczy mocy*. XLIII Otwarte Seminarium z Akustyki OSA'96, Ustroń–Zawodzie, 17–20 wrzesień 1996, Materiały, tom 2, s. 611–616
10. J. Stanclik: *Prosty makromodel akustycznych wzmacniaczy mocy o zasilaniu niesymetrycznym*. XLIII Otwarte Seminarium z Akustyki OSA'96, Ustroń–Zawodzie, 17–20 wrzesień 1996, Materiały, tom 2, s. 617–622
11. J. Stanclik: *Modelowanie elementów na potrzeby wspomaganego komputerem projektowania układów elektronicznych*. Prace Naukowe Inst. Telekom. i Akustyki Polit. Wrocław nr 41, Seria Monografie nr 20, Wyd. PWr, Wrocław 1982

J. STANCLIK

MACROMODELS OF INTEGRATED CIRCUIT AUDIO POWER AMPLIFIERS

Summary

This paper deals with two macromodels of integrated circuit audio power amplifiers intended for PSpice computer program (the more accurate and complex model and the less complex and accurate model). The both macromodels are much less complex than the original amplifiers but accurate enough for typical applications. The more accurate model, intended for amplifiers supplied from two symmetrical supply sources, is derived from [8]. It ensures the proper values of currents taken from the supply sources, the supplied power, output power and dissipated power as well as the power efficiency and total harmonic distortion factor depending on the input voltage, load impedance and supply voltages. The less complex macromodel refers to amplifiers supplied from one supply source [10]. The macromodel structures, identification algorithms and identification examples of typical integrated circuit audio power amplifiers are also presented in the paper. Macromodel parameters are evaluated from typical data sheet. Performance characteristics of the macromodels have been verified by their analysis and the comparison of the results with the data sheets. Finally the macromodel applications to the computer simulation of the audio power amplifiers are also shown.

Key words: macromodel, audio power amplifier, PSpice computer program.

The generation of learning data points for circuit performance modelling by means of fuzzy system

RYSZARD HORBOWSKI, MICHAŁ BIAŁKO

*Wydział Elektroniki, Informatyki i Telekomunikacji Zakład Układów Elektronicznych
Politechnika Gdańska*

Received 1998.07.03

Authorised 1998.09.23

In this paper the application of fuzzy system to generation of learning data points for electronic circuit performance modelling is presented. Under some assumptions the proposed approach helps to generate the learning data point base that possesses two advantages: it is free of superfluous redundancy and sufficiently represents behaviour of the circuit. Using simple examples we show, how properly chosen learning points may influence model errors and structure.

Key words: fuzzy and neurofuzzy modelling, fuzzy logic, electronic circuit modelling, knowledge engineering, expert systems.

1. INTRODUCTION

Modelling of electronic circuits performance enables us to predict the behaviour of an investigated circuit without its analysis on the level of individual elements. Applications of such models have been presented in some papers [1, 2] significantly speeding up a process of circuit optimisation.

The input-output data pair set can be obtained in two ways: on the basis of measurements of a circuit or on the basis of simulation. In the first case it is usually not possible to obtain any additional data pairs, whereas in the latter one we can generate as many additional points as we need.

An appropriate selection of these points is essential in circuit modelling, and in the remaining part of this paper we will consider only points, which were generated by the circuit simulator. In this paper we will consider two types of circuit parameters: input (independent) parameters, called as circuit variables, and output (dependent) parameters called as circuit functions. In general, finding appropriate

model structure from data base points is iterative and time consuming task. The time of the model generation and its accuracy are heavily dependent on a quality of learning data points. Thus, the selection of these points should be carried out as skillful as possible. The large and redundant data base causes long time of a model structure identification, whereas not sufficiently covered input space of the circuit variables can cause large model errors.

The simplest technique of generating data points is randomly sampling of the input space of the circuit. If each of the circuit variables (e.g. resistances, capacitances, currents, voltages, technology parameters) x_i belongs to the interval $[L_i, U_i]$ then the input space X of the circuit is defined by the Cartesian product: $X = [L_1, U_1] \times [L_2, U_2] \times \dots \times [L_N, U_N]$, where N is the number of circuit variables. In this method, we randomly choose the base points from the input space: $P'_j = (p_1^{(j)}, p_2^{(j)}, \dots, p_N^{(j)})$, where $p_i \in [L_i, U_i]$. For each P'_j the circuit simulator determines the values of the selected circuit function $y^{(j)}$ (e.g. voltage gain, cut-off frequency, power dissipation, SR, CMRR, input and output resistances, etc.). Finally, we obtain the learning data point set P consisting of learning data points $P_j = (p_1^{(j)}, p_2^{(j)}, \dots, p_N^{(j)}, y^{(j)})$. Despite its simplicity, randomly sampling of the input space has one disadvantage: it tends to produce clusters and therefore is not well suited for the model construction.

The second technique that is often utilized is uniform covering of the input space. It ensures that input space can be sufficiently covered by the samples. The main disadvantage of this technique is known as a *curse of dimensionality*: the number of required points grows exponentially with the dimension of the input space. If we assert k_i points uniformly distributed in the interval $[L_i, U_i]$ for each variable x_i , then for N variables the total number of data points is determined by $K = \prod_{i=1}^N k_i$. If, for example, $k_i = 5$ (5 points per variable) for each i , and $N = 5$ (5 variables), then $K = 5^5 = 3125$.

In [1, 2] Latin Hypercube sampling technique (briefly described in [1]) has been used to generate of data points; this method is a specific kind of the uniform sampling. Described above two techniques are used if no *a priori* knowledge about the circuit is available. However, often such a knowledge is available and should be used to aid the generation process. This paper tries to show, how it can be done.

The paper is organized as flows. In section 2 we try to get to know, what kind of circuit functions we can deal with. In section 3, on the base of section 2, an expert prediction of the circuit function is discussed. Section 4 describes in details the proposed approach to the model error prediction (MEP) performed by the fuzzy system, and section 5 describes the way of its learning. The algorithm that makes use of the designed fuzzy system is proposed in section 6. Section 7 provides us with two examples of the proposed algorithm and finally, section 8 contains some conclusions and possible future research.

2. TYPE

In or
circuits d

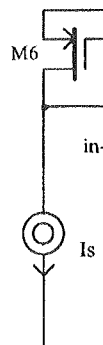


Fig. 1.

We have
current s
calculated

Parameter
transistor
capacitan

Based
the circuit
conducta
Derived f

The time
quality of
ed out as
of a model
the circuit

ing of the
capacitan-
[U_i] then
product:
variables. In
ut space:
etermines
Frequency,
inally, we
ta points
the input
s not well

t space. It
The main
number of
ace. If we
ole x_i , then
 k_i . If, for
oles), then

) has been
e uniform
dge about
should be
be done.

at kind of
an expert
details the
the fuzzy
makes use
s with two
onclusions

2. TYPES OF NONLINEARITIES OCCURED IN ELECTRONIC CIRCUITS; BASIC ASSUMPTIONS

In order to get to know what kind of nonlinearities can be met in the electronic circuits domain let us consider the simple two-stage op-amp presented in Fig. 1.

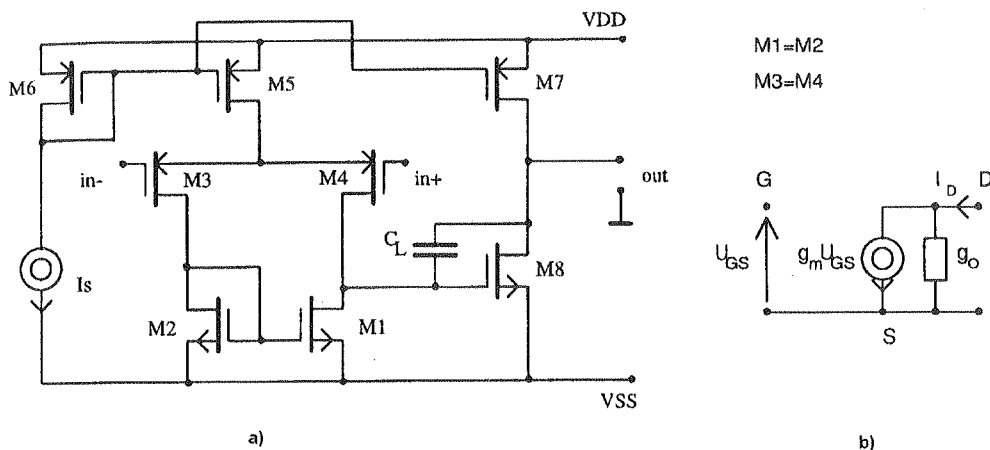


Fig. 1. Simple two-stage operational amplifier (a) and assumed model of the MOS transistor (b)

We have assumed simple model of the MOS transistor as the voltage controlled current source with output conductance g_o . Individual parameters of the model are calculated as follows [3]:

$$g_m = \sqrt{\frac{k'}{2} \frac{W}{L} I_D} \quad g_o = \lambda I_D \quad (1)$$

Parameters k' and λ are technological parameters, and we assume that for all transistors they are the same. W/L of each transistor are circuit parameters as well as capacitance C_L and current I_S .

Based on equations (1) we have derived formulas that determine 6 parameters of the circuit. These are: differential voltage gain K_{ud} , CMRR, output resistance R_o or conductance G_o , total power dissipation P_{tot} , cut-off frequency f_p , and slew rate SR. Derived formulas are presented below.

$$K_{ud} \approx K_{ud0} \sqrt{\frac{(W/L)_{3,4} (W/L)_8 (W/L)_6}{(W/L)_5 (W/L)_7} \frac{1}{I_S}}$$

$$K_{ud}[\text{dB}] \approx 20 \log K_{ud0} + 10 \log ((W/L)_{3,4} (W/L)_8) + 20 \log (W/L)_6 - 10 \log ((W/L)_5 (W/L)_7) - 20 \log I_S$$

$$G_o \approx G_{o0} I_s \frac{(W/L)_7}{(W/L)_6} \quad R_o = \frac{1}{G_o} \approx R_{o0} \frac{(W/L)_6}{I_s (W/L)_7}$$

$$P_{tot} = (V_{CC} - V_{SS}) I_s \frac{(W/L)_5 (W/L)_7}{(W/L)_6^2} \quad SR \approx 2 I_s \frac{(W/L)_5}{C_L (W/L)_6}$$

$$f_i \approx f_{i0} \frac{1}{2\pi} \frac{1}{C_L} \sqrt{(W/L)_{3,4} I_s \frac{(W/L)_5}{(W/L)_6}}$$

Parameters: K_{ud0} , G_{o0} , R_{o0} and f_{i0} depend only on technological parameters, e.g. on k' and λ , and for established technology are constant. CMRR parameter can be expressed by a formula similar to the K_{ud} .

Looking at the listed formulas we can derive some conclusions:

- all circuit functions depend monotonically on each circuit variable with no inflexion points;
- often occurred nonlinearities are: $\log(x)$, $\sqrt{x}$, $1/x$ and $1/\sqrt{x}$;
- by appropriate choice of the circuit parameter we can minimize its nonlinearity with respect to the circuit variables, e.g. R_o is nonlinear with respect to two variables whereas G_o with respect to only one;
- by finding the logarithm of some nonlinear functions which are more nonlinear than $\log(x)$ we can decrease their nonlinearity to $\log(x)$ (e.g. K_{ud} i $K_{ud}[\text{dB}]$).

Taking into account given above statements we assume that considered circuit functions are monotonous with respect to each variable for all possible combinations of the remaining ones. Such an assumption may not be true even for the circuit shown in Fig. 1. Our circuit is determined by much more parameters that, in some cases, can not be even expressed analytically, and can be non-monotonous. As an example we can get nonlinearity error that for OTA has been defined in [1]. In this case our approach may not always perform well. Although, the proposed algorithm has been developed on the basis of the above assumptions, it performs well in some other cases. Of course it is possible to develop other algorithm starting from other, less restricted assumptions.

3. THE PREDICTION OF A BEHAVIOUR OF CIRCUIT FUNCTIONS

Let us assume that for fixed values of all circuit variables except one, x_i , we have three values of any circuit function, y_k , of any considered circuit. In other words we evaluate three values of any circuit function for three values of any circuit variable for fixed values of the remaining ones. Additionally, we assume that the values of circuit variables mentioned above are equidistant. These values are presented in Fig. 2. as black points.

Let us have a look on Fig. 2. (a). What can we tell about the whole characteristics $y_k = f(x_i)$ if we have only three points presented in Fig. 2. (a)? Taking into account considerations performed in section 2 we expect that this characteristic will

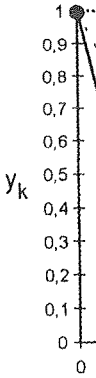


Fig. 2

be similar behavior if we have likely random function knowledge

4. FU

Let function between Our could be tic $y(x)$

We error wi In orde assumed

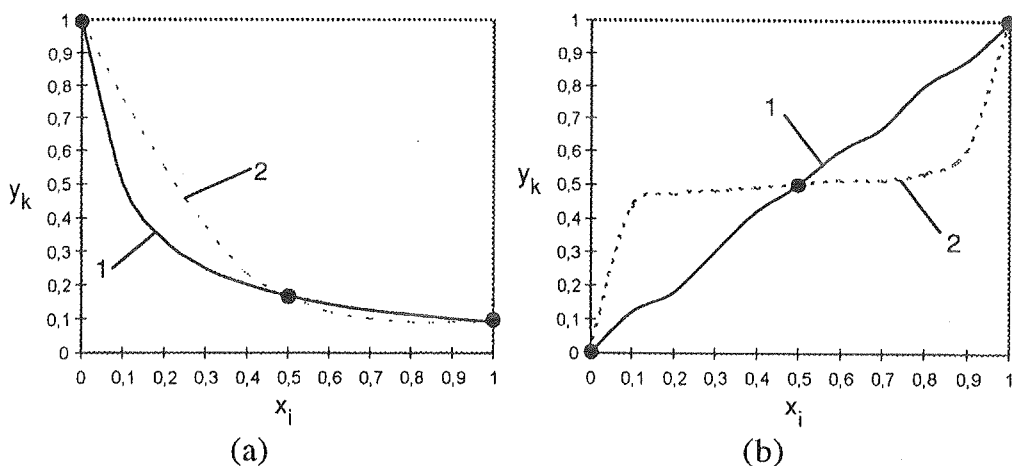


Fig. 2. Examples of distribution of three points representing the behaviour of a circuit function

nonlinearity
ect to two

nonlinear
[dB]).

red circuit
ible com-
en for the
ers that, in
onous. As
d in [1]. In
proposed
performs
m starting

be similar to 1 or 2. We can not indicate which of them better approximate behaviour of an actual function. Both characteristics are equally possible. However, if we have points presented in Fig. 2. (b) we will expect, that characteristic 1 is more likely rather than characteristic 2. So, being given only three points of circuit function we try to foresee the actual character of this function using our *a priori* knowledge about it.

4. FUZZY SYSTEM AS A PREDICTOR OF APPROXIMATION ERROR

Let us assume, that we perform piecewise linear approximation of the actual function which is shown in Fig. 3. In fact it can be treated as linear interpolation between three data points.

Our goal is to estimate how large the expected error of such an approximation could be. We define this error as maximum relative error between actual characteristic $y(x)$ and its approximation $\hat{y}(x)$:

$$\text{MEP} = \text{model error prediction} = \max_{x \in [L, U]} \left| \frac{y(x) - \hat{y}(x)}{y(x)} \right|. \quad (2)$$

We can easily notice, that for characteristic presented in Fig. 3.(a) the expected error will be large whereas for approximation presented in Fig. 3.(b) it will be small. In order to express some general rules related to the expected error value of the assumed approximation we introduce a linearity coefficient (LC) defined as

$$\text{LC} = \min \left\{ \left| \frac{y(x_2) - y(x_3)}{y(x_1) - y(x_2)} \right|, \left| \frac{y(x_1) - y(x_2)}{y(x_2) - y(x_3)} \right| \right\} = \min \left\{ \left| \frac{\Delta_1}{\Delta_2} \right|, \left| \frac{\Delta_2}{\Delta_1} \right| \right\}. \quad (3)$$

TIONS

x_p , we have
words we
it variable
e values of
esented in

characteris-
aking into
eristic will

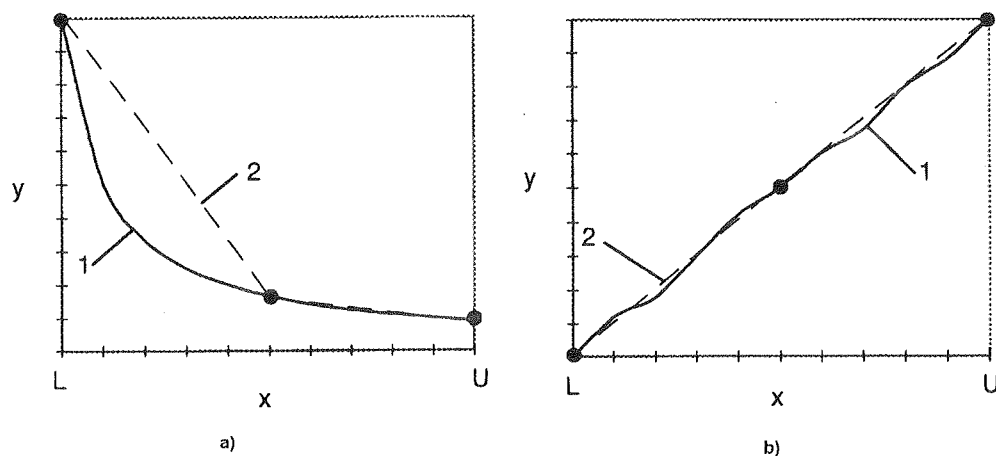


Fig. 3. Actual characteristics (1) and its assumed piecewise linear approximation (2)

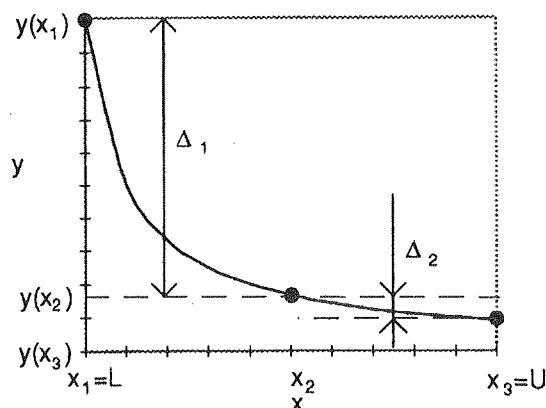


Fig. 4. The interpretation of notations for definition (3)

All notations existing in definition (3) are explained in Fig. 4.

Now we can express some natural statements (rules) related to the MEP:

IF (LC IS "large") THEN (MEP IS "small")

IF (LC IS "middle") THEN (MEP IS "small")

IF (LC IS "small") THEN (MEP IS "middle")

...

Such terms of linguistic variables (LC and MEP) like "middle" or "large" will be in our system represented by the appropriate fuzzy sets.

As it can be seen, only LC is used in the error. For the function $y(x_3)$ the error is much smaller. Thus, we can define

where $y(x)$ is the function. Such a definition can be calculated

For a

IF (Δ IS "small") THEN (LC IS "large")
IF (Δ IS "middle") THEN (LC IS "middle")
IF (Δ IS "large") THEN (LC IS "small")

...

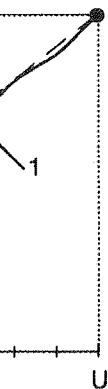
This rule can be used to do it

IF ((LC IS "large") AND (MEP IS "small")) THEN (MEP IS "small")
IF ((LC IS "middle") AND (MEP IS "small")) THEN (MEP IS "small")
IF ((LC IS "small") AND (MEP IS "middle")) THEN (MEP IS "middle")

...

The system has 12 rules.





As it can be noticed, the rules presented above, in which antecedents there appear only LC, are not sufficient to determine an expected value of the approximation error. For example, for two sets of the circuit functions of the form $(y(x_1), y(x_2), y(x_3))$ having values: (1, 2, 10) and (101, 102, 110) the values of the parameter LC are the same whereas we expect that for the second set the approximation error will be much lower than for the first one.

Thus, we need to introduce second parameter. It will be denoted as Δ , and its definition is as follows:

$$\Delta = \frac{y_{\max} - y_{\min}}{y_{\max} + y_{\min}} \quad (4)$$

where $y_{\max} = \max\{y(x_1), y(x_3)\}$ and $y_{\min} = \min\{y(x_1), y(x_3)\}$.

Such a defined parameter Δ is equal to one half of relative change of the circuit function in the interval $[L, U]$ corresponding to the mean value of this function calculated as $y_{\text{mean}} = (y_{\max} + y_{\min})/2$.

(2)

For defined parameter Δ we can derive some linguistic rules, e.g.:

IF (Δ IS "large") THEN (MEP IS "large")
 IF (Δ IS "middle") THEN (MEP IS "middle")
 IF (Δ IS "small") THEN (MEP IS "very small")

...

This rule set alone is not able to determine the expected error value, as well. In order to do it we have derived the folded rule base consisting of rules of the form:

IF ((LC IS "large") AND (Δ IS "large")) THEN (MEP IS "middle")
 IF ((LC IS "large") AND (Δ IS "middle")) THEN (MEP IS "small")

...

The total number of rules depends on the number of fuzzy sets partitioning the domains of LC and Δ . We have assumed 3 fuzzy sets for LC and 4 for Δ , so we have 12 rules. Example of parameter spaces partition is presented in Fig. 5.

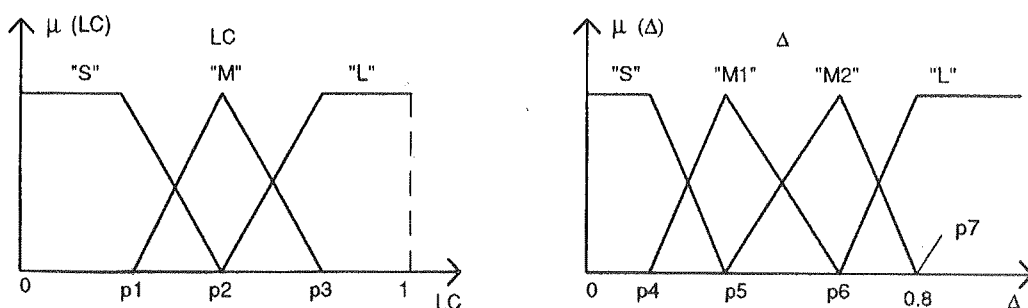


Fig. 5. Assumed partition of spaces of LC and Δ

All rules may be presented in the form of a look-up table:

We want to remark, that fuzzy sets corresponding to LC, Δ and to expected error in Fig. 6 are quite different. Moreover, in practice, all fuzzy sets asserted to MEP will be different (there will be 12 different singleton sets). These singletons will be fixed during a learning process. Now, the structure of our "error predictor" is defined. It has been derived on the basis of two assumptions related to the behaviour of circuit function and to the assumed approximation. The task of our system is to determine (approximately) the expected approximation error of piecewise linear approximation based on 3 data points only (see Fig. 3). Unfortunately, our system is not able to predict the value of this error, because it is not easy to determine exactly the locations of fuzzy sets appearing in consequences of our rules. In order to do it we have performed the learning process which is described in the next section.

$\Delta \backslash LC$	S	M	L
S	M2	S	S
M1	L	M1	S
M2	L	M2	S
L	L	L	M1

Fig. 6. Example of the look-up table of the constructed fuzzy system

5. LEARNING PROCESS OF THE FUZZY SYSTEM

For the fuzzy system being constructed we have chosen the following fuzzy strategies: singleton-type fuzzyfication, max-product fuzzy inference, singleton-type fuzzy sets associated with system output and center of gravity defuzzification method [4, 5]. In order to identify parameters of our fuzzy system, that is locations of all fuzzy sets, we have used genetic algorithm [5]. Genetic algorithms belong to the group of stochastic algorithms, that can effectively search a space of all possible solutions in order to find the best of them. Parameter identification has been performed on the basis of learning data points. To generate these points we have chosen a test function of the form $f(x) = a + 1/x$, $x \in [0.1, 1]$. For such a defined function it is possible to determine analytically the error of approximation (MEP) as a function of parameter a . There has been chosen 81 data pairs of the form $(y(x_1), y(x_2), y(x_3))$ such that $|x_2 - x_1| = |x_3 - x_2|$, $x_1 < x_2 < x_3$, $x_1, x_2, x_3 \in [L, U]$. For each pair $(y(x_1), y(x_2), y(x_3))$ and for various values of the parameter a we have computed values of the related LC and Δ parameters as well as the error of

approx...
been co...
this cas...
corresp...
location...
(p1, p2...
p7). Th...
such, th...
has bee...
on the...
chosen...
paramet...
generati...
mutatio...
have ch...



6
As
a mode...
consum...
control...
It is e...
prepro...
case is...
sophist...
a circu...
predict...

ected error
to MEP
as will be
dictor" is
behaviour
stem is to
ise linear
system is
ne exactly
r to do it
tion.

approximation (MEP). In this manner almost entire input domain of our system has been covered. We have additionally added 9 learning points for which $LC = 0$ — in this case the expected error equals 0 independently of the value of Δ . Since fuzzy sets corresponding to LC and Δ form the fuzzy partition of related domains, thus the location of the fuzzy sets corresponding to LC is determined by only 3 parameters (p_1, p_2, p_3 in Fig. 5) whereas corresponding to Δ by 4 parameters (p_4, p_5, p_6 and p_7). The task of genetic algorithm has been to find values of the parameters $p_1, \dots, p_7$ such, that the mean square error of the system calculated across all training points has been minimized. For each individual in one generation of the genetic algorithm, on the basis of the data points, the initial values of output singletons have been chosen and trained in some iterations of conjugate gradient method [4]. Main parameters of the genetic algorithm have been: number of individuals in each generation: 100, number of generations: 1000, crossover probability: 0.5 and mutation probability: 0.005. We have fired the genetic algorithm several times and have chosen the best result. It is (partially) presented in Fig. 7.

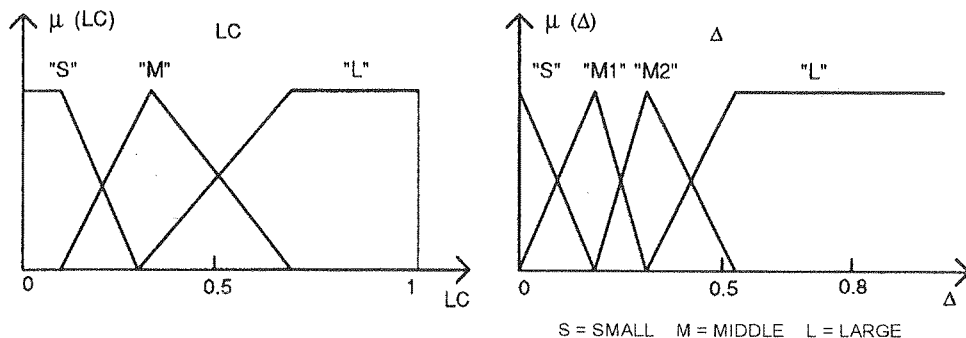


Fig. 7. Locations of fuzzy sets found by genetic algorithm

6. POSSIBILITIES OF APPLICATIONS OF THE FUZZY SYSTEM TO A POINT BASE GENERATION

ing fuzzy
eton-type
zification
locations
ong to the
d possible
has been
we have
a defined
(MEP) as
m ($y(x_1)$,
U]. For
y we have
error of

As it has been mentioned, properly chosen base points enable to generate a model possessing good quality. But often generation process alone can be very time consuming, especially for large number of circuit parameters. Thus, algorithm that controls the process of model generation can significantly influence generation time. It is easy to generate as much data pairs as possible and then perform their preprocessing rejecting superfluous ones; but such an approach in high dimensional case is not possible to perform. Therefore, we have to search for other, more sophisticated methods. These methods arise usually from our *a priori* knowledge about a circuit or a system. One of such an approach, based on the developed fuzzy error predictor is presented below.

Let us assume that we have to generate data base points for any electronic circuit function, still maintaining all previous assumptions about a modelled circuit function. Circuit possesses N parameters denoted as x_i , $i = 1, \dots, N$; each of them is limited to interval $[L_i, U_i]$. Proposed algorithm is as follows:

1. All interval $[L_i, U_i]$ is sampled for $x_i^{(1)} = L_i$, $x_i^{(2)} = L_i + (U_i - L_i)/2$ and $x_i^{(3)} = U_i$.
2. For all possible vectors $X_{a_1, a_2, \dots, a_N} = (x_1^{(a_1)}, x_2^{(a_2)}, \dots, x_N^{(a_N)})$, $a_1, \dots, a_N \in \{1, 2, 3\}$ evaluate the circuit function. In this way we get 3^N data vectors of the form $P_{a_1, a_2, \dots, a_N}(x_1^{(a_1)}, x_2^{(a_2)}, \dots, x_N^{(a_N)}, y^{(a_1, \dots, a_N)})$.
3. $(a_2, a_3, \dots, a_N) = (1, 1, \dots, 1)$;
 - a) take the vector $Y = (y^{(1, a_2, \dots, a_N)}, y^{(2, a_2, \dots, a_N)}, y^{(3, a_2, \dots, a_N)})$ and calculate for it the value of expected error of the approximation (MEP);
 - b) if $\text{MEP} > \text{threshold value}$ then for $x_1^{(1')} = x_1^{(1)} + \frac{x_1^{(2)} - x_1^{(1)}}{2}$ and $x_1^{(2')} = x_1^{(2)} + \frac{x_1^{(3)} - x_1^{(2)}}{2}$ evaluate circuit function: $y^{(1', a_2, \dots, a_N)}$ and $y^{(2', a_2, \dots, a_N)}$, else take the next combination of $(a_2, a_3, \dots, a_N)$ and go to (a);
 - c) calculate MEP for $(y^{(1, a_2, \dots, a_N)}, y^{(1', a_2, \dots, a_N)}, y^{(2, a_2, \dots, a_N)})$ as well for $(y^{(2, a_2, \dots, a_N)}, y^{(2', a_2, \dots, a_N)}, y^{(3, a_2, \dots, a_N)})$;
 - d) if $\text{MEP}(y^{(1, a_2, \dots, a_N)}, y^{(1', a_2, \dots, a_N)}, y^{(2, a_2, \dots, a_N)}) > \text{MEP}(y^{(2, a_2, \dots, a_N)}, y^{(2', a_2, \dots, a_N)}, y^{(3, a_2, \dots, a_N)})$ add to the point base $P_{1', a_2, \dots, a_N}(x_1^{(1')}, x_2^{(a_2)}, \dots, x_N^{(a_N)}, y^{(1', a_2, \dots, a_N)})$ else add to the point base $P_{2', a_2, \dots, a_N}(x_1^{(2')}, x_2^{(a_2)}, \dots, x_N^{(a_N)}, y^{(2', a_2, \dots, a_N)})$;
 - e) if all combinations of $(a_2, a_3, \dots, a_N)$ has been chosen go to (4) else take the next combination $(a_2, a_3, \dots, a_N)$ and go to (a).
4. Repeat 3 for remaining circuit variables.

Fig. 8 explains how for the first variable individual points are calculated.

The described above algorithm is only one of possible, and can be easily changed. The main element in the algorithm is the previously designed fuzzy

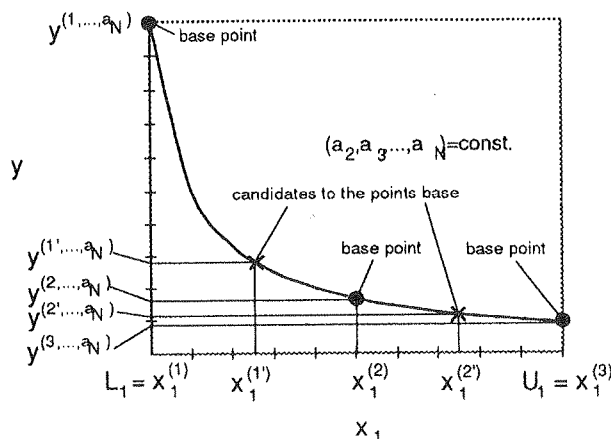


Fig. 8. Example of calculation of individual points for the first circuit variable in step (3) of proposed algorithm

predict
whereas
obvious
our fuzz
property
very im
large. I
will co
mante
a mode
system

As

for x_1 ,
We ha
a mod
the sec
points
between
 $x_1^{(1)}$ or

Resu

Error
thresho
value
[%]

10
5
2
1

Looki

— the
— in

ic circuit
d circuit
f them is

$x_i^{(3)} = U_i$
1, 2, 3}
the form

the value

(1)
and

$\dots, a_N)$, else

$(2, a_2, \dots, a_N)$,

$(3, a_2, \dots, a_N)$

the point

the nest

be easily
ed fuzzy

predictor of the expected error. This predictor has been tuned for one function only whereas we apply it for all possible functions satisfying our assumptions. It is obvious, that in this case error prediction can not be performed precisely. However, our fuzzy system has been designed on the basis of very general rules, so it possesses property of generalisation. Besides, in practical cases exact value of the error is not very important, because for strongly nonlinear functions the value of MEP will be large. In this case the exact value of MEP will play no role because the value of MEP will considerably exceed the threshold value. We have to add, that skilled performante models perform smooth circuit approximation, thus actual error of such a model will be lower than for the piecewise linear one. In this respect our fuzzy system can be treated as upper limiter of the approximation error.

7. EXAMPLES

As the first example we have chosen 5-variable nonlinear function.

$$f(x_1, x_2, x_3, x_4, x_5) = 1 + x_1 + 2\sqrt{(x_1 + x_2)} + 3 \log_{10}(1 + x_2^2 + x_3^3) + 10 \sin\left(\frac{\pi}{2} x_4\right) + \frac{1}{0.1 + x_5} \quad (5)$$

for $x_1, x_2, x_3 \in [0, 10]$, $x_4, x_5 \in [0, 1]$.

We have generated data base points representing this function. We have used a modified version of the algorithm proposed in section 6. Namely, we have added the second step to the algorithm: we have enabled our algorithm to insert additional points located strictly between $x_i^{(1)}$ and $x_i^{(1)}$, $x_i^{(1)}$ and $x_i^{(2)}$, $x_i^{(2)}$ and $x_i^{(2)}$ as well as between $x_i^{(2)}$ and $x_i^{(3)}$ (see Fig. 8). These points are taken into account only if one of $x_i^{(1)}$ or $x_i^{(2)}$ has been added. Table 1 summarises obtained results:

Table 1
Results obtained for 5-variable function by applying modified version of the proposed algorithm

Error threshold value [%]	Number of added data points with respect to each variable					Total number of added points in the first step	Total number of added points in the second step	Total number of added points
	1 st step / 2nd step							
	x1	x2	x3	x4	x5			
10	0/0	0/0	0/0	0/0	54/20	54	20	74
5	0/0	7/2	7/0	14/0	81/81	109	81	190
2	0/0	32/16	26/0	81/2	81/81	220	99	319
1	0/0	45/27	29/2	81/43	81/81	236	153	389

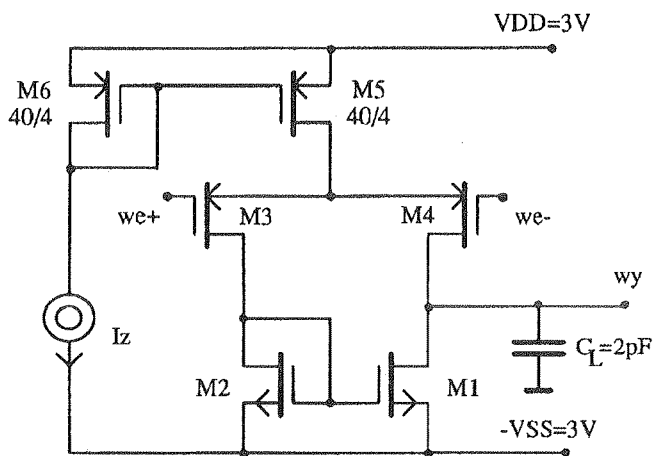
Looking at the Table 1 we can derive some conclusions:

- the number of added points grows up if the error threshold value decreases;
- investigated function is mostly nonlinear with respect to variable x_5 , next to x_4 ;

- for threshold error = 5% in both steps maximal number of added points for x_5 occurs. It means, that only in two steps it is not possible to approximate (in the assumed above manner) the function with error value lower then 5%.

Obtained results are consistent with analytical prediction: the most nonlinear function in (5) (in sense of assumed approximation) is $1/(a+x)$ while the least nonlinear is $x + 2\sqrt{(a+x)}$. Nonlinearity of strongly nonlinear function x^3 is decreased by finding the logarithm of it.

As the second example we have taken a simple CMOS differential pair shown in Fig. 9. As a circuit function we have chosen output resistance R_o and as circuit variable — current I_z as well as (W/L) of M1, M2 and M3, M4. The ranges of circuit



Ranges of circuit variables:

$$I_z = 10 \mu A \dots 100 \mu A$$

$$(W/L)_{M1,2} = 5/4 \dots 50/4$$

$$(W/L)_{M3,4} = 40/4 \dots 400/4$$

Fig. 9. Schematic diagram of the differential pair

parameters are presented in Fig. 9. First we have generated 27 ($3 \times 3 \times 3$) data pairs and next we have run the modified (as in previous example) algorithm with threshold value 10%. It has been added 28 new data points. In order to demonstrate the influence of data point distribution we have assumed very simple fuzzy model of R_o . We have namely partitioned spaces of circuit parameters using fuzzy sets as it is shown in Fig. 10.

Next we have trained such a model on the basis of various learning sets that have uniformly covered parameters space of the model. After each learning process (we have learned only output singletons of the model in fixed number of iterations) we have tested model accuracy in 1000 testing points uniformly distributed in the circuit parameter space. Obtained results are presented in Table 2.

Since the structure of the model is probably one of the simplest thus the model errors are large. However, the aim of this paper is not to construct accurate models, but to show how the point location may influence the accuracy and the time of model training. Results contained in Tab. 2 show, that the mean model error for 55 adequate chosen learning points is better than for 216 uniformly distributed ones.

Learning
the mod
125 poin
Since
model st
structure
changes,

nts for x_5
te (in the

nonlinear
the least
on x^3 is

shown in
as circuit
of circuit

variables:

A

4 .. 50/4

/4 .. 400/4

ata pairs
hm with
onstrate
model of
ts as it is

that have
ocess (we
ions) we
ne circuit

ne model
e models,
e time of
or for 55
ed ones.

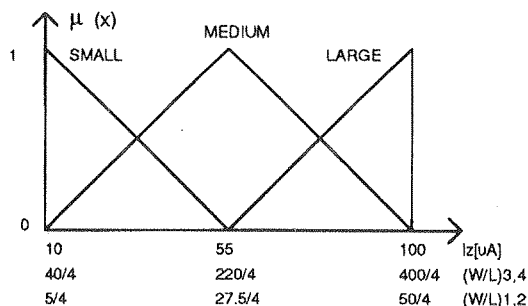


Fig. 10. Fuzzy sets dividing spaces of the differential pair parameters

Table 2

Results obtained for the simple model of the differential pair

Learning set	Mean error [%]	Max. error [%]	Number of learning points
$3 \times 3 \times 3$	15.4	57.2	27
$4 \times 4 \times 4$	13.2	38.6	64
$5 \times 5 \times 5$	9.4	38.7	125
$6 \times 6 \times 6$	10.5	32.0	216
$3 \times 3 \times 3$ + added points	9.3	36.4	55

Learning time in this second case is of course 4 times longer. The maximum error of the model learned on the basis of 55 points is better than for uniformly distributed 125 points.

Since the location of learning points added by the algorithm can indicate the model structure, thus on the basis of locations of added points we have changed the structure of our model. The partitioning of (W/L) of transistors we have left without changes, but the partition of I_z we have assumed as in Fig. 11.

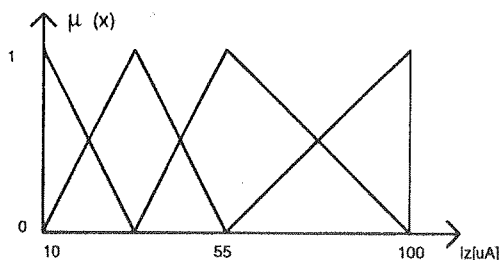


Fig. 11. The modified partition of the I_z space

Table 3

Results obtained for modified model of the simple differential pair

Learning set	Mean error [%]	Max. error [%]	Number of learning points
$4 \times 4 \times 4$	4.8	26.5	64
$5 \times 5 \times 5$	5.1	27.4	125
$6 \times 6 \times 6$	4.1	21.0	216
$3 \times 3 \times 3 + \text{added points}$	4.6	23.2	55

For this model we have performed the same tests as for the previous. Obtained results summarises Table 3.

In this case a significant improvement in the model errors occurs, especially for the mean error. Additionally the accuracy of the model trained on the basis of 55 learning points is better than for 125 uniformly distributed ones.

8. CONCLUSIONS AND FUTURE RESEARCH

Based on presented examples we have shown how *a priori* knowledge about behaviour of an electronic circuit can be used, and how properly chosen learning points can influence the model accuracy and structure. However, our approach will perform well only for restricted class of functions. Moreover, our fuzzy system has been trained for one kind of function only, so for other kinds of functions it works only approximately. In order to overcome its drawbacks some changes are required. First, we can start from other, less restricted assumptions about circuit functions. Second, for the investigated function an error predictor the best suited for it should be used. Since we do not know much about a modelled function, thus we propose to determine some classes of functions. At the beginning it is necessary to derive some formulas for some representatives of a given circuit class. In order to determine the class, to which investigated function belong, we propose to use a classifier, as is presented in Fig. 12.

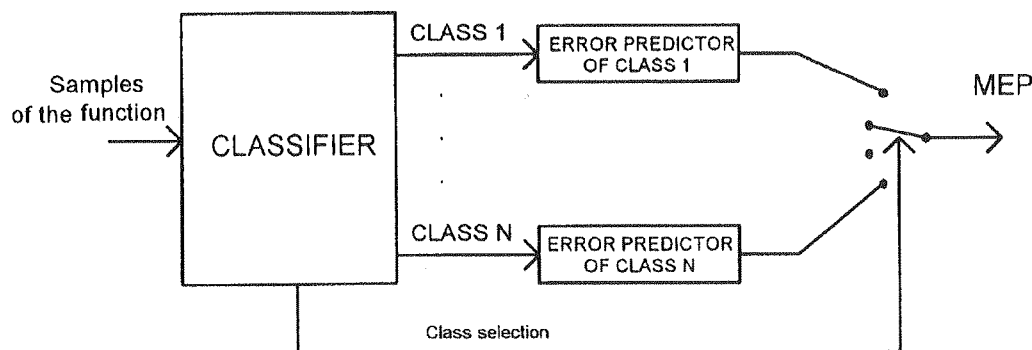


Fig. 12. Proposition of the usage of classifier

Samples of the circuit function could be taken in some points of circuit input space, and on the basis of these samples classifier would decide about appropriate class. The error predictor of each class would be learned on the basis of a function representative for this class. Such an approach would ensure, that MEP value could be calculated more exactly.

REFERENCES

1. M.A. Stybliński, S. Aftab: *Combination of Interpolation and Self-Organizing Approximation Techniques* — A New Approach to Circuit Performance Modelling. IEEE Trans. on CAD, Nov. 1993, Vol. 12, No. 11, pp. 1775–1785
2. A. Torralba, J. Chavez, L. Franquelo: *Circuit Performance Modelling by Means of Fuzzy Logic*. IEEE on CAD, Nov. 1996, Vol. 15, No. 11, pp. 1391–1398
3. A. Guziński: *Liniowe elektroniczne układy analogowe*. WNT, W-wa, 1992
4. R.R. Yager, D.P. Filev: *Podstawy modelowania i sterowania rozmytego*. WNT, W-wa, 1995
5. D. Rutkowska, M. Piliński, L. Rutkowski: *Sieci neuronowe, algorytmy genetyczne i systemy rozmyte*. PWN, W-wa, 1997

R. HORBOWSKI, M. BIAŁKO

GENEROWANIE PUNKTÓW UCZĄCYCH DLA MODELI FUNKCJI UKŁADOWYCH
UKŁADÓW ELEKTRONICZNYCH ZA POMOCĄ SYSTEMU ROZMYTEGO

Streszczenie

W pracy zaprezentowano sposób konstrukcji i uczenia systemu rozmytego, zadaniem którego jest określenie błędu założonej aproksymacji funkcji układowych układów elektronicznych. Konstrukcja systemu oparta została na wiedzy a priori o zachowaniu się tych funkcji oraz na podstawowym założeniu wynikającym z przykładowo przeprowadzonej analizy wybranego układu elektronicznego. Wyrażone lingwistycznie reguły rozmyte, określają jednocześnie strukturę systemu, który dodatkowo był uczony na podstawie danych numerycznych.

Bazując na zaprojektowanym systemie rozmytym zaproponowano algorytm generacji punktów uczących. Jego działanie zaprezentowano na dwóch prostych przykładach pokazując, jak właściwie dobrane punkty uczące mogą wpływać na czas jego uczenia jak i na jego dokładność. Przedstawiono również możliwość wykorzystania klasyfikatora w systemie predykcji błędu aproksymacji.

Słowa kluczowe: modelowanie rozmyte i neurorozmyte, logika rozmyta, modelowanie układów elektronicznych, inżynieria wiedzy, systemy ekspertowe.

MEP

mak
para
takt
tych
nej.
para
pods
należ
ratow
licze
W a
wart

znac
błęd
opisa
wart
w ty
znac
stand
częśc
czasie
z me
najm

ekstr
MTI
wyzn

Wyznaczanie maksymalnego błędu przedziału czasu sygnałów taktowania

ANDRZEJ DOBROGOWSKI, MICHAŁ KASZNIA

Instytut Elektroniki i Telekomunikacji, Politechnika Poznańska

Otrzymano 1998.07.17

Autoryzowano 1998.09.14

W artykule przedstawione zostały zagadnienia związane z wyznaczaniem wartości maksymalnego błędu przedziału czasu. Maksymalny błąd przedziału czasu (MTIE) jest parametrem charakteryzującym sygnały taktowania sieci synchronizacyjnej. Sygnały te taktują operacjami realizowanymi w cyfrowych sieciach i urządzeniach. Właściwa jakość tych sygnałów jest warunkiem prawidłowego procesu synchronizacji sieci telekomunikacyjnej. Istotnym zagadnieniem w eksploatacji cyfrowej sieci telekomunikacyjnej jest pomiar parametrów związanych z długookresowymi zmianami fazy sygnałów taktowania. Do podstawowych parametrów charakteryzujących wolne zmiany fazy sygnałów taktowania należą dewiacja Allana, zmodyfikowana dewiacja Allana, dewiacja czasu, średniokwadratowy błąd przedziału czasu oraz maksymalny błąd przedziału czasu. Pomiar i obliczenia tych parametrów wymagają znacznego zaangażowania sprzętowego i czasowego. W artykule skupiono się na metodach umożliwiających skrócenie czasu wyznaczania wartości maksymalnego błędu przedziału czasu.

W pierwszej części pracy przedstawione zostały definicje wielkości o podstawowym znaczeniu dla charakteryzowania sygnału taktowania, a mianowicie błędu czasu oraz błędu przedziału czasu oraz definicja maksymalnego błędu przedziału czasu. Następnie opisane zostały warunki dokonywania pomiaru błędu czasu. Na podstawie pomierzonych wartości błędu czasu wyznacza się wartości innych parametrów sygnałów taktowania, w tym także maksymalnego błędu przedziału czasu. Przedstawiona została zasada wyznaczania estymaty MTIE zgodnie z definicją podaną przez międzynarodowe instytucje standaryzacyjne. Zastosowanie takiej metody jest jednak bardzo czasochłonne. W następnej części pracy przedstawiono szczegółowo trzy metody umożliwiające w znacznie krótszym czasie wyznaczanie wartości MTIE charakteryzującej badany sygnał taktowania. Pierwsza z metod, nazwana przez autorów metodą większego skoku, pozwala na znalezienie z co najmniej dobrym przybliżeniem rzeczywistej wartości MTIE charakteryzującej ciąg danych.

Pozostałe dwie metody — metoda decyzji brzegowych oraz metoda położenia wartości ekstremalnych, pozwalają w sposób efektywny czasowo na znalezienie dokładnej estymaty MTIE. Przedstawiono także wyniki eksperymentu porównującego efektywność czasową wyznaczania wartości MTIE poszczególnymi metodami dla przykładowego ciągu wyników

pomiaru. Potwierdzona została efektywność czasowa zaproponowanych metod. Czas wyznaczania był od kilkunastu do kilkuset razy krótszy, niż czas wyznaczania metodą opartą bezpośrednio na definicji parametru. Przytoczona została także opisana w literaturze metoda zgrubnego szacowania MTIE. Zwrócono uwagę na wady tej metody.

Zastosowanie opisanych metod w praktyce jest pomocne w eksploatacji sieci telekomunikacyjnej. Umożliwi szybkie wyznaczanie wartości parametru badanego sygnału taktowania sieci synchronizacji. Porównanie otrzymanych wartości z zaleceniami (normami) instytucji standaryzacyjnych pozwoli na dokonanie oceny działania sieci telekomunikacyjnej. Pracę wykonano w ramach projektu TB-44-536/98/DS.

Słowa kluczowe: sieć synchronizacji, sygnał taktowania, parametry sygnału taktowania, błąd czasu, błąd przedziału czasu, maksymalny błąd przedziału czasu.

1. WSTĘP

Istotnym zagadnieniem w eksploatacji cyfrowej sieci telekomunikacyjnej jest pomiar parametrów sygnałów synchronizacji sieci. Sygnały te taktują operacjami zachodzącymi w sieci, stąd zwane są także sygnałami taktowania. Właściwa jakość sygnałów taktowania jest warunkiem prawidłowego procesu synchronizacji sieci telekomunikacyjnej. Sygnały taktowania transmitowane w sieci poddawane są różnym mechanizmom skażeniowym powodującym pogorszenie ich jakości. Z punktu widzenia eksploatacji sieci ważnym zagadnieniem jest pomiar parametrów związanych z wolnozmiennymi fluktuacjami fazy sygnałów taktowania. Fluktuacje te, inaczej określane jako wędrowanie lub pełzanie fazy (ang. *wander*), są długookresowymi wahaniami chwil znamienych sygnału taktowania wokół ich nominalnego usytuowania w czasie. Według standardów wahania długookresowe mają częstotliwość mniejszą niż 10 Hz [1, 2].

Do podstawowych parametrów charakteryzujących wolne zmiany fazy sygnału taktowania należą dewiacja Allana, zmodyfikowana dewiacja Allana, dewiacja czasu, średniokwadratowy błąd przedziału czasu oraz maksymalny błąd przedziału czasu [1]. Pomiary i obliczenia tych parametrów wymagają znacznego zaangażowania sprzętowego i czasowego. Niezbędne są specjalistyczne urządzenia pomiarowe oraz sterujące nimi dostatecznie szybkie komputery, które wykonują także analizę wyników pomiaru i ostateczne obliczenia parametrów. Czas wykonywania pomiaru, niezbędny do określenia poszczególnych parametrów, sprecyzowany jest przez międzynarodowe standardy. Istotny jest także czas wykonywania obliczeń parametrów, który zależy od długości ciągu danych oraz od skomplikowania matematycznej formuły definiującej dany parametr.

W literaturze [3, 4, 8] przedstawione zostały podstawowe zagadnienia związane z pomiarami parametrów sygnałów synchronizacji. Opisane zostały wymagania dotyczące sprzętu niezbędnego w procesie pomiarowym. W pracy [3] zasygnalizowane zostały możliwości modyfikacji metod wyznaczania parametrów sygnałów synchronizacji, umożliwiające skrócenie czasu oczekiwania na wyniki. Efekty badań nad tymi metodami, w odniesieniu do dewiacji Allana, dewiacji czasu i maksymalnego błędu przedziału czasu, zawarte zostały w pracach [5, 6]. Dobre wyniki

uzyskan
stawiony
przedzia
czasu bę
Następn
tego par
paramet
nano me
obliczeń

Sygn
stosowan
niach. S

gdzie A j
całkowit
fazowego

Częst
przebieg
Wielk
jest błąd
różnicą p

Zakłada
podstawi
TIE (ang
podstawi

Funkcja
błędu cza

gdzie t je
Błąd
początko

uzyskano dla dewiacji Allana i dewiacji czasu. Przedmiotem rozważań przedstawionych w niniejszej pracy będą metody wyznaczania maksymalnego błędu przedziału czasu. Najpierw więc zostanie zdefiniowany maksymalny błąd przedziału czasu będący jedną z możliwych charakterystyk funkcji czasu sygnału taktowania. Następnie w pracy zawarta będzie pogłębiona analiza metod wyznaczania wartości tego parametru. Przedstawiona zostanie metoda oparta bezpośrednio na definicji parametru oraz metody umożliwiające znaczne skrócenie czasu obliczeń. Porównano metody wyznaczania maksymalnego błędu przedziału czasu ze względu na czas obliczeń dla pozyskanego w eksperymencie ciągu danych.

2. MAKSYMALNY BŁĄD PRZEDZIAŁU CZASU

Sygnał taktowania jest nominalnie okresowym sygnałem generowanym przez zegar, stosowanym do taktowania operacji realizowanych w cyfrowych sieciach i urządzeniach. Sygnał taktowania (określany także jako sygnał zegara) może mieć postać:

$$s(t) = A \sin[\varphi(t)], \quad (1)$$

gdzie A jest stałym współczynnikiem (amplitudą), t jest czasem idealnym, a $\varphi(t)$ jest całkowitą fazą chwilową. Funkcja czasu danego zegara jest stosunkiem procesu fazowego $\varphi(t)$ i kątowej częstotliwości nominalnej tego zegara:

$$T(t) = \frac{\varphi(t)}{2\pi\nu_{nom}}. \quad (2)$$

Częstotliwość nominalna jest to postulowana, stała w czasie, częstotliwość przebiegu wytwarzanego przez zegar.

Wielkością o podstawowym znaczeniu dla scharakteryzowania sygnału taktowania jest błąd czasu (ang. *Time Error*, TE). Błąd czasu, oznaczany jako $TE(t)$ lub $x(t)$, jest różnicą pomiędzy czasem $T(t)$ badanego zegara i czasem $T_{ref}(t)$ zegara odniesienia:

$$x(t) = T(t) - T_{ref}(t) \quad (3)$$

Zakłada się, że jakość zegara odniesienia jest lepsza od jakości zegara badanego. Na podstawie pomierzonych wartości błędu czasu można określić błąd przedziału czasu TIE (ang. *Time Interval Error*) będący różnicą miar przedziału czasu, uzyskanych na podstawie wskazań badanego zegara i zegara odniesienia:

$$TIE(t, \tau) = [T(t + \tau) - T(t)] - [T_{ref}(t + \tau) - T_{ref}(t)] \quad (4)$$

Funkcja błędu przedziału czasu TIE(t, τ) może zostać wyrażona za pomocą funkcji błędu czasu $x(t)$ w postaci:

$$TIE(t, \tau) = x(t + \tau) - x(t) \quad (5)$$

gdzie t jest chwilą początkową obserwacji a τ przedziałem obserwacji.

Błąd przedziału czasu TIE(t, τ) jest wielkością zmienną w czasie i zależy od początkowej chwili obserwacji t . Z punktu widzenia charakterystyki sygnału tak-

towania istotne jest określenie średniej oraz granicznej wartości błędu przedziału czasu w badanym sygnale taktowania. Średniokwadratowy błąd przedziału czasu TIErms definiowany jest następująco:

$$\text{TIErms}(\tau) = \sqrt{\langle [x(t+\tau) - x(t)]^2 \rangle} \quad (6)$$

gdzie nawiasy kątowe $\langle \rangle$ oznaczają średnią po nieskończonym czasie. Maksymalny błąd przedziału czasu (ang. *Maximum Time Interval Error*, MTIE) został pierwotnie zdefiniowany jako maksymalna międzyszczytowa wartość błędu czasu sygnału taktującego względem sygnału wzorcowego w pewnym przedziale obserwacji τ [8]:

$$\text{MTIE}_{t_0}(\tau) = \max_{t_0 \leq t \leq t_0 + \tau} [x(t)] - \min_{t_0 \leq t \leq t_0 + \tau} [x(t)], \quad (7)$$

gdzie chwila t_0 wyznacza początek przedziału obserwacji.

Międzynarodowe instytucje standaryzacyjne rozszerzyły następnie definicję MTIE na wszystkie możliwe przedziały obserwacji o tej samej długości τ występujące w czasie pomiaru T definiując ten parametr następująco:

$$\text{MTIE}(\tau, T) = \max_{0 \leq t_0 \leq T - \tau} [\text{MTIE}_{t_0}(\tau)]. \quad (8)$$

W definicji MTIE sformułowanej w standardzie ETSI [1] z listopada 1995 r. przyjęto, że czas T jest tożsamy z całą osią czasu:

$$\text{MTIE}(\tau) = \max_{-\infty \leq t_0 \leq \infty} \left(\max_{t_0 \leq t \leq t_0 + \tau} [x(t)] - \min_{t_0 \leq t \leq t_0 + \tau} [x(t)] \right) \quad (9)$$

Zalecenie ITU-T G.810 [2] w wersji z sierpnia 1996 r. w definicji MTIE uwzględnia explicite zarówno oczywisty fakt, że czas pomiaru T musi być skończony jak i losowy charakter procesu błędu czasu. Według tego zalecenia maksymalny błąd przedziału czasu MTIE(τ) definiuje się jako percentyl β zmiennej losowej X określonej wzorem:

$$X = \max_{0 \leq t_0 \leq T - \tau} \left(\max_{t_0 \leq t \leq t_0 + \tau} [x(t)] - \min_{t_0 \leq t \leq t_0 + \tau} [x(t)] \right). \quad (10)$$

Jeśli wyniki pomiaru błędu czasu zapisane są w postaci ciągu próbek wziętych z odstępem próbkowania τ_0 , to maksymalny błąd przedziału czasu estymowany jest według wzoru:

$$\hat{\text{MTIE}}(n\tau_0) = \max_{1 \leq k \leq N - n} \left(\max_{k \leq i \leq k + n} x_i - \min_{k \leq i \leq k + n} x_i \right), \quad (11)$$

gdzie: $\{x_i\}$ — ciąg N próbek błędu czasu $x(t)$ wziętych z odstępem τ_0 ;

$\tau = n\tau_0$ — przedział obserwacji;

$n = 1, 2, \dots, N - 1$.

Zalecenie G.810 [2] stwierdza, że wzór (11), podany także jako estymator MTIE w standardzie ETSI [1] dostarcza punktowej estymaty w oparciu o dane uzyskane

w pojed
przedstaw
maksyma
błąd prz
stawie p
dobrej ja
nej wpro
(ang. Ma
błąd prze
czasu sy
jego wejś

Aby
należy d
wartości
pomiaro
w skład
wiązanie

Zgod
najmniej
a więc c
wartości
parametr
i maksym
tru. Zale
wacji, a
obliczeń
12 razy c

W ba
łagodniej
czas pom
 $\tau_{\max}$, zaś
MTIE m
obserwac
sygnałów
że czas p
malne. W
położon
dłuższy c
W cel
wzorem

przedziału
czasu

(6)
maksymalny
erwotnie
sygnału
ci τ [8]:

(7)

definicję
stępujące

(8)

1995 r.

(9)

względnie
i losowy
przedziału
wzorem:

(10)

wziętych
wany jest

(11)

or MTIE
zyskane

w pojedynczym sensie pomiarowym o czasie trwania $T = (N - 1) \tau_0$. W artykule przedstawiono efektywne ze względu na czas obliczeń metody wyznaczania estymaty maksymalnego błędu przedziału czasu na podstawie estymatora (11). Maksymalny błąd przedziału czasu dla pewnego sygnału taktowania wyznaczany jest na podstawie pomiaru błędu czasu tego sygnału względem sygnału odniesienia o bardzo dobrej jakości. W analizie parametrów sygnałów taktowania sieci telekomunikacyjnej wprowadzone zostało pojęcie maksymalnego względnego błędu przedziału czasu (ang. *Maximum (Relative) Time Interval Error*, $M(R)TIE$). Maksymalny względny błąd przedziału czasu definiuje się jako maksymalną międzyszczytową wartość błędu czasu sygnału taktowania na wyjściu urządzenia względem sygnału taktowania na jego wejściu [1, 2].

3. POMIAR BŁĘDU CZASU I WYZNACZANIE MTIE

Aby dokonać obliczeń wartości parametrów sygnałów synchronizacji najpierw należy dokonać pomiaru błędu czasu. Pomierzone z odstępem próbkowania τ_0 wartości (próbki) błędu czasu $x_i = x(i\tau_0)$ należy zapisać w pamięci urządzenia pomiarowego. W tym celu może zostać wykorzystany dysk komputera wchodzącego w skład systemu pomiarowego. W przypadku długich serii pomiarów takie rozwiązanie jest najwygodniejsze.

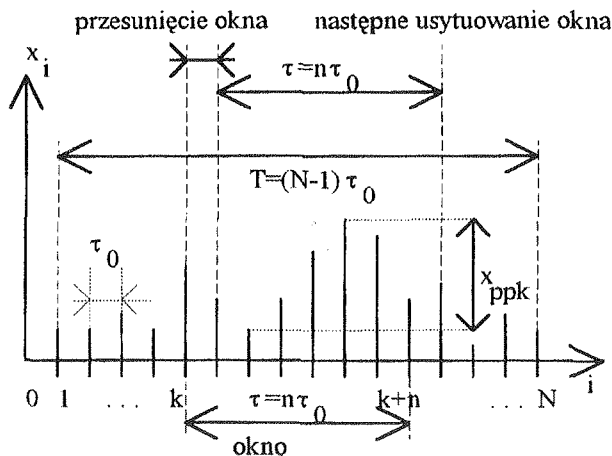
Zgodnie ze standardem [1] minimalny przedział obserwacji τ musi być co najmniej trzy razy większy od odstępów próbkowania τ_0 . Dla przedziału $\tau = 0,1$ s, a więc dla minimalnego przedziału obserwacji, dla którego standardy określają wartości graniczne, maksymalny odstęp próbkowania wynosi $\tau_0 = 1/30$ s. Estymując parametry sygnału synchronizacji należy rozważyć relację pomiędzy czasem pomiaru i maksymalnym przedziałem obserwacji, dla którego liczona będzie wartość parametru. Zalecenia nie precyzują zależności pomiędzy maksymalnym przedziałem obserwacji, a długością ciągu wyników pomiaru dla obliczeń MTIE. Natomiast dla obliczeń dewiacji czasu zalecane jest [1], aby czas trwania pomiaru był co najmniej 12 razy dłuższy od maksymalnego przedziału obserwacji $\tau_{\max}$.

W badaniach nad efektywnością różnych algorytmów obliczeń można przyjąć łagodniejsze kryteria. Z postaci estymatora dewiacji czasu wynika, że minimalny czas pomiaru musi być trzykrotnie dłuższy od maksymalnego przedziału obserwacji $\tau_{\max}$, zaś dla dewiacji Allana ten czas jest 2 razy dłuższy od $\tau_{\max}$. Natomiast wartość MTIE można obliczyć, gdy czas pomiaru jest równy maksymalnemu przedziałowi obserwacji. W praktyce pomiarowej przyjmuje się, że wyznaczanie parametrów sygnałów taktowania bazuje na jednakowym ciągu danych. Przyjmuje się również, że czas pomiaru powinien być większy, niż przedstawione powyżej wartości minimalne. W artykule, przy eksperymencie weryfikacji metod wyznaczania MTIE, posłużono się danymi uzyskanymi w pomiarze, którego czas trwania był cztery razy dłuższy od rozważanego maksymalnego przedziału obserwacji.

W celu znalezienia wartości MTIE w przedziale obserwacji τ , należy, zgodnie ze wzorem (11), dokonać przejrzenia wszystkich przedziałów szerokości τ występują-

cych w ciągu N próbek błędu czasu. W tym celu utworzone zostaje „okno” o szerokości $\tau = n\tau_0$, obejmujące $n + 1$ kolejnych próbek błędu czasu, które umieszczone zostaje na początku ciągu wartości $\{x_i\}$. Utworzone w ten sposób okno przesuwane jest następnie o odstęp τ_0 , a więc o jedną próbkę, do końca ciągu $\{x_i\}$. Dla każdego usytuowania okna znajdowana jest wartość międzyszczytowa błędu czasu, czyli maksymalny w tym usytuowaniu okna błąd przedziału czasu. Wartość $MTIE(\tau)$ jest maksymalną wartością międzyszczytowa błędu czasu obliczoną po wszystkich możliwych usytuowaniach okna o szerokości τ . Dla ciągu próbek o długości N i przedziału szerokości $\tau = n\tau_0$, istnieje $N - n$ usytuowań okna. Jeśli, na przykład, błąd czasu mierzony był z odstępem $\tau_0 = 1/30$ sekundy przez czas 4000 sekund, a wartość przedziału obserwacji wynosi $\tau = 10$ s, to pojedyncza operacja przeglądania okna obejmującego 301 próbek ($n + 1 = 301$) wykonywana jest 119701 razy, bowiem $N = 120001$. Zasada wyznaczania maksymalnego błędu czasu przedstawiona jest na rysunku 1. Algorytm wykonywania obliczeń przedstawia diagram na rysunku 2. Linia przerywaną otoczona została najbardziej pracochłonna procedura przeglądania okna.

W celu sprawdzenia jakości badanego sygnału ciąg wartości $MTIE$ wyznaczonych dla przedziałów obserwacji τ z określonego zakresu (np. od 0,1 s do 1000 s) należy przedstawić w postaci wykresu o logarytmicznej skali wartości i logarytmicznej skali czasu. Tak przedstawione wyniki porównuje się następnie z nałożonymi wartościami granicznymi (normami) dla $MTIE$ w określonym punkcie sieci, zdefiniowanymi przez instytucje standaryzacyjne. Niezbędne jest więc wyznaczenie ciągu wartości maksymalnego błędu przedziału czasu. Zastosowanie do obliczeń opisanej wyżej metody opartej bezpośrednio na definicji $MTIE$, polegającej na przesuwaniu i przeglądaniu okien jest niezwykle czasochłonne. Wraz ze wzrostem wielkości przedziału obserwacji τ maleje wprawdzie liczba usytuowań okna, a więc liczba

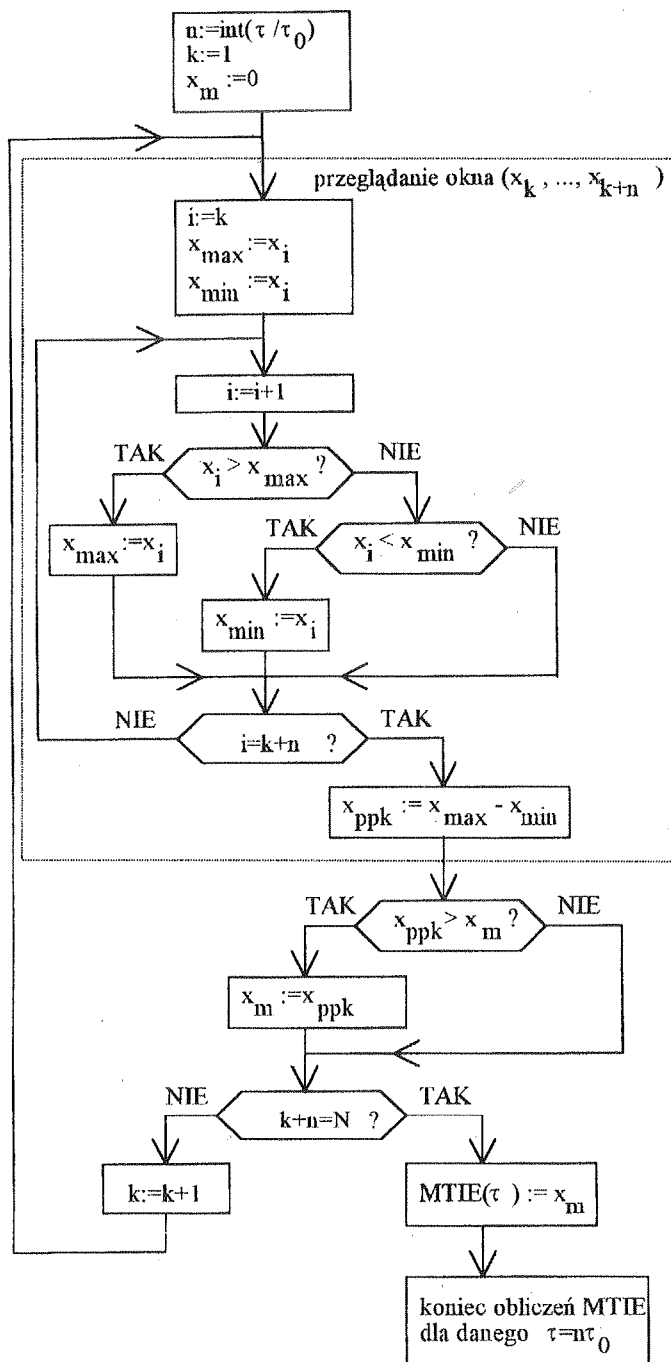


Rys. 1. Zasada wyznaczania $MTIE$

Rys. 2. AL
 $n + 1$ —
 $x_{\min}$ — w

„okno”
e umiesz-
ób okno
agu $\{x_i\}$.
wa błędu
Wartość
czoną po
próbek
na. Jeśli,
czas 4000
operacja
st 119701
su przed-
diagram
na proce-

yznacz-
o 1000 s)
arytmicz-
łożonymi
ci, zdefi-
nie ciągu
opisanej
esuwaniu
wielkości
ęc liczba



Rys. 2. Algorytm wykonywania obliczeń wartości MTIE dla danego przedziału obserwacji τ . Oznaczenia: $n+1$ — liczba próbek w oknie, k — kolejny numer próbki, i — numer próbki w oknie, $x_{\max}$, $x_{\min}$ — wartość maksymalna i minimalna dla danego usytuowania okna, x_m — bieżąca maksymalna wartość międzyszczytowa, x_{ppk} — wartość międzyszczytowa dla danego okna

Tabela 1

Zależność między szerokością przedziału obserwacji (okna), liczbą próbek w oknie i liczbą przeglądanych okien

szerokość przedziału obserwacji	liczba próbek w oknie	liczba przeglądanych okien
τ [s]	$n + 1$	$N - n$
0,1	4	119998
1	31	119971
10	301	119701
100	3001	117001
1000	30001	90001

przeglądanych okien. Rośnie natomiast liczba próbek n określająca szerokość okna, czyli liczba wartości, które należy przejrzeć dla każdego położenia okna. W Tabeli 1 przedstawiona jest zależność między liczbą próbek w przedziale, liczbą operacji przeglądania okna i szerokością przedziału obserwacji τ dla pomiaru błędu czasu z odstępem $\tau_0 = 1/30$ s przez czas 4000 sekund.

Długość czasu wyznaczania wartości MTIE zwiększa się wraz ze wzrostem wielkości przedziału obserwacji τ . Spowodowane to jest zwiększeniem liczby próbek w oknie i jednoczesnym nieadekwatnym zmniejszeniu liczby przeglądanych okien. (Oczywiście w przypadku, gdy wielkość przedziału obserwacji τ będzie dostatecznie duża, czas wyznaczania będzie maleł. Ze względu jednak na przyjęte założenie, że czas pomiaru jest cztery razy większy od maksymalnego przedziału obserwacji, taka sytuacja nie będzie miała miejsca). Zależności przedstawione w Tabeli 1 dotyczą pomiaru błędu czasu wykonanego przez 4000 sekund, a maksymalny przedział obserwacji wynosi 1000 s. Sprawdzenie jakości sygnału taktowania zgodnie ze standardami międzynarodowymi [1] wymaga obliczeń dla większych przedziałów obserwacji (do 100000 s) oraz wykonania pomiaru w znacznie dłuższym czasie. Prowadzi to do zwiększenia całkowitej liczby wartości błędu czasu czyli wydłużenia ciągu próbek, który należy przejrzeć. Wzrośnie zatem liczba operacji, a więc również czas obliczeń. Konieczne staje się zatem znalezienie metody, która pozwoli w istotnym stopniu skrócić ten czas i nie spowoduje wprowadzenia znaczących błędów do wyników obliczeń.

Bardzo prostym i jednocześnie skutecznym sposobem jest zwiększenie odstępu próbkowania τ_0 . Zmniejszy się przez to ogólna liczba danych uzyskanych w pomiarze oraz liczba próbek przypadających na przedział obserwacji o danej wielkości. Wszystko to w efekcie spowoduje skrócenie czasu obliczeń. Poprzez zwiększenie odstępu próbkowania τ_0 zwiększy się jednak minimalny przedział obserwacji, co uniemożliwi obliczanie MTIE dla małych wartości τ . Aby temu zapobiec wyniki

pomiaru
a do ob
 $m = 10$),
może je
zapropo
wykorzy
i dewiac
maksym
paramet
kach obl
W ta
modyfik

W p
metodę
większeg
próbkow
metodą,
większyc
okna, kt
 τ większ
wartości
położen
Kolejne
usytuow
wartości
błędów
w pewn
i zauwa
Zasada
została
rysunku
graniczn
(skoku).
przedstaw
wywoływ
z definicj

1
ie

miaru należałoby zapisywać z dużą częstotliwością (a więc małym odstępem τ_0), a do obliczeń MTIE dla dużych wartości τ wybierać co m -tą próbkę (np. dla $m = 10$), a więc wziąć krótszy ciąg wartości błędu czasu. Takie opuszczanie próbek może jednak spowodować istotne błędy w obliczaniu MTIE. W pracach [5, 6] zaproponowano opartą na uśrednianiu metodę przekształcania zbioru danych, którą wykorzystano z dobrym skutkiem do przyspieszenia obliczeń dewiacji Allana i dewiacji czasu. Metoda ta nie może być jednak zastosowana do wyznaczania maksymalnego błędu przedziału czasu. Ze względu na graniczny charakter tego parametru, uśrednianie próbek błędu czasu spowodowałoby istotne błędy w wynikach obliczeń.

W takim przypadku metoda skracająca czas obliczeń musi opierać się na modyfikacji algorytmu, a nie modyfikacji danych uzyskanych z pomiaru.

4. PRZYSPIESZONE METODY WYZNACZANIA MTIE

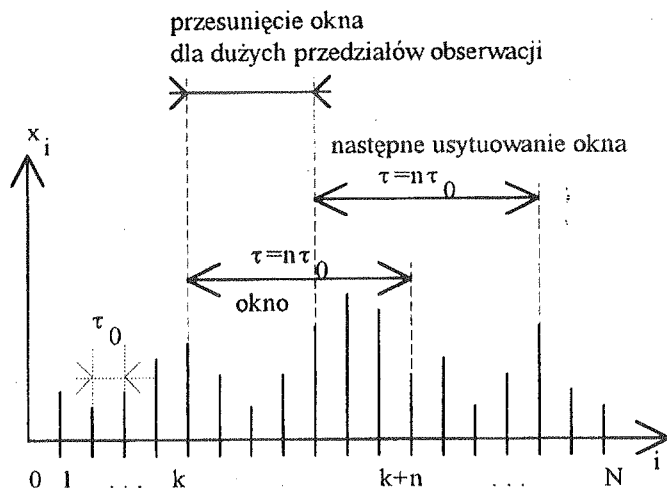
4.1. Metoda większego skoku

W pracy [3] zaproponowano i modyfikowano następnie w pracach [5, 6] metodę szacowania maksymalnego błędu przedziału czasu, nazwaną metodą większego skoku. Wyznaczanie MTIE dla małych, w stosunku do odstępów próbkowania τ_0 , przedziałów obserwacji odbywa się zgodnie z opisaną wyżej metodą, opartą bezpośrednio na definicji parametru. Dla przedziałów obserwacji większych od pewnej granicznej wartości opuszczana jest pewna liczba usytuowań okna, które przesuwa się o więcej, niż jedną próbkę. Dla przedziału obserwacji τ większego od określonej wartości (np. $300 \tau_0$) przesunięcie (skok) okna w ciągu wartości błędu czasu $\{x_i\}$ następuje o odstęp $10 \tau_0$, czyli o 10 próbek. Liczba położeń okna jest więc 10 razy mniejsza, niż przy przesunięciu o odstęp τ_0 . Kolejne zwiększenie skoku okna (do $100 \tau_0$), i co za tym idzie zmniejszenie liczby usytuowań okna, następuje dla przedziałów obserwacji większych od kolejnej wartości granicznej (np. $3000 \tau_0$). Wadą tej metody jest możliwość wprowadzenia błędów do wyników obliczeń. Wartość międzyszczytowa błędu czasu istniejąca w pewnym przedziale może zostać pominięta ze względu na skok okna i zauważona dopiero przy obliczeniach dla większego przedziału obserwacji. Zasada wyznaczania MTIE zgodnie z metodą większego skoku przedstawiona została na rysunku 3. Algorytm wykonywania obliczeń przedstawia diagram na rysunku 4. W zależności od wartości przedziału obserwacji τ i wartości granicznych τ_{g1} , τ_{g2} , na początku obliczeń ustalana jest wielkość przesunięcia (skoku). Dla stosunkowo dużych wartości τ procedura przeglądania okna, przedstawiona na diagramie w postaci bloku otoczonego linią przerywaną, wywoływana jest 10 lub 100 razy rzadziej, niż w przypadku obliczania zgodnie z definicją.

ość okna,
W Tabeli
operacji
ędu czasu

wzrostem
m liczby
glądanych
 τ będzie
dnak na
ymalnego
Zależności
go przez
rawdzenie
wymi [1]
0 s) oraz
większenia
ek, który
obliczeń.
n stopniu
wyników

e odstęp
w pomia-
wielkości.
większenie
wacji, co
ec wyniki



Rys. 3. Wyznaczanie wartości MTIE metodą większego skoku

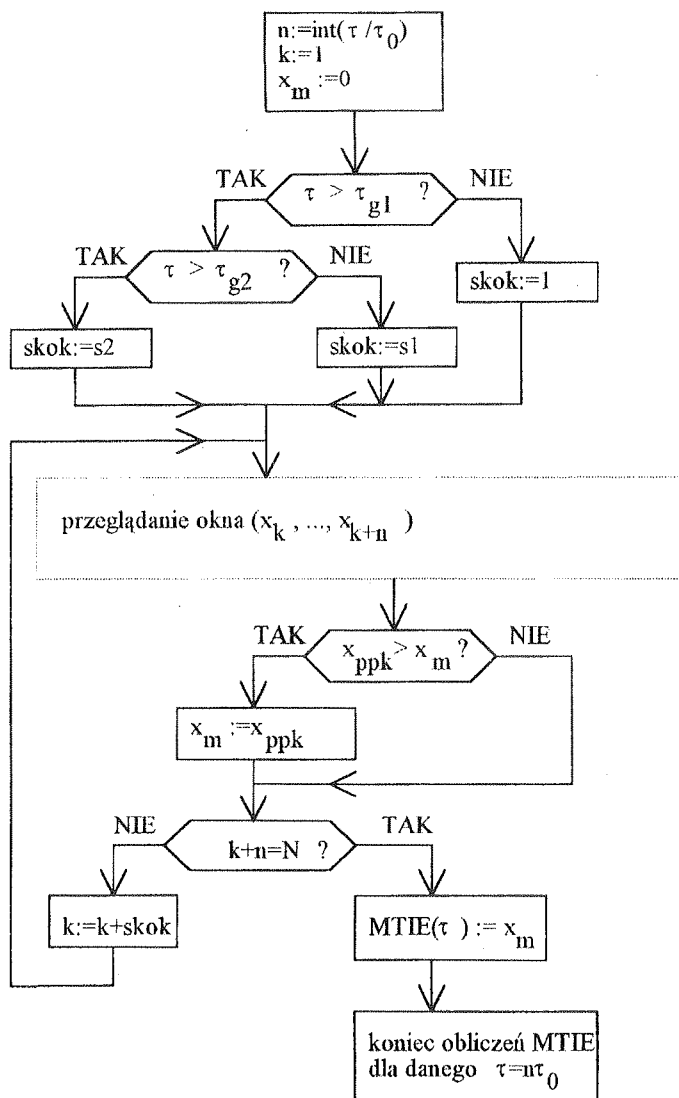
4.2. Metoda decyzji brzegowych

W pracy [7] zauważono, że w obliczeniach MTIE możliwe jest znaczące ograniczenie liczby operacji przeglądania okna tylko do koniecznych sytuacji. Trzeba natomiast poddać analizie wartości brzegowej okna, a więc pierwszej i ostatniej wartości w oknie. Na tej podstawie opracowana została metoda, którą autorzy niniejszego artykułu nazwali metodą decyzji brzegowych i która zostanie opisana poniżej. Można zauważyć, że w procesie przesuwania okna w ciągu próbek $\{x_i\}$, przy przesunięciu o jedną próbkę, zmiana zawartości okna następuje tylko na dwóch pozycjach: pierwsza wartość opuszcza okno, a pojawia się nowa wartość na końcu okna. Przeszukiwanie okna w nowym położeniu nie jest konieczne w każdym przypadku. Wystarczy porównać wartości brzegowe (nowa i stara wartość) z wartością maksymalną i minimalną dla poprzedniego położenia okna i na tej podstawie podjąć decyzję o poszukiwaniu w oknie wartości maksymalnej lub minimalnej. Na końcu okna, a więc rozpatrując nową wartość w oknie, możliwe są następujące sytuacje:

- 1) nowa > max;
- 2) nowa = max;
- 3) max > nowa > min;
- 4) nowa = min;
- 5) nowa < min.

Wartości min i max w powyższych relacjach odnoszą się do starego usytuowania okna. Z przedstawionych relacji wynika, że na końcu okna zdarzenia modyfikujące wartości ekstremalne w nowym usytuowaniu okna to zdarzenia 1) i 5). Jeśli wystąpi

Rys. 4. A
próbek w
okna, τ_{g1} ,



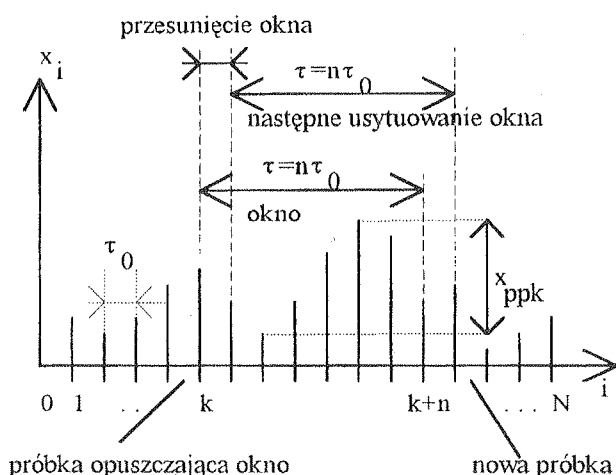
Rys. 4. Algorytm poszukiwania wartości MTIE metodą większego skoku. Oznaczenia: $n+1$ — liczba próbek w oknie, k — kolejny numer próbki, skok — przesunięcie okna, s_1, s_2 — wartości przesunięcia okna, τ_{g1}, τ_{g2} — wartości graniczne przedziału obserwacji, x_m — bieżąca maksymalna wartość między-szczytowa, x_{ppk} — wartość międzyszczytowa dla danego usytuowania okna,

jedno z tych zdarzeń, znana jest nowa wartość ekstremum. O przeszukiwaniu okna decyduje analiza wartości na drugim krańcu (początku) okna. Dla wartości, która właśnie opuściła okno, możliwe są następujące zdarzenia:

- 1) opuszczająca = max;
- 2) max > opuszczająca > min;
- 3) opuszczająca = min.

W nowym usytuowaniu okna wartość maksymalna i minimalna zmieniają się odpowiednio dla zdarzeń 1) i 3). W tych przypadkach jedna z wartości ekstremalnych opuściła okno i konieczne staje się znalezienie nowej wartości maksymalnej lub minimalnej. W przypadku wartości międzyszczytowa dla danego położenia okna nie zmieniła się. Biorąc pod uwagę obydwa krańce okna, konieczność przeszukiwania całego okna pojawia się tylko wtedy, gdy zajdą zdarzenia 1) i 3) na początku okna. W pozostałych przypadkach albo nie zachodzi zmiana wartości międzyszczytowej dla danego położenia okna, albo zmiana występuje, lecz znane są nowe wartości minimalne i maksymalne (zdarzenia 1) i 5) na końcu okna). We wszystkich tych przypadkach przeszukiwanie okna jest niepotrzebne. Na rysunku 5 przedstawione zostało umieszczenie analizowanych wartości w tej metodzie obliczeń. Algorytm obliczeń przedstawia diagram na rysunku 6. Wywołanie procedury przeglądania okna, oznaczonej na diagramie linią przerywaną, następuje tylko w przypadku opuszczenia okna przez wartość minimalną lub maksymalną.

Wadą metody decyzji brzegowych jest zależność jej efektywności od danych. W skrajnym przypadku, przy dużym odchyleniu częstotliwości badanego zegara, efekt tego odchylenia będzie wyraźnie dominował nad zjawiskami szumowymi w procesie błędu czasu. Wówczas ostatnia wartość prawie zawsze będzie wartością



Rys. 5. Umieszczenie analizowanych wartości w metodzie decyzji brzegowych wyznaczania MTIE

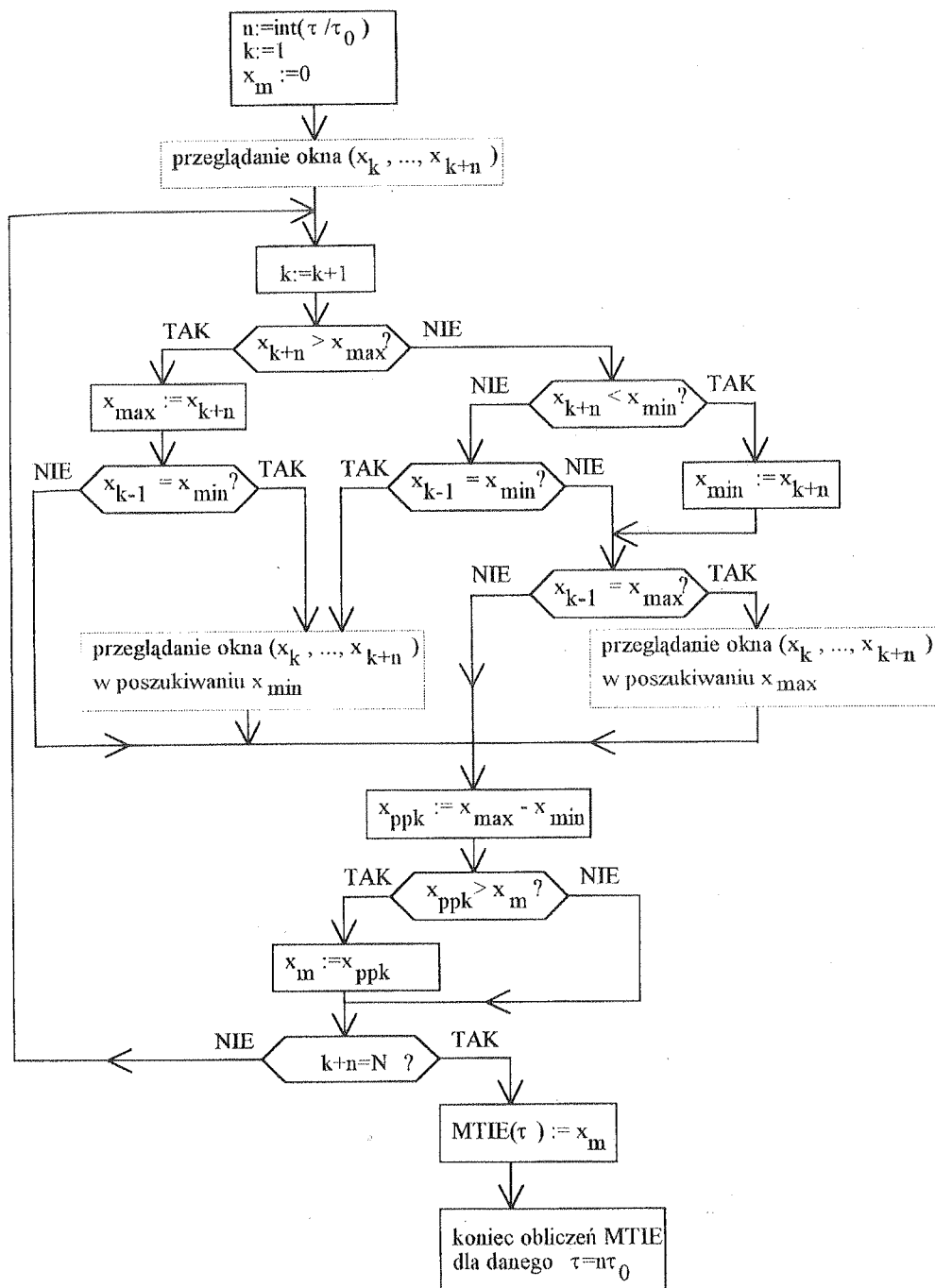
NIE

Rys. 6. Algorytm umieszczenia próbek w oknie usytuowania

ni okna
ści, która

eniają się
kstremal-
symalnej
położenia
nieczność
ia 1) i 3)
wartości
cz znane
kna). We
rysunku
metodzie
ywołanie
aną, na-
alną lub

danych.
o zegara,
mowymi
wartością

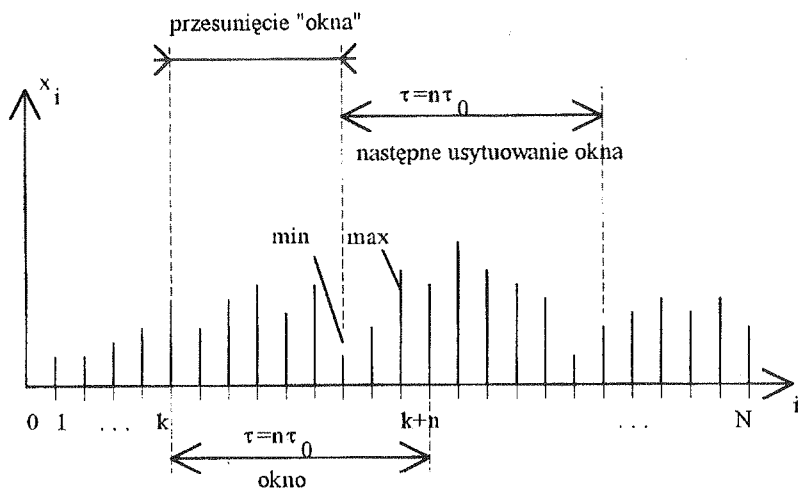


Rys. 6. Algorytm poszukiwania wartości MTIE metodą decyzji brzegowych. Oznaczenia: $n+1$ — liczba próbek w oknie, k — kolejny numer próbki, $x_{\max}$, $x_{\min}$ — wartość maksymalna i minimalna dla danego usytuowania okna, x_{k+n} — nowa wartość w oknie, x_{k-1} — wartość opuszczająca okno, x_m — bieżąca wartość międzyszczytowa, x_{ppk} — wartość międzyszczytowa dla danego usytuowania okna

maksymalną lub minimalną, co zmusi algorytm do poszukiwania nowego ekstremum dla danego usytuowania okna. W praktyce, przy relatywnie wysokiej jakości zegarów sieci synchronizacyjnej, sytuacja taka nie powinna mieć miejsca.

4.3. Metoda położenia wartości ekstremalnych

Inna metoda przyspieszająca wyznaczanie MTIE zasygnalizowana w pracy [7] opiera się na określeniu położenia wartości minimalnej i maksymalnej znalezionej dla danego usytuowania okna. Także w tej metodzie przeszukiwanie okna ograniczone jest tylko dla niektórych usytuowań okna. Na odcinku między najwcześniejszą (pierwszą) próbką w oknie (o indeksie k), a najbliższą jej próbką o wartości maksymalnej lub minimalnej znalezionej dla danego usytuowania okna, nie wystąpi już tego typu wartość dla nowego położenia okna. Nie jest zatem konieczne przeglądanie tego odcinka. Można więc wykonać przesunięcie okna w ten sposób, że pierwszą próbką w nowym położeniu będzie właśnie najbliższa próbka o wartości ekstremalnej z poprzedniego położenia okna. Przesunięcie takie przedstawione zostało na rysunku 7.



Rys. 7. Przesunięcie okna w metodzie położenia wartości ekstremalnych

Po znalezieniu próbek o wartości maksymalnej i minimalnej dla danego położenia okna, oprócz zapisu ich wartości należy także zapisać ich położenie w postaci kolejnego numeru próbki w ciągu wyników pomiaru lub chwili pobrania próbki w procesie pomiaru. Najbliższe k -tej próbce bieżącego usytuowania okna (najwcześniejsze) położenie wartości ekstremalnej wyznacza nowe położenie pierwszej próbki dla nowego usytuowania okna. W przypadku, gdy najbliższą próbką o wartości ekstremalnej jest właśnie próbka o indeksie k , przesunięcie okna wykony-

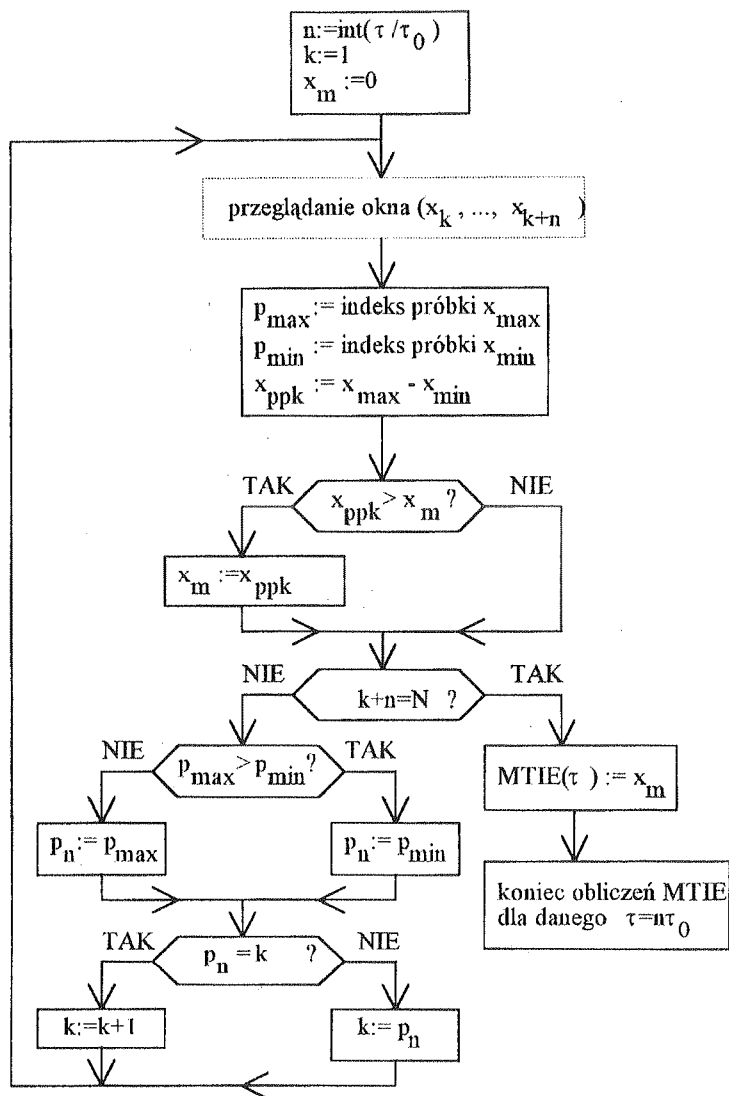
wane je-
szukane
szczytow-
ku 8.

Rys. 8. Alg-
 $n + 1$ — lic-
dla danego
pierwszego

ego eks-
wysokiej
iejsca.

wane jest o odstęp τ_0 . Po wykonaniu przesunięcia dla danego położenia okna szukane są wartości minimalna i maksymalna oraz obliczana jest wartość międzyszczytowa. Algorytm obliczania MTIE tą metodą przedstawiony jest na rysunku 8.

pracy [7]
alezionej
a ograni-
żeńszą
wartości
e wystąpi
nieczne
posób, że
wartości
stawione



położenia
w postaci
ia próbki
kna (naj-
pierwszej
ą próbką
wykony-

Rys. 8. Algorytm poszukiwania wartości MTIE metodą położenia wartości ekstremalnych. Oznaczenia: $n + 1$ — liczba próbek w oknie, k — kolejny numer próbki, $x_{\max}$, $x_{\min}$ — wartość maksymalna i minimalna dla danego usytuowania okna, $p_{\max}$, $p_{\min}$ — indeks próbki maksymalnej i minimalnej, p_n — położenie pierwszego w oknie ekstremum, x_m — bieżąca wartość międzyszczytowa, x_{ppk} — wartość międzyszczytowa dla danego usytuowania okna

Podobnie jak w przypadku metody decyzji brzegowych, efektywność tej metody jest zależna od danych. W przypadku dużego odchylenia częstotliwości prawie każda próbka o indeksie k może być próbką minimalną lub maksymalną. Wówczas przesunięcie okna w ciągu danych będzie odbywało się o jedną próbkę, co znacznie spowolni wykonywanie obliczeń. Jednak jak zauważono wcześniej, w praktyce taka sytuacja nie powinna się wydarzyć.

Metoda położenia wartości ekstremalnych sprawdza się szczególnie dla dużych przedziałów obserwacji obejmujących ponad kilkaset próbek. Przesunięcie do najbliższego ekstremum wykonywane przez okno może być porównywalne z jego szerokością. Ma to miejsce w przypadku, gdy próbki o wartościach ekstremalnych (lokalne minima i maksima) są oddalone od siebie. Pomijane jest wówczas od kilkuset do kilku tysięcy usytuowań okna. W działającej podobnie metodzie większego skoku przesunięcie okna nie jest większe, niż 3,3% szerokości okna.

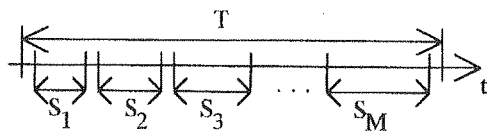
4.4. Metoda przedziałów rozłącznych

Opisane wyżej metody obliczania maksymalnego błędu przedziału czasu odnosiły się do definicji estymatora parametru podanej we wzorze (11). W pracy [8] zaproponowane zostało alternatywne podejście do zagadnienia pomiaru MTIE. W tej metodzie nie jest znajdowana maksymalna wartość międzyszczytowa poszukiwana po wszystkich przedziałach o szerokości τ istniejących w ciągu wyników pomiaru błędu czasu. W tym przypadku przeprowadzane jest M oddzielnych, niezależnych pomiarów wartości międzyszczytowej błędu czasu w pewnym czasie T obejmującym zarówno okresy pomiaru, jak i przerwy między nimi. Każdy z M pomiarów przeprowadzany jest dla rozłącznych przedziałów obserwacji S , ułożonych w ciąg geometryczny z ilorazem q danym wzorem:

$$q = \sqrt[M-1]{S_{\max}/S_{\min}} \quad (12)$$

gdzie $S_{\max}$ i $S_{\min}$ są odpowiednio maksymalnym i minimalnym przedziałem obserwacji. Rozkład przedziałów S przedstawiony został na rysunku 9.

Celem tej metody jest nie tyle wyznaczenie wartości maksymalnego błędu przedziału czasu dla określonego przedziału obserwacji τ , ile sprawdzenie, czy mierzony sygnał spełnia zalecone wymagania na MTIE. Wyniki obliczone w prze-



Rys. 9. Rozkład przedziałów S w metodzie przedziałów rozłącznych

Rys. 10.
punktó

działac
rzedną
wartoś
nanosi
spełnie
Przykła
bardzo
wartoś
Dotycz
wartos
przedzi
przedzi

Rys. 11.
stępująca

ej metody
wie każda
Wówczas
o znacznie
tytce taka

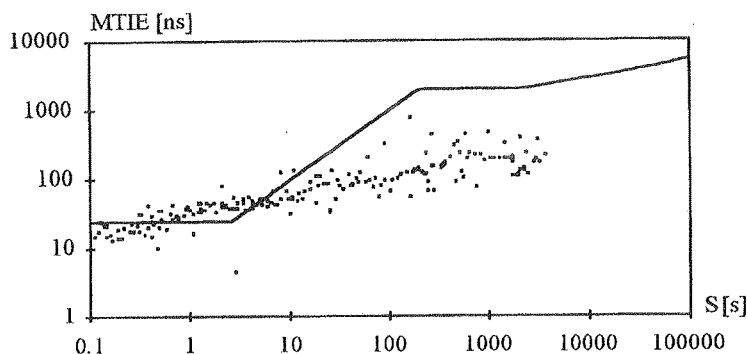
ła dużych
nięcie do
wnywalne
wartościach
Pomijane
działające
ksze, niż

łu czasu
rze (11).
gadnienia
wartość
ości τ is-
łku prze-
zczytowej
pomiaru,
zany jest
metryczny

(12)

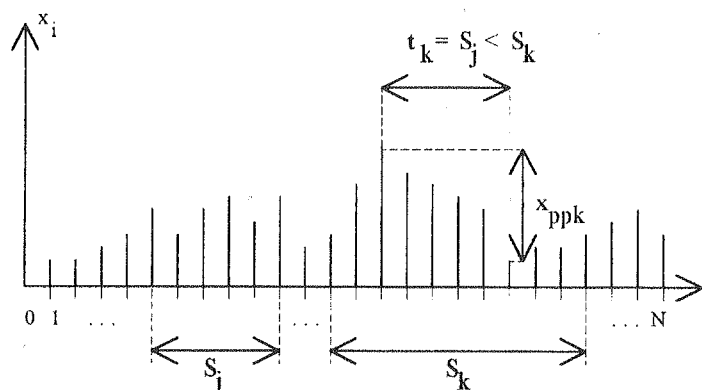
em obser-

ego błędu
zenie, czy
e w prze-



Rys. 10. Wyniki wyznaczania wartości MTIE metodą przedziałów rozłącznych w postaci „chmury punktów” z nałożoną wartością graniczną (normą) dla wtórnego źródła sygnału synchronizacji (SSU)

działach S naniesione są na wykres w postaci „chmury” punktów, których współrzędną na osi czasu jest szerokość przedziału S , a na osi wartości maksymalna wartość międzyszczytowa błędu czasu w tym przedziale. Na tak otrzymany wykres nanosi się wartości graniczne (normy) dla badanego zegara. Warunkiem koniecznym spełnienia norm jest pozostawianie wszystkich punktów poniżej wykresu normy [8]. Przykładowy wykres przedstawiony został na rysunku 10. Wadą tej metody jest bardzo duże przybliżenie otrzymanych wyników w stosunku do rzeczywistych wartości maksymalnego błędu przedziału czasu właściwych dla mierzonego sygnału. Dotyczy to szczególnie małych przedziałów, gdyż ewentualne pojawienie się dużej wartości międzyszczytowej na krótkim odcinku czasu podczas pomiaru dla dużego przedziału S spowoduje, że wartość ta nie zostanie uwzględniona dla krótszego przedziału S . Przypadek ten przedstawiony został na rysunku 11.

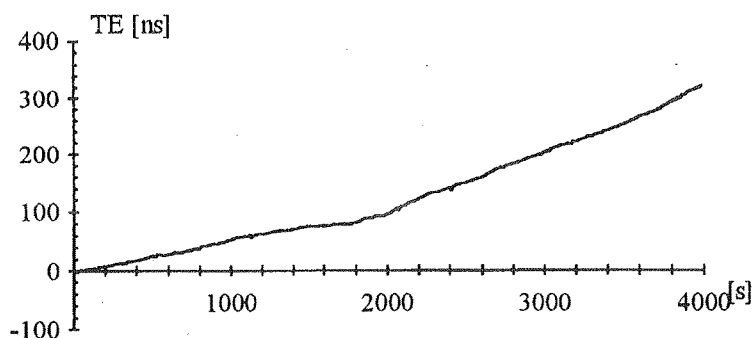


Rys. 11. Ilustracja błędu w metodzie przedziałów rozłącznych. Wartość międzyszczytowa x_{ppk} występująca dla pomiaru w przedziale S_k nie zostanie uwzględniona dla przedziału szerokości S_j , dla którego w rzeczywistości wystąpiła

Niewątpliwą zaletą metody przedziałów rozłącznych, podkreślaną przez autora pracy [8], jest możliwość wykonywania szacowania w czasie rzeczywistym w trakcie pomiaru, bez konieczności zapisywania długich ciągów próbek błędu czasu.

5. WYNIKI EKSPERYMENTU

Opisane w poprzednim rozdziale metody wyznaczania maksymalnego błędu przedziału czasu zostały wykorzystane do znalezienia wartości MTIE dla przykładowego pomiaru błędu czasu. Obiektem pomiaru był sygnał generatora wzorcowego systemu pomiarowego SP-2000 [3, 4], który to system został zaprojektowany i wykonany w Instytucie Elektroniki i Telekomunikacji Politechniki Poznańskiej. Pomiar wykonywany był z odstępem $\tau_0 = 1/30$ s przez okres 4000 sekund. Przebieg błędu czasu przedstawiony został na rysunku 12.



Rys. 12. Przebieg błędu czasu wykorzystanego do wyznaczania wartości MTIE

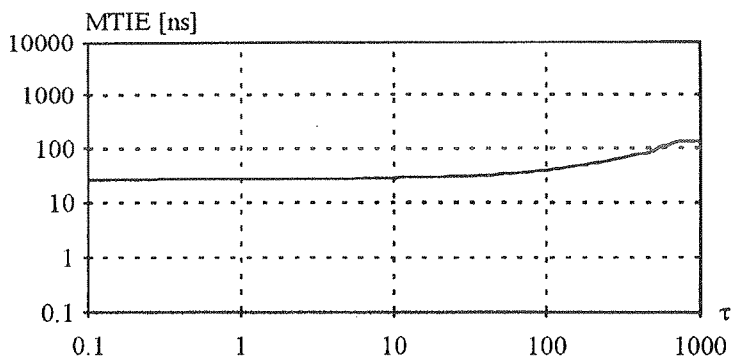
Wyniki pomiaru błędu czasu wykorzystane zostały do wyznaczenia wartości maksymalnego błędu przedziału czasu. Minimalny przedział obserwacji miał wartość $\tau = 0,1$ s, natomiast przedział maksymalny wartość $\tau = 1000$ s. Znaleziono zostało 80 wartości MTIE (po 20 na jedną dekadę), które przedstawione zostały w postaci wykresu o logarytmicznej skali wartości i logarytmicznej skali czasu na rysunku 13. Czas trwania obliczeń wykonanych metodą wynikającą bezpośrednio z definicji z wykorzystaniem komputera PC z procesorem Pentium i zegarem 200 MHz przekroczył 41 godzin. Przeprowadzone zostało także poszukiwanie wartości MTIE przy pomocy trzech metod uproszczonych: metody większego skoku, metody decyzji brzegowych oraz metody położenia wartości ekstremalnych. Wykres na rysunku 14 przedstawia przebieg błędu względnego dla obliczeń metodą większego skoku.

W Tabeli 2 porównany został czas trwania obliczeń przeprowadzonych opisanymi metodami. Pod względem długości czasu obliczeń najlepszą metodą jest metoda położenia wartości ekstremalnych. Wyznaczenie tą metodą 80 wartości MTIE zgodnie z przyjętymi założeniami zajęło nieco ponad 24 minuty. Jest to ponad czterokrotnie mniej, niż czas wyznaczania metodą większego skoku oraz ośmiokrot-

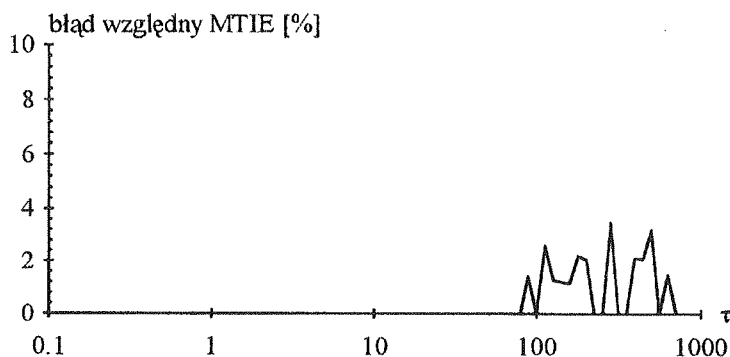
n
metoda
metoda
metoda
metoda
ekstrem

autora
prakcie

błędu
przy-
zorc-
owany
ńskiej.
zebieg



Rys. 13. Wyniki wyznaczania maksymalnego błędu przedziału czasu



Rys. 14. Przebieg błędu względnego dla wyznaczania MTIE metodą większego skoku

wartości
wartość
zostało
postaci
nku 13.
definicji
0 MHz
MTIE
decyzji
nku 14
u.
pisany-
metoda
MTIE
ponad
mikrot-

Tabela 2

Czas trwania obliczeń dla poszczególnych metod wyznaczania wartości MTIE

metoda wyznaczania MTIE	czas obliczeń dla $\tau \in (0, 1, 10)$	czas obliczeń dla $\tau \in (10, 1000)$	łączy czas trwania obliczeń
metoda klasyczna	41 min 24 s	40 h 25 min 36 s	41 h 7 min
metoda większego skoku	41 min 24 s	1 h 1 min 45 s	1 h 43 min 9 s
metoda decyzji brzegowych	7 min 17 s	3 h 10 min 1 s	3 h 27 min 18 s
metoda położenia wartości ekstremalnych	3 min 8 s.	21 min 17 s	24 min 15 s

nie mniej, niż metodą decyzji brzegowych. W porównaniu z czasem obliczeń przy użyciu metody klasycznej, wszystkie metody uproszczone uzyskały kilkudziesięciokrotne skrócenie czasu obliczeń. Czas wyznaczania wartości MTIE dla przedziałów obserwacji z zakresu od 0,1 s do 10 s metodą większego skoku jest identyczny, jak dla metody klasycznej, gdyż w tym zakresie poszukiwania prowadzone są bez modyfikacji. Wprowadzenie większego skoku dla przedziałów obserwacji z zakresu od 10 s do 1000 s spowodowało blisko czterdziestokrotne skrócenie czasu obliczeń w tym zakresie. Należy dodać, że wyniki obliczeń metodą większego skoku obarczone są błędami, nie większymi jednak, niż 4% dla przedziałów obserwacji powyżej 100 s. Wyznaczanie metodą decyzji brzegowych przebiegało bardzo szybko (ponad pięć razy szybciej niż dla metody klasycznej) dla pierwszego zakresu wartości przedziału obserwacji. Czas wyznaczania wzrósł jednak zdecydowanie dla przedziałów obserwacji od 10 s do 1000 s. Spowodowane zostało to koniecznością przeglądania dużej liczby próbek w przypadku opuszczenia okna przez próbkę o wartości ekstremalnej. Zastosowanie kombinowanej metody obliczeń — wyznaczanie metodą decyzji brzegowych dla przedziałów obserwacji w pierwszym zakresie i metodą większego skoku dla przedziałów obserwacji w zakresie drugim dałoby łączny czas obliczeń wynoszący 1 h 9 min 2 s. Czas ten jest krótszy, niż w przypadku obliczeń osobno każdą metodą, lecz ponad dwa razy dłuższy, niż w przypadku wyznaczania metodą położenia wartości ekstremalnych. Obliczenia z zastosowaniem tej metody były najszybsze zarówno dla pierwszego zakresu przedziałów obserwacji, jak i dla zakresu drugiego.

W przedstawionych w artykule metodach wyznaczania maksymalnego błędu przedziału czasu skrócenie czasu obliczeń uzyskuje się przez ograniczenie liczby usytuowań okna w ciągu danych, które poddaje się pełnemu badaniu. W artykule nie zajmowano się optymalizacją procesu badania okna. Zastosowano metodę pełnego przeglądu okna, która w przypadku okna o wielu lokalnych maximach i minimach może być jedyną metodą pozwalającą na wyznaczenie globalnych ekstremów okna niezależnie od właściwości badanego łańcucha losowego. Eksperyment obliczeniowy przeprowadzono na komputerze. Wykorzystanie układów ASIC zmieniłoby wyniki, lecz skutki rozważonych w artykule algorytmów byłyby takie same.

BIBLIOGRAFIA

1. ETSI draft ETS DE/TM-3017: *Generic requirements for synchronization networks*, Helsinki 1995
2. ITU-T Rec G.810. *Considerations on Timing and Synchronization Issues*
3. A. Dobrogowski, M. Jessa, M. Kasznia, K. Lange: *Pomiary i obliczenia parametrów sygnałów synchronizacji sieci telekomunikacyjnej*. Poznańskie Warsztaty Telekomunikacyjne PWT'96, materiały konferencyjne str. 132–139, Poznań 11–12 grudnia 1996
4. A. Dobrogowski, M. Jessa, M. Kasznia: *Podstawy pomiarów parametrów sygnałów i urządzeń synchronizacji sieci telekomunikacyjnej*. Przegląd Telekomunikacyjny i Wiadomości Telekomunikacyjne, nr 6 str. 315–325
5. M. Kasznia: Some Approach to Computation of ADEV, TDEV and MTIE. Proc. of 11th European Frequency and Time Forum, pp. 544–548, Neuchatel 4–6 March 1997

6. M. K. Krajo
Bydgo
7. A. D. Pozna
Pozna
8. S. B. Chara
pp. 90

CA

In the
interval e
provide a
signals is
topics in
describing

The p
time inter
equipmen
make tim

First
— time e
condition
parameter
calculation
presented
of the art
First met
approxima
method of
results of
calculation
time of ca
parameter

Using
results wit

Key words
interval er

6. M. Kasznia: *Aspekty obliczeniowe wyznaczania wartości parametrów sygnałów synchronizacji*. Krajowe Sympozjum Telekomunikacji KST'97, materiały konferencyjne tom A, str. 178—187, Bydgoszcz, 10—12 września 1997
7. A. Dobrogowski, M. Kasznia: *Metody obliczania maksymalnego błędu przedziału czasu*. Poznańskie Warsztaty Telekomunikacyjne PWT'97, materiały konferencyjne str. IV.3—1—IV.3—5, Poznań 10—11 grudnia 1997
8. S. Bregni: *Measurement of Maximum Time Interval Error for Telecommunications Clock Stability Characterization*. IEEE Trans. on Instrumentation and Measurement, October 1996 vol. 45, No. 5, pp. 900—906

A. DOBROGOWSKI, M. KASZNIA

CALCULATION OF MAXIMUM TIME INTERVAL ERROR OF TIMING SIGNALS

S u m m a r y

In the article the problems of maximum time interval error calculation are described. Maximum time interval error MTIE is a parameter describing timing signals in synchronization network. The signals provide all operations in digital network and network equipment with time scale. Proper quality of the signals is necessary for correct synchronization process of telecommunication network. One of the main topics in the exploitation of digital telecommunication network is a measurement of parameters describing long term changes in timing signals phase.

The parameters are: Allan deviation, modified Allan deviation, time deviation, root mean square of time interval error and maximum time interval error. To measure and calculate the parameters a lot of equipment and time must be involved. The article proposes and describes the methods, which enable to make time saving calculation of maximum time interval error.

First part of the article presents the basic definitions of the quantities characterizing timing signals — time error and time interval error and the definition of maximum time interval error. Then, the conditions of time error measurement are described. Based on results of time error measurement the other parameters are calculated, among them the maximum time interval error. The method of MTIE calculation based directly on the definition of the estimator given by international standards is presented. Unfortunately this method of MTIE calculation is dramatically time consuming. The next part of the article presents the details of the methods enabling the MTIE calculation in much shorter time. First method, called by authors the "bigger jump" method, enables to find with at least good approximation the real value of MTIE estimate. The other two methods — method of edge decision and method of extremum values position enable time effective calculation of MTIE estimate. In the article the results of an experiment are included. The experiment compares the time effectiveness of MTIE calculations using different methods. The results confirm the time effectiveness of proposed methods. The time of calculation was shorter from ten to hundred times than the time of calculation directly based on parameter definition. The alternative method of rough MTIE calculation was also included.

Using these methods makes in-service measurement relatively fast. Comparison of the obtained results with international standards makes possible to check the performance of the network.

Key words: synchronization network, timing signal, the parameters of timing signals, time error, time interval error, maximum time interval error.

the
place
char
syste
can
quas
appl

Key

One
started i
duced to
in part
of stron
notice a
have bee
developm
in the co
inexpensi
matic me

Concept of WWW-based distributed measurement system

TOMASZ STARECKI

Instytut Podstaw Elektroniki, Politechnika Warszawska

Received 1998.09.09

Authorised 1998.10.30

A concept of WWW-based distributed measurement system is presented. Due to use of the internet as an interface, virtually any number of the instruments can be located in any place without any distance limit. Under proper software, structure of the system can be changed while the system is running. Due to use of the WWW and Java language, the system controlling software is platform independent. The concept presented in the paper can be used in a variety of applications, of which the most important are wide-area quasi-real-time data acquisition systems, eg. country-wide pollution monitoring system. An application of the concept will be likely presented in a separate paper.

Key words: distributed measurement systems, automatic measurement systems, WWW.

INTRODUCTION

One can say that the history of automatic measurements and control was started in late 60's, when the standard of CAMAC interface [1] was introduced to the market. The technology has changed quite a lot since then. This is in part due to rapid progress in the field of electronics and in part a result of strong market demands for such automatic systems. Recently one can notice a fast progress in the field of wide-area data acquisition systems, that have been usually considered as extremely expensive. However, some substantial developments in communications (especially networking) have finally resulted in the concept of distributed measurement systems, making such systems relatively inexpensive and hence significantly broadening the range of applications of automatic measurements and control.

DISTRIBUTED MEASUREMENT SYSTEM AS A TERM

The term of "distributed measurement system" is being used more and more often nowadays, although it is difficult to find a precise definition of this term. In such a system we can consider the distribution of:

- a) measuring instruments and control devices,
- b) data processing,
- c) functions responsible for controlling the operation of the whole system.

In modern systems virtually any device is based on a microprocessor or microcontroller that implements at least some basic functions mentioned in (b) and (c). Hence, conditions given as (b) and (c) should be treated as necessary but not sufficient, otherwise virtually any modern automatic measurement system should be considered as a distributed system. This means, however, that there is still need for a factor that could be used while distinguishing between distributed and non-distributed measurement system. And certainly, the distance between the devices can be used as such an indicator. The problem is the value that should be assumed as the threshold. At first it can be set at a few tens of meters, cause it is quite easy to give an example of a system fulfilling this requirement that would be intuitively recognised as a distributed one. However, some popular interface standards allow maximum distance between the devices to be a few hundred meters or more (eg. for RS422, RS423, RS485 the value is set at 1200 m [2]). So it is possible to give an example of another system fulfilling conditions (b) and (c), that would be rather considered as a non-distributed one, despite the distance between some of its elements is over 100 m. This leads to the conclusion, that in the definition of distributed measurement system, the structure of the system is more important than the distance between its elements. Finally, distributed measurement systems could be defined as a system with distributed data processing and distributed functions responsible for controlling the operation of the whole system, with additional assumption that structure of the system and means of data transfer used in the system do not limit the distance between the devices.

COSTS OF THE SYSTEM

Standard elements of a distributed measurement system are:

- measuring instruments and control devices,
- system controller,
- network (means of data transfer between the elements of the system),
- software used in the process of measuring and control.

Until not so long ago, the main factor limiting the range of applications of distributed measurement systems was very high costs of establishing and operation of the system network. So that practical applications of such distributed systems

were li
systems
telecom
strong
any are
exchang
control
system
softwar
costs o
possible
real-tim
money
softwar
Moreov
Netscap
ting WV
importa
real-tim
process
WWW,
only cor
a Web b

There
recently
designs
based on

In th
is relativ
some da
the prob
server. 7

to be b
possibilit
or Java
case of
serving C

In the
bidirectio
implemen

and more
term. In

tem.

cessor or
ed in (b)
ssary but
at system
t there is
istributed
etween the
should be
cause it is
would be
interface
hundred
[2]). So it
) and (c),
distance
n, that in
system is
istributed
ted data
eration of
stem and
etween the

m),

cations of
operation
d systems

were limited to very few sophisticated and very expensive systems (eg. military systems) or some small, local area systems. However, strong development in telecommunication and global "explosion" of the Internet in last years resulted in strong decrease of costs and increase of availability of networks covering virtually any area (LANs, WANs, etc.). Hence, if one takes into consideration that data exchange in the system can be done by means of the Internet and the system controller can be implemented on an ordinary PC computer, the cost of the system is defined by the costs of measuring-and-control devices and the cost of software. Thus, a way of decreasing the final costs of the system is to reduce the costs of the system software by implementing it in WWW technology. It is possible, of course, to have some other software solutions — eg. based on real-time operating systems like QNX [3]. But use of WWW means no additional money spent on operating systems and similar software, as many standard software components used by WWW technology are available as freeware. Moreover, the freeware is not only browsers (including most popular ones like Netscape Navigator or MS Internet Explorer), but also some software implementing WWW servers [4, 5] and tools for WWW software development [6]. The most important argument for using WWW technology is platform independency, while real-time operating system mechanisms are usually limited to some particular processors — eg. X86 in the case of QNX [7]. What is more — in the case of WWW, the system can be controlled from virtually any place in the words. The only condition is to have a computer connected to the Internet and equipped in a Web browser.

SOME ASPECTS OF WWW-BASED CONTROL

There are many single Web-based measuring-and-control devices introduced recently to the public [8, 9], some companies offer development tools for such designs [10], but there is still no information on distributed measurement systems based on WWW technology.

In the case of a single device, Web-based remote control of such a device is relatively easy. If it is a free-running device, that has to periodically output some data only (with no option of remote programming of its setup, etc.), the problem resolves itself into forcing the device to act as a local WWW server. The task is a little bit more complicated if the device is assumed to be back-programmable, because if using HTML language [11] the only possibility of invoking an action on a server is to run a CGI script [12] or Java applet [13]. Usually it is easier if done with CGI script, so in the case of bidirectional data exchange, the device should be also capable of serving CGI scripts [9].

In the case of a whole WWW-based system, it is not enough just to support bidirectional data transfer (from and to the terminal devices). The software must be implemented in such a way that it allows simultaneous management (reading the

measurement results and controlling the devices) of multiple devices and in a way leading to the highest reliability and speed of the system. In particular one must take into consideration possible effects of temporary unaccessibility of some URLs and transmission dropouts — these effects are not rare in the case of wide area systems and they can cause a hangup of the software running on the device calling the server. The main problem, however, is implementing an automatic control of the system, as HTML is a language used for Web pages definition only, and it does not contain any direct mechanisms that would allow creation of automatic Web pages, i.e. the pages that would automatically (eg. periodically) update their contents and invoke some actions on other URLs. In order to create such pages one must use some indirect mechanisms allowed by HTML, as:

- a) CGI scripts [12],
- b) JavaScript inlines [14],
- c) Java applets [13].

CGI script is used for invoking an action on the server, so it can be treated as a form of procedure located and run on the sever, but called (with parameters) by any computer connected to the network. In particular CGI script can be implemented as a batch or EXE file, thus making possible to perform any possible operation on the hardware.

JavaScript inline is a kind of program enclosed in a Web page in the form of a human readable code, which is run (interpreted) on the computer on which the Web page containing JavaScript inline is browsed. JavaScript is a simplified programming language (at least from the system programming point of view) and offers very poor graphic support (quite important for the tasks as presentations of the measurement results, etc.) and making interaction with the system resources relatively difficult if possible. In particular JavaScript does not allow write file operations, not mentioning any operations on serial (COM) or parallel ports (LPT), that are quite often used as interfaces by measuring instruments and control devices.

Java applets, just as JavaScript inlines, are a kind of code run on the computer on which the Web page containing the applet is browsed. But, on contrary to JavaScript, Java applets contain code that has been already compiled and this results in faster execution. Comparing both languages from the system programming point of view, the main advantage of Java over JavaScript is that Java is an entire object programming language with many libraries, including many system mechanisms making interaction with system resources relatively easy. That is why Java applets are much stronger software tool than JavaScript inlines. Probably the most important advantage of Java in the case of WWW-based distributed measurement systems, is RMI (Remote Method Invocation) library. The library is dedicated to distributed network applications and contains many mechanisms which make such applications easy to implement and strongly resistant to problems caused by transmission errors, etc. [15].

Assu
control

Fig. 1. S
 U_i — mea

easily no
tasks:

- a) th
fu
- b) th
re
- c) th
to
(Z
Z
di

In th
devices,
P between
used due
controlled
devices, t
devices t
devices a
The inter
system. 7

STRUCTURE OF THE SYSTEM AND DATA FLOW

Assuming there is a system of a relatively simple structure, consisting of a system controller K and number of measuring-and control devices U (Fig. 1a), one can

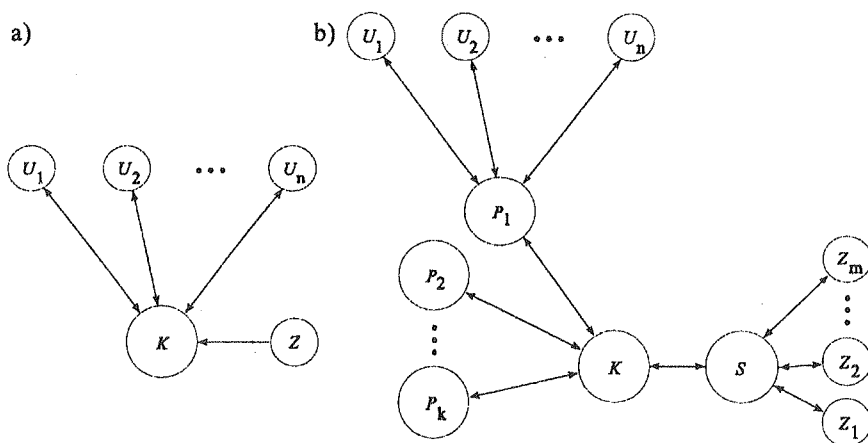


Fig. 1. Structure of a distributed measurement system: a) simple structure, b) complex structure; U_i — measuring-and-control devices, K — controller, Z, Z_i — external computer(s), P_j — intermediate servers, S — server

easily notice that the whole process of the system control consists of at least three tasks:

- the devices U_i must implement some individual measuring and control functions,
- the controller K must be able to control the devices by sending setup data and reading measurement results,
- the controller K must be also capable of working as a server in order to make the final results of data processing available to another computer (Z) that can be also used to control the system (usually the computer Z is external to the system, but the function can be also implemented directly on the controller K).

In the case of biggersystems with great number of measuring-and-control devices, the more likely is the structure with one or more intermediate servers P between the controller K , and the devices U (Fig. 1b). The intermediate servers are used due to limited number of devices that can be simultaneously served by a single controller. Thus, instead of direct data exchange with a great number of terminal devices, the controller K can send a command resulting in service of some terminal devices to an intermediate server P . Then the server can exchange the data with the devices and send the results (usually after some data processing) to the controller. The intermediate servers can be mirrored in order to achieve higher reliability of the system. The main role of the optional server S is to release the controller from

simultaneous servicing of many computers Z , that is possible in some public systems (eg. weather forecast system). A relatively simple but effective separation of the controller K from the computers Z can be achieved if the final data are periodically copied from the controller K to the server S (rare transmissions between the controller and the server), while leaving the service of the computers Z_i to the server S .

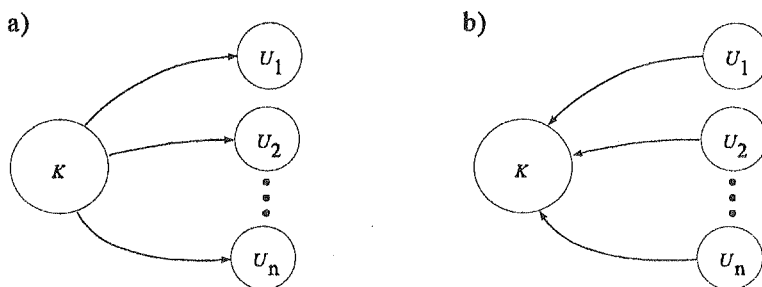


Fig. 2. Data exchange between the controller K and the devices U_i : a) initiated by the devices (bottom-to-top data flow), b) initiated by the controller (top-to-bottom data flow)

Even in the case of a simple structure as shown at Fig. 1a, one can notice that data exchange between the controller K and the devices U can be implemented in two different ways:

- by polling the devices by the controller (top-to-bottom data transfer) — Fig. 2a,
- by automatic (periodical) data transfer from the devices to the controller (bottom-to-top data exchange) — Fig. 2b.

Both methods have some good points and drawbacks. In the case of the solution from Fig. 2a number of data transfers is smaller, as the device is polled only when the controller needs data from the particular device and only the data that are needed are sent. The drawback of the method is that the system is sensitive to transmission dropouts, etc. Such events can cause serious disruption of service and slow down or even hangup the system. The main advantage of the solution is that any data exchange is started by the controller which requires only the data necessary and exactly at the moment they are needed. Some potential problems caused by transmission breaks, etc. can be reduced by running a few copies of the browser simultaneously, so that each copy will service only a few (eg. single) connections.

In the case of bottom-to-top data exchange (Fig. 2b) the worst result of connection break is that a single device is stopped, not the whole system. However, the method is featured with excessive data flow (usually not all the data are required), strong limitation in the maximum number of devices that can be simultaneously serviced by the controller (by the WWW server software) and much higher difficulties in controlling the devices (especially in real-time). Thus structure

The
systems

- u
- l
- c
- c
- s
- t
- e
- u
- d

In the
meteorological
stations
point of
top data
ces relat
station s
continuo
most im
data flow
in the sy

In de
applets t
RMI m
distribut

Due
systems
quasi re
system. A
Work su
tion — M

Special ack
many valu

ne public
eparation
inal data
missions
omputers

with the bottom-to-top data flow seems to be more reliable, but as the transmissions are initiated by the devices (not by the controller) the structure is suitable rather for the system with the devices of equal priority.

CONCLUSIONS

There is no doubt that implementation of WWW-based distributed measurement systems is possible. The main advantages of such systems are:

- unlimited number of the devices and the distance between them,
- low software costs,
- due to use of the WWW and Java language the system is platform independent (it is possible to run the software on different platforms within the same system),
- the system can be controlled from any place in the world by a computer equipped in a Web browser and connected to the Internet,
- under proper software, the structure of the system can be changed (eg. some devices can be added or disconnected) while the system is running.

ttom-to-top

notice that
mented in

transfer)

controller

e solution

only when

that are

nsitive to

ervice and

on is that

necessary

caused by

e browser

ections.

result of

However,

data are

t can be

and much

structure

In the case of huge data acquisition systems that have regular structure (eg. meteorological data measurement system) in which dropout of a few terminal stations (measuring instruments) or even intermediate servers is not critical from the point of view of the whole system functionality, it is better choice to use bottom-to-top data flow structure (Fig. 2b). This makes the system more resistant to disturbances related to data transmission problems, and the need of changing the terminal station setup (reverse data flow) is very rare. In relatively local systems in which continuous operation of all devices is vital to proper operation of the system, the most important feature is a highly reliable network. In such a case top-to-bottom data flow (Fig. 2a) can be used, as the system hangup (or another strong disturbance in the system operation) caused by the device or network failure is very unlikely.

In development of the software used for the system control it is better to use Java applets than other mechanisms as JavaScript inlines or CGI scripts. The reason is RMI mechanisms embedded in Java, that allow easy development of reliable distributed network applications.

Due to the advantages mentioned above, WWW-based distributed measurement systems can be used in variety of applications of which the most important are global quasi real-time data acquisition systems eg. whole-country pollution monitoring system. An application of the concept will be likely presented in a separate paper. Work supported in great part by NIME (National Institute of Multimedia Education — Makuhari, Japan).

ACKNOWLEDGMENTS

Special acknowledgments to dr Ronald Oliver from Edith Cowan University for serious thinking and many valuable remarks during several evening sessions at NIME.

REFERENCES

1. W. Nowakowski: *Systemy interfejsu w miernictwie*. Warszawa, WKiŁ 1987, pp. 214–243
2. *Overview of the Interface Standards in: Data Transmission Circuits Data Book Vol. 1*, Texas Instruments 1995, pp. 4–11 – 4–12
3. D. Hildebrand: *An Architectural Overview of QNX*,
<http://www.qnx.com/literature/whitepapers/archoverview.html>
Last update: 10 Sep 1998
4. R.B. Denny: *Windows httpd 1.4 World Wide Web Server*,
<http://tech.west.ora.com/win-httpd/>
Last update: 01 Aug 1997
5. *Freeware HTTP Server for Windows NT*,
<http://emvac.ed.ac.uk/html/internet/toolchest/https/software.htm>
Last update: 11 Aug 1997
6. D. Connolly: *World Wide Web and HTML Tools*,
<http://www.w3.org/Tools/Overview.html>
Last update: 23 Jan 1997
7. *QNX System Architecture; Chapter 2 — Microkernel*,
<http://www.qnx.com/literature/qnx-sysarch/microkernel.html>
Last update: 25 Mar 1998
8. T. Starecki: *WWW-based video camera remote control*, Technical Reports of IEICE, ET96-121 (1997-03), pp. 59–64; also <http://www.ipe.pw.edu.pl/~tomi/papers/>
9. Real-time data from Phar Lap's miniature weather station can be found at:
<http://smallest.pharlap.com/>
10. N. Nicolaisen: *Let's get small*, Byte, April 1997, pp. 28–32
11. D. Raggett, A. Le Hors, I. Jacobs: *HTML 4.0 Specification (W3C Recommendation, revised on 24 April 1998)*
<http://www.w3.org/TR/REC-html40/>
12. *The Common Gateway Interface* (interface specification and some examples),
<http://hoo.hoo.ncsa.uiuc.edu/cgi/>
Last update: 07 Mar 1996
13. *The Java™ Tutorial: A practical guide for programmers*,
<http://java.sun.com/docs/books/tutorial/>
Last update: 07 Sep 1998
14. *JavaScript Guide*,
<http://developer.netscape.com/docs/manuals/communicator/jsguide4/index.htm>
Last update: 26 Nov 1997
15. *Java™ Remote Method Invocation (RMI)*:
<http://java.sun.com/products/jdk/1.2/docs/guide/rmi/index.html>
Last update: 02 Mar 1998

T. STARECKI

KONCEPCJA ROZPROSZONEGO SYSTEMU POMIAROWEGO
STEROWANEGO PRZEZ WWW

Streszczenie

W artykule przedstawiono koncepcję rozproszonego systemu pomiarowego sterowanego za pomocą WWW. Dzięki wykorzystaniu Internetu, jako interfejsu łączącego poszczególne urządzenia, praktycznie nie występują ograniczenia liczby urządzeń występujących w systemie, ani odległości między nimi. Przy odpowiedniej konstrukcji oprogramowania struktura systemu może być modyfikowana podczas pracy

14 – 243
I, Texas

systemu. Dzięki wykorzystaniu WWW i języka Java oprogramowanie sterujące pracą systemu jest niezależne od platformy sprzętowej. Przedstawiona w artykule koncepcja systemu pomiarowego może być wykorzystana do wielu zastosowań, z których najważniejszymi wydają się bardzo rozległe systemy zbierania danych pracujące w czasie prawie-rzeczywistym, jak np. ogólnokrajowy system monitorowania zanieczyszczeń środowiska. Przykład praktycznej realizacji koncepcji rozproszonego systemu pomiarowego sterowanego za pomocą WWW zostanie przedstawiony w oddzielnym artykule.

Słowa kluczowe: rozproszone systemy pomiarowe, automatyczne systemy pomiarowe, WWW.

Work supported in great part by NIME (National Institute of Multimedia Education — Makuhari, Japan).

ET96-121

mendation,

za pomocą
praktycznie
nimi. Przy
czas pracy

tr

to
m
wy
na
to:
W
EM
zy:

St
ter

Tran
mocy z
tranzyst
traktow
W el
elektrycz
tego cie
w strukt
elementu
romodel
powinier

Modelowanie nieizotermicznych stałoprądowych charakterystyk tranzystora Darlingtona mocy za pomocą programu APLAC

JANUSZ ZARĘBSKI, KRZYSZTOF GÓRECKI

Katedra Radioelektroniki Morskiej, Wyższa Szkoła Morska w Gdyni

Otrzymano 1998.09.10

Autoryzowano 1998.11.03

W pracy sformułowano trzy różne postacie elektrotermicznego makromodelu tranzystora Darlingtona mocy (EMTD) dla programu APLAC. Pierwszy z opracowanych makromodeli opiera się na wbudowanych elektrotermicznych modelach elementów składowych występujących w strukturze tranzystora Darlingtona mocy. Przy konstruowaniu dwóch następnych, zastosowano tzw. hybrydowe oraz globalne elektrotermiczne modele tranzystora bipolarnego stanowiące obwodową modyfikację wbudowanych modeli tego elementu. Wyniki symulacji, zweryfikowane doświadczalnie, wykazały wyraźną poprawę dokładności EMTD w przypadku wykorzystywania zmodyfikowanych elektrotermicznych modeli tranzystora bipolarnego.

Słowa kluczowe: tranzystor Darlingtona, makromodel tranzystora mocy, zjawiska elektrotermiczne.

1. WSTĘP

Tranzystor Darlingtona mocy jest popularnym półprzewodnikowym elementem mocy zawierającym typowo kilka elementów półprzewodnikowych takich jak: tranzystory bipolarne, rezystory i diody p-n. Z tego też powodu może on być traktowany jako układ scalony małej skali integracji.

W elemencie takim, pracującym typowo w warunkach wydzielania dużych mocy elektrycznych, zamienianych na ciepło przy nieidealnych warunkach odprowadzania tego ciepła do otoczenia, można oczekiwać znacznych nadwyżek temperatury w strukturze. Ponieważ temperatura w sposób zasadniczy wpływa na właściwości elementu (jego charakterystyki), dlatego model elementu składowego, czy też makromodel całego układu elektrycznego — reprezentującego tranzystor Darlingtona, powinien uwzględniać efekt samonagrzewania oraz wzajemne sprzężenia termiczne

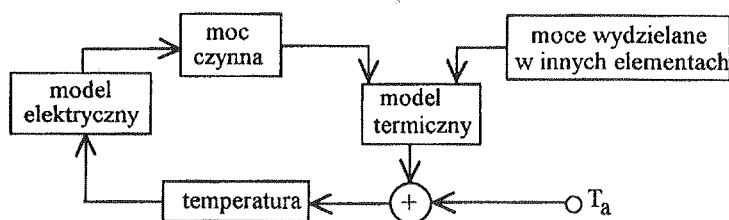
między elementami rozważanej struktury scalonej. Taki model/makromodel nazywa się modelem/makromodelem elektrotermicznym, a otrzymane stąd charakterystyki nazywane są charakterystykami nieizotermicznymi. Oczywiście postać elektrotermicznego modelu elementu, czy też makromodelu układu (tranzystor Darlingtona mocy) musi uwzględniać ograniczenia programu, np. ograniczenia wykorzystywanych algorytmów obliczeniowych, w którym ten model/makromodel będzie zaimplementowany.

Celem pracy jest przedstawienie możliwości zastosowania programu APLAC opracowanego na uniwersytecie w Helsinkach [1, 2] do modelowania nieizotermicznych charakterystyk stałoprądowych tranzystora Darlingtona mocy. Sformułowano trzy różne postacie elektrotermicznego makromodelu tranzystora Darlingtona (EMTD) dla programu APLAC wykorzystując do tego celu wbudowane modele elementów składowych, modele stanowiące ich modyfikację (modele hybrydowe) oraz modele własne (autorskie) opisane w pracy [3], pozwalające na wyznaczenie zarówno nieizotermicznych charakterystyk zaciskowych, co jest ważne z punktu widzenia inżyniera – projektanta układów elektronicznych lub energoelektronicznych jak i ocenę rozkładu mocy i temperatury wnętrza struktury. Wyniki analiz tranzystora Darlingtona BU323A zweryfikowano doświadczalnie przez pomiar prądów i napięć zaciskowych.

2. MODELE ELEMENTÓW W PROGRAMIE APLAC

Program APLAC przeznaczony jest głównie do analizy (stałoprądowej, stanów przejściowych, małosygnałowej, Fouriera wrażliwościowej, szumowej itp.) układów elektronicznych, optymalizacji parametrów układów elektronicznych oraz sterowania magistralą pomiarową IEEE 488.

Istotne znaczenie w procesie symulacji odgrywają wbudowane modele elementów półprzewodnikowych oraz możliwość tworzenia przez użytkownika własnych modeli elementów. W programie APLAC występują dwa typy modeli elementów: modele elektryczne (z parametrami zależnymi od temperatury) i modele termiczne. Łączne wykorzystanie tych modeli (ich połączenie) tworzy tzw. model elektrotermiczny elementu. Model taki (rys. 1) pozwala na wyznaczanie prądów i napięć zaciskowych elementu oraz temperatury wnętrza w tym elemencie.



Rys. 1. Schemat blokowy elektrotermicznego modelu elementu

Prog
wodnik
(3 pozio
tranzyst
wymieni
mi (opc
w złącza
cyfrowy
gacji po
nikowy
i źródł
ności, in
rezystan
może by
źródło m

W p
(model
densator
ny w p
elementa
pozwała
peraturę
z punktu
formułow
modele e

Każd
tranzysto
nych war

l nazywa
terystyki
elektroter-
rlingtona
zysywa-
zie zaim-

APLAC
otermicz-
ułowano
rlingtona
e modele
rydowe)
znaczenie
z punktu
ktronicz-
ki analiz
z pomiar

j, stanów
układów
sterowa-

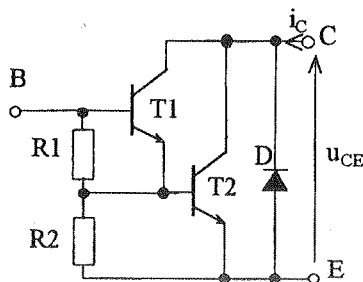
elementów
ch modeli
v: modele
e. Łączne
termiczny
iskowych

Program APLAC zawiera modele elektryczne następujących elementów półprzewodnikowych: dioda, tranzystor bipolarny, tranzystor JFET, tranzystor MOSFET (3 poziomy modelu), tranzystor DMOS, tranzystor MESFET (4 poziomy modelu), tranzystor BSiM, heterozłączowy tranzystor bipolarny, tranzystor IGBT. Spośród wymienionych elementów jedynie tranzystor bipolarny opisany jest dwoma modelami (opcjonalnie): Ebersa-Molla (z uwzględnieniem efektu powielania lawinowego w złączach) i Gummela-Poona. Dostępne są także wbudowane modele układów cyfrowych, które realizują zadaną funkcję logiczną oraz uwzględniają czasy propagacji poszczególnych funktorów. Modele wymienionych elementów półprzewodnikowych stanowią zawsze połączenie modelu złącza oraz elementów biernych i źródeł sterowanych. Modele elementów biernych uwzględniają pasożytnicze pojemności, indukcyjności i upływności. Źródła napięciowe i prądowe mają różną od zera rezystancję i konduktancję wewnętrzną, natomiast wydajność źródeł sterowanych może być opisana za pomocą złożenia funkcji elementarnych, jednak konkretne źródło może być sterowane tylko prądami lub tylko napięciami.

W programie APLAC jedynie wybrane elementy: dioda, tranzystor bipolarny (model Ebersa-Molla), tranzystor JFET, tranzystor MESFET oraz rezystory, kondensatory, induktry i źródła sterowane, mają wbudowany skupiony model termiczny w postaci obwodowej. Model taki uwzględnia własne i wzajemne (między elementami) sprzężenia termiczne, a także zjawiska konwekcji i radiacji oraz pozwala na wyznaczenie nadwyżki temperatury wnętrza elementu ponad temperaturę otoczenia przy znanej mocy czynnej wydzielanej w elemencie. Ważną z punktu widzenia modelowania i analizy jest opcja DefModel pozwalająca na formułowanie własnych modeli w postaci podukładów, wykorzystując do tego celu modele elementów wbudowane w programie APLAC.

3. SPOSÓB FORMUŁOWANIA ELEKTROTERMICZNEGO MAKROMODELU TRANZYSTORA DARLINGTONA

Każdy układ tranzystora Darlingtona można przedstawić jako połączenie kilku tranzystorów bipolarnych (co najmniej dwóch) oraz opcjonalnie rezystorów o różnych wartościach, bocznikujących złącza baza – emiter poszczególnych tranzystorów



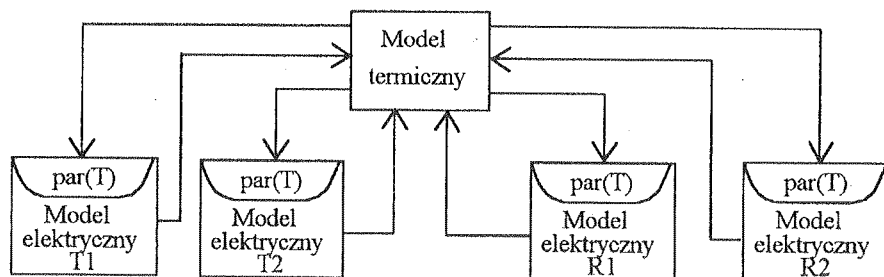
Rys. 2. Struktura wewnętrzna tranzystora BU323A

oraz diod usprawniających i diody zabezpieczające. W pracy, będzie rozpatrywana struktura tranzystora Darlingtona mocy przedstawiona na rys. 2, odpowiadająca tranzystorowi BU323A. Zatem elektryczny makromodel tranzystora Darlingtona mocy można łatwo sformułować w postaci pliku połączeń (zgodnie z topologią struktury — rys. 2) modeli elementów występujących w strukturze, których parametry zależą od ustalonej temperatury. Jednakże w rzeczywistości parametry modeli elementów składowych w strukturze tranzystora Darlingtona zależą od temperatury wnętrza danego elementu T_j , określonej zależnością

$$T_j = T_a + \sum_i R_{thij} \cdot p_i \quad (1)$$

gdzie T_a jest temperaturą otoczenia, R_{thij} to rezystancja termiczna między elementem i -tym oraz j -tym (własna rezystancja termiczna gdy $i = j$), p_i — moc wydzielana w elemencie j -tym.

Tak więc elektrotermiczny makromodel tranzystora Darlingtona (EMTD) zawiera 4 elementy (2 tranzystory i 2 rezystory, natomiast znaczenie diody w obszarze aktywnym normalnym jest nieistotne), a każdy element posiada swój model elektryczny oraz model termiczny. Stąd ogólna struktura EMTD pokazana na rys. 3 zawiera 4 bloki reprezentujące obwodowe modele elektryczne tranzystorów i rezystorów z parametrami zależnymi od temperatury wnętrza (T_j) oraz obwodowy model termiczny uwzględniający własne i wzajemne rezystancje termiczne oraz moce elektryczne wydzielane w elementach składowych struktury. Jak zatem widać, podstawowe znaczenie w procesie modelowania (analizy) mają dokładne modele elektryczne tranzystora bipolarnego i rezystora, poprawnie uwzględniające zjawiska fizyczne w elemencie, mające wpływ na właściwości zaciskowe elementu (prądy i napięcia) w rozpatrywanym zakresie pracy elementu. Również ważne jest prawidłowe sformułowanie modelu termicznego uwzględniającego samonagrzewanie w poszczególnych elementach jak i wzajemne oddziaływania termiczne między tymi elementami.

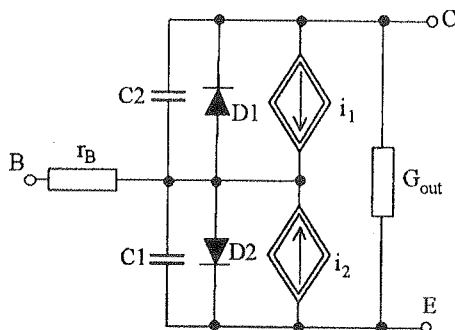


Rys. 3. Ogólna struktura EMTD

Przy formułowaniu obydwu typów modeli (elektrycznych i termicznych) należy pamiętać o tym, aby parametry występujące w modelach były łatwe do „pozyskiwania” przez użytkownika, np. z katalogu, pomiarów itp.

4. EMTD OPARTY NA MODELACH ELEMENTÓW WBUDOWANYCH W PROGRAMIE APLAC

Wbudowany w programie APLAC model elektryczny tranzystora bipolarnego do analizy elektrotermicznej jest zmodyfikowanym modelem Ebersa-Molla w postaci obwodowej przedstawionej na rys. 4.



Rys. 4. Reprezentacja obwodowa modelu elektrycznego tranzystora bipolarnego w programie APLAC

W omawianym modelu składowe prądów wstrzykiwanych przez złącza (reprezentowane przez diody D1, D2) są modelowane w identyczny sposób jak diody w programie SPICE (jest uwzględniona składowa idealna prądu, składowa prądu generacyjno-rekombinacyjnego, powielanie lawinowe oraz rezystancja szeregową). Składowe prądów zbieranych są reprezentowane przez źródła sterowane o wydajności i_1 , i_2 , przy czym

$$i_1 = \alpha \cdot I_S \cdot \left(\exp\left(\frac{u_{BE}}{V_T}\right) - 1 \right) \quad (2)$$

$$i_2 = \alpha_R \cdot I_S \cdot \left(\exp\left(\frac{u_{BC}}{n_V \cdot V_T}\right) - 1 \right), \quad (3)$$

gdzie α , α_R , I_S , n_V , V_T to parametry modelu.

Występujące w modelu rezystancje modelują odpowiednio rezystancję szeregową bazy (rezystor r_B) oraz efekt Early'ego reprezentowany przez nieliniową konduktancję G_{out} , przy czym

$$G_{out} = G_0 + n_g \cdot \left(\frac{dI_{BE}}{dV_{BE}} + \frac{dI_{BC}}{dV_{BC}} \right), \quad (4)$$

gdzie G_0 i n_g są parametrami modelu tranzystora.

Prąd każdej z diod i_D przedstawiono jako sumę prądów reprezentujących model idealny (i_{D1}) z rezystorem szeregowym oraz współczynnikiem emisji N oraz prąd przebicia (i_{D2}), opisane odpowiednio wzorami:

$$i_{D1} = I_S \cdot \left[\exp\left(\frac{u_D - r_S \cdot i_D}{N \cdot V_t}\right) - 1 \right] \quad (5)$$

$$i_{D2} = IBV \cdot \exp\left(-\frac{u_D - r_S \cdot i_D + V_{BV}}{NBV \cdot V_t}\right) + IBNL \cdot \exp\left(-\frac{u_D - r_S \cdot i_D + V_{BV}}{NBVL \cdot V_t}\right), \quad (6)$$

gdzie N , IBV , NBV , $IBVL$, $NBVL$ oznaczają parametry modelu diody.

Zależności temperaturowe parametrów V_t , I_S , r_S , V_{BV} , zaczerpnięte z pracy [2], dane są wzorami.

$$V_t = \frac{k}{q} \cdot T_j \quad (7)$$

$$I_S = I_{S0} \left(\frac{T_j}{T_{nom}}\right)^{p_t/\eta} \exp\left[\frac{qE_{g0}}{\eta k T_{nom}} - \frac{qE_{g0}}{\eta k T_j}\right] \quad (8)$$

$$V_{BV} = V_{BV0} (1 + t_{BV1} (T_j - T_{nom}) + t_{BV2} (T_j - T_{nom})^2) \quad (9)$$

$$R_S = R_{S0} (1 + t_{RS1} (T_j - T_{nom}) + t_{RS2} (T_j - T_{nom})^2), \quad (10)$$

gdzie k , q są stałymi fizycznymi, T_{nom} jest temperaturą nominalną, w której wyznaczono wartości parametrów, natomiast I_{S0} , p_t , η , E_{S0} , V_{BV0} , t_{BV1} , t_{BV2} , R_{S0} , t_{RS0} , t_{RS2} są parametrami modelu tranzystora bipolarnego. Modele pojemności C1 i C2 nie będą tu omawiane, gdyż elementy te nie występują w wersji stałoprądowej omawianego modelu.

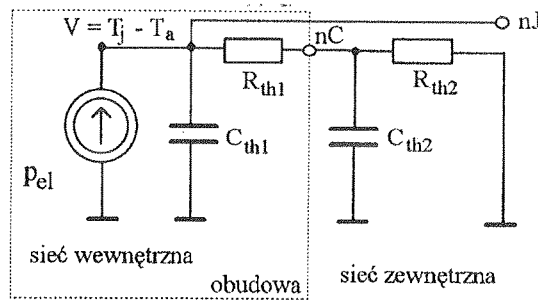
Do poważnych mankamentów omawianego modelu tranzystora bipolarnego można zaliczyć przede wszystkim: stałą wartość współczynnika wzmocnienia prądowego α , pominięcie prądów generacyjnych, modelowanie rezystancji szeregowych obszarów emitera i kolektora za pomocą rezystancji szeregowych diod D1 i D2, jak również przyjęcie modelu zjawiska przebicia lawinowego złączy danego wzorem (6), uniemożliwiającego uzyskanie efektu zróżnicowania wartości napięć przebicia przy zmianie konfiguracji pracy tranzystora. Możliwe jest tylko uzyskanie zmiany „twardości” charakterystyki. Z kolei model termiczny tranzystora o postaci obwodowej pokazanej na rys. 5 jest tworzony automatycznie przez program APLAC, jeśli użytkownik wprowadzi różną od zera wartość rezystancji termicznej R_{th1} .

W przypadku stałoprądowym pojemności termiczne C_{th1} , C_{th2} nie występują w omawianym modelu. Wydajność źródła prądowego p_{el} odpowiada wartości mocy elektrycznej w elemencie, natomiast wartość R_{th1} , określa wartość rezystancji termicznej złącze—obudowa. Dołączenie do zacisku nC termicznej sieci zewnętrznej pozwala na uwzględnienie w modelu termicznym rezystancji termicznej obudowa—otoczenie. Węzły obudów różnych tranzystorów można łączyć ze sobą, co pozwala na modelowanie wzajemnych sprzężeń termicznych. Przy braku sieci

zewnętr
modelo
nie wy
(obwod
Dru
tona je
tranzys
rys. 6, g

natomiast
Paramet
nej oraz
Sfor
sprowad
tranzyst
modelu
pomocą

DefMod
Trans T
+ BV =
+ etc =
Trans T2
+ BV =
+ etc =
Diode D
+ isr =
Res R1 b
Res R2 2
Res rt1 n
EndMod



Rys. 5. Model termiczny tranzystora w programie APLAC

zewnętrznej punkt nC jest zwierany automatycznie do masy. Oczywiście takie modelowanie oznacza, że pomiędzy obudowami elementów sprzężonych termicznie nie występuje różnica temperatur. Węzeł nJ pozwala na dołączenie własnego (obwodowego) modelu termicznego elementu.

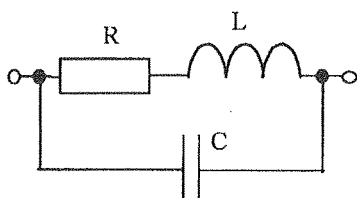
Drugim z rozważanych elementów składowych struktury tranzystora Darlingтона jest rezystor. Model termiczny rezystora jest identyczny jak w przypadku tranzystora bipolarnego, natomiast model elektryczny rezystora przedstawiono na rys. 6, gdzie rezystancja R opisana jest wzorem

$$R = R_{nom} \cdot \left[1 + t_{c1} \cdot (T_j - T_{nom}) + t_{c2} \cdot (T_j - T_{nom})^2 \right] \quad (11)$$

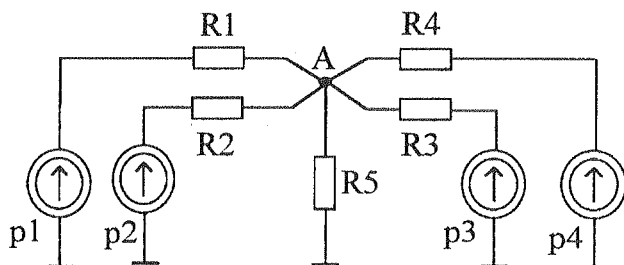
natomiast elementy inercyjne (L , C) można pominąć w analizie stałoprądowej. Parametry R_{nom} , t_{c1} , t_{c2} oznaczają odpowiednio rezystancję w temperaturze nominalnej oraz liniowy i kwadratowy temperaturowy współczynnik zmian rezystancji.

Sformułowanie elektrotermicznego makromodelu tranzystora Darlingтона mocy sprowadza się zatem do połączenia obwodów reprezentujących modele elektryczne tranzystora bipolarnego i rezystora zgodnie z topologią z rys. 2, natomiast postać modelu termicznego przedstawiono na rys. 7 [4]. Zatem EMTD zdefiniowany jest za pomocą opcji DefModel o postaci:

```
DefModel Darl 3 e b c param 4 ro1 ro2 rth11 rth12
Trans T1 c b 2 IS = 1.46e-13 rb = 1m alpha = 0.987 alphas = 0.846 ng = 0.00008 DIODE_BC
+ BV = 400 ibv = 1m nbv = 38.75 DIODE_BE BV = 8 ibv = 1m nbv = 38.75 rth = 7.4 etj = nt1
+ etc = nx
Trans T2 c 2 e IS = 7.31e-13 rb = 1m alpha = 0.9846 alphas = 0.846 ng = 0.00008 DIODE_BC
+ BV = 400 ibv = 1m nbv = 38.75 DIODE_BE BV = 8 ibv = 1m nbv = 38.75 rth = 7.4 etj = nt2
+ etc = nx
Diode D1 e c IS = 1.389e-12 bv = 400 ibv = 1m nbv = 27.3 rs = 2 tbv1 = 5e-4 vj = 1 xti = 2 nr = 2
+ isr = 6.05e-9 trs1 = 3.36e-3 rth = 7.4 etj = nt3 etc = nx
Res R1 b 2 ro1 rth = 7.4 etj = nt4 etc = nx tc1 = 1.25e-3
Res R2 2 e ro2 rth = 7.4 etj = nt5 etc = nx tc1 = 4.3e-3
Res rt1 ntx gnd rth11
EndModel
```



Rys. 6. Model elektryczny rezystora



Rys. 7. Model termiczny tranzystora Darlingtona

Wydajności źródeł p_1 , p_2 , p_3 , p_4 reprezentują moce wydzielane w poszczególnych elementach składowych tranzystora Darlingtona. Rezystory $R_1 \div R_4$ reprezentują własne rezystancje termiczne elementów składowych między złączem a obudową, reprezentowaną w tym przypadku przez podkładkę montażową (punkt A na schemacie), natomiast rezystor R_5 modeluje zarówno wzajemną rezystancję termiczną między dowolnie wybraną parą elementów (R_{thij} między elementem i -tym oraz j -tym) jak i rezystancję termiczną między obudową (podkładką) a otoczeniem (R_{thc-a}).

Jak wynika z koncepcji omawianej struktury modelu termicznego, wzajemne rezystancje termiczne między elementami są takie same. Powyższe spostrzeżenie wynika bezpośrednio z analizy schematu z rys. 7, gdyż temperatura wnętrza wybranego elementu i -tego wynikająca z wydzielania mocy w elemencie j -tym jest zawsze równa iloczynowi wydajności źródła prądowego p_j oraz rezystora R_5 . Problem identyfikacji wartości rezystorów w modelu termicznym rozwiązano na drodze eksperymentalnej. Zmierzono wartości własnej rezystancji termicznej diody $R_{th33} = 37,4 \text{ K/W}$ oraz wzajemnej rezystancji termicznej między tranzystorem T2 a diodą $R_{th23} = 30 \text{ K/W}$. Z uwagi na małe rozmiary struktury krzemowej, przyjęto, że dla wszystkich pozostałych elementów wartości własnych rezystancji termicznych są równe R_{th33} , natomiast wartości wzajemnych rezystancji termicznych są równe R_{th23} [5].

Po

wynika
muszą
otrzym
 $R_5 = 3$
następu

• d

• d

• d

• d

Za

znaczon
wartośc
ciągłe)
puterow
główną
pomiar
wzmocn
ich war
zatem o
tru alph
„optycz
przebiec
znak z c
dla wart
T2 wyn
kterystyl
istotną
w progr
alpha.

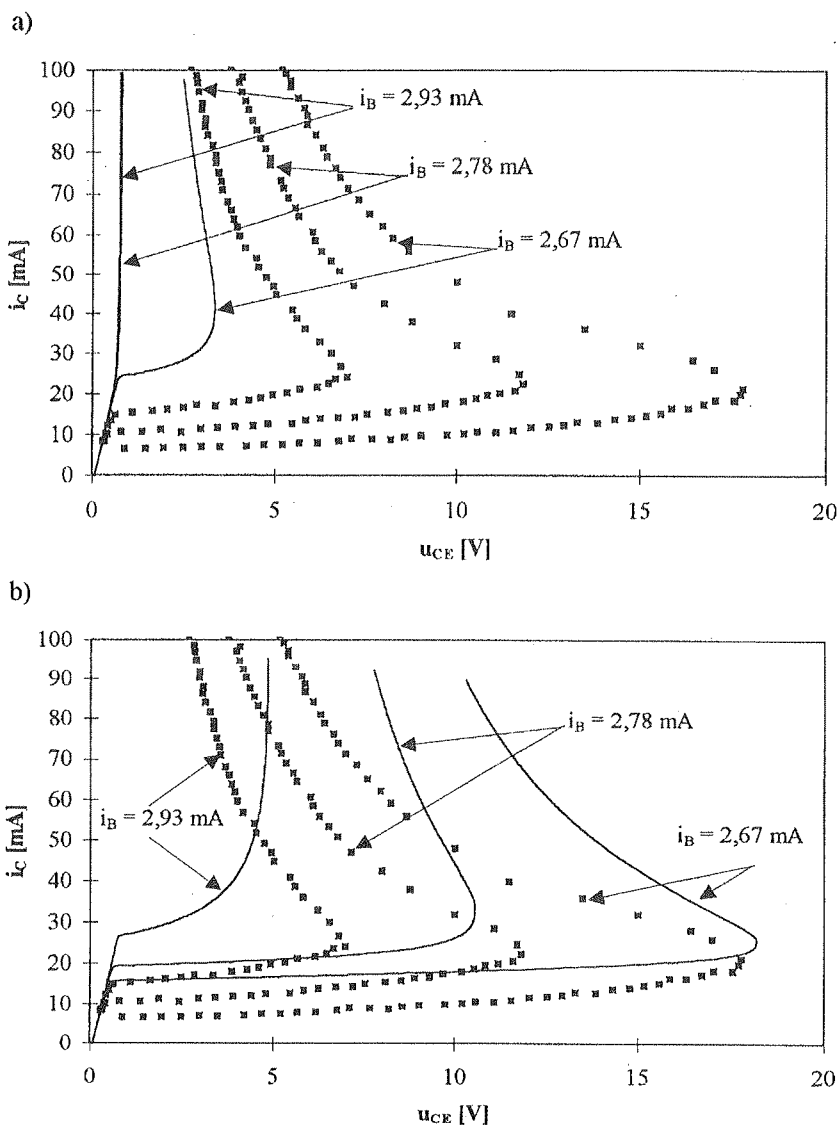
Po rozwiązaniu układu równań

$$\begin{cases} p_1 \cdot R_{thii} + (p_2 + p_3 + p_4) \cdot R_{thij} = p_1 \cdot (R1 + R5) + (p_2 + p_3 + p_4) \cdot R5 \\ p_2 \cdot R_{thii} + (p_1 + p_3 + p_4) \cdot R_{thij} = p_2 \cdot (R1 + R5) + (p_1 + p_3 + p_4) \cdot R5 \\ p_3 \cdot R_{thii} + (p_1 + p_2 + p_4) \cdot R_{thij} = p_3 \cdot (R1 + R5) + (p_1 + p_2 + p_4) \cdot R5 \\ p_4 \cdot R_{thii} + (p_1 + p_2 + p_3) \cdot R_{thij} = p_4 \cdot (R1 + R5) + (p_1 + p_2 + p_3) \cdot R5 \end{cases} \quad (12)$$

wynikającego z warunku, że wartości temperatur wnętrza obliczone z równania (1) muszą być równe wartościom otrzymanym z analizy sieci termicznej z rys. 7, otrzymano następujące wartości rezystorów: $R1 = R2 = R3 = R4 = 7,4 \, \Omega$ oraz $R5 = 30 \, \Omega$. Parametry modelu elektrycznego, otrzymane również z pomiarów mają następujące wartości:

- dla tranzystora T1: $IS = 1,46 \cdot 10^{-13} \, A$, $RB = 1m\Omega$, $\alpha = 0,9846$, $\alpha_{har} = 0,846$, $ng = 0,00008$, parametry diody D1: $BV = 400 \, V$, $IBV = 1mA$, $NBV = 38,75$, parametry diody D2: $BV = 8V$, $IBV = 1mA$, $NBV = 38,75$
- dla tranzystora T2: $IS = 7,31 \cdot 10^{-13} \, A$, $RB = 1m\Omega$, $\alpha = 0,9846$, $\alpha_{har} = 0,846$, $ng = 0,00008$, parametry diody D1: $BV = 400 \, V$, $IBV = 1mA$, $nbv = 38,75$, parametry diody D2: $BV = 8V$, $IBV = 1mA$, $NBV = 38,75$
- dla rezystora R1: $R_{nom} = 281 \, \Omega$, $tc1 = 1,25 \cdot 10^{-3} \, K^{-1}$
- dla rezystora R2: $R_{nom} = 21,18 \, \Omega$, $tc1 = 4,3 \cdot 10^{-3} \, K^{-1}$.

Za pomocą programu APLAC, wykorzystując przedstawiony EMTD, wyznaczono nieizotermiczne charakterystyki tranzystora BU323A dla podanych wartości parametrów. Na rys. 8 przedstawiono charakterystyki obliczone (linie ciągłe) oraz otrzymane z pomiarów (punkty) z wykorzystaniem mikrokomputerowego systemu pomiarowego opisanego w pracy [6]. Przyjęto hipotezę, że główną przyczyną uzyskanej na rys. 8a dużej rozbieżności pomiędzy wynikami pomiarów i obliczeń jest przyjęta w makromodelu stała wartość współczynników wzmocnienia prądowego tranzystorów (parametr α) podczas, gdy w praktyce ich wartości silnie zależą od wartości prądu kolektora tranzystora. Powtórzono zatem obliczenia z podanymi wyżej wartościami parametrów, z wyjątkiem parametru α . Wartość tego parametru dobierano tak, aby uzyskać jak najlepszą „optyczną” zgodność między wynikami obliczeń i pomiarów w obliżu punktów przebiecia termicznego, tzn. punktów, w których rezystancja różniczkowa zmienia znak z dodatniego na ujemny. Przedstawione na rys. 8b wyniki obliczeń, uzyskane dla wartości współczynników wzmocnienia prądowego α tranzystorów T1 oraz T2 wynoszących odpowiednio 0,978 oraz 0,92, wykazują lepszą zgodność z charakterystykami pomierzonymi i potwierdzają słuszność przyjętego założenia, że istotną przyczyną małej dokładności EMTD jest przyjęcie we wbudowanym w programie APLAC modelu tranzystora bipolarnego, stałej wartości parametru α .



Rys. 8. Nieizotermiczne charakterystyki wyjściowe tranzystora BU323A: pomierzone (punkty) i obliczone (linie ciągłe) przy wykorzystaniu pomierzonych a) oraz zmodyfikowanych b) wartości parametrów

5. MODYFIKACJE EMTD

Jak zatem wynika z rozdziału 4, na dokładność EMTD duży wpływ ma postać elektrotermicznego modelu tranzystora bipolarnego. Opracowano zatem dwa typy zmodyfikowanych EMTD, przy formułowaniu których wykorzystano autorskie

elektrote
nością m
niające
temperat
wykorzy
Zgodnie
źródłami
oraz nap
zmianam
danych v

gdzie a , b
[3]. Spos
Prakt
lu tranz
wodow r
odpowie
sterowan
ny mode
źródło pr
sterowan
termiczna
Darlington
rezystor o

elektrotermiczne modele tranzystora bipolarnego, legitymujące się większą dokładnością niż modele tego elementu wbudowane w programie APLAC, a uwzględniające zjawiska zmian współczynnika wzmocnienia prądowego tranzystora od temperatury jego wnętrza i prądu kolektora. Pierwsza z omawianych modyfikacji wykorzystuje ideę formułowania modeli hybrydowych dla programu SPICE [7]. Zgodnie z tą ideą wbudowany model APLACa (WMA) uzupełniony jest dwoma źródłami sterowanymi (rys. 9): prądowym — modelującym zmianę prądu kolektora oraz napięciowym — modelującym zmiany napięcia baza—emiter, spowodowane zmianami temperatury wnętrza ponad temperaturę otoczenia, o wydajnościach danych wzorami

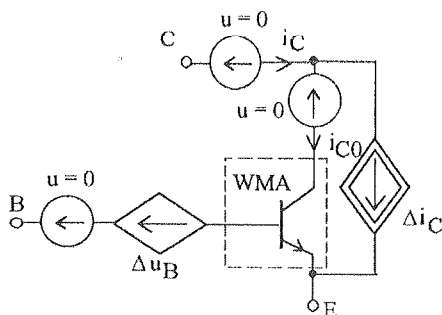
$$\Delta i_C = -i_{C0} \cdot \sum_i \gamma_i \cdot \exp\left(\alpha_i \cdot \sqrt{i_C^2 + d_i}\right) \cdot (1 + a \cdot (T_j - T_0)) +$$

$$+ i_{C0} \cdot a \cdot (T_j - T_0) - b \cdot (1 + c \cdot (T_j - T_0)) i_C^2$$
(13)

$$\Delta u_B = V_t \cdot \ln\left(\frac{i_C}{i_{C0}}\right),$$
(14)

gdzie $a, b, c, d_i, \alpha_i, \gamma_i$ to parametry modelu współczynnika wzmocnienia prądowego [3]. Sposób wyprowadzenia wzorów (13), (14) przedstawiono w Dodatku A.

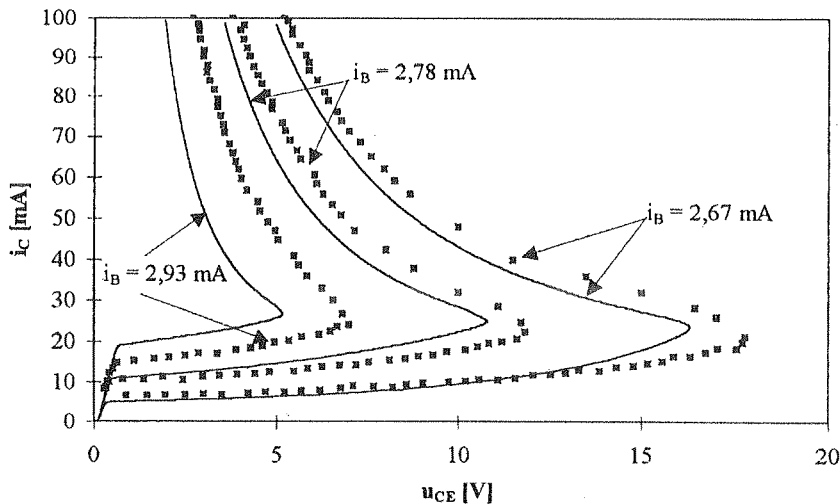
Praktyczna implementacja omawianego elektrotermicznego hybrydowego modelu tranzystora bipolarnego wymaga sformułowania dodatkowych prostych obwodów realizujących konwersję wartości prądów w wybranych gałęziach modelu na odpowiednie napięcia, co wynika ze składni programu APLAC (wydajność źródeł sterowanych może być funkcją tylko napięć lub tylko prądów). Również wbudowany model termiczny rozważanego elementu wymaga uzupełnienia o sterowane źródło prądowe o wydajności odpowiadającej sumie mocy wydzielanej w źródłach sterowanych występujących na rys. 9. Modyfikacji uległa również wewnętrzna sieć termiczna z rys. 5 w odniesieniu do wszystkich elementów składowych tranzystora Darlingtona. Węzeł nC jest teraz zwarty do masy, natomiast do węzła nJ poprzez rezystor o wartości odpowiadającej R_{thi} dołączone jest sterowane źródło napięciowe



Rys. 9. Struktura obwodowa hybrydowego modelu tranzystora bipolarnego

o wydajności równej $\sum_{i=1}^3 (p_j \cdot R_{thij} \cdot R_{thii}^{-1})$, reprezentujące wzajemne sprzężenia termiczne.

Na rys. 10 pokazano wyjściowe charakterystyki tranzystora Darlingtona BU323A otrzymane przy użyciu hybrydowego EMTD. Wartości parametrów występujących w opisie źródeł sterowanych otrzymano z pomiarów, natomiast parametry modelu wbudowanego wynoszą: $\alpha = \beta_0^{-1}$ oraz $I_S = I_S(T_a)$. Jak widać uzyskano znaczną poprawę zgodności charakterystyk obliczonych z pomierzonymi.



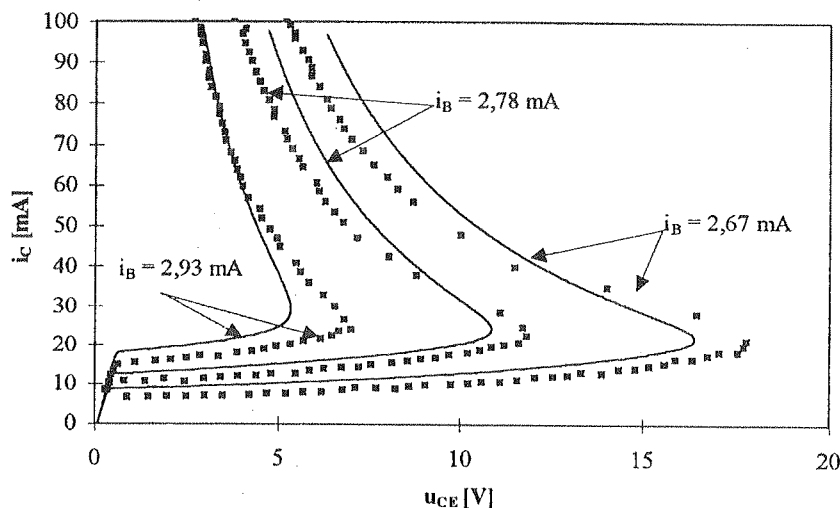
Rys. 10. Wyniki pomiarów i symulacji BU323A przy użyciu hybrydowego EMTD

Dруга z zapowiadanych modyfikacji elektrotermicznego makromodelu tranzystora Darlingtona opiera się na globalnym elektrotermicznym modelu tranzystora bipolarnego (GETM) zaimplementowanym do programu APLAC w postaci podukładu z wykorzystaniem opcji DefModel. dokładny opis modelu GETM tranzystora bipolarnego zamieszczono w pracy [3], natomiast opis globalnego elektrotermicznego makromodelu tranzystora Darlingtona zamieszczono w pracy [5].

Na rys. 11 przedstawiono, zweryfikowane doświadczalnie, wyniki symulacji nieizotermicznych charakterystyk wyjściowych tranzystora BU323A otrzymane z wykorzystaniem globalnego EMTD. Wartości parametrów, podobnie jak w przypadku hybrydowego EMTD otrzymano z pomiarów. Jak widać, zgodność obliczeń i pomiarów jest nieco lepsza niż w przypadku makromodelu hybrydowego.

Rys. 11

W
tora D
elektro
elektry
miczny
rów otr
towałn
Stwierd
modelu
prądow
tranzyst
elektro
z wyko
globaln
w przyp
APLAC



Rys. 11. Wyniki pomiarów i symulacji BU323A wykonane z wykorzystaniem elektrotermicznego modelu globalnego EMTD

UWAGI KOŃCOWE

W pracy przedstawiono trzy postacie elektrotermicznego makromodelu tranzystora Darlingtona mocy. Przy jego formułowaniu wykorzystano wbudowane elektrotermiczne modele elementów składowych, stanowiące połączenie ich modeli elektrycznych i termicznych w postaci obwodowej. Wyznaczono rodzinę nieizotermicznych charakterystyk wyjściowych tranzystora BU323A z wartościami parametrów otrzymanymi z pomiarów charakterystyk tego elementu. Uzyskano nieakceptowalnie dużą rozbieżność między charakterystykami obliczonymi i pomierzonymi. Stwierdzono, że główną przyczyną tych rozbieżności jest przyjęcie we wbudowanym modelu tranzystora bipolarnego stałej wartości współczynnika wzmocnienia prądowego. Zaproponowano dwie modyfikacje elektrotermicznego makromodelu tranzystora Darlingtona formułując elektrotermiczny makromodel hybrydowy oraz elektrotermiczny makromodel globalny tego elementu. Wyniki symulacji otrzymane z wykorzystaniem makromodeli zmodyfikowanych, a zwłaszcza makromodelu globalnego, pozostają w dużo lepszej zgodności z wynikami pomiarów niż w przypadku EMTD sformułowanego w oparciu o modele wbudowane w programie APLAC.

BIBLIOGRAFIA

1. T. Veijola, M. Valtonen: *An Object-Oriented Approach to Combined Electrical and Thermal Circuit Analysis*, 12th European Conference on Circuit Theory and Design, Istambul 1995, Vol. 2, p. 1055
2. T. Veijola i inni: APLACTM 7.0 Reference Manual Vol. I. Helsinki University of Technology, Circuit Theory Laboratory & Nokia Corporation Research Center, 1996
3. J. Zarębski: *Modelowanie, symulacja i pomiary przebiegów elektrotermicznych w elementach półprzewodnikowych i układach elektronicznych*. Prace Naukowe Wyższej Szkoły Morskiej w Gdyni, Gdynia 1996
4. K. Górecki, J. Zarębski: *Możliwości zastosowania programu APLAC do modelowania nieizotermicznych charakterystyk stałoprądowych tranzystora Darlingtona mocy*. III Konferencja Naukowo-Techniczna Zastosowanie Komputerów w Elektrotechnice ZKwE'98, Kiekrz—Poznań 1998, s. 337
5. W.J. Stepowicz, J. Zarębski, K. Górecki: *Electrothermal Macromodel of the Power Darlington Transistor*. 1997 European Conference on Circuit Theory and Design — ECCTD'97, Budapest 1997, Vol. 1, p. 191
6. K. Górecki, J. Zarębski: *Układ do automatycznego pomiaru stałoprądowych nieizotermicznych charakterystyk wybranych elementów półprzewodnikowych*. XXVIII Międzyuczelniana Konf. Metrologów MKM'96, Częstochowa 1996, t. 1, s. 192
7. J. Zarębski, K. Górecki: D.C. Electrothermal Hybrid BJT Model for SPICE. Proc. VI Int. Conf. Simulation of Semiconductor Devices and Processes — SISDEP'95, Erlangen (RFN) 1995, Vol. 6, p. 306

DODATEK A

Wyprowadzenie zależności opisujących wydajności źródeł sterowanych występujących w hybrydowym modelu tranzystora bipolarnego

Model tranzystora bipolarnego w zakresie aktywnym normalnym można przedstawić za pomocą równań

$$i_C = \beta_F \cdot i_B + I_S \cdot \left(1 + \frac{\beta_F + 1}{\beta_R} \right) \quad (A1)$$

$$u_{BE} = V_T \cdot \ln \left(\frac{i_C}{I_S} + 1 \right), \quad (A2)$$

gdzie współczynnik wzmacnienia prądowego β_F opisany jest zależnością

$$\beta(T_j) = \frac{\beta_0 \cdot (1 + a \cdot (T_j - T_0))}{1 + b \cdot (1 + c \cdot (T_j - T_0)) \sqrt{i_C^2 + d_1}} \cdot \left[1 - \sum_i \gamma_i \cdot \exp \left(\alpha_i \cdot \sqrt{i_C^2 + d_1} \right) \right] \quad (A3)$$

łącąc zależności (A1) z (A3) otrzymujemy

$$i_C = \frac{\beta_0 \cdot i_B (1 + a \cdot (T_j - T_0)) \left[1 - \sum_i \gamma_i \cdot \exp \left(\alpha_i \cdot \sqrt{i_C^2 + d_1} \right) \right]}{1 + b \cdot (1 + c \cdot (T_j - T_0)) \sqrt{i_C^2 + d_1}} + I_S \cdot \left(1 + \frac{1}{\beta_R} + \frac{\beta_{0F}}{\beta_{0R}} \right). \quad (A4)$$

W modelu wbudowanym APLAC-a prąd kolektora w zakresie aktywnym normalnym jest opisany wzorem

$$i_{C0} = \beta_0 \cdot i_B + I_S \cdot \left(\frac{1}{\beta_{0R}} + \frac{\beta_{0F}}{\beta_{0R}} \right) \quad (A5)$$

natomiast napięcie baza-emiter wzorem

$$u_{BE} = V_T \cdot \ln \left(\frac{i_{C0}}{I_S} + 1 \right). \quad (A6)$$

Odejmując stronami zależności (A4) i (A5) oraz zakładając, że $\beta_{0F} \gg \beta_{0R}$ oraz $\beta_{0F} \gg 1$ otrzymuje się wzór opisujący wydajność źródła prądowego Δi_C postaci

$$\begin{aligned} \Delta i_C = & -i_{C0} \cdot \sum_i \gamma_i \cdot \exp(\alpha_i \cdot \sqrt{i_C^2 + d_1}) \cdot (1 + a \cdot (T_j - T_0)) + \\ & + i_{C0} \cdot a \cdot (T_j - T_0) - b \cdot (1 + c \cdot (T_j - T_0)) \cdot i_C^2. \end{aligned} \quad (A7)$$

Z kolei, odejmując stronami zależności (A2) i (A6) i zakładając, że $i_C \gg I_S$ oraz $i_{C0} \gg I_S$ otrzymuje się wzór opisujący wydajność źródła napięciowego Δu_B postaci

$$\Delta u_B = V_T \cdot \ln \left(\frac{i_C}{i_{C0}} \right). \quad (A8)$$

J. ZARĘBSKI, K. GÓRECKI

MODELLING OF THE NONISOTHERMAL DC CHARACTERISTICS OF THE POWER DARLINGTON TRANSISTOR USING APLAC

S u m m a r y

In the paper three various forms of the electrothermal macromodel of the power Darlington transistor (EMDT) for APLAC user were formulated. First of them is based on the built-in electrothermal models of the component devices of Darlington transistor structure. In the designing process of the next two, so called, hybrid and global electrothermal models of the bipolar transistor were formulated and utilized. In the case of using of the modified versions of the electrothermal BJT models, the simulation results verified experimentally have shown a visible improvement of accuracy of EMTD.

Key words: semiconductor devices, Darlington transistor, modelling, electrothermal phenomena.

m
na
tel
ko

SA

Jedn

Są one
kolorow
legające
kształt.
nym pr
obrazu.
histogra
dostosc
nych w
(przedn
W arty
wykorz
W p
istnieją
celów p

Metoda specyfikacji histogramu wykorzystująca macierz współwystępowania i jej zastosowanie dla transformacji obrazu

JERZY SIUZDAK, TOMASZ CZARNECKI

Instytut Telekomunikacji, Politechnika Warszawska

Otrzymano 1998.05.22

Autoryzowano 1998.08.15

W pracy zaprezentowano nową, globalną metodę specyfikacji histogramu obrazu monochromatycznego wykorzystującą własności macierzy współwystępowania. Opracowana metoda pozwala, w zależności od potrzeby, uwydatnić na obrazie większe szczegóły lub tekstury. Metoda może być również zastosowana do przekształcenia luminancji obrazów kolorowych.

Słowa kluczowe: obrazy monochromatyczne TV, histogramy, macierz współwystępowania.

1. WSTĘP

Jednymi z ważnych metod obróbki obrazu są przekształcenia jego histogramu. Są one stosowane głównie do obrazów czarno-białych lub do luminancji obrazów kolorowych. Metody te polegają na takiej transformacji jaskrawości obrazu podlegającego obróbce, iż histogram przekształconego obrazu przyjmuje z góry zadany kształt. Takie przekształcenie nosi nazwę specyfikacji histogramu [1]. W szczególnym przypadku, jeśli żądany jest równomierny kształt histogramu przekształconego obrazu, mamy do czynienia z wyrównywaniem (linearyzacją) histogramu (ang. *histogram equalization*) [1]. Celem specyfikacji histogramu jest na ogół lepsze dostosowanie przedstawienia zawartości treściowej obrazu do możliwości percepcyjnych wzroku ludzkiego. Pozwala to na przykład uwydatnić pewne fragmenty obrazu (przedmioty, tekstury), które są słabo widoczne na obrazie nie przetworzonym. W artykule przedstawiono nową, globalną metodę specyfikacji histogramu, która wykorzystuje pewne własności macierzy współwystępowania.

W pracy można wydzielić dwie części. W pierwszej z nich opisano niektóre, istniejące, globalne metody specyfikacji histogramu. Będą one później użyte dla celów porównawczych. W drugiej części przedstawiono proponowaną metodę specy-

fikacji histogramu i porównano rezultaty stosowania zaproponowanej metody z metodami znanymi wcześniej.

2. NIEKTÓRE METODY SPECYFIKACJI HISTOGRAMU

Rozpatrzmy specyfikację histogramu od strony matematycznej. Oznaczmy przez $h(i)$, $i = 0 \dots K - 1$ (K — liczba poziomów jaskrawości) wartości histogramu obrazu początkowego, znormalizowane do całkowitej liczby piksli, tzn.

$$h(i) = \frac{\text{liczba piksli o poziomie jaskrawości } i}{\text{liczba wszystkich piksli}}. \quad (1)$$

Zdefiniujmy funkcję $H(i)$, będącą analogiem dystrybuanty w rachunku prawdopodobieństwa, jako

$$H(i) = \sum_{m=0}^i h(m). \quad (2)$$

W analogiczny sposób zdefiniujmy funkcje $g(i)$ i $G(i)$ dla przekształconego obrazu. W szczególnym przypadku wyrównywania (linearyzacji) histogramu mamy:

$$g(i) = \frac{1}{K}; \quad G(i) = \frac{i+1}{K}; \quad i = 0 \dots K-1. \quad (3)$$

Żądana transformacja polega na przyporządkowaniu poziomowi jaskrawości i pierwotnego obrazu, poziomu jaskrawości j , który spełnia równość

$$H(i) = G(j). \quad (4)$$

Ostatecznie, żądane przekształcenie dane jest wzorem

$$j = G^{-1}[H(i)]. \quad (5)$$

Omawiana procedura daje dokładne wartości tylko w przypadku funkcji ciągłych; w przypadku dyskretnych wartości jaskrawości równość (5) może być spełniona tylko w sposób przybliżony. W szczególności, w przypadku wyrównywania histogramu, jeśli „słupek” histogramu o danej jaskrawości i jest k razy większy niż średnia, to tę jaskrawość należy odwzorować w k różnych poziomów na przykład pomiędzy j_1 a j_k . Potrzebna jest więc reguła takiego odwzorowania. W pracy [2] zaproponowano kilka takich reguł, jednak w praktyce okazało się, że wszystkie dają bardzo podobne obrazy.

Specyfikacja histogramu nie zachowuje w ogólnym przypadku średniej jaskrawości obrazu, co może być istotną wadą w pewnych warunkach (np. obraz TV). W pracy [3] podano pewną odmianę metody linearyzacji histogramu, pozwalającą zachować średnią jasność obrazu. Wymaga ona obliczenia średniej jaskrawości obrazu k_m zgodnie z wzorem

$$k_m = \sum_{i=0}^{K-1} h(i) i, \quad (6)$$

metody

a następnie dokonaniu niezależnej linearyzacji histogramów jaskrawości dla poziomów jaskrawości $i \leq k_m$ oraz dla poziomów jaskrawości $i > k_m$. W tej metodzie definiuje się następujące funkcje

$$h_L(i) = \frac{\text{liczba piksli o poziomie jaskrawości } i}{\text{liczba wszystkich piksli o jaskrawości } i \leq k_m} \quad \text{dla } i \leq k_m \quad (7)$$

y przez
obrazu

$$h_U(i) = \frac{\text{liczba piksli o poziomie jaskrawości } i}{\text{liczba wszystkich piksli o jaskrawości } i > k_m} \quad \text{dla } i > k_m \quad (8)$$

(1)

oraz

a praw-

$$H_L(i) = \sum_{m=0}^i h_L(m), \quad i \leq k_m, \quad (9)$$

(2)

$$H_U(i) = \sum_{m=k_m+1}^i h_U(m), \quad i > k_m. \quad (10)$$

obrazu.

Ostatecznie, jaskrawość o poziomie i transformowanego obrazu jest przekształcana na jaskrawość o poziomie j , zgodnie z wzorem [3]:

(3)

$$j = \begin{cases} k_m H_L(i) & \text{dla } i \leq k_m \\ k_m + 1 + (N - k_m - 2) H_U(i) & \text{dla } i > k_m. \end{cases} \quad (11)$$

ci i pier-

(4)

Ta metoda pozwala zwiększyć kontrast zachowując jednocześnie średnią jasność obrazu [3]. Problem transformacji histogramu może być również potraktowany jako pewne zagadnienie wariacyjne [4], polegające na minimalizacji łącznego odstępów między „słupkami” histogramu z pewną zadaną funkcją wagową $w(m)$. W praktyce zastosowanie tej metody polega na usunięciu z histogramu pierwotnego wszystkich zerowych „słupków” i utworzeniu nowego histogramu wyłącznie z niezerowymi „słupkami”. Takich niezerowych „słupków” jest w ogólnym przypadku L , gdzie $L \leq K$ (K — liczba możliwych poziomów jaskrawości pierwotnego obrazu). Jednocześnie musi zostać zapamiętane przekształcenie starej jaskrawości i na miejsce j ($j = 1, 2, \dots, L$) w nowym histogramie

(5)

$$j = D(i). \quad (12)$$

ciągłych;
na tylko
ogramu,
nia, to tęzy j_1 a j_k .
no kilka

e obrazy.

niej jask-

raz TV).
walającą

krawości

Do wielkości j stosuje się przekształcenie, dane w przypadku dyskretnym w ogólnej postaci przez [4]

$$k = k(j) = k[D(i)] = \frac{K-1}{L-1} \frac{\sum_{m=1}^{j-1} w(m)}{\sum_{m=1}^L w(m)}. \quad (13)$$

(6)

Funkcja wagowa $w(m)$ generalnie zależy od postaci histogramu obrazu wejściowego. Warto dodać, że niektóre poprzednio omówione przekształcenia można uzyskać jako przypadki szczególne tej metody, zmieniając jedynie funkcję wagową $w(m)$. Dla $w(m) = 1$ mamy

$$k(j) = \frac{K-1}{L-1}(j-1), \quad (14)$$

(liniowe rozciągnięcie skali), zaś dla $w(m) = h(m)$ otrzymujemy

$$k(j) = (K-1) \sum_{m=1}^{j-1} h(m), \quad (15)$$

(linearyzacja histogramu).

W dwóch powyższych metodach specyfikacji histogramu nie brano pod uwagę przestrzennej korelacji jaskrawości piksli. W pracy [5] zaproponowano metodę linearyzacji histogramu, która wykorzystuje informacje o tej korelacji. Obliczona jest względna częstość współwystępowania poziomów jaskrawości sąsiednich piksli i w ten sposób można policzyć warunkowe gęstości prawdopodobieństw wystąpienia każdego poziomu jaskrawości względem innych poziomów jaskrawości obrazu. Generowane tą metodą obrazy poprawiają jakość obrazu wyjściowego, lepiej oddając jego charakter aniżeli tradycyjne metody specyfikacji histogramu. Załóżmy, że mamy dany piksel o jaskrawości a i jego otoczenie 5×5 . Prawdopodobieństwo (łącznie) występowania w tym otoczeniu piksela o jaskrawości b obliczane jest zgodnie z wzorem [5]:

$$P(b, a) = \frac{\frac{1}{25} \sum L_b(a)}{\text{liczba wszystkich piksli w obrazie}}, \quad (16)$$

gdzie $L_b(a)$ jest liczbą piksli o jaskrawości b (b może przyjmować wartości z zakresu $b = a-2, a-1, a, a+1, a+2$) w otoczeniu 5×5 piksela o jaskrawości a , zaś sumowanie rozciąga się na wszystkie otoczenia piksli o jaskrawości a .

Następnie obliczane jest prawdopodobieństwo warunkowe

$$P(b|a) = \frac{P(b, a)}{P(a)} = \frac{P(b, a)}{h(a)}, \quad (17)$$

gdzie $h(a)$ dane jest wzorem (1).

Ostatecznie nowy poziom jaskrawości j związany jest ze starym poziomem jaskrawości i przekształceniem

$$j = T(i) \cdot (K-1), \quad (18)$$

gdzie $(K-1)$ jest maksymalnym poziomem jaskrawości obrazu, zaś $T(i)$ dane jest wzorem

$$T(i) = \frac{\sum_{a=0}^i \sum_{b=a-2}^{b=a+2} P(b|a)}{\sum_{a=0}^{K-1} \sum_{b=a-2}^{b=a+2} P(b|a)}, \quad i = 0, 1, \dots, K-1. \quad (19)$$

(14) Wszystkie omówione do tej pory metody mogą być stosowane globalnie i lokalnie. W przypadku zastosowania lokalnego do obliczania odnośnych parametrów oraz histogramów wykorzystuje się nie cały obraz, a jedynie pewne, określone otoczenie danego piksla.

(15)

3. ZASTOSOWANIE MACIERZY WSPÓŁWYSTĘPOWANIA

W dalszych rozważaniach bardzo istotną rolę pełni tzw. macierz współwystępowania (ang. *cooccurrence matrix*) [1]. Dla jej wprowadzenia założmy najpierw, że mamy dany pewien obraz czarno-biały, którego jaskrawość J jest określona zależnością

$$J = J(i, j), \quad J = 0, 1, 2, \dots, K-1, \quad i = 1, 2, \dots, N, \quad j = 1, 2, \dots, M \quad (20)$$

Tutaj K określa liczbę poziomów szarości, zaś N i M rozmiary obrazu. Macierz współwystępowania jest określona zawsze dla danej konfiguracji pary piksli. Jeśli pierwszy z piksli ma współrzędne (i, j) , to drugi piksel pary ma współrzędne $(i-r, j-q)$, gdzie r, q są liczbami całkowitymi, które nie są jednocześnie równe zero (gdyby tak było, to wówczas para degenerowałaby się do pojedynczego piksla). Ponieważ liczby r, q mogą przybierać różne wartości, dla danego obrazu istnieje wiele macierzy współwystępowania. Macierz współwystępowania jest zdefiniowana jako macierz kwadratowa o wymiarach $K \times K$, której każdy element $W_{rq}(x, y)$ jest równy liczbie opisanych powyżej par piksli, z których pierwszy ma jaskrawość x , zaś drugi jaskrawość y . Można więc powiedzieć, że macierz współwystępowania jest dwuwymiarowym histogramem jaskrawości określonej konfiguracji pary piksli. Zapisując powyższe warunki w sposób analityczny mamy:

(16)

zakresu
i a , zaś

$$W_{rq}(x, y) = \sum_i \sum_j \delta_{J(i,j),x} \delta_{J(i-r,j-q),y}, \quad x = 0, 1 \dots K-1, \quad y = 0, 1 \dots K-1. \quad (21)$$

Sumowanie rozciąga się po całym obrazie, zaś symbol δ oznacza deltę Kroneckera, zdefiniowaną następująco

(17)

$$\delta_{k,l} = \begin{cases} 1 & \text{dla } k = l \\ 0 & \text{dla } k \neq l \end{cases} \quad (22)$$

skrawo-

(18)

lane jest

(19)

Ponieważ przy przekształceniach obrazu chodzi często o poprawę subiektywnie odbieranej jakości obrazu, przy określeniu metody specyfikacji histogramu należy wykorzystać własności wzroku ludzkiego. Niektóre z tych własności są omówione w pracach [6, 7]. W szczególności należy podkreślić, że minimalny rozróżnialny przez człowieka kontrast dla małych natężeń oświetlenia maleje wraz ze wzrostem natężenia oświetlenia. Dla dużych i średnich wartości natężeń oświetlenia minimalny rozróżnialny kontrast nie zależy od tego natężenia. Ponadto czułość wzroku ludzkiego zależy od częstości przestrzennych obrazu. Dla każdego natężenia oświetlenia istnieje pewna częstotliwość przestrzenna, dla której wzrok jest najczulszy: na ogół rośnie ona wraz z poziomem oświetlenia. Maksimum czułości przypada na częstości przestrzenne od około 1 cykla/stopień dla małych wartości luminancji do około

4 cykli/stopień dla dużych luminancji. co więcej wzrok ludzki jest bardziej wrażliwy na szumy w „gładkich” rejonach obrazu aniżeli w rejonach o dużej liczbie szczegółów.

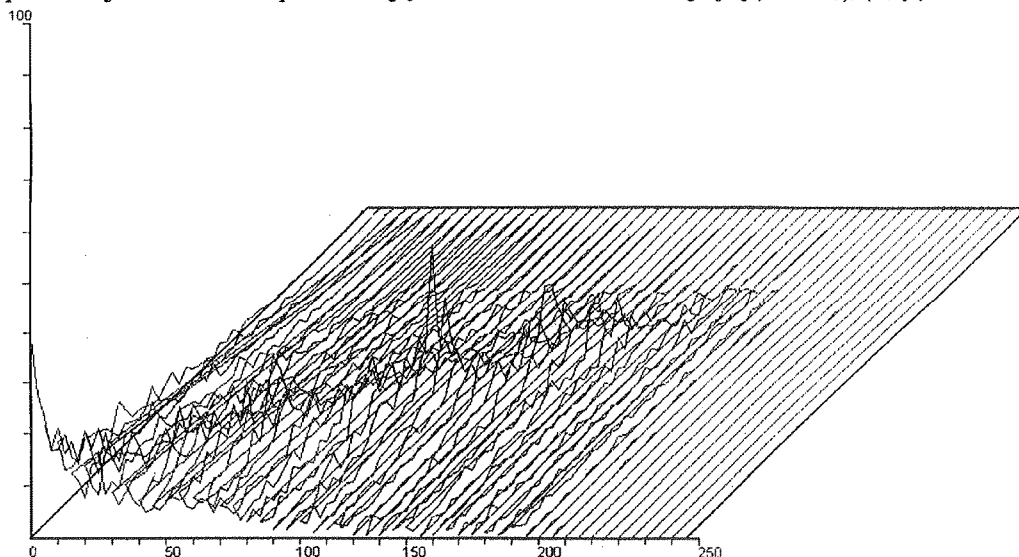
Po to, aby powiązać macierz współwystępowania z własnościami wzroku ludzkiego, należy określić, jakiej częstotliwości przestrzennej odpowiada dana odległość między parą piksli w macierzy współwystępowania. W tym celu należy oczywiście znać fizyczną odległość na ekranie monitora odpowiadającą jednemu pikselowi, jak również odległość, z której operator obserwuje monitor. Dla typowych wartości: monitor 15 calowy o aspekcie 4:3, 640 piksli w linii, przybliżone przeliczenie dyskretnej odległości między pikslami na częstość przestrzenną wyrażoną w cyklach

Odległość między pikslami	Częstość przestrzenna przy odległości obserwacji 0.5 m	Częstość przestrzenna przy odległości obserwacji 1 m
5	4 cykle/deg	8 cykli/deg
10	2 cykle/deg	4 cykle/deg
20	1 cykl/deg	2 cykle/deg

na stopień podaje poniższa tabelka.

Przykład macierzy współwystępowania typu $W_{r,r}(x, y)$ dla $r = 15$ ($r = q$) pokazano na rys. 1.

Dobór wartości r jest nieprzypadkowy. Jak wynika z powyższej tabelki i prac [6, 7] odpowiada on częstościom przestrzennym, na które oko ludzkie jest najbardziej wrażliwe. Dla ustalonej wartości x , zależność $W_{r,r}(x, y)$ określa, jakiego rodzaju piksele o jaskrawości x przeważają na obrazie. Jeśli funkcja $f(y) = W_{r,r}(x, y)$ koncen-



Rys. 1. Przykład macierzy współwystępowania $W_{r,r}(x, y)$ dla obrazu z rys. 5 dla wartości parametru $r = 15$. Skala na osi pionowej jest dowolna

truje s
w sasi
Ewentu
wane s

Z k
wokół
różnyc
Wreszc
dalszyc
wokół
sumow
współw
funkcji
prowad

Prze
malizac
takie m

Pierwsza
modułu
te miary
kilku m
podstaw
innym o
Miar
bardzo
znowu
wartości
Taki cha
tylko jed

truje się w pobliżu wartości x , wówczas piksele o tej jasności znajdują się w sąsiedztwie pikseli o podobnej jasności, należą zatem do jednolitych obszarów. Ewentualne różnice jasności w tych obszarach są najprawdopodobniej spowodowane szumem lub innymi zakłóceniami.

Z kolei jeśli funkcja $f(y)$ rozciąga się do wartości odległych (lub koncentruje się wokół wartości odległych) od x , to wtedy sąsiednie piksele należą do wyraźnie różnych obszarów, czyli jasność x określa zarys przedmiotów na obrazie. Wreszcie należy sądzić, że przypadek pośredni odpowiada tekstuom obrazu. Dla dalszych przekształceń numerycznych należy wybrać miarę koncentracji funkcji $f(y)$ wokół wartości x . W tym celu znormalizowano wartości funkcji $f(y)$, tak aby przy sumowaniu po y dawały jedność. Pozwala to uniezależnić się od wartości funkcji współwystępowania dla x . Następnie można obliczyć klasyczne miary rozpiętości funkcji (rozkładu), takie jak np. odchylenie standardowe. W ujęciu matematycznym prowadzi to do następujących przekształceń

$$F(y) = W_{r,r}(x, y) = f(y)/\eta, \quad \text{dla } y = 0, 1, 2, \dots, K-1, \quad (23)$$

$$\text{gdzie } \eta = \sum_{k=0}^{K-1} f(k) = \sum_{k=0}^{K-1} W_{r,r}(x, k).$$

Przez $F(y)$ oznaczono znormalizowaną zgodnie z (23) funkcję $f(y)$. Po normalizacji można policzyć miary koncentracji funkcji $F(y)$ wokół x . Zbadano trzy takie miary, zdefiniowane poniżej

$$m_1(x) = \sum_{k=0}^{K-1} (k-x)^2 F(k) \quad (24)$$

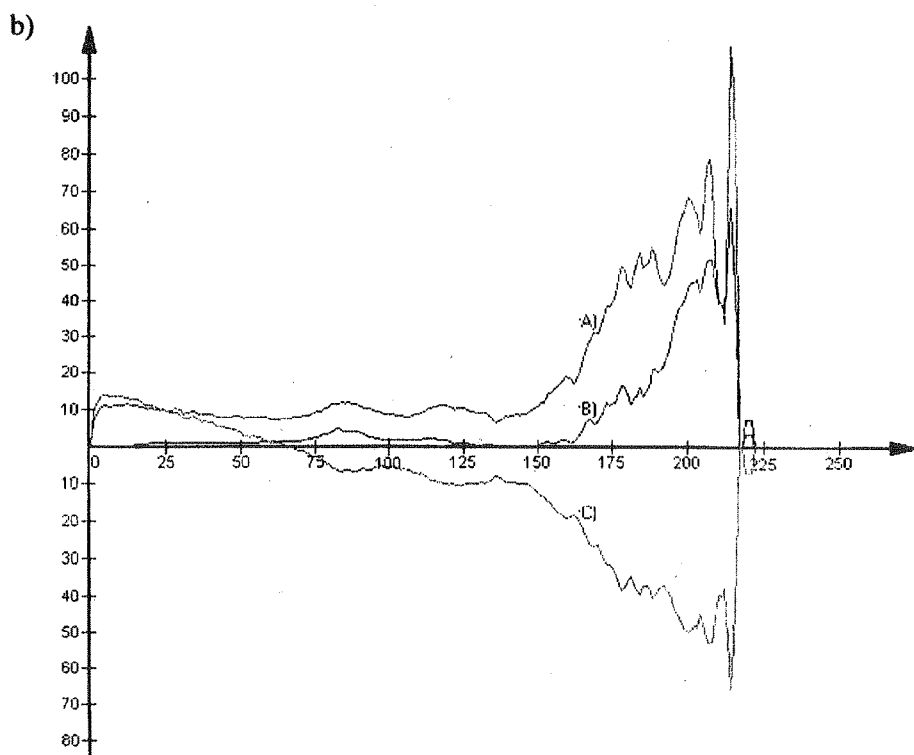
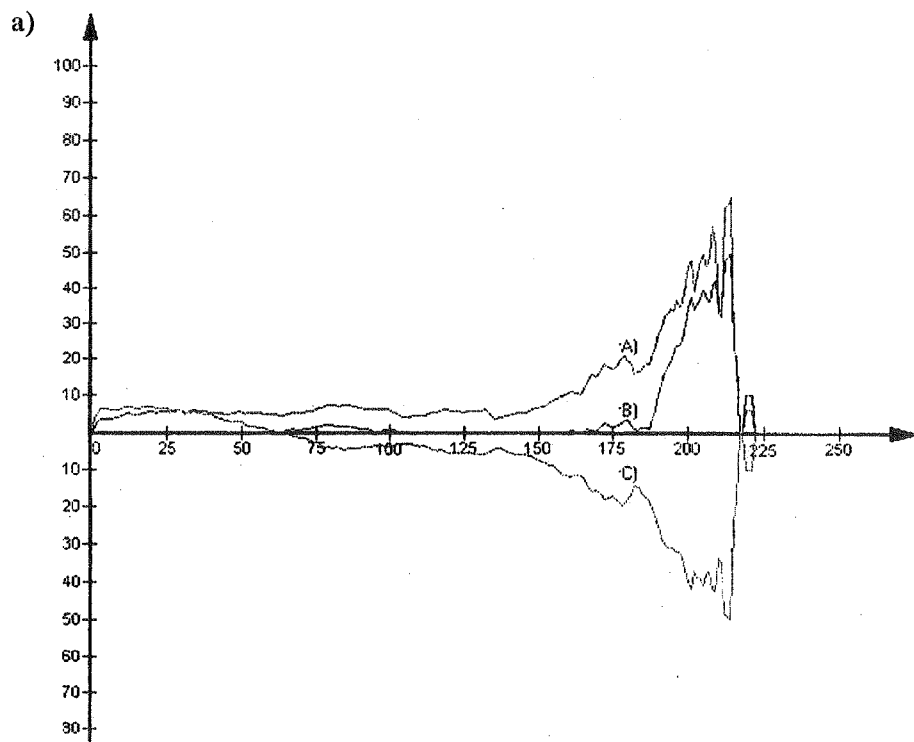
$$m_2(x) = \sum_{k=0}^{K-1} |k-x| F(k) \quad (25)$$

$$m_3(x) = \sum_{k=0}^{K-1} k F(k) - x. \quad (26)$$

Pierwsza z tych miar jest wariancją rozkładu $F(x)$, druga — średnią wartością modułu $y - x$, wreszcie trzecia — różnicą między wartością średnią y a x . Wszystkie te miary zależą od wartości x . Na rys. 2 pokazano wszystkie miary koncentracji dla kilku macierzy współwystępowania (różnych wartości r) dla obrazu z fot. 5. Na podstawie obserwacji tych wykresów i analogicznych wykresów odpowiadających innym obrazom, można wysnuć następujący wniosek.

Miary koncentracji (m_1 , m_2 , m_3) dla danego obrazu i danego r mają bardzo podobny przebieg (z dokładnością do stałego mnożnika). Ponadto, znowu z dokładnością do stałego mnożnika, przebieg ten mało zależy od wartości r .

Taki charakter zmian pozwala ograniczyć się w dalszych rozważaniach do obliczeń tylko jednej miary koncentracji i tylko dla jednego r . Pozostałe funkcje również dla

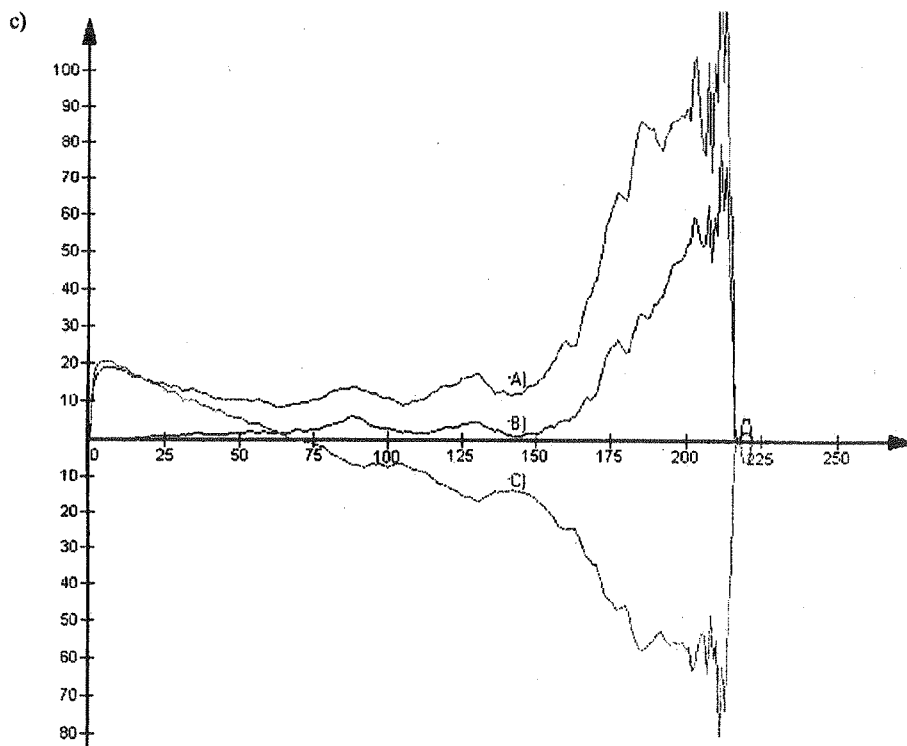


Rys. 2.

innych
dotycz
liwość

4.

Ide
w histo
zgodni
poziom
innych
w któ
o dany
obrazu
etapem
obrabia
(linear



Rys. 2. Miary m_1 , m_2 , m_3 (A, B, C) koncentracji wokół x dla macierzy współwystępowania dla obrazu z rys. 5 i różnych wartości r : a) $r = 5$, b) $r = 15$, c) $r = 25$

innych r mają podobny przebieg. Podkreślimy jeszcze raz, że powyższy wniosek dotyczy wartości r , odpowiadającym częstościom przestrzennym, dla których wrażliwość wzroku ludzkiego jest maksymalna.

4. PROPONOWANA METODA PRZEKSZTAŁCENIA HISTOGRAMU

Idea proponowanego przekształcenia obrazu polega na dokonaniu takich zmian w histogramie obrazu, aby uwydatnić bądź tekstury, bądź większe szczegóły obrazu, zgodnie z przebiegiem funkcji $m_i(x)$. Ma to sens jedynie wtedy, gdy wszystkie poziomy szarości występują na obrazie równomiernie. Gdyby pominąć aspekt innych częstości występowania różnych jasności, dochodziłoby do sytuacji, w której (ze względu na znormalizowanie funkcji $F(k)$) niewielkie ilości pikseli o danym poziomie jasności, nie mające znaczenia przy ogólnej ocenie jakości obrazu, decydowałyby o przebiegu transformacji jasności. Stąd pierwszym etapem proponowanej metody jest dokonanie wstępnej transformacji jasności obrabianego obrazu tak, aby histogram obrazu wynikowego był równomierny (linearyzacja histogramu).

Od strony matematycznej omawiane przekształcenia są opisane wzorami (1)–(5). Jeśli przez K oznaczymy liczbę poziomów jaskrawości, to żądane przekształcenie dane jest wzorem

$$j = G^{-1}[H(i)] = KH(i) - 1, \quad (27)$$

gdzie $H(i)$ jest dane wzorem (2). Jak już wspomniano, omawiana procedura daje dokładne wartości tylko w przypadku funkcji ciągłych; w przypadku dyskretnych wartości jaskrawości, równość (27) może być spełniona tylko w sposób przybliżony.

Wstępne przekształcenie jaskrawości obrazu zgodnie z wzorem (27) pozwala wyeliminować aspekt związany z inną częstotliwością występowania pikseli o różnych jaskrawościach. Dalsze przekształcenia są dokonywane na wstępnie zlinearyzowanym obrazie. Obliczana jest macierz współwystępowania (21), a następnie funkcja $m_2(x)$ (25). W celu wyeliminowania przypadkowych fluktuacji, funkcja $m_2(x)$ jest uśredniana w pięcioelementowym otoczeniu poziomu x , zgodnie z zależnością

$$m_2(x) = \frac{1}{5} \sum_{i=x-2}^{x+2} m_2(i). \quad (28)$$

Po obliczeniu uśrednionej funkcji $m_2(x)$, definiowane jest następujące przekształcenie $A(x)$:

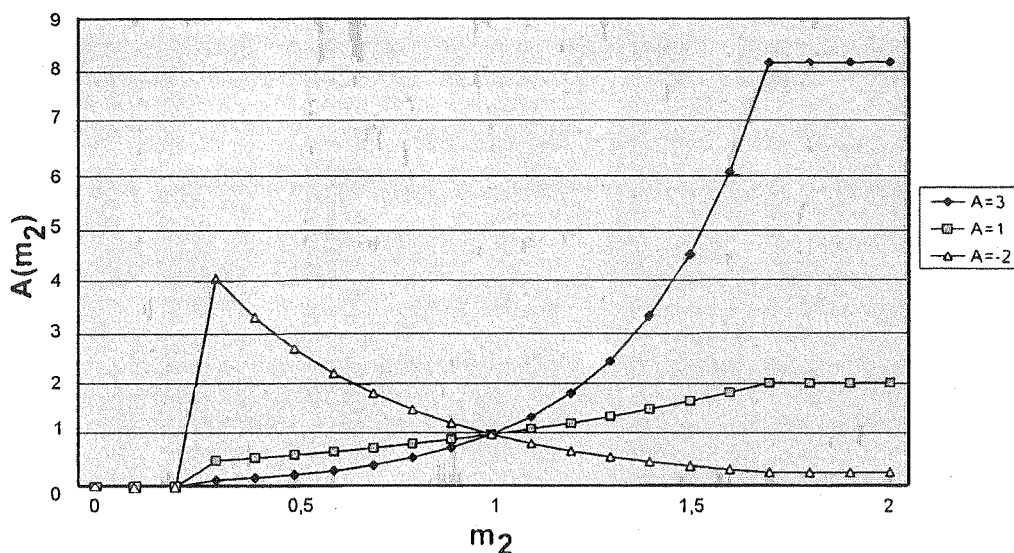
$$A(x) = \begin{cases} 0 & \text{dla } \frac{m_2(x)}{\langle m \rangle} < 0.3 \\ \exp \left[\alpha \left(\frac{m_2(x)}{\langle m \rangle} - 1 \right) \right] & \text{dla } 1.7 > \frac{m_2(x)}{\langle m \rangle} \geq 0.3 \\ \exp(\alpha \cdot 0.7) & \text{dla } \frac{m_2(x)}{\langle m \rangle} \geq 1.7 \end{cases} \quad (29)$$

gdzie $\langle m \rangle$ jest średnią wartością $m_2(x)$ określoną przez

$$\langle m \rangle = \frac{1}{K} \sum_{i=0}^{K-1} m_2(i). \quad (30)$$

Przykładowy przebieg funkcji $A(x)$ w zależności od $m_2(x)$ dla $\alpha > 0$ i $\alpha < 0$ pokazano na rys. 3. W przypadku parametru $\alpha > 0$ funkcja $A(x)$ przyjmuje wartości zerowe dla małych wartości funkcji $m_2(x)$, co odpowiada jednorodnym obszarom, wartości mniejsze od jedności dla wartości funkcji $m_2(x)$ mniejszych od średniej (tekstury, drobniejsze szczegóły obrazu) oraz wartości większe od jedności dla funkcji $m_2(x)$ większej od średniej (większe szczegóły obrazu).

W przypadku, kiedy $\alpha < 0$, funkcja $A(x)$ przyjmuje wartości zerowe dla małych wartości funkcji $m_2(x)$ (jednorodne obszary), wartości większe od jedności dla wartości funkcji $m_2(x)$ mniejszych od średniej (tekstury, drobniejsze szczegóły obrazu) oraz wartości mniejsze od jedności dla funkcji $m_2(x)$ większej od średniej (większe szczegóły obrazu). W obydwu przypadkach wartości $A(x)$ są ograniczone dla dużych $m_2(x)$, aby bardzo duże wartości $m_2(x)$ nie wypaczały przekształcenia.

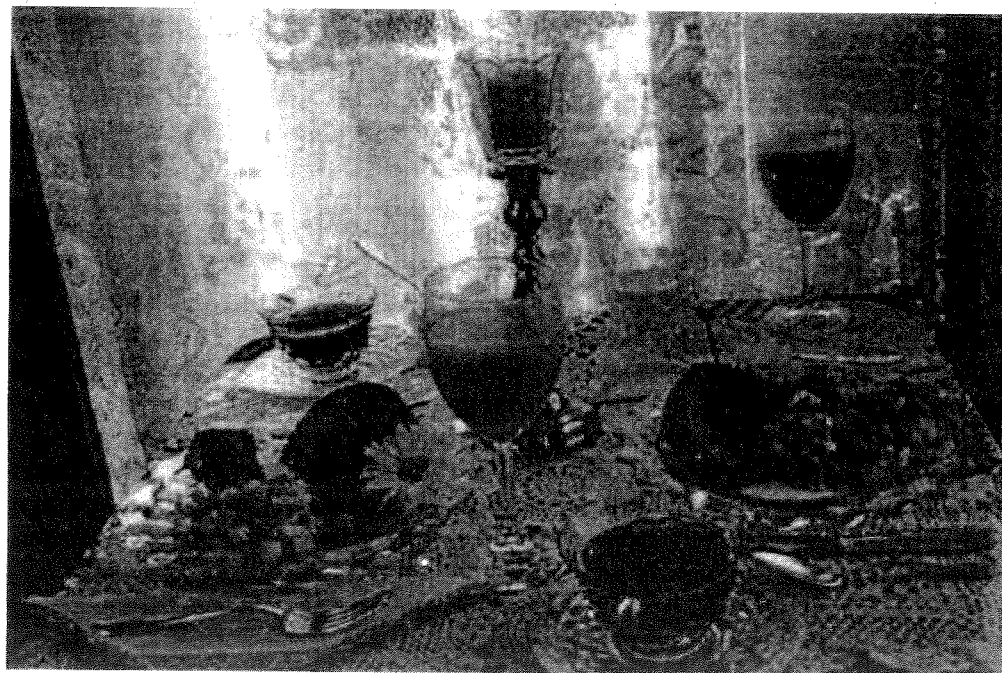
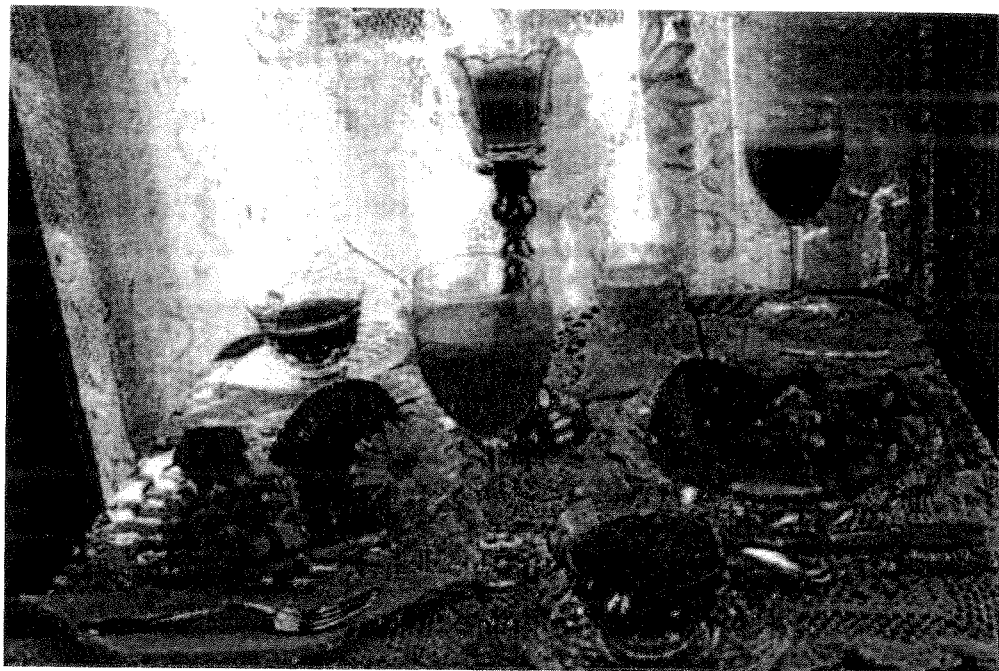


Rys. 3. Przebieg $A(x)$ w funkcji $m_2(x)$ w zależności od parametru α : a) $\alpha = 3$, b) $\alpha = 1$, c) $\alpha = -2$

Postać funkcji $A(x)$ może być inna niż tutaj podana, byle by został zachowany opisany powyżej jej charakter. W trakcie badań stosowano również i inne formy funkcji $A(x)$, jednakże najlepsze rezultaty udało się osiągnąć dla postaci opisanej wzorem (29), co nie znaczy, że jest to postać w jakimś sensie optymalna.

W następnym etapie przekształcenia dokonywana jest linearyzacja (wyrównywanie) obrazu, w którym histogram jaskrawości został zastąpiony funkcją $A(x)$. Pozwala to na rozciągnięcie (podkreślenie) tych poziomów jaskrawości obrazu, dla których funkcja $A(x)$ osiąga dużą wartość. Jeśli $\alpha > 0$, funkcja $A(x)$ osiąga duże wartości dla większych szczegółów obrazu i zostaną one w ten sposób uwydatnione. Dzieje się to kosztem tekstur i drobniejszych szczegółów obrazu. Przeciwnie jest w przypadku $\alpha < 0$; wówczas uwydatnione zostaną tekstury, co dzieje się kosztem istotniejszych szczegółów obrazu. Dobrze to widać na fot. 4, gdzie pokazano transformację obrazu przy obydwu znakach parametru α . Przy $\alpha > 0$ uwydatnione zostają większe szczegóły obrazu (obrus i filiżanki), co dzieje się kosztem tekstury na firankach. Przy $\alpha < 0$ uwydatnione zostają tekstury na firankach kosztem obrusu i filiżanek. Postać matematyczna omawianej linearyzacji jest podobna do przekształcenia danego wzorem (27) z tym, że histogram $h(x)$ został zastąpiony funkcją $A(x)$. Przekształcenie to zmienia jaskrawość piksla z jaskrawości j na jaskrawość k , zgodnie z wzorem

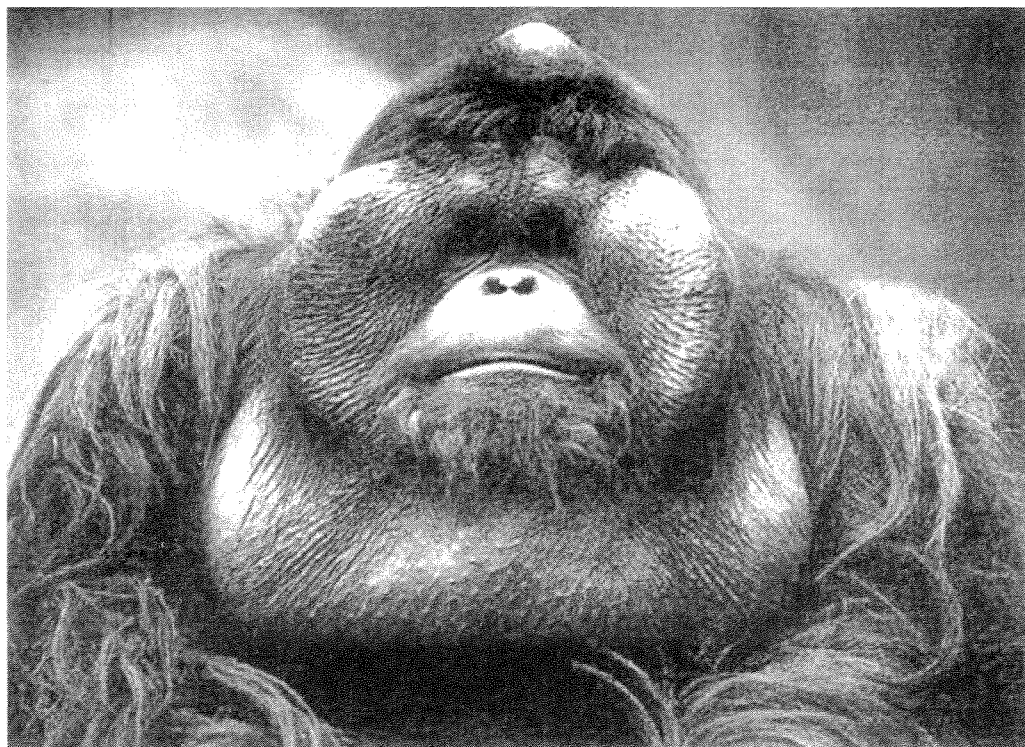
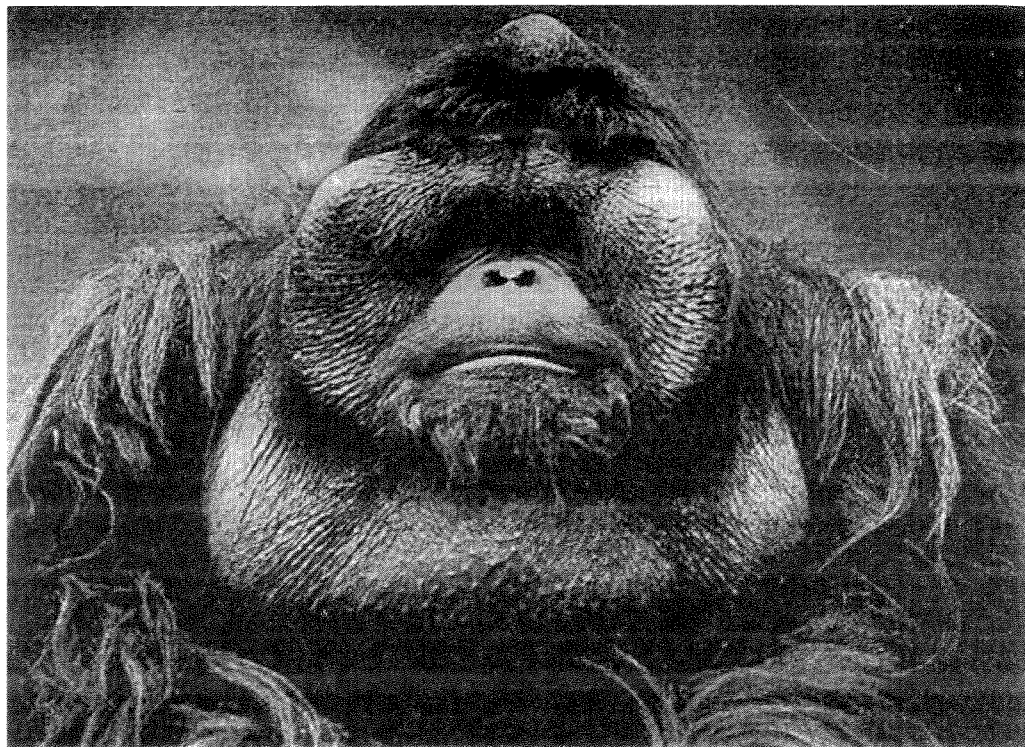
$$k = \frac{\sum_{n=0}^j A(n)}{\sum_{n=0}^{K-1} A(n)} - 1. \quad (31)$$

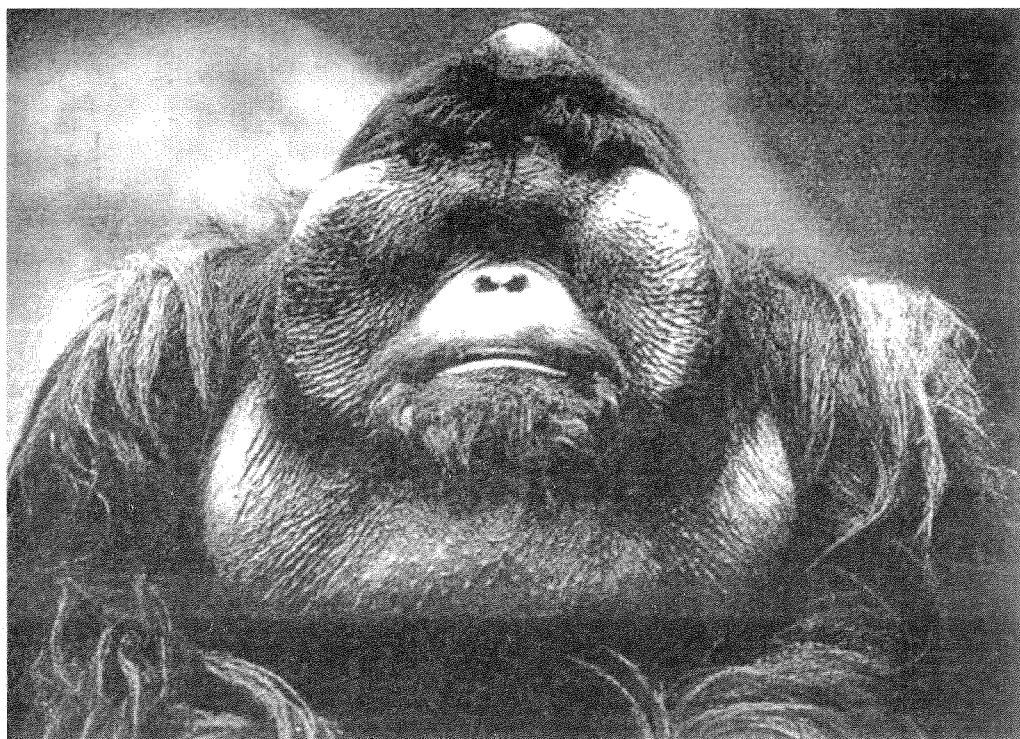
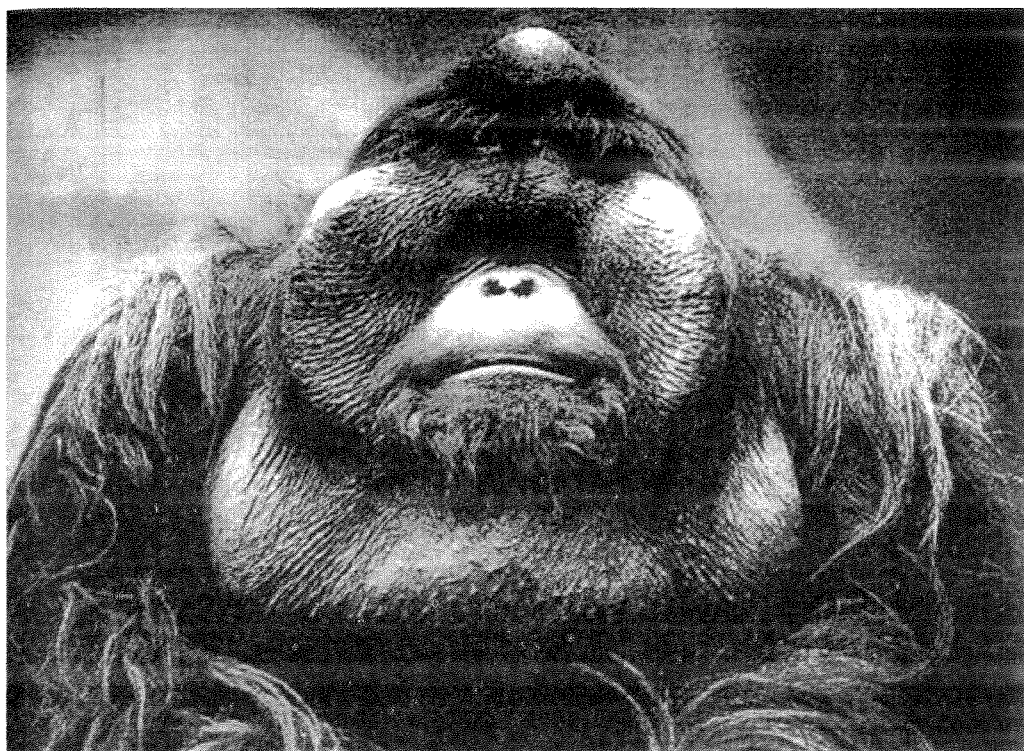


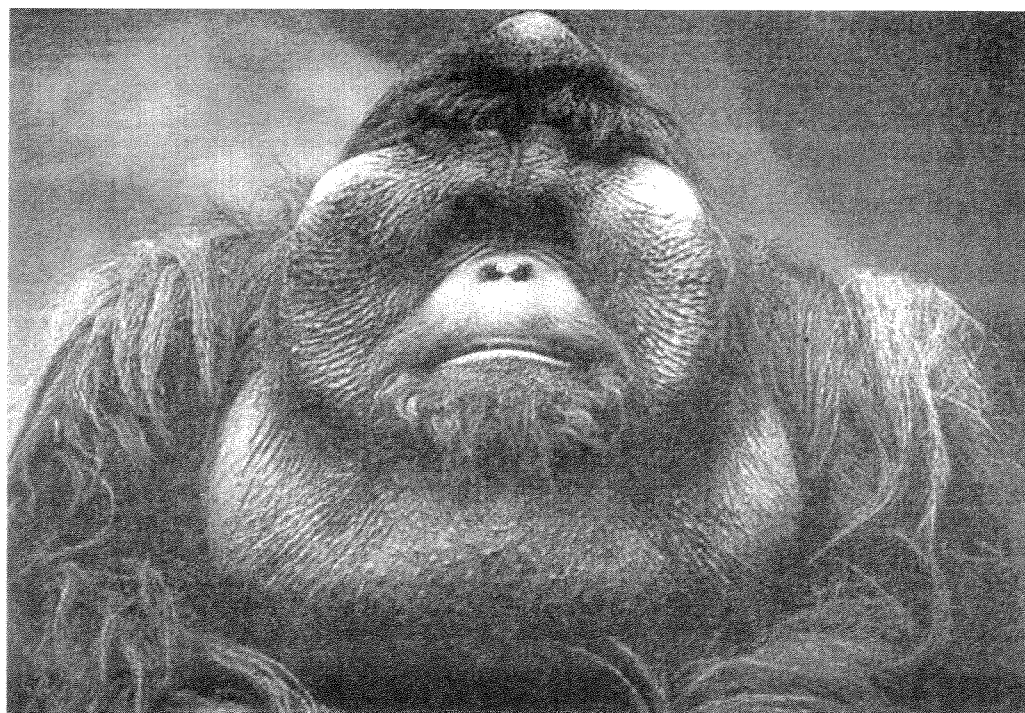
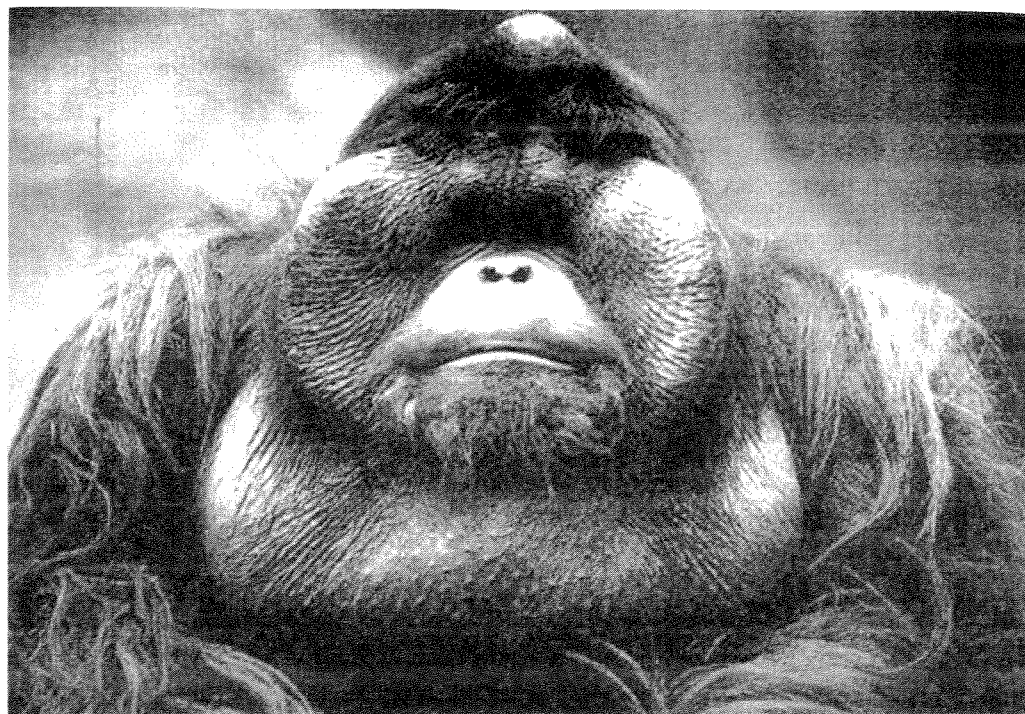
Rys. 4. Pr



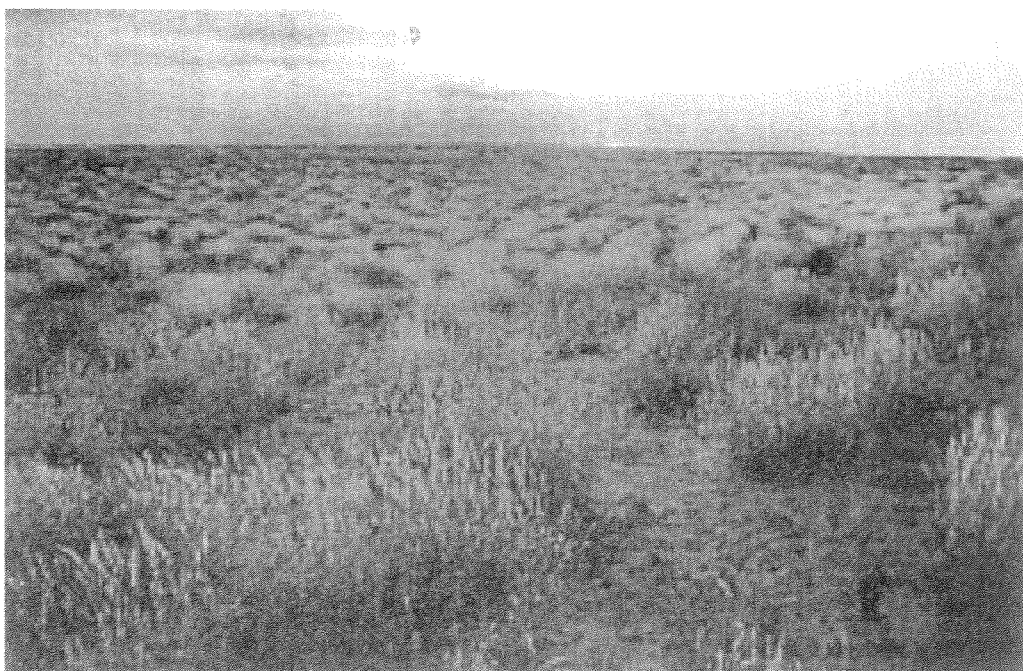
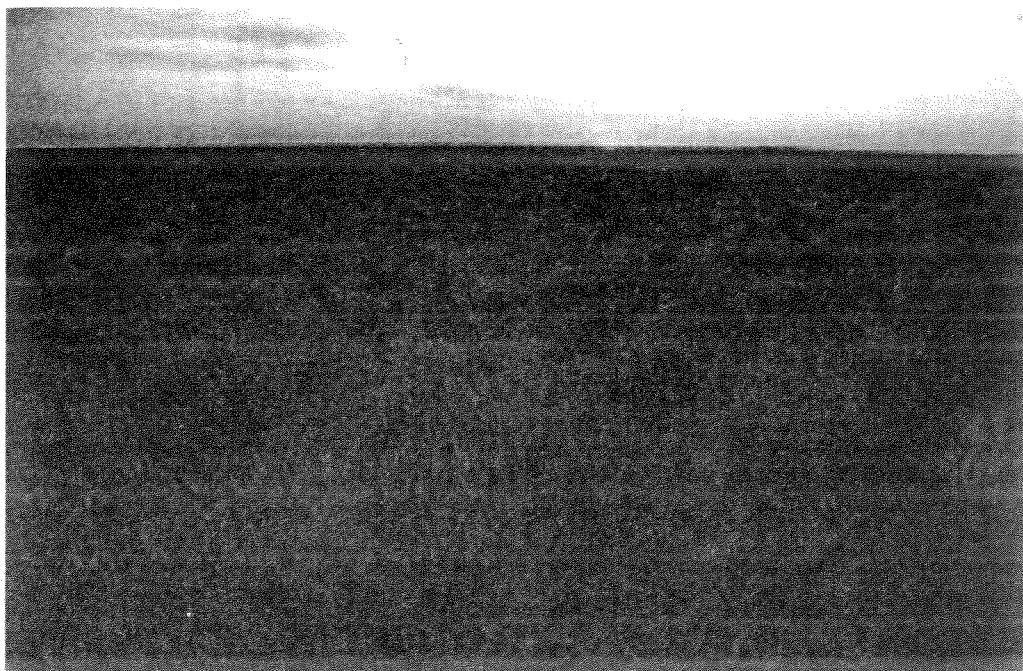
Rys. 4. Przykład transformacji obrazu: a) obraz oryginalny, b) obraz o wyrównanym histogramie, c) obraz przekształcony $\alpha = 3$, d) obraz przekształcony $\alpha = +3$





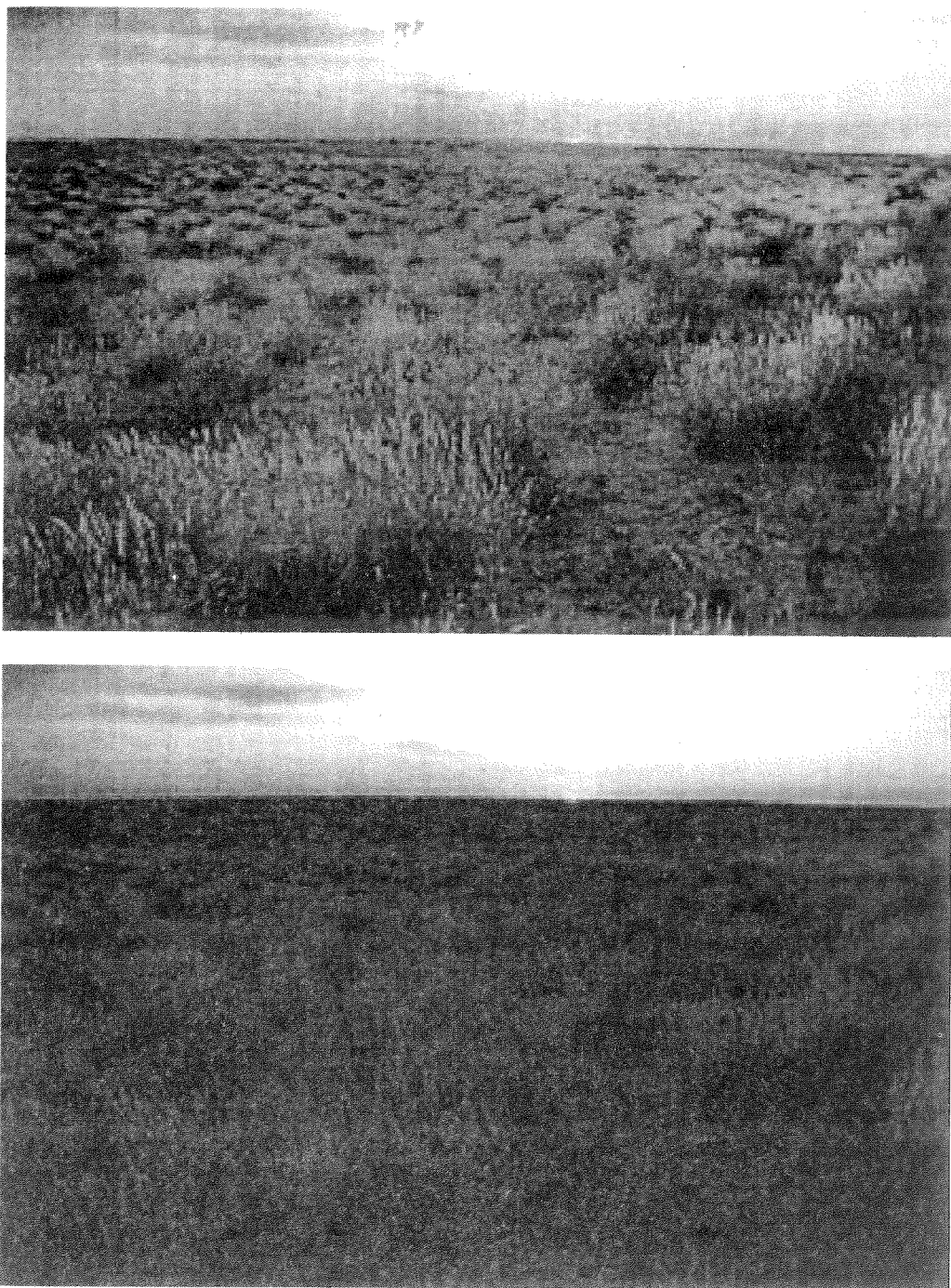


Rys. 5. Obraz przekształcony zgodnie z różnymi algorytmami specyfikacji histogramu: a) oryginał, b) obraz o wyrównanym histogramie c) opisywana metoda $\alpha = +2$, d) metoda [3], e) metoda [4], f) metoda [5]





Rys. 6
o wy



Rys. 6. Obraz przekształcony zgodnie z różnymi algorytmami specyfikacji histogramu: a) oryginał, b) obraz o wyrównanym histogramie c) opisywana metoda $\alpha = +1$, d) metoda [3], e) metoda [4], f) metoda [5]

PODSUMOWANIE

W niniejszej pracy zaprezentowano nową, globalną metodę specyfikacji histogramu obrazu monochromatycznego wykorzystującą własności macierzy współwystępowania. Opracowana metoda pozwala, w zależności od potrzeby, uwydatnić na obrazie większe szczegóły lub tekstury. Metoda może być również zastosowana do przekształcenia luminacji obrazów kolorowych. Można zaryzykować stwierdzenie, że opisywana metoda daje dla większości obrazów możliwość uzyskania subiektywnej poprawy jakości poprzez odpowiedni dobór parametru α . Zwróćmy uwagę, że wyrównywanie histogramu jest właściwie szczególnym przypadkiem opisywanej metody dla $\alpha = 0$ występującego w wyrażeniu na funkcję $A(x)$ (29).

BIBLIOGRAFIA

1. R.C. Gonzales, P. Wintz: *Digital Image Processing*. Addison-Wesley, Reading, 1987
2. T. Pavlidis: *Grafika i przetwarzanie obrazów*. WNT, Warszawa 1987
3. Y.T. Kim: *Contrast Enhancement Using Brightness Preserving Bi-Histogram Equalisation*. IEEE Trans. Consumer Electronics, vol. 43, no. 1, February 1997, pp. 1–8
4. I. Altas, J. Louis, J. Belward: *A Variational Approach to the Radiometric Enhancement of Digital Imagery*. IEEE Trans. Image process., vol. 4, no. 6, June 1995, pp. 845–849
5. M. Kamel, L. Guan: *Histogram equalization utilizing spatial correlation for image enhancement*. Proceed. SPIE, vol. 1199, 1989, pp. 712–721
6. S. Daly: *The visible differences predictor: an algorithm for the assessment of image fidelity*. SPIE Proceed., 1992, vol. 1666, pp. 2–15
7. P.G.J. Barten: *Physical model for the contrast sensitivity of the human eye*. SPIE Proceed., 1992, vol. 1666, pp. 57–72

J. SIUZDAK, T. CZARNECKI

A METHOD OF HISTOGRAM SPECIFICATION EMPLOYING THE COCURRENCE MATRIX AND ITS APPLICATION TO IMAGE TRANSFORMATION

Summary

Presented is a method of histogram specification employing the cocurrence matrix properties. Its application to monochrome image transformation is discussed. The method enables one to enhance either greater details or texture depending on the requirements. The method may be adapted also for colour images.

Key words: image processing, histogram specification, image filtering.

T
a sim
teleph
popu
netwo
netwo
the ha
uses s
tal no

Joint system for acoustic echo and noise control in hands-free communication devices

KRZYSZTOF BIELAWSKI, ALEXANDER A. PETROVSKY

Instytut Informatyki, Politechnika Białostocka

Received 1998.10.20

Authorised 1999.01.19

This paper addresses the problem of speech enhancement in the hands-free telephony where the signal to be transmitted is corrupted by ambient noise and echo signal. First, the most common issues which strongly influence the hands-free communication and commonly used algorithms, approaches to the adaptive control as well the discrete time model of the electro-acoustic path in hands-free communication and new approaches to the acoustic echo canceller and noise reduction systems are presented. Then advantages of both solutions for echo and noise control joined together in one combined system are discussed in order to increase the speech communication performance. Paper shows the leading and most promising combined systems, introduces new classification, discusses the novel adaptive echo cancellation systems for the realisation as a low cost and high performance devices on the digital signal processor.

Key words: hands-free communication, acoustic echo cancellation, noise reduction, adaptive algorithms, front-end processing.

1. INTRODUCTION

The voice communication is with humankind from the beginning. First, as a simple oral conversation then through loudhailer to the television and the mobile telephone, which are the most common used communication devices nowadays. The popularity and simplicity of the distance communication provided by the telephone network as also technical progress make this device evaluated from the analogue network telephone to the mobile cellular telephones, from the hand-held device to the hands-free device and moves to the new environment in which they are not being uses so far i.e. car. This makes the modern telephones very sensitive to environmental noises and acoustic phenomenon. This is especially in case of the small mobile

telephones. Because of the network mismatches speech signal can become even more disturbed before it reaches the recipient. For this reason the cancelling echo line in the telephone network has been successfully used since 1960's. In the echo line canceller, the adaptive filter is used to model the unknown echo path in communication channel and then cancel the echo signal from the conversation. An adaptive system is required because the echo path is not only unknown, but it is also different for each network path.

This problem of the speech transmission has been satisfying solved until the hands-free telephones appeared and acoustic properties of the environment start playing important role in signal transmission quality. In hands-free devices the conventional receiver has been replaced by the separate microphone(s) and loudspeaker, such a solution provides what we are already used with the conventional telephones, but without any restriction upon one's physical location. This is a great advantage of hands-free telephone, but also brings some additional problems in signal's quality transmission control. Additionally to the line transmission control there is acoustic signal control needed before the transmission to achieve as best as it is possible signal quality. Every realisation of the hands-free device has to deal with the problems, which appear in situation when microphone(s) is located some significant distance from the mouth of the speaker and it is in field of the loudspeaker. It means that there is a coupling present between the microphone(s) and loudspeaker, reverberation and the ambient noise, which can drown out the communication and reduce the intelligibility of the speech. However, controlling the signal quality before transmission is also solve by the adaptive systems but not so easily as in network echo line cancellation problem. Adaptive solution requirements are much higher for a much longer environment impulse response and for a much faster convergence since the echo path is time-varying during the conversation. Unfortunately, acoustic echo is not the only problem meet in hands-free device. Another one is background noise, which should be reduce as well as the residual far-end and near-end speaker's echo to provide as best as it possible signal quality.

Note that both problems of acoustic echo control and noise reduction in hands-free devices have been investigated separately from many years but only few works [1-7] have been made on system which combine this two solution in one. This paper shortly presents the recently used acoustic echo cancellers and noise reduction systems and then presents the state-of-art and perspectives to the combined system for echo and noise control in hands-free communication devices.

2. NOISE AND ECHO CONTROL IN HANDS-FREE DEVICES

There are two possible methods for echo elimination in hands-free devices. One requires the use of a unidirectional microphone, which is carefully positioned so that it picks up little sound generated by loudspeaker and heavy padded walls. This

TOM

techn
restr
15 d
not s
doub
cing
comb

M
electr
and
electr
is old
becau
Canc

E
(DSP
Coupl
are el
backg

A
the s
noise
simila
Canc
it rec
replac
the si

M
11, 1
first i
respo
non-li
two-si
input
altere
comp
drive
presen
filter.

An
modif

) s

technique requires extensive and expensive treatment of the environment and still restricted position of the speaker. But, it can reduce the direct pickup by about 15 dB¹⁾, which is not enough to prevent noticeable echo for far-end speaker and do not satisfy the ITU-T Recommendation for ordinary telephone in both situation of double and single talk condition. This technique is used in old fashion teleconferencing broadcasting and it is too expensive for use in a car and after all must be combined with some other device to echo control.

More effective technique is to use an acoustic echo canceller [8, 9, 10], which electronically removes audio coupling and allows almost free placement of speakers and microphones without any special acoustical treatment needed. There are two electronic techniques to eliminate the echo in hands-free telephone. Echo Suppression is older of these two echo control methods and will not be presented in this paper because it is now rapidly falling out of favour to the second one called Echo Cancellation [8, 9].

Echo Cancellation techniques employ the high speed digital signal processor (DSP) which electronically compares the microphone signal to the speaker signal. Coupling or echoing causes any similarities in these two signals. These similarities are electronically subtracted from microphone signal leaving only local speech and background signals.

As said before, echo is not only one distortion in hand-free communication the second one are background noises. There is much technique for control noise level in the signal contaminated with additive noise. One method is similar to Echo Cancellation technique. In this method the adaptive Noise Cancellation [11, 12] is used to clean the microphone signal from noises, it requires at least two sensors, but in real live this technique is mostly replace by the some statistical signal estimates and a priori assumptions of the signals used in one channel solution.

Most general adaptive system employs Adaptive Signal Processing (ASP) [10, 11, 13, 14] element presented in Fig. 1. The ASP element has two components: first is the adjustable or programmable filter. For this purpose the finite impulse response (FIR) filter, infinite impulse response (IIR) filter or some other type of non-linear filter can be used. If delay is allowed, the adjustable filter can be two-sided or non-casual. The adjustable filter has auxiliary signal input x and also input from the adaptive algorithm, which causes the filter parameters to be altered. The adjustable or programmable filter produces an outputs which is compared with the primary signal y to generate the error signal e that is used to drive the second component of ASP element — adaptive algorithm. In general, the presence of the adaptive algorithm makes the ASP element a *non-linear system or filter*.

An adaptive algorithms which adjusting the adaptive filter to optimum, are modification of standard iterative procedures for solving minimisation problems in

¹⁾ See section 7 for definition of performance measure (ERLE).

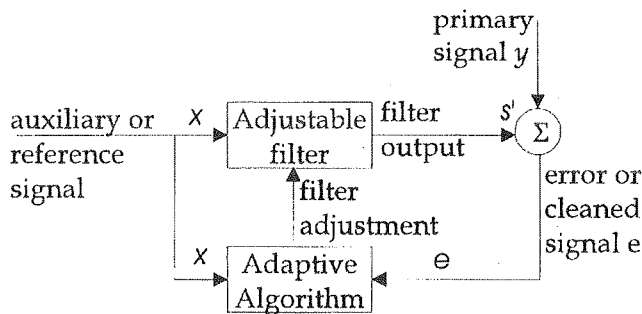


Fig. 1. Adaptive signal processing element

real time setting. the most common criterion used in adaptive minimisation is error criterion as Mean Square Error (MSE) [14] namely,

$$E = E(w) = E[e^2(w)] = E[(y - s')^2] = R_y - 2R_{xy}^T w + w^T R_x w \quad (1)$$

where, $w^T(n) = [w_1, w_2, w_3, \dots, w_N]$ is filter's coefficient vector, y is desired output, s' is ASP element estimation product for output y .

$$R_y = E[y^2(n)] = \text{variance of } y(n)$$

$$R_x = E[x(n) x^T(n)] = \text{covariance of matrix } x(n) \quad (2)$$

$$R_{yx} = E[y(n) x^T(n)] = R_{xy}^T \text{ covariance between } y(n), x(n).$$

The two basic iterative methods are Steepest Descent method (SDM) and Newton's method (NM). Both methods [10, 13] are based on assumption that filter's coefficients are updated according to the equation

$$w(n+1) = w(n) + \delta w(n). \quad (3)$$

The SD algorithm has the form

$$\delta w(n) = w(n) - w(n-1) = -\frac{1}{2}\mu \nabla (w(n-1)) \quad (4)$$

where $\nabla(w)$ is the MSE gradient defined in (5)

$$\nabla(w) = dE/dw = -2R_{yx} + 2R_x w. \quad (5)$$

The NM has the form

$$\delta w(n) = -\mu H^{-1}(w(n-1)) \nabla(w(n-1)) \quad (6)$$

where $\nabla(w)$ is the MSE gradient and

$$H(w) = d^2 E/dw dw^T = 2R_x \quad (7)$$

is the Hessian.

As could be notice the MSE criteria required the covariance matrix to be known, but in real setting statistical information about signal is usually not available, therefore approximations are used to construct the working algorithm.

In very general sense almost all adaptive algorithms have the same form, namely

$$\mathbf{w}_{new} = \mathbf{w}_{old} + \text{gain} \times \text{gradient} \times \text{error}.$$

The error is a scalar signal, the gradient is a vector signal, the gain is a scalar or a matrix. In order to calculate the gain, gradient and error at each time, it is only required to store a fixed span of previous data. The size of the gain or step size controls the speed of adaptation.

Priori to adaptive algorithm is the specification of the filter structure. Mainly the adaptive filters can be realised as:

- transversal — finite impulse response (FIR)
- recursive — ionfinite impulse response (IIR)
- lattice filter
- transform-domain filters.

3. ADAPTIVE ALGORITHMS USED IN ECHO AND NOISE CONTROL

In adaptive digital signal processing, discreet-time filters of a transversal structure with N variable coefficient are commonly used. An algorithm used to update the coefficient is the Normalised Least Mean Square (NLMS) algorithm [13, 15, 16] derived from Steepest Descent method. This algorithm has the low complexity, simple implementation structure, but the convergence characteristic of the filter using this algorithm are influenced by the statistical properties of the input signal. But so far this algorithm is widely used in time-domain solution for adaptive echo and noise control, resulting in many modification and advance step size control mechanisms [15, 17, 18]. At the same time, Recursive Least Square (RLS) algorithm [13, 18, 19] based on Newton method are the starting point to the other more promising algorithms. To avoid convergence dependencies of the signal's statistic the RLS perform decorrelation in time-domain, where decorrelation order is equal the filter order. However RLS algorithm can not be used in real-time application because of the autocorrelation matrix inversion involved, which makes it very large computational. High performance delivered by RLS algorithm, and one sample interval delay makes it the starting point to construct the new time-domain algorithm with reduce complexity, as Block Orthogonal Projection (BOP) algorithm [13, 18].

Another method to increase the performance and reduce the complexity is to perform adaptive filter operations in frequency-domain [20, 22]. This gives following algorithm — Block Frequency Domain Adaptive Filter (BFDAF) [13, 23] which complexity and long processing delay is improved in Partitioned BFDAF (PBFDAF) or Decoupled PBFDAF (DPBFDAF) [13, 17, 18, 22—26].

The Discrete Fourier Transform used in the frequency domain algorithms presented earlier can be interpreted as an analysis bank splitting the input signal in

different frequency bands. By using a better analysis bank than the DFT, the Subband Adaptive Filter (SAF) is created [18, 27–30], which is the next possible approach to the adaptive filtering. Every frequency band is adapted separately in the hope that there is only a small amount of correlation between adjacent bands. Bands adaptation is done by the most common used adaptive algorithms but the popularity have NLMS, BFDAF, Affine Projection Algorithm (APA) [13, 22]. In SAF usually an (almost) uniform filter bank is used but this is a rather arbitrary choice. In acoustic echo cancellation a multiresolution approach [18, 28, 31] involving a tree structure of filter banks, offers many advantages. In this case, the signal is split in two bands, the lowest band is again split in two bands, and so on. Since the quality of echo suppression is judged by the human ear, which analyses sounds with a filter bank that is uniform on the Bark ($\sim$ logarithmic) scale, it is desirable to adapt the filtering criterion to this scale which can be done by the Wavelett transform and multiresolution filtering [30–33].

The optimal algorithm should have “no delay”, “low” complexity and “optimal” performance, let's named it IDEAL. The graphical comparison of the mentioned algorithms is presented in Fig. 2.

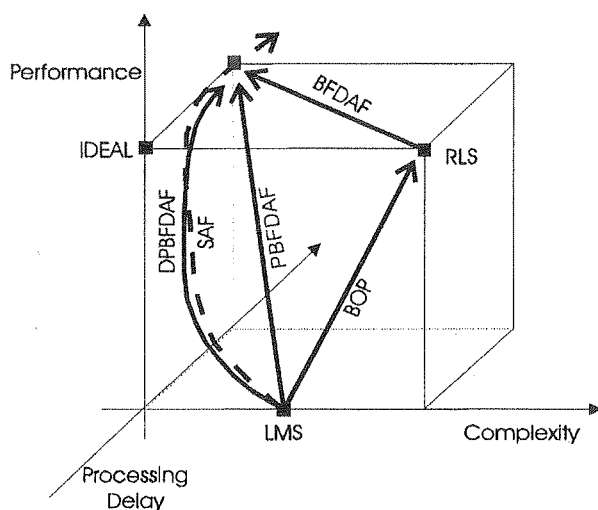


Fig. 2. Comparison of the adaptive algorithms

The LMS algorithm family has indeed (almost) no delay, but complexity and especially performance are not too good. An efficient RLS algorithm, introduces more complexity for a better performance, with an efficient BOP implementation lying somewhere between LMS and RLS. The frequency domain transformation in BFDFAF is only efficient for large block-sizes, but then a lot of delay is introduced. By taking an arbitrary PBFDAF algorithm a DPBFDFAF algorithm can be constructed with lower complexity and better decorrelation properties. A SAF with a comp-

lexity and performance comparable to a DPBFDAF can have a larger delay but this depends on the implementation properties.

4. GENERAL ACOUSTIC ECHO CANCELLER STRUCTURE

The idea behind acoustic echo cancellation is to predict the echo signal value and thus subtract it out. Putting an adaptive signal processing element on the side of the transmission link does this removal of re-occurring signal, when we have acoustic loop. Adjustable filter of ASP element estimates the impulse response (example is shown in Fig. 3) of the echo path(s) and subtracts this estimation from the microphone signal, to obtain an estimate of the input without echoes.

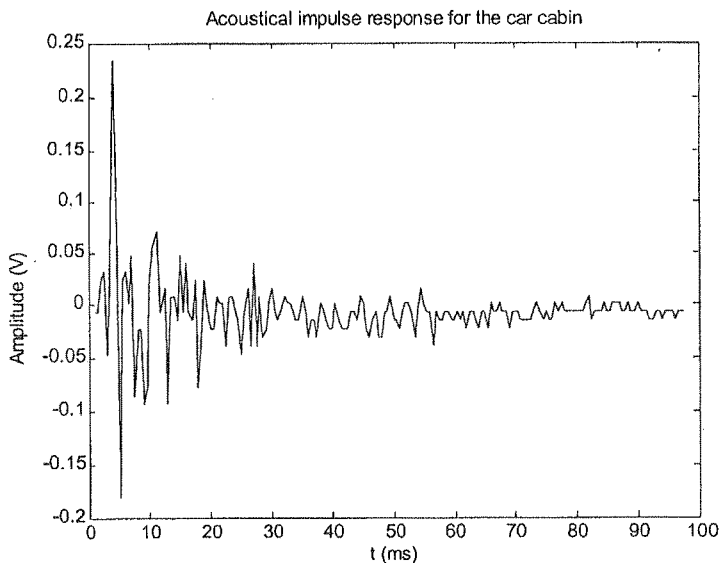


Fig. 3. Example of the impulse response in mid-size car

The first task of the AEC is to suppress the direct path echo to avoid that the closed loop gain is larger than one and the loop becomes unstable [34]. The second task is to suppress the indirect echoes in such a way that the participants do not hear their own echoes. To achieve this the adaptive filter must model a large part of the room's echoes and the number of required coefficient is linearly related to the reverberation time. AEC application requiring filters with long impulse response, it is not uncommon that thousands of FIR coefficients are needed to achieve the desired performance. The level of performance is directly influenced by the statistical properties of the adaptive filter input signals.

For this reason the ASP element is used in order to evaluate the correlation between far-end signal $x(t)$ and echo signal $x'(t)$ residual in signal picked up by microphone $s(t)$ to estimate the impulse response of the echo path [29]. During ASP

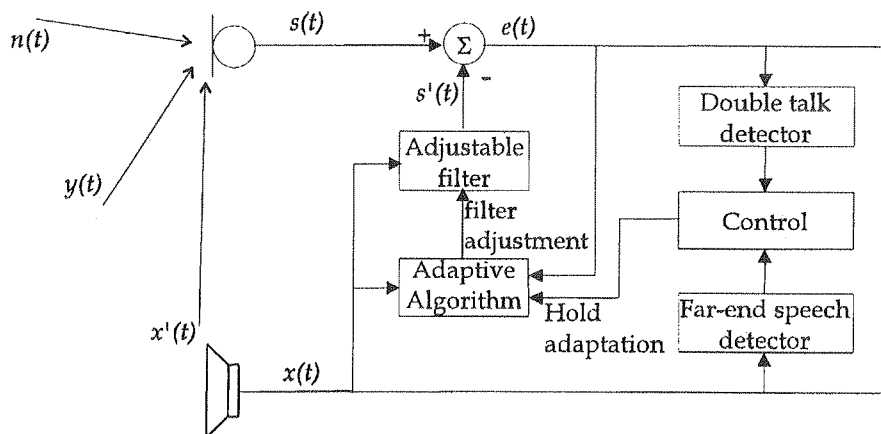


Fig. 4. Discrete time model of acoustic echo cancellation

element training mode, it is assumed that the local user is not speaking through the microphone, and the far-end user is assumed to be talking. To meet this requirements the control mechanism are used (see Fig. 4). System control is based on two detectors: Double Talk Detector (DTD) [8, 9, 35, 36] and far-end speech activity detector more often replaced by Silence Detector (SD) [8, 9, 35, 36]. SD detects the silence period at the far-end signal. Adaptation process must be held, if silence or double talk are detected to avoid the corruption of the adjustable filter coefficients. First, because there is no echoes if there is no far-end signal. Second, because the optimal filtering using steepest descent algorithm can not be done during double talk periods.

5. GENERAL ADAPTIVE NOISE CANCELLER STRUCTURE

There are two approaches to control the noise: passive and active. The traditional approach uses passive techniques such as enclosures, barriers, silencers to attenuate the undesired noise. But when we analyse the subject in the context of hands-free communication this technique is less interesting. Therefore, active noise control is more liked because the characteristic of an acoustic noise source and environment are not constant and undesired noise are nonstationary so control system must be adaptive in order to cope with these changing characteristics.

Adaptive noise control technique is based on subtracting noise from a received signal, which is controlled in adaptive manner for the purpose of improved signal to noise ratio (SNR). Adaptive Noise Canceller (ANC) presented in Fig. 5 is a dual-input adaptive system [10–13, 37]. The inputs of the system are derived from pair of sensors: a primary sensor and a reference (auxiliary) sensor. The primary sensor receives information bearing signal corrupted by additive noise. The signal and noise are uncorrelated with each other. The reference sensors receive a noise that is

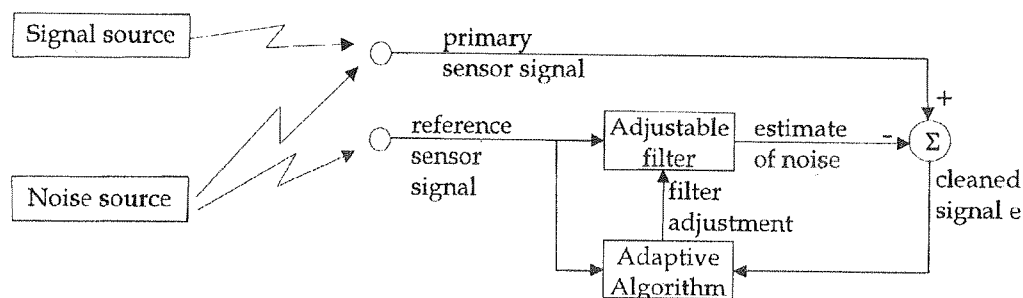


Fig. 5. Adaptive noise cancellation structure

uncorrelated with the signal but correlated with the noise in the primary sensor in unknown way.

As can be seen in Fig. 5 the ANC structure has the same properties as AEC structures and as the echo cancellers can be implemented with the same wide of structures and algorithm. But the noise control in adaptive manner based on two sensors extend the single microphone acoustic echo canceller approach, besides the car environment is very restricted, microphones can not be fixed in such a way that they form the reference input which will be free of speech signal. For this reason the developers are mostly concern on the single sensor noise estimation using low-cost method as Spectral Subtraction (SS) [38] performed in frequency domain and estimate the noise spectrum during non-speech periods of the conversation. The other one is Adaptive Predictor (AP) totally performed in time domain. AP systems introduce delay in transmitted signal equal speech pitch period, while the noise is much less uncorellated.

The most popular adaptive single sensor speech enhancement technique is based on the Wiener solution. Unfortunately, there is a conflict between ratio of the noise reduction and quality of the resulting speech signal, which is degraded by the unpleasant tonal noises. For this reason the modification of the Wiener filter to noise reduction is made resulting in Minimum Mean Square Error (MMSE) estimation performed in frequency domain with Short Time Spectral Amplitude (MMSE-STSA) suppression rule [3], or its derivative the Logarithmic Spectral Amplitude Estimator (MMSE-LSA).

6. APPROACHES TO COMBINED SYSTEMS FOR ECHO AND NOISE REDUCTION

The previous paragraphs explain that the hands-free communication systems require the solution of the two fundamental problems. First, coupling and echoing require use of an echo control device to guarantee the stability of the electro-acoustic loop and to supply sufficient echo reduction. The second problem to be solve is the reduction of noise which becomes necessary due to the relative large distance from

the microphone to the mouth of the speaker especially in mobile communication environment where the signal to noise ratio can be very low. The above problems were investigated independently for many years resulting in spread wide of algorithms and structures, but only combined approach can deliver the satisfying performance for both echo and noise reduction.

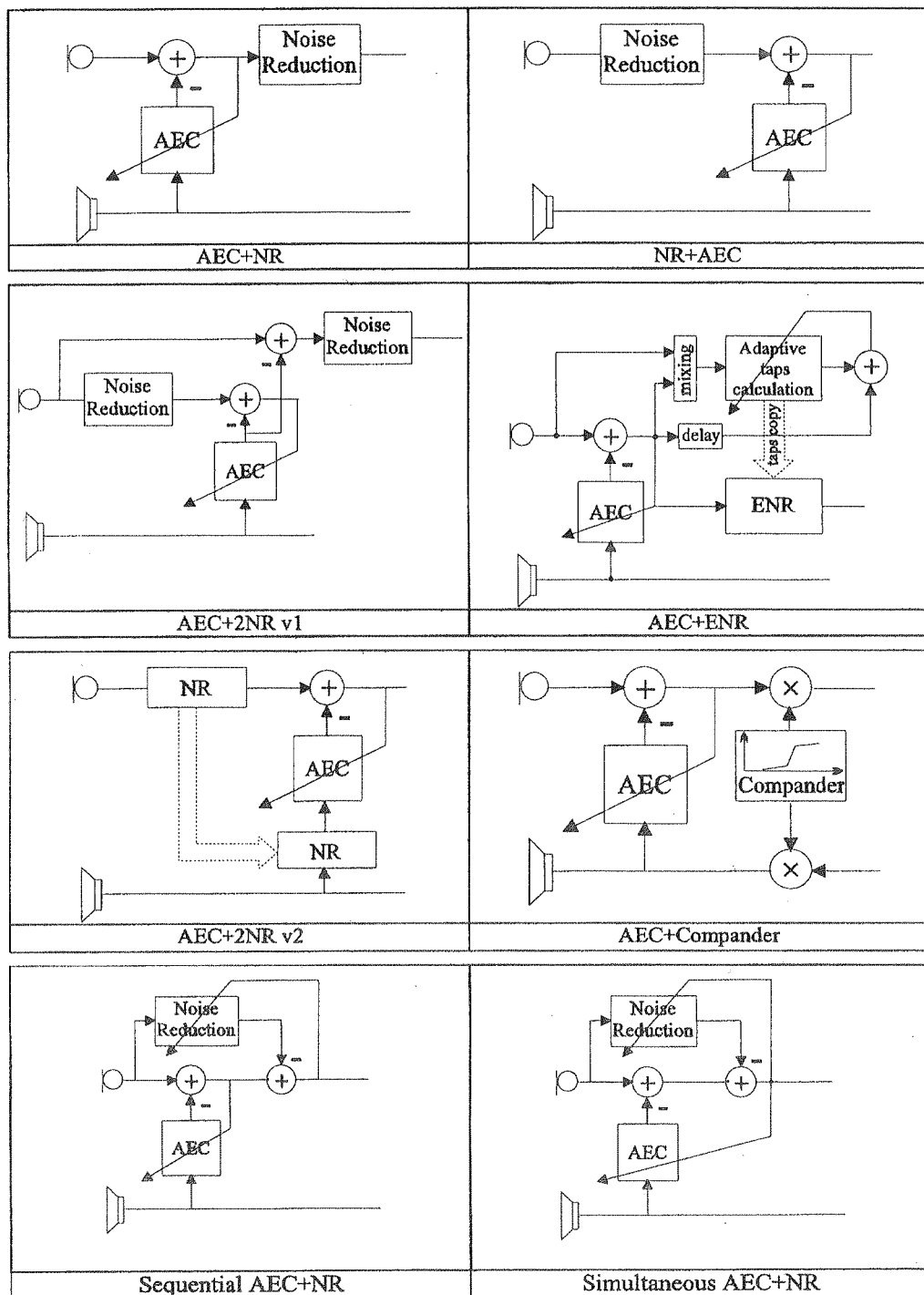
As mentioned above the hands-free telephone set-up makes that far-end speaker signal highly correlated with echoes at the microphone signal. Therefore, the combined system will have the echo canceller, which almost ideally solves the echo reduction problem. But from insufficient canceller convergence or insufficient compensator order there is need to employ some additional device to reduce the residual echo as well the ambient noises. For this, the combine systems with echo cancellation and noise reduction must be consider in case of these two processing operation order. Table 1 depicts principal cases. The optimal structure is composition of the two cascaded optimal filters where the Acoustic Echo Canceller (AEC) systems proceeds the Noise Reduction (NR) system as shown in AEC + NR (see Table 1). This structure order is used in most time and frequency domain solution. The inverted order structure (NR + AEC) could be also used to reduce the noise influence on the AEC systems. However processing done by the NR system reduces the noise, but disturbs the original echo signal. The change of echo done by NR in NR + AEC structure, introduce the non-linear changes into the input signal to the AEC. This occurrence can slow or even stop the convergence at the not satisfying level, and yet it gives the better performance in speech detection during the double-talk period and less disturbance in coefficient vector optimisation process.

Unfortunately, the disturbing noise is less reduced, if it do not disturb the AEC, and compensator almost always leaves the residual echo signal due to the echo change in NR or insufficient filter order. Therefore, to exploit advantages of NR + AEC structure order with the additional noise reduction filter at systems output, which can be combine with additional echo reduction is needed resulting in AEC + 2NR v1 and AEC + NR v2 combined system structure. The NR operation enhances the Signal to Noise ratio, but it also introduces the non-linear distortion to the original echo, so the copy of the NR filter in the AEC branch is aimed to reduce this potential disturbance.

According to the [2] the AEC + NR with second order Affine Projection Algorithm (APA 2) it is much more efficient compered to the NR + AEC even with the additional NR as it is depict in AEC + 2NR v2 (see Table 1). Quite new modification made to the AEC + NR structure are proposed in [49] resulting in Sequential AEC + NR and Simultaneous AEC + NR, where the NR unit is assumed to be an Acoustic Noise Canceller (ANC). Basis of this approach is made at assumption that echo and near-end speech have the similar power spectral densities, so the Wiener solution to filter noise from the noisy echo and from the noisy near-end speech are not very different. The advantages of the Sequential AEC + NR structure is that the noise and echo canceling filters are completely decoupled, after

Table 1

The main structures of the mono-channel combined echo cancellation and noise reduction systems



they converge, but systems reach low ERLE²⁾ in noisy environment. The Simultaneous AEC + NR achieve the better ERLE, because adaptation of echo canceling filter does not suffer due to the noisy error signal. However, in this structure NR unit introduces the distortion to the near-end speech.

The structure mentioned so far can be classify as classical cascaded structure of the combined system AEC + NR, NR + AEC, AEC + 2NR, as well Sequential AEC + NR and Simultaneous AEC + NR assumed that NR system performs reduction based on compensated signal or else microphone signal, also noise attenuation is done (due to the single microphone system) using frequency domain MMSE noise reduction technique for speech enhancement or most common algorithm like Wiener or Power Spectrum Subtraction (SS) suppression rule. But there are few novel approaches which used the modelled (estimated) echo signal and microphone signal to estimate the noise signal to be cancelled or masked, as in AEC combined with Compander (AEC + Compander structure, see Table 1). Other solution to the noise and residual echo attenuation has been proposed in [1, 2, 42, 43] where the selective frequency attenuation technique is used for echo and noise reduction in sending path. The mentioned algorithm used the mixture of microphone and compensated signal as a input to additional noise and echo adaptive filter with compensated delayed signal as a reference. The approach is sketch in Table 1 and named AEC + ENR. Of course the reverse configuration can be realised where combined echo and noise reduction filter (ENR) proceeds the AEC, but authors of experiments clearly show that the configuration AEC + ENR is preferable.

The short composition of a available adaptive algorithms used with the presented structure (see Table 2) gives the situations overview done in order to evaluate the idea of the combined system for echo and noise reduction. Not every structure was implemented as a real time application many of them are on the theoretical and experimental stage of research. But there are some structures, which are successfully realised as a real time application. The main structure is AEC + ENR, which performs calculation in time-domain, introduces one sample period processing delay, and achieves satisfying performance. Unfortunately in noisy environment single microphone AEC + NR solution do not attenuate the noise.

So far the main effort is concentrated to make the single microphone noise and echo reduction system more effective and design the algorithm which has the opitmised and reduce processing delay. But if only echo compensation is concerned as a main task of the system there is no need to employ more then one microphone. However, if we want to reduce the background noises in the very noisy environment where useful signal reached the low signal-to-noise ratio the multiple microphone systems should be concerned. Even if this solution introduces the additional hardware, more complex algorithms and control mechanism it is worth the effort, because multimicrophone algorithms can exploit the spatial properties of the sound

²⁾ See section 7 for definition.

Table 2

Overview of adaptive algorithms used in combined system structure

Combined System Structure	Algorithm of Echo Compensator	Algorithm of Noise Reduction	Additional control Devices
AEC + NR	• NLMS • Fast RLS • Subband APA	• MMSE-LSA • MMSE-STA • SS	• DTD • SD
NR + AEC	• NLMS • BFDAP • APA	• APA 2 • MMSW-STA • SS	• DTD • SD
AEC + 2NR	• BFDAP • NLMS	• APA • MMSE-STA	• DTD
Sequential AEC+NR and Simultaneous AEC + NR	• NLMS • SNLMS* • NLMS/SNLMS	• NLMS • SNLMS • NLMS/SNLMS	• DTD • SD
AEC + ENR	• NLMS	• NLMS	• DTD • SD

* SNLMS- Subband NLMS

field, which give the most robust algorithms in presents of the nonstationary noises. Besides the knowledge of spatial properties are useful in reduction of potentially annoying reverberation of the near-end speaker, which is very important in case of long reverberant environment with the relatively high noise level. But there are some restriction which are not easy to overcome in the multimicrophone systems, as:

- how to compensate the echo in each microphone signal, and reduce its complexity
- how combine several signal in one combined and compensated signal
- how explore the delivered spatial properties of the sound field and make system able to be run in real time
- time delay compensation if the near speaker is not in symmetric position with respect to the microphones.

The great advantage of using a multimicrophone noise reduction system has been exploit in the work of Allen [39], Kaneda [40] and Zielinski [37]. They introduced a microphone array with subsequent Wiener filtering which can reduce the received noise signal considerably. This approach exploit fact that the speech signal is received as a correlated signal in the microphones whereas the noise signals appear to be mutually uncorrelated. There are only few works, so far which employ multimicrophone technique in combined system for echo and noise reduction. The concepts of single microphone system presented earlier can be transferred into two multi-channel approaches where the main structures are presented in Fig. 6. These two structures have the acoustic echo cancellation realised on each microphone channel. The noise reduction filtering is applied to the compensated signals (see structure multi AEC + NR v1) but in structure multi AEC + NR v2 the noise reduction filter taps computation are based also on input signal before compensation on which the analysis for noise reduction performed on compensated signal is

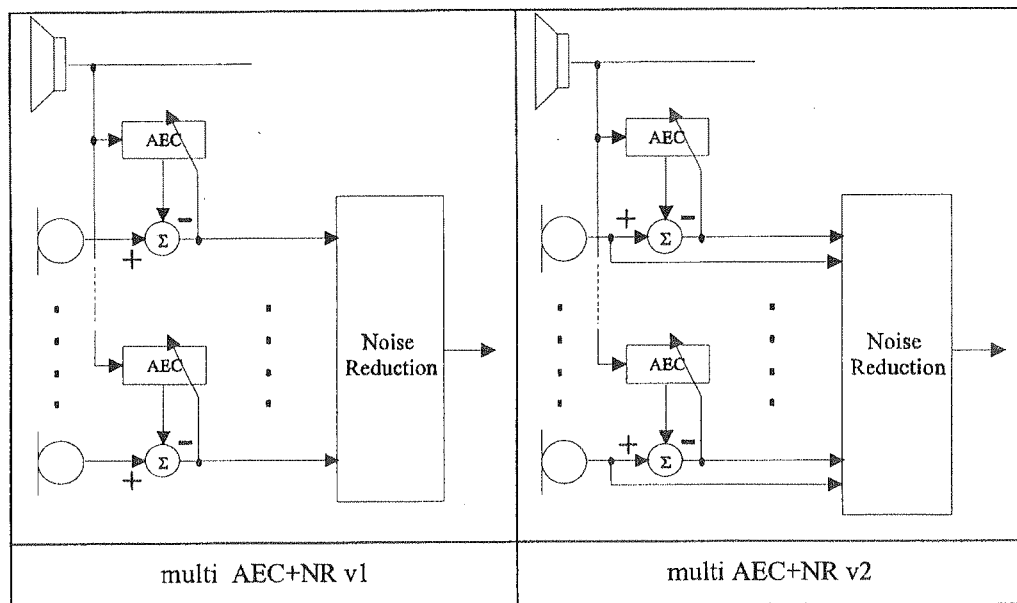


Fig. 6. Multi-channel structure of the combined systems

done. The advantage of these combinations is that the coherent components of the far end speech are removed before the signals are passed to the noise reduction. The noise reduction has to deal only with the reverberated sound which is mainly not coherent and noise. The obvious disadvantage of this structure is that systems need an echo canceller at each microphone channel. These increase the complexity of the system and may not meet the real time systems expectations. Experiments show that this approach is quite reasonable and meets expectations for two microphone designs. For more than two microphones this approach is impractical. However, modification of structure multi AEC + NR v1 according to the principle of the of echo and noise reduction depict in AEC + ENR single microphone order structure can be applied to the multimicrophone approach resulting in two microphone systems presented in Fig. 7.

Due to the difficulties found in multi-channel systems with AEC in each channel the combined system can be stated using noise reduction system before passing any signals to the echo canceller [1]. The idea is presented in Fig. 8. The noise and the reverberation reduction now introduce a time varying filter into the ech path. The copy of that filter is placed before the far-end input of the echo canceller. Because of the time variant behaviour the AEC filter must be recomputed each time the coefficient of the noise reduction filter has changed. It increases complexity, because each time when delay in channels occurs it must be compensated before the microphone signals are mixed. Due to complex synchronisation control to keep

Fig. 8

AEC
comb
S
reduc
proc

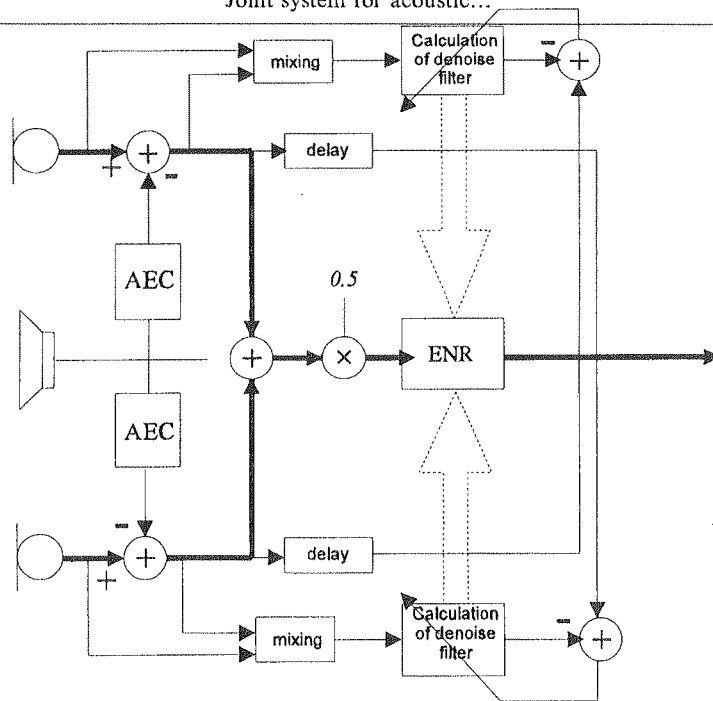


Fig. 7. Two microphone system based on the AEC + ENR concept

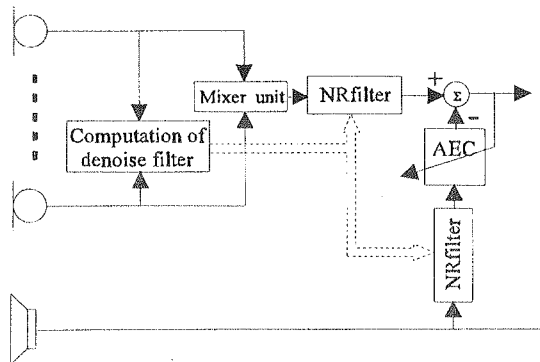


Fig. 8. Block diagram of combined system with the noise reduction before passing the signal to the echo canceller

AEC tracking system response, multi-channel configuration is not often used in combined system.

Simulation done for each system clearly shows its advantage to echo and noise reduction. The real time implementation of multi AEC + ENR system on DSP processor exploits higher properties for noise reduction than in single microphone

AEC + ENR system. But for more than two microphones systems the mixing function as well as the time delay compensation algorithms are currently investigated.

The multi microphone systems presented so far operate mainly on the same range of adaptive algorithms such as single microphone structures. Microphone array in frequency domain is quite promising technique for solving the problem of reducing the noise and reverberation. Also beamforming microphone arrays for speech application which can be classified in three categories: conventional beamforming, adaptive beamforming, microphone arrays with adaptive postfiltering are promising. For now, the combining time-varying beamformers with acoustic echo cancellers for full-duplex hands-free communication remains an open area for research. However, it seems logical that future hands-free communications system will combine the beamforming an echo cancellation function in one integrated full duplex system.

Also psychoacoustic employed in hands-free device can result in efficient and reasonable complexity systems as it has been presented in [6]. Another good starting point for further research is combining the AEC with the codec in cellular telephone or filtering the microphone signal with switching matrix of band for better SNR in noisy environment.

7. REALISATION OF AN ACOUSTIC ECHO AND NOISE CONTROL SYSTEMS ON DSP PROCESSOR

This paragraph will consider examples of two systems intended for echo and noise cancellation realised on single DSP processor. The performance judgement of the presented system is based on output signals $e(t)$ and echo and noise corrupted input signal $s(t)$ (see Fig. 4), which are normalised to obtain the Echo Return Loss Enhancement (ERLE) [8, 9, 46] defined by:

$$ERLE = 10 \cdot \log_{10} (E\{s^2(t)\} / E\{e^2(t)\}),$$

where $E\{\}$ — designates mathematical expectation.

System based on Echo Shaping idea with AEC + ENR structure can be realised as a on-chip program using low cost floating point Texas Instruments TMS320C3x DSP processor [41] with additionally, four channel A/D converter and at least one D/A converter to perform signal conversion. Single microphone and two microphone version of system are based on AEC + ENR concept using the transversal NLMS algorithm for echo compensation and Echo Shaping algorithm [42, 43] in ENR part based also on transversal NLMS algorithm. The full realisation and implementation description can be found in [44, 45]. Figure 10 shows the systems work and achieved echo and noise attenuation measured in ERLE for the same testing signal. Concerning single microphone system the echo shaping introduces the additional attenuation at 5 dB during speech activities which finally gives echo attenuation at 30 dB. But, if there is a silence (see Fig. 9 (left) about 3 sec.), that means that microphone picks up only noise, the shaping unit does not introduce any

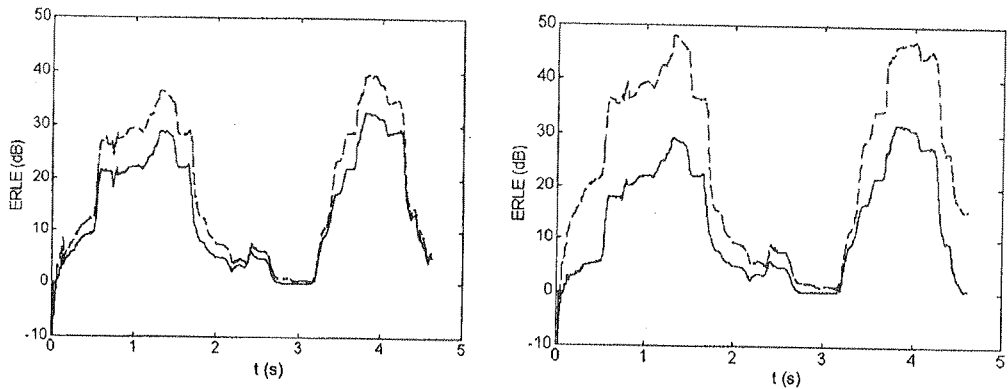


Fig. 9. Time flow of ERLE of the single microphone (left) and two microphone combined system for a five seconds testing signal in the cabin of the mid size car with engine (on [44])

improvement to compensator performance. A two microphone system gives better results than previous one. It has echo attenuation at about 40 dB and noise attenuation at 5 dB (see Fig. 9 (right) about 3 sec.). Echo shaping unit almost doubles the compensator echo attenuation, which is clearly visible.

Table 3
Parameters of the single and two microphones AEC + ENR system implemented on DSP processor

	single microphone system			two-microphones system	
AEC filter length	256	256	512	256	512
ENR filter length	32	64	64	64	64
ERLE	15 dB	20 dB	30 dB	24 dB	38 dB
Efficient echo path length cancellation	20 ms	20 ms	100 ms	20 ms	100 ms
One sample processing time	29,99 μ s	33,159 μ s	50,055 μ s	60,041 μ s	93,833 μ s

Another example of Echo Cancellation system but this time intended for cancellation of a long acoustic echo was offered in work [47, 48]. An acoustic echo canceller is realized also on a single DSP, the Texas Instruments TMS320C30 and based on the Decoupled Partitioned Block Frequency Domain Adaptive Filter. The system is mounted on a Loughborough Sound Images prototyping board, inserted in PC-solt. The hardware configuration is depicted in Fig. 10 [48].

The program of the Acoustic Echo Canceller consists of several modules [47] example is presented in Tab. 4 [47, 48].

Table 4

Maximum value of the system parameter and examples for two different echo length

Variable			Maximum	Example 1	Example 2
Filter length		N	8192	2016	2560
Filter part	Block length	B	256	8	64
	Partition length	Q	256	56	192
	FFT size	M	256	64	256
Update part	Block length	A	512	504	512
	Partition length	P	1024	504	512
	FFT size	L	1024	1024	1024
Sample frequency		f_s		7 kHz	13 kHz
Echo path length				288 ms	192 ms
Processing delay				1.6 ms	6.5 ms

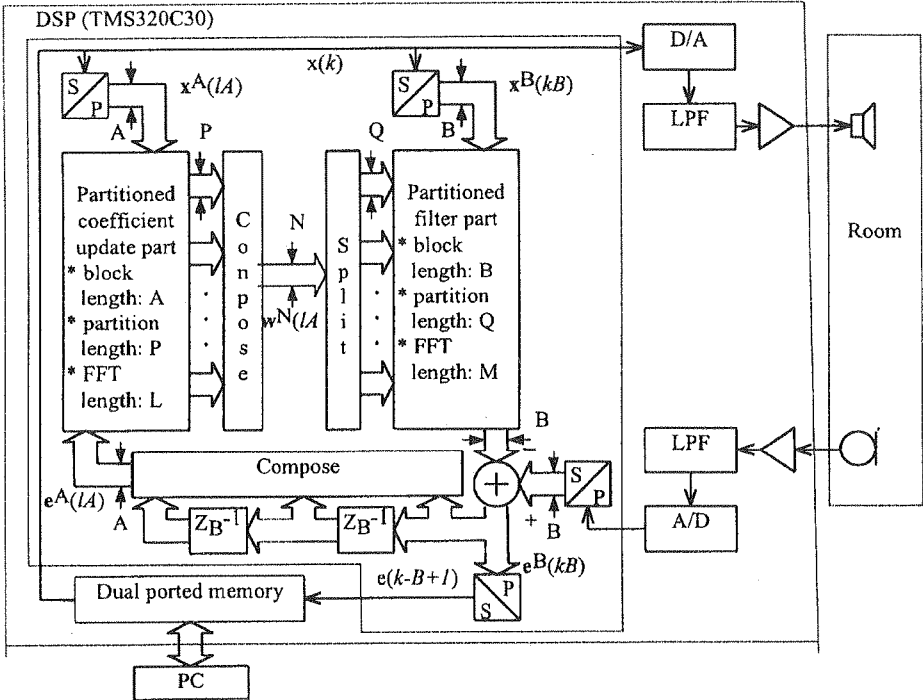


Fig. 10. Implementation scheme of Acoustic Echo Canceller

8. RECENT STUDIES

Recent studies made on combined acoustic echo and noise reduction systems use the psychoacoustic to improve the noise reduction algorithms [50–52]. Also residual echo attenuation are now investigated in order to employ psychoacoustically motivated algorithms [51]. There are two approaches investigated: adaptation signal analysis to the human ear, and to use the phenomenon of auditory masking. First idea is base on assumption that analysis filter (filters bank) should process the signals as the human ear do. In this solution the Cochlea model is used and the sub-bands are distributed non-linearly in frequency domain, with an equalised power distribution, according to the human cochlea function for spacing of the filters. This approach seems to be similar to the ordinary sub-band adaptive filters but initial experiment [53] shows that even when number of filters is limited by the size of the modifying FFT (or DCT) the results have been very encouraging and will be investigated further. Second approach made with use of the psychoacoustic is build upon the fact that the human ear will not perceive any additive signal components as long as their power spectral density lies completely below the so called masking threshold [51]. Because echo compensator do not necessary imposes a complete cancellation of the echo and also additive background noise is present the masking effect is used to reduce the noise and mask the residual echo. This idea is successfully realised based on structure AEC + NR, where the AEC is a conventional echo compensator but NR unit realises the psychoacoustical algorithm, as it has been presented in [51].

CONCLUSION

The acoustic echo in enclosed environment degrades the communication quality in cellular telephone at hands-free operation. This paper presents the authors approach to the nowadays acoustic echo cancellation solutions on which the practical systems are build. Authors realise that the discussion on the front-end processing in hands-free devices is not over yet and it is most advanced in single full-duplex audio channel voice communication. However, spatial realism desired in tele- and videoconferencing, as well as high quality mobile communication requirements forced researches for effective audio/video multichannel communication making great progress at multimicrophone base solution and it's practical realisation which is mentioned in this work. We hope that presented systematisation will put some light on the present stage of research in field of front-end processing in hands-free devices.

REFERENCE

1. R. Martin, P. Vary: *Combined Acoustic Echo Cancellation, Two Microphone Approach*. Annales des Telecommunications, 1994, Vol. 49, No. 7—9, pp. 429—438
2. R. Martin, J. Altenhoner: *Coupled Adaptive Filters for Acoustic Echo Control and Noise Reduction*. Proc. Int. Conf. Acoustic, Speech and Signal Processing Vol. 5, Detroit, May 1995, pp. 3042—3046
3. B. Ayad, G.F. & R. Le Bouquin-Jeannès: *Optimization of a Noise Reduction Preprocessing in an Acoustic Echo and Noise Controller*. Proceedings of ICASSP, Atlanta, May 1996, pp. 953—956
4. P. Scalart, A. Benamar: *A system for Speech enhancement in the context of hands-free radiotelephony with combined noise reduction and acoustic echo cancellation*. Speech Communication, 1996, Vol. 20, No. 3—4, pp. 203—214
5. R. Le Bouquin-Jeannès, G. Faucon, B. Ayad: *How to improve acoustic echo and noise canceling using a single talk detector*. Speech Communication, 1996, Vol. 20, No. 3—4, pp. 191—202
6. H. Matt, M. Walker: *Handsfree speaking for communication terminals*. VIII European Signal Processing Conference (EUSIPCO-96), Trieste, Italy, Sept. 1996, pp. 1139—1142
7. G.W. Elko: *Microphone array systems for hands-free telecommunication*. Speech Communication, 1996, Vol. 20, No. 3—4, pp. 229—240
8. E. Heansler: *The hands-free telephone problem — An annotated bibliography*. Signal Processing, 1992, Vol. 27, No. 3, pp. 259—271
9. E. Heansler: *The hands-free telephone problem — An annotated bibliography update*, Annales des Telecommunication, 1994, Vol. 49, No. 7—8, pp. 360—367
10. M. Bellanger: *Signal processing in mobile communication systems*. Proc. Of the 13th Int. Conf. On Digital Signal Processing, Grece, 1997, pp. 7—12
11. B. Widrow, S. Stearus: *Adaptive signal processing*, Prentice-Hall, 1996
12. E. Bataillou, H. Rix: *A new method for noise cancellation*. Signal Processing 1994, Vol. 7, pp. 1046—1049
13. V. Solos, X. Kong: *Adaptive signal processing algorithm. Stability and performance*. Prentice Hall, 1995
14. S. Haykin: *Adaptive filter theory*, Third edition, Prentice-Hall, 1996
15. A.M. Kuzminsky: *A robust step size adaptation scheme for LMS adaptive Filters*. IEEE-97
16. E.R. Ferrara: *Implementation of LMS Adaptive Filters*. IEEE trans., Vol. ASS-28, No. 4, pp. 474—475, August 1980
17. G.A. Clark, S.K. Mitra, S.R. Parker: *Block Implementation of Adaptive Digital Filters*. IEEE Trans., Vol. CAS-28, No. 6, pp. 584—592, June 1981
18. G.P.M. Egelmerts: *Adaptive Filtersystems: Complexity-performance Optimisation*. NERG 1993, part 58, nr 2, pp. 59—64
19. S.T. Alexander: *Fast adaptive filters: a geometrical approach*. IEEE ASSP Magazine, 1986, pp. 18—28
20. F. Capman, J. Boudy, P. Lockwood: *Acoustic Echo Cancellation and Noise Reduction in the Frequency-Domain: A Global Optimisation*. VIII European Signal Processing Conference (EUSIPCO-96), Trieste, Italy, Sept. 1996, pp. 27—32
21. G.A. Clark, S.R. Parker, S.K. Mitra: *Unifilied Approach to Time- and Frequency-Domain Realization of FIR Adaptive Digital Filter*. IEEE trans., Vol. ASSP-31, No. 5, pp. 1073—1083, October 1983
22. M. Dentino, J. Mc Cool, B. Widrow: *Adaptive Filtering in the Frequency Domain*. Proceedings IEEE. Vol. 66. No. 12, pp. 1658—1659, December 1978
23. N.J. Bershad, P.L. Feintuch: *Analysis of the Frequency Domain Adaptive Filter*. Proceedings IEEE. Vol. 67. No. 12, pp. 1658—1659, December 1979

24. G.P.M. Egelmeers: *Decoupling of partition factors in Partitioned Block FDAF*. ECCTD, Davos, Switzerland, August 1993, pp. 323–329
25. P.C.W. Sommen: *Partitioned Frequency Domain Adaptive Filter*. Proc. Asilomar Conf. On Signals and Systems, Pacific Grove, CA, USA, 1989, pp. 679–681
26. P.C.W. Sommen: *On the Convergence Properties of Partitioned Block Frequency Domain Adaptive Filter (PBFDAF)*. Proc. V European Signal Processing Conference (EUSIPCO-90), Barcelona, Spain, 1990, pp. 201–204
27. W.L. Yeung, C.P. Kwon: *Subband domain Wiener filtering*. Proc. Of the 13th Int. Conf. On Digital Signal Processing, Grece, 1997, pp. 483–486
28. K. Eneman, M. Moonen: *A Relation Between Subband and Frequency-Domain Adaptive Filtering*. Proc. Proc. Of the 13th Int. Conf. On Digital Signal Processing, Grece, 1997, pp. 25–28
29. L.R. Rabiner, B. Gold: *Theory And Application of Digital Signal Processing*. Prentice-Hall, Inc. Englewood cliffs, New Jersey, 1975
30. P.L. De León: *Subband and Wavelet Transforms: Design and applications*. chap. 11, edited by A. Akansk, M.J.T.A. Smith, Kluwer Academic Publishers, 1995
31. S. Fischer, K.U. Simmer: *Beamforming microphone arrays for speech acquisition in noisy environments*. Speech Communication, 1996, Vol. 20, No. 3–4, pp. 115–227
32. F. Amano, H.P. Meana, A de Luca, G. Duchen: *A multirate Acoustic Echo Canceller structure*. IEEE Transaction on Communication, 1995, Vol. 43, No. 7, pp. 2172–2176
33. N. Jayant, P. Noll: *Digital coding of waveforms*. Prentice-Hall, 1984
34. K. Bielawski, A.A. Petrovsky: *Speaker's acoustic environment in hands-free communication device*. Proceedings of the 7th Structural Acoustics Conference, Zakopane, Poland, April 1998, pp. 21–27
35. C. Carlemalm: *A novel low complexity algorithm for fast detection of echo path changes an double talk*. The First European Conf. On Signal Analysis And Prediction (EURASIP), ICT Press, Prague, June 1997, pp. 311–314
36. K. El-Malah, P. Kabal: *Comparision of voice activity algorithms for wireless personal communication systems*. Proc. IEEE Conf. on Electrical and Computer Engineering, May 1997, pp. 470–473
37. R. Zelinski: *Noise Reduction Based on Microphone Array with LMS Adaptive Post-Filtering*. Elect. Lett., Vol. 26, No. 24, pp. 2036–2037, November 1990
38. R. Martin: *Spectral Substration Based on Minimum Statistics*. VII European Signal Processing Conference (EUSIPCO-94), Edinburg, 1994, pp. 1182–1185
39. J. Allen, D. Berkley, J. Blauert: *Multimicrophone signals-processing technique to remove room reverberation from speech signals*. J. Acoust. Soc. Am., Oct. 1977, Vol. 62, No. 4
40. Y. Kaneda, M. Tohyama: *Noise suppresionsignal processing using 2-point received signals*. Electronics and Commun. In Japan, 1984, 67-A, (12), pp. 19–28
41. Texas Instruments: *TMS320C3x User's Guide*. Texas Instruments Incorporated, 1996
42. R. Martin, S. Gustafson: *The Echo Shaping approach to acoustic echo control*. Speech Communication, Vol. 20, No. 3–4, pp. 181–190, 1996
43. R. Martin, S. Gustafsson: *An Improved Echo Shaping Algorithm for Acoustic Echo Control*. VIII European Signal Processing Conference (EUSIPCO-96), Trieste, Italy, Spet. 1996, pp. 25–28
44. K. Bielawski, A.A. Petrovsky: *Adaptacyjny system eliminacji echa i szumów w samochodowych instalacjach głośnomówiących*. Elektronizacja nr 5'98, pp. 13–16
45. A.A. Petrovsky, K. Bielawski: *A DSP-based hands-free radiotelephony system with combined noise reduction and acoustic echo cancellation for mobile (car) environment*. Int. Conf. Modern Telecommunication systems Naroch, Belarus, September 1997, pp. 12–17
46. H.W. Gierlich: *The auditory perceived quality of hands-free telephones: Auditory judgments, instrumental measurements and their relationship*. Speech Communication, 1996, Vol. 20, No. 3–4, pp. 215–227

47. G.P.M. Egelmeers, P.C.W. Sommen, G.M. Alexiu, C.P.A. Vervoort: *Real-time implementation of low delay acoustic echo-canceller on one TMS320C30 DSP*. Proc. ICSPAT, Dallas USA, October 1994
48. G.P.M. Egelmeers, P.C.W. Sommen, J. de Boer: *Realization of an acoustic echo canceller on a single DSP*. VIII European Signal Processing Conference (EUSIPCO-96), Trieste, Italy, Sept. 1996, pp. 33–36
49. S.G. Sankaran: *Implementation and realisation of echo cancellation algorithms*, M.Sc. thesis, Electrical Eng. Virginia Polytechnic Institute, Blacksbur, Virginia, Dec, 1996
50. P. Dreiseitel, E. Hänsler, H. Puder: *Acoustic echo and noise control — a long lasting challenge*. IX European Signal Processing Conference (EUSIPCO'98), Rhodes-Greece, Sept. 1998, pp. 945–952
51. S. Gustafsson, P. Jax: *Combined residual echo and noise reduction: a novel psychoacoustically motivated algorithm*. IX European Signal Processing Conference (EUSIPCO'98), Rhodes-Greece, Sept. 1998, pp. 961–964
52. T. Haulick, K. Linhard, P. Schrögmeier: *Residual noise suppression using psychoacoustic criteria*. Proc. Of the 13th Int. Conf. On Digital Signal Processing, Greece, 1997
53. A. Hussain, R. Campbell: *Non-linear processing in Cochlear spaced sub-bands using artificial neural networks for multimicrophone adaptive speech enhancement*. IX European Signal Processing Conference (EUSIPCO'98), Rhodes-Greece, Spet. 1998, pp. 1209–1212

K. BIELAWSKI, A.A. PETROVSKY

POŁĄCZONE SYSTEMY ELIMINACJI ECHA I REDUKCJI HAŁASU W URZĄDZENIACH GŁOŚNOMÓWIĄCYCH

Streszczenie

Przedstawiono zagadnienia związane z użytkowaniem i konstruowaniem urządzeń głośnomówiących w systemach komunikacyjnych. Omówiono wpływ takich zjawisk akustycznych jak echo, pogłos, i hałas na jakość komunikacji podczas korzystania z telefonicznych instalacji hands-free. Zaprezentowano rodzinę algorytmów adaptacyjnych stosowanych w systemach kontroli zjawisk akustycznych, jak również rozwiązania wykorzystujące te algorytmy do kontroli echa i zakłóceń odwołując się do literatury. Dokonano zestawienia i porównania struktur i algorytmów stosowanych w połączonych systemach eliminacji echa i zakłóceń co stanowi istotę niniejszej pracy. Zestawienie to jest wynikiem badań autorów w dziedzinie systemów łączących w sobie rozwiązania kontroli echa i hałas, w celu poprawy jakości komunikacji z wykorzystaniem telefonów komórkowych z instalacją głośnomówiącą. Przedstawiono także złożoność i możliwości implementacyjne wybranych systemów na procesorze sygnałowym.

Słowa kluczowe: eliminacja echa, algorytmy adaptacyjne, redukcja zakłóceń, urządzenia głośnomówiące.

Praca finansowana z grant'u W/II/6/96.

SPIS TREŚCI

J. Stanclik: Makromodele scalonych wzmacniaczy mocy	445
R. Horbowski, M. Białko: Generowanie punktów uczących dla modeli funkcji układowych układów elektronicznych za pomocą systemu rozmytego	471
A. Dobrogowski, M. Kasznia: Wyznaczanie maksymalnego błędu przedziału czasu sygnałów taktowania	489
T. Tarecki: Koncepcja rozproszonego systemu pomiarowego sterowanego przez WWW	509
J. Zarębski, K. Górecki: Modelowanie nieizotermicznych stałoprądowych charakterystyk tranzystora Darlingtona mocy za pomocą programu APLAC	519
J. Siuzdak, T. Czarnecki: Metoda specyfikacji histogramu wykorzystująca macierz współwystępowania i jej zastosowanie dla transformacji obrazu	535
K. Bielawski, A.A. Petrovsky: Połączone systemy eliminacji echa i redukcji hałasu w urządzeniach głośnomówiących	555

CONTENTS

J. Stanclik: Macromodels of integrated circuits audio power amplifiers	445
R. Horbowski, M. Białko: The generation of learning data points for circuit performance modelling by means of fuzzy system	471
A. Dobrogowski, M. Kasznia: Calculation of Maximum Time Interval Error of timing signals ..	487
T. Starecki: Concept of WWW-based distributed measurement system	509
J. Zarębski, K. Górecki: Modelling of the nonisothermal DC characteristics of the power Darlington transistor using APLAC	519
J. Siuzdak, T. Czarnecki: A method of histogram specification employing the cocurrence matrix and its application to image transformation	535
K. Bielawski, A.A. Petrovsky: Joint system for acoustic echo and noise control in hands-free communication devices	555

„Zapowiedzi wydawnicze”

Kwartalnik Elektroniki i Telekomunikacji w 1999 roku rozpoczyna 45 tom swoich publikacji.

Kwartalnik jest czasopismem Komitetu Elektroniki i Telekomunikacji Polskiej Akademii Nauk, dociera do wszystkich krajowych ośrodków naukowych oraz technicznych, a także szeregu instytucji zagranicznych zajmujących się jego podstawową tematyką.

W zeszycie 1/1999 r. na uwagę Czytelników zasługuje kilka artykułów dotyczących wyników prac doświadczalnych autorów z krajowych Politechnik. Są to tytuły: „Wpływ pola magnetycznego pierścieniowych magnesów trwałych na wzrost mocy jonowego lasera argonowego”, „Pojemność kodowa programowalnych transkoderów opartych na strukturze typu PAL z różną liczbą termów dołączonych do poszczególnych wyjść”.

Zamieszczono również dwie prace z Politechniki Wrocławskiej dotyczące transmitancji opisujących obiekty przemysłowe. W wypadku występowania zmian ich parametrów w toku eksploatacji można stosować do sterowania układy adaptacyjne.

Treść zeszytu uzupełniają prace dotyczące projektowania struktur sieci magistral miejscowych dla rozproszonych systemów sterowania.

LISTA RECENZENTÓW KWARTALNIKA ELEKTRONIKI I TELEKOMUNIKACJI

Dr inż. *Jrōslaw Arabas*
Politechnika Warszawska,
Instytut Telekomunikacji

Doc. dr inż. *Zenon Baran*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Jerzy Baranowski*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Prof. dr hab. inż. *Daniel Jōzef Bem*
Politechnika Wrocławska,
Instytut Telekomunikacji i Akustyki

Dr hab. inż. *Wojciech Burakowski*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Andrzej Czajkowski*
Politechnika Łōdzka,
Instytut Automatyki

Prof. dr hab. inż. *Zbigniew Czyż*
Przemysłowy Instytut
Telekomunikacji, Warszawa

Dr hab. inż. *Andrzej M. Dąbrowski*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Janusz Dobrogowski*
Politechnika Poznańska,
Instytut Elektroniki i Telekomunikacji

Prof. dr hab. inż. *Janusz Dobrowolski*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Dr inż. *Jarōslaw Domaszewicz*
Politechnika Warszawska,
Instytut Telekomunikacji

Dr inż. *Przemysław Dymarski*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Adam Fioł*
Politechnika Warszawska,
Instytut Radioelektroniki

Dr inż. *Zbigniew Gajo*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Prof. dr hab. inż. *Anatōł Gosiewski*
Politechnika Warszawska,
Instytut Automatyki i Informatyki Stosowanej

Prof. dr hab. inż. *Krzysztof Grabowski*
Politechnika Gdańska,
Katedra Techniki Mikrofalowej
i Telekomunikacji Optycznej

Dr inż. *Antoni Grzanka*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Prof. dr hab. inż. *Wojciech Gwarek*
Politechnika Warszawska,
Instytut Radioelektroniki

emer. Prof. dr hab. inż. *Jan Hennel*
Politechnika Warszawska,
Instytut Mikroelektroniki i Optoelektroniki

Prof. dr hab. inż. *Krzysztof Holejko*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Ryszard Jachowicz*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Dr inż. *Andrzej Jakubiak*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Włodzimierz Janke*
Politechnika Gdańska,
Katedra Elektroniki Ciała Stałego

Prof. dr hab. inż. *Janusz M. Jaworski*
Politechnika Warszawska,
Instytut Elektroniki Teoretycznej
i Miernictwa Elektrycznego

Prof. dr hab. inż. *Romuald Józwicki*
Instytut Konstrukcji Przyrządów
Precyzyjnych i Optycznych,
Politechnika Warszawska

Prof. dr hab. inż. *Zbigniew Kaczkowski*
Polska Akademia Nauk,
Instytut Fizyki, Warszawa

Doc. dr hab. inż. *Andrzej Kasprzak*
Politechnika Wrocławska,
Instytut Sterowania i Techniki Systemów

Dr inż. *Jerzy Kęsik*
Politechnika Warszawska,
Instytut Mikroelektroniki i Optoelektroniki

Prof. dr hab. inż. *Jerzy Klamka*
Instytut Technologii
Elektronowej, Warszawa

Prof. dr hab. inż. *Andrzej Kobus*
Instytut Technologii
Elektronowej, Warszawa

Doc. dr hab. inż. *Adam Kowalewski*
Akademia Górniczo-Hutnicza,
Katedra Automatyki, Kraków

Prof. dr hab. inż. *Czesław Królikowski*
Politechnika Poznańska,
Instytut Elektroenergetyki

Dr inż. *Zdzisław Kozłowski*
Politechnika Warszawska,
Instytut Radioelektroniki

Prof. dr hab. inż. *Andrzej Kraśniewski*
Politechnika Warszawska,
Instytut Telekomunikacji

Doc. dr hab. inż. *Jerzy Kruszewski*
Politechnika Warszawska,
Instytut Mikroelektroniki i Optoelektroniki

Prof. dr hab. inż. *Adam Kujawski*
Politechnika Warszawska,
Instytut Fizyki

Prof. dr hab. inż. *Małgorzata Kujawińska*
Politechnika Warszawska,
Instytut Mikromechaniki i Fotoniki

Prof. dr hab. inż. *Juliusz Lech Kulikowski*
Polska Akademia Nauk, Warszawa,
Instytut Biocybernetyki
i Inżynierii Biomedycznej

Dr inż. *Zbigniew Lisik*
Politechnika Łódzka,
Instytut Elektroniki

Prof. dr hab. inż. *Józef Lubacz*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Jerzy Luciński*
Politechnika Łódzka,
Instytut Elektroniki

Prof. dr hab. inż. *Tadeusz Łuba*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Jerzy Majewski*
Wyższa Szkoła Morska, Gdynia,
Wydział Elektryczny

Prof. dr inż. *Władysław Majewski*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr inż. *Jan Malachowski*
Wojskowa Akademia Techniczna,
Instytut Fizyki Technicznej, Warszawa

Prof. dr hab. inż. *Bogdan Majkusiak*
Politechnika Warszawska,
Instytut Mikroelektroniki i Optoelektroniki

Prof. dr inż. *Marek Matczak*
Uniwersytet Szczeciński,
Wydział Mat.-Fiz., Katedra Fizyki

Prof. dr hab. inż. *Marian Milek*
Politechnika Zielonogórska,
Instytut Metrologii Elektrycznej

Prof. dr hab. inż. *Tadeusz Morawski*
Politechnika Warszawska,
Instytut Radioelektroniki

Prof. dr hab. inż. *Bogdan Mroziewicz*
Instytut Technologii Elektronowej,
Warszawa

Prof. dr hab. inż. *Jan Mulawka*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Prof. dr hab. inż. *Andrzej Napieralski*
Politechnika Łódzka,
Instytut Elektroniki

Prof. dr hab. inż. *Stanisław Nowak*
Akademia Górniczo-Hutnicza, Kraków

Dr hab. inż. *Jan Ogrodzki*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Prof. dr inż. *Jerzy Osiowski*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Dr inż. *Andrzej Pacut*
Politechnika Warszawska,
Instytut Automatyki
i Informatyki Stosowanej

Prof. nzw., dr hab. inż.
Zdzisław Papir
Akademia Górniczo-Hutnicza,
Kraków

Prof. dr hab. inż. *Władysław Paszek*
Politechnika Warszawska,
Instytut Maszyn i Urządzeń Elektrycznych

Dr hab. inż. *Jan Petykiewicz*
Politechnika Warszawska,
Instytut Fizyki

Dr inż. *Andrzej Pfitzner*
Politechnika Warszawska,
Instytut Mikroelektroniki i Optoelektroniki

Prof. dr hab. inż. *Marian S. Piekarski*
Politechnika Wrocławska,
Instytut Telekomunikacji i Akustyki

Prof. dr hab. inż. *Jan Purczyński*
Politechnika Szczecińska,
Instytut Elektroniki i Informatyki

Doc. dr inż. *Olgierd Przesmycki*
Politechnika Warszawska,
Instytut Telekomunikacji

Dr hab. inż. *Roman Rykaczewski*
Politechnika Gdańska,
Katedra Systemów Informacyjnych

Dr inż. *Franciszek Seredyński*
Polska Akademia Nauk, Warszawa
Instytut Podstaw Informatyki

Prof. dr inż. *Stanisław Sławiński*
Politechnika Warszawska,
Instytut Telekomunikacji

Prof. dr hab. inż. *Stanisław Skoczowski*
Politechnika Szczecińska,
Wydział Elektryczny

Dr inż. *Jerzy Skwarczyński*
Akademia Górniczo-Hutnicza, Kraków
Instytut Maszyn i Sterowania Układów
Elektroenergetycznych

Prof. dr hab. inż. *Tadeusz Sobczyk*
Akademia Górniczo-Hutnicza, Kraków
Instytut Maszyn i Sterowania Układów
Elektroenergetycznych

Prof. dr hab. inż. *Marek Stabrowski*
Politechnika Warszawska,
Instytut Elektrotechniki Teoretycznej
i Miernictwa Elektrycznego

Dr inż. *Jerzy Szabatin*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Prof. dr hab. inż. *Mieczysław Szustakowski*
Wojskowa Akademia Techniczna,
Warszawa,
Instytut Fizyki Technicznej

Dr inż. *Jacek Szymanowski*
Politechnika Warszawska,
Instytut Automatyki
i Informatyki Stosowanej

Dr inż. *Lech Śliwa*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Dr inż. *Tadeusz Świdorski*
Wojskowy Instytut Łączności,
Zakład Radiokomunikacji, Zegrze

Prof. dr inż. *Michał Tadeusiewicz*
Politechnika Łódzka,
Instytut Podstaw Elektroniki

Prof. dr hab. inż. *Leszek Trybus*
Politechnika Rzeszowska,
Katedra Automatyki i Informatyki

Dr hab. inż. *Aleksander Urbaś*
Politechnika Warszawska,
Instytut Podstaw Elektrotechniki

Dr hab. inż. *Krzysztof Wesolowski*
Politechnika Poznańska,
Instytut Elektroniki i Telekomunikacji

Dr hab. inż. *Tadeusz Więckowski*
Politechnika Wrocławska,
Instytut Telekomunikacji i Akustyki

Prof. dr hab. inż. *Jacek Wojciechowski*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Prof. nzw.,
Dr hab. inż. *Andrzej Wojtkiewicz*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Dr hab. inż. *Tomasz Woliński*
Politechnika Warszawska,
Instytut Fizyki

Dr inż. *Włodzimierz Wolski*
Politechnika Wrocławska,
Instytut Telekomunikacji i Akustyki

Dr hab. inż. *Zygmunt Wróbel*
Uniwersytet Śląski, Sosnowiec
Instytut Problemów Techniki

Dr inż. *Kornel B. Wydro*
Politechnika Warszawska,
Instytut Automatyki
i Informatyki Stosowanej

Prof. dr hab. inż. *Jan Zabrodzki*
Politechnika Warszawska,
Instytut Informatyki

Dr inż. *Krzysztof Zamłyński*
Politechnika Warszawska,
Instytut Podstaw Elektroniki

Dr inż. *Piotr Zwierko*
Politechnika Warszawska,
Instytut Telekomunikacji

Subscription for external subscribers:

The promotional subscription price in 1997 is \$ 90 including postage for institutions. Subscriptions should be sent to the publisher, Polish Scientific Publishers PWN Ltd, Journal Division, Miodowa 10, 00-250 Warsaw, POLAND, fax (48) (22) 826 09 50, (48) (22) 826 71 63, with a copy by fast to Electronics and Telecommunications Quarterly 00-665 Warsaw, ul. Nowowiejska 15/19, p. 470. Subscription is occupied after showing a cheque or transfer documents. Our bank account is as follows: Bank account: PBK VIII O/Warszawa nr 11101037-1052-2700-1-43. At subscriber's request this journal will be air mailed at additional postage to 50% gross price to European countries and 65% overseas.

Subscription orders available through the local press distribution or through the Foreign Trade Enterprise ARS Polona, 00-068 Warszawa, Krakowskie Przedmieście 7, Poland.

Ruch S.A. fulfils foreign customers orders, starting from any issue in the calendar year: tel. (48) (22) 620 120 39; fax (48) (22) 620 17 62.