

Ulepszone rozwiązanie układowe konweytorów prądowych w technice CMOS

STANISŁAW KUTA, WITOLD MACHOWSKI, JACEK JASIELSKI

Akademia Górniczo-Hutnicza, Kraków

Otrzymano 1994.01.18

Autoryzowano do druku 1994.07.12

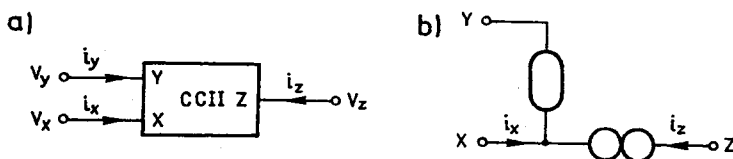
W pracy dokonano krótkiego przeglądu i oceny podstawowych rozwiązań układowych konweytorów prądowych CCII w technice CMOS. Porównania tych układów dokonano w oparciu o ich symulacje w programie HSPICE. Przedstawiono projekty ulepszonych rozwiązań układowych CCII z wewnętrzną pętlą sprzężenia zwrotnego i kaskadowymi lustrami prądowymi o zwiększonym zakresie dynamicznym napięć wejściowych. Proponowane układy mogą pracować przy niskich napięciach zasilających ± 3.3 V oraz charakteryzują się dużą dokładnością i liniowością przetwarzania sygnałów.

1. WPROWADZENIE

W 1968 roku Smith i Sedra [1] opracowali koncepcję trójwrotnika ze wspólnym punktem masy, który nazwano konweytor prądowym, a następnie ci sami autorzy w 1970 roku przedstawili ulepszoną wersję tego układu, nazwanego konweytor prądowym drugiej generacji CCII (Second-Generation Current Conveyor).

Układ CCII znalazł bardzo szerokie zastosowanie jako podstawowy blok w syntezie układów pracujących w modzie prądowym. Należy podkreślić, że układy w modzie prądowym wykazują wiele zalet w porównaniu do układów napięciowych. W szczególności układy te zapewniają większe pasmo częstotliwości, większą liniowość przetwarzania sygnałów, szerszy zakres dynamiczny oraz mogą pracować przy mniejszych napięciach zasilających, co ma szczególne znaczenie w układach VLSI i ULSI zawierających część analogową i cyfrową, bowiem pozwala to na ograniczenie mocy rozpraszanej w układzie scalonym.

Schemat blokowy konweytor prądowy przedstawiono na rys. 1.a.



Rys. 1. Konweyor prądowy CCH: a) schemat blokowy; b) schemat zastępczy w postaci nullora

Chwilowe wartości napięć i prądów konweyora prądowego, odpowiednio pierwszej i drugiej generacji, opisują następujące równania hybrydowe:

a) CCI

$$\begin{bmatrix} I_y \\ U_x \\ I_z \end{bmatrix} = \begin{bmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & \pm 1 & 0 \end{bmatrix} \begin{bmatrix} U_y \\ I_x \\ U_z \end{bmatrix} \quad (1)$$

b) CCII

$$\begin{bmatrix} I_y \\ U_x \\ I_z \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \\ 0 & \pm 1 & 0 \end{bmatrix} \begin{bmatrix} U_y \\ I_x \\ U_z \end{bmatrix} \quad (2)$$

W konweytorze CCI impedancja widziana z zacisku X przenosi się na zacisk Y tak, że przy $U_x = U_y$, $I_x = I_y$. W konweytorze CCII, który będzie przedmiotem dalszych rozważań, impedancja widziana z zacisku Y jest zawsze równa nieskończoności ($I_y = 0$).

Jeżeli współczynnik transmisji prądowej z wejścia X na wyjście Z wynosi $+1$ (zgodnie z oznaczeniami na rys. 1a, $I_x = I_z$), to taki konweytor oznaczamy CCII^+ . W przypadku, gdy współczynnik ten wynosi -1 , to konweytor oznaczamy CCII^- .

Dla ułatwienia syntezy schematów funkcjonalnych zawierających konweyory CCII mogą być one przedstawione przy pomocy elementów zdegenerowanych. Na rys. 1b przedstawiono schemat zastępczy konweyora CCII w postaci nullora, zawierającego nulator na wejściu i norator na wyjściu. Jedna para zacisków układu z rys. 1b działa jednocześnie jak zwarcie i rozwarcie, podczas gdy prąd i napięcie drugiej pary zacisków są całkowicie dowolne.

W ostatnich trzech latach opublikowano wiele prac [3 ÷ 9], w których opisano ulepszone wersje układowe konweytorów CCII w technice CMOS, jednakże rozwiązania te nie są jeszcze w pełni zadowalające. W niniejszej pracy przedstawiono projekt konweyora CCII w technice CMOS z wewnętrzną pętlą sprzężenia zwrotnego i kaskadowymi lustrami prądowymi o zwiększonym zakresie dynamicznym napięć wyjściowych [10], co pozwoliło na stosowanie małych napięć zasilających ± 3 V oraz uzyskano dużą dokładność przetwarzania sygnałów w układzie.

Porównania różnych rozwiązań układowych konweytorów prądowych dokonano w oparciu o ich symulacje w programie HSPICE, przy czym parametry modeli tranzystorów na poziomie LEVEL = 2 odpowiadają 2 μm komercyjnej technologii CMOS z AMS.

2. PRZEGLĄD I OCENA PODSTAWOWYCH ROZWIĄZAŃ UKŁADOWYCH KONWEYTORÓW CCII W TECHNICIE CMOS

Pierwsze układy konweytorów CCII⁺ i CCII⁻ w technice CMOS opisano w pracy [3]. Schemat ideowy konweytoru CCII⁺ opisanego w tej pracy, przedstawiono na rys. 2a. Para różnicowa na tranzystorach M_1 i M_2 , zakończona wtórnikiem źródłowym na tranzystorze M_5 i wewnętrzną pętlą ujemnego sprzężenia zwrotnego z wyjścia wzmacniacza X na wejście odwracające pary różnicowej tworzy wtórnik napięciowy zapewniający $U_x = U_y = U_{in}$ oraz małą rezystancję wyjściową z zacisku X.

Ponieważ wyjście wtórnika napięciowego, a więc i wejście X konweytoru jest punktem o małej impedancji, zatem łatwo może nastąpić przepływ prądu z lub do tego punktu. Jeżeli $U_x = U_y > 0$, to prąd $I_x = U_y/R_x$ wypływa z punktu X. Prąd drenu tranzystora M_5 równy $I_2 + I_x$ zostaje przeniesiony przy pomocy lustra prądowego na tranzystorach M_6 i M_7 do punktu Z tak, że przy $I_2 = I_3$, $I_z = I_1 + I_x - I_3 = I_x$.

Dla realizacji ujemnego konweytoru CCII⁻ należy zastosować dodatkowe lustro prądowe dostarczające prąd $I_2 + I_x$ wypływający z punktu Z oraz prąd $I_2 = I_3$ wpływający do punktu Z. Wtedy prąd $I_z = I_3 - (I_2 + I_x) = -I_x$.

Jak wykazano w pracy [3], w układzie na rys. 2a spełnione są następujące zależności:

$$V_x \simeq V_y \frac{g_{m2}}{g_{m2} + g_{d2} + g_{d4}} \Big|_{R_x = \infty} \quad (3)$$

$$r_x = \frac{V_x}{I_x} \simeq \frac{(g_{m1} + g_{m2})(g_{d2} + g_{d4})}{g_{m1}g_{m2}g_{m5}} \Big|_{V_x = 0} \quad (4)$$

$$r_z \simeq \frac{1}{g_{d7} + g_{d11}} \quad (5)$$

gdzie:

- r_x — rezystancja wejściowa widziana z punktu X,
- r_z — rezystancja wyjściowa widziana z punktu Z,
- g_{mi} — transkonduktancje tranzystorów,
- g_{di} — konduktancje drenowe tranzystorów.

Jak wynika z powyższych zależności napięcie U_x jest mniejsze od napięcia wejściowego U_y , a dokładność powtarzania tego napięcia w punkcie X szybko maleje dla małych wartości R_x z uwagi na stosunkowo dużą wartość rezystancji wyjściowej

r_x wtórnika napięciowego. Należy również podkreślić, że układ posiada niezbyt dużą rezystancję wyjściową r_z w punkcie Z.

Wady tych prostych rozwiązań mogą być zilustrowane przy pomocy charakterystyk statycznych układu z rys. 2a otrzymane w wyniku jego symulacji w programie HSPICE. Jak już wspomniano wcześniej do symulacji wykorzystano modele tranzystorów na poziomie LEVEL = 2 i parametrach odpowiadających $2\ \mu\text{m}$ w technologii CMOS z AMS. Wymiary W/L tranzystorów zestawiono w tabelce na rys. 2a.

Na rys. 2b przedstawiono charakterystyki $(U_y - U_x) = f(U_y)$ dla trzech wartości R_x . Jak widać, już dla rezystancji $R_x = 1\text{ k}\Omega$ występują duże błędy powtarzania napięcia wejściowego U_y w punkcie X, ponieważ rezystancja wyjściowa układu wynosi $r_x \simeq 75\ \Omega$. Na rys. 2c przedstawiono charakterystyki transmisji prądowej z wejścia X na wyjście Z: $I_x = f(U_y)$, $I_z = f(U_y)$ dla trzech różnych wartości rezystancji $R_x = R_y = R$. Tranzystory M_6, M_7 lustra prądowego posiadają jednakowe wymiary W/L , lecz tranzystor M_6 pracuje zawsze przy mniejszym napięciu drenowym niż M_7 . Na skutek zjawiska modulacji długości kanału prąd drenu I_{d7} jest większy niż prąd drenu I_{d6} , stąd w układzie $I_x > I_z$. Zmniejszając szerokość W kanału tranzystora M_7 można doprowadzić do zrównania się prądów $I_x = I_z$ w wybranym punkcie pracy, lecz zjawisko modulacji długości kanałów oraz efekty podłoża tranzystorów w lustrach prądowych są źródłem dużych błędów w transmisji prądu z wejścia X na wyjście Z. Błędy transmisji prądowej są jeszcze większe, gdy $R_x \neq R_z$, ponieważ jeszcze silniej uwidatnia się zjawisko modulacji długości kanałów.

Dużą dokładność transmisji prądowej z wejścia X na wyjście Z ($I_x = I_z$) oraz znaczne zwiększenie rezystancji wyjściowej r_z może zapewnić kaskadowe lustro prądowe. Schemat ideowy konweyora CCII⁺ z kaskadowymi lustrami prądowymi, opisanego w pracy [4], przedstawiono na rys. 3.

Rezystancja wyjściowa wtórnika napięciowego, widziana w punkcie X, wynosi

$$r_x = \frac{r_{out}}{1 + LG}, \quad (6)$$

gdzie:

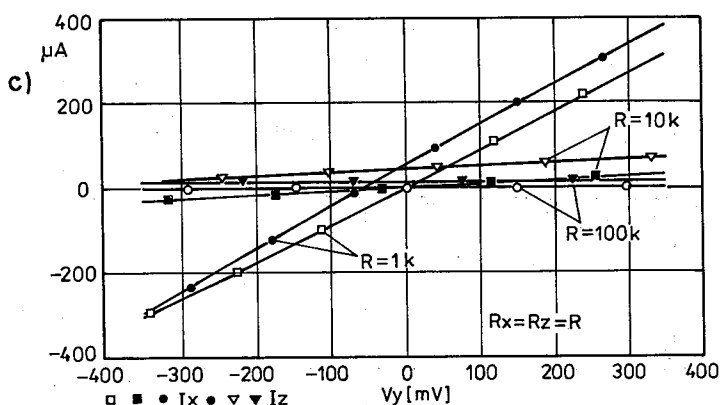
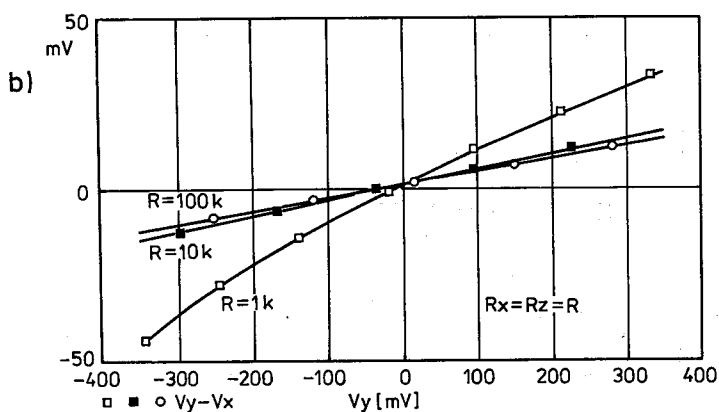
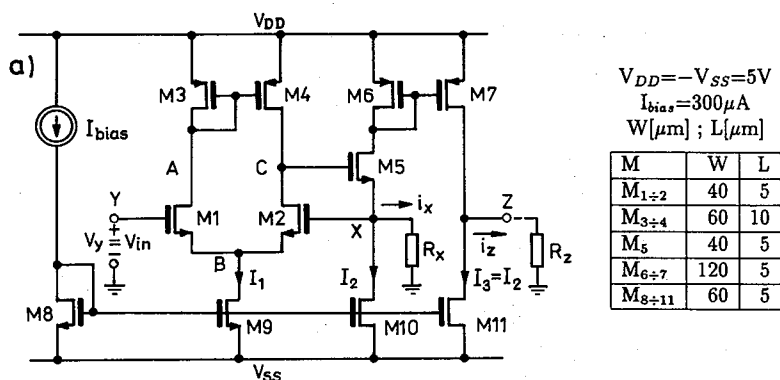
r_{out} — rezystancja wyjściowa wzmacniacza (w punkcie X) przy otwartej pętli sprzężenia zwrotnego,

LG — wzmocnienie otwartej pętli ujemnego sprzężenia zwrotnego.

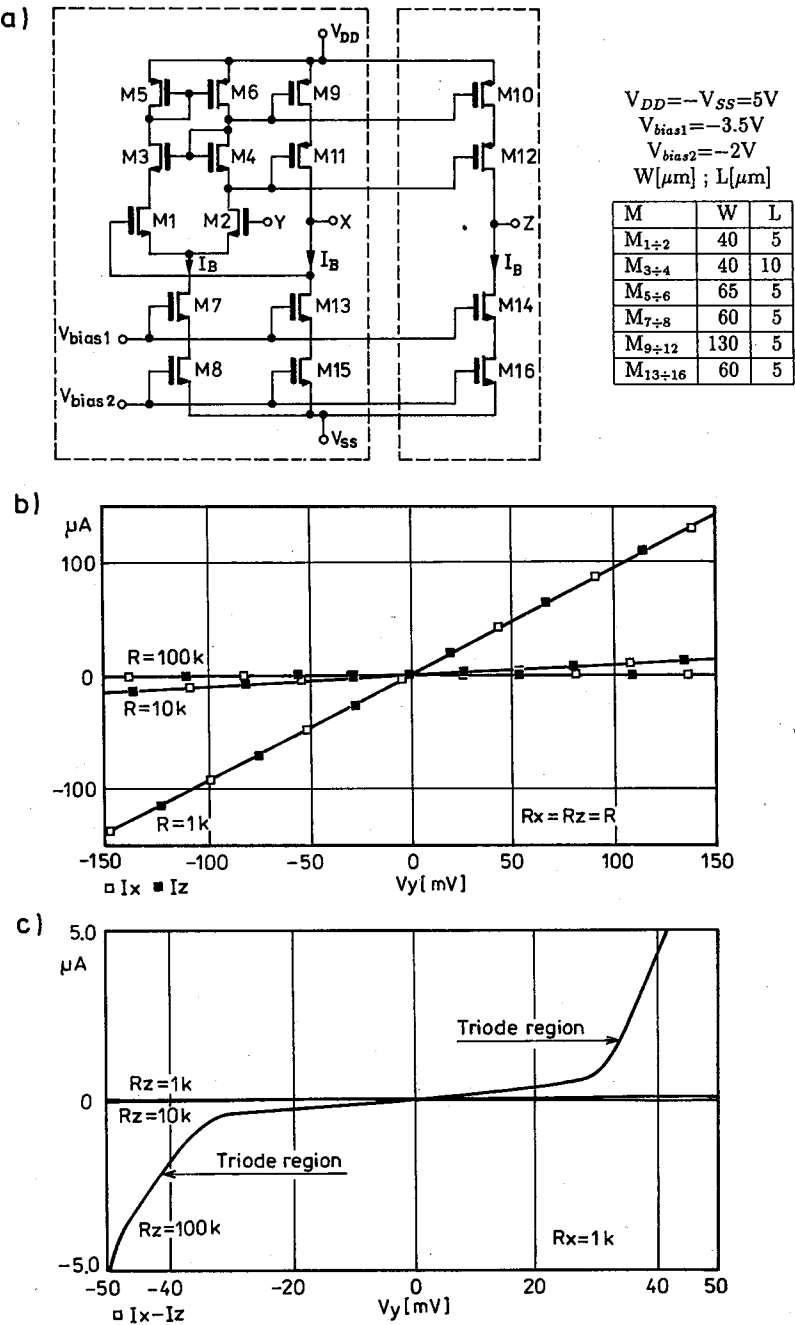
Układ ten nie pozwala na uzyskanie dostatecznie małej rezystancji wyjściowej r_x , ponieważ r_{out} jest bardzo duże, a wzmacniacz jest jednostopniowy.

Gdy rezystancja obciążenia R_x nie spełnia warunku $R_x \gg r_x$, to występuje błąd powtarzania napięcia U_y w punkcie X ($U_x < U_y$).

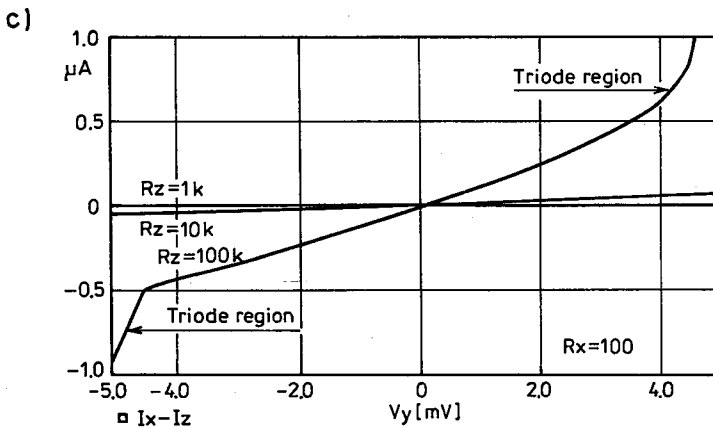
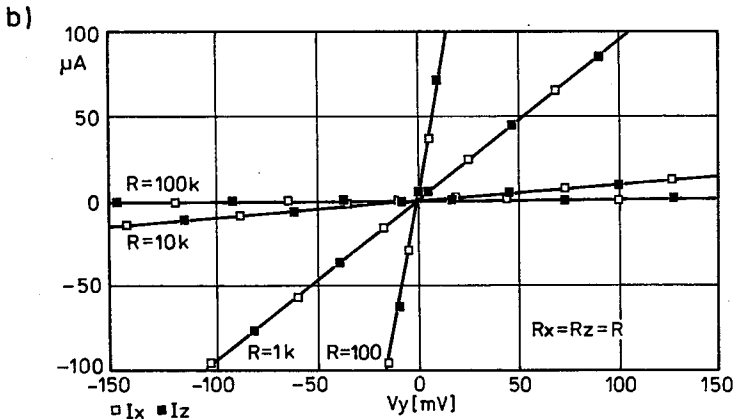
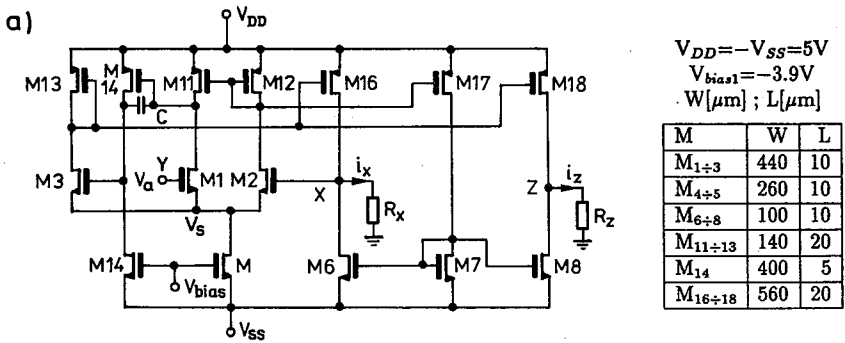
Jak wynika z przeprowadzonej analizy, rezystancja wyjściowa r_x rozważanego układu wynosi $56\ \Omega$, co już przy rezystancji $R_x = 1\text{ k}\Omega$ daje 5-cio procentowy błąd powtarzania napięcia wejściowego w punkcie X.



Rys. 2. Prosty układ konwejora CCII⁺ w technice CMOS: a) schemat ideowy; b) błąd bezwzględny powtarzania napięcia wejściowego U_y w punkcie X: $(U_y - U_x) = f(U_y)$ dla różnych wartości R_x ; c) charakterystyki transmisji prądowej z wejścia X na wyjście Z: $I_x = f(U_y)$, $I_z = f(U_y)$ dla różnych wartości $R_x = R_z = R$



Rys. 3. Konweyor prądowy CCII⁺ z kaskadowym lustrem prądowym: a) schemat ideowy; b) charakterystyki transmisji prądowej: $I_x = f(U_y)$, $I_z = f(U_y)$; c) błąd bezwzględny transmisji prądowej: $(I_x - I_z) = f(U_y)$ dla $R_x \neq R_z$



Rys. 4. Konwejer prądowy CCII⁺ z ulepszonym wtórnikiem napięciowym: a) schemat ideowy; b) charakterystyki transmisji prądowej: $I_x = f(U_y)$, $I_z = f(U_y)$; c) błąd bezwzględny transmisji prądowej: $(I_x - I_z) = f(U_y)$ dla $R_x \neq R_z$

Kaskadowe lustra prądowe zapewniają dużą dokładność powtarzania prądu I_x na wyjściu Z (tzn. $I_z \simeq I_x$), co ilustruje rys. 3b przedstawiający zależności $I_x = f(U_y)$, $I_x = f(U_y)$ dla trzech wartości $R_x = R_y = R$ oraz rys. 3c przedstawiający zależności $I_x - I_z = f(U_y)$ dla $R_x = 1 \text{ k}\Omega = \text{const}$ i trzech różnych wartości rezystancji R_z . Jak wynika z rys. 3c największa różnica prądów ($I_x - I_z$), spowodowana zjawiskiem modulacji długości kanałów tranzystorów, nie przekracza 2%, gdy rezystancja R_x i R_z różnią się o dwa rzędy wielkości (tzn. $R_x = 1 \text{ k}\Omega$, $R_z = 100 \text{ k}\Omega$).

W pracy [5] opisano konweyor prądowy CCII⁺ w którym wtórnik napięciowy został zrealizowany przy użyciu dwustopniowego wzmacniacza (rys. 4a).

Zgodnie z zależnością (4) większe wzmocnienie układu przy otwartej pętli ujemnego sprzężenia zwrotnego pozwala na zmniejszenie rezystancji wyjściowej r_x wtórnika napięciowego. Opisany układ posiada rezystancję wyjściową r_x około 34Ω , tak że dla rezystancji obciążenia $R_x > 1 \text{ k}\Omega$ błąd powtarzania napięcia U_y na wyjściu X jest mniejszy niż 4%, lecz dla $R_x = 100 \Omega$ błąd ten wynosi prawie 30%.

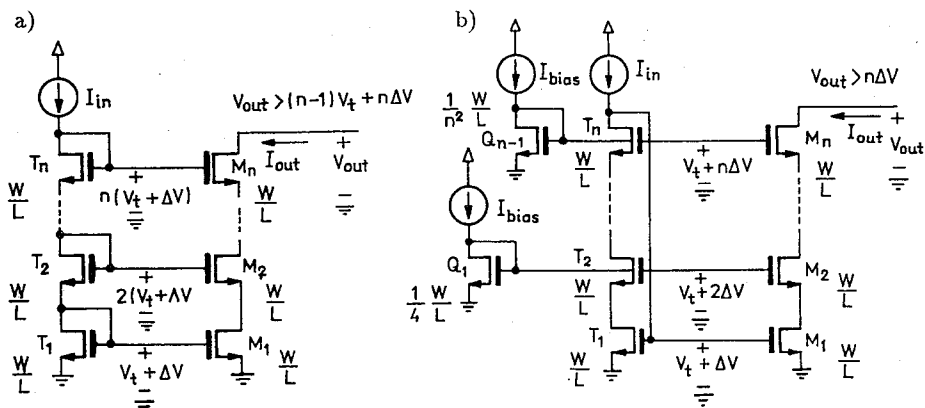
Przy jednakowych rezystancjach $R_x = R_z$, pary tranzystorów M_6, M_8 oraz M_{16}, M_{18} pracują przy zbliżonych napięciach drenowych dlatego w tym przypadku prąd I_z dokładnie nadąża za prądem I_x (rys. 4b). Gdy $R_z \gg R_x$, na skutek zjawiska modulacji długości kanałów tranzystorów, występuje niewielka różnica pomiędzy prądami I_x oraz I_z , co ilustruje rys. 4c.

3. ULEPSZONE UKŁADY KONWEJORÓW PRĄDOWYCH CCII

Żaden z opisanych układów w punkcie 2 nie zapewnia równocześnie małej rezystancji wyjściowej w punkcie X oraz dużej dokładności transmisji prądu z wejścia X na wyjście Z. Dla realizacji profesjonalnych układów konwejorów prądowych należy zastosować kaskadowe lustra prądowe, co pozwala na znaczne ograniczenie szkodliwego zjawiska modulacji długości kanałów tranzystorów, a tym samym zapewnia dużą dokładność transmisji prądowej. Wykorzystanie kaskadowych lusterek prądowych o zwiększonym zakresie dynamicznym napięć wyjściowych [10] pozwala na realizację układów pracujących przy małych napięciach zasilających $\pm 3,3 \text{ V}$.

Jak pokazano na rys. 5a, w klasycznym n -kaskadowym lustrze prądowym zbudowanym z identycznych tranzystorów, napięcie na n szeregowo połączonych tranzystorach $M_1 \div M_n$ nie może być mniejsze od wartości:

$$V_{out} > n\Delta V + (n-1)V_t \quad (7)$$



Rys. 5. n -kaskadowe lustro prądowe: a) układ klasyczny; b) układ o zwiększonym zakresie dynamicznym napięcia wyjściowego

przy czym V_t jest napięciem progowym tranzystorów MOSFET

$$\Delta V = \sqrt{\frac{I_{in}}{K \frac{W}{L}}} \quad (8)$$

$K = \frac{\mu C_{ox}}{2}$ jest parametrem konduktancji, zaś $I_{in} = I_{out}$.

Spełnienie warunku (5) zapewnia pracę wszystkich tranzystorów $M_1 \div M_n$ w obszarze nasycenia, czyli napięcie dren — źródło każdego z tranzystorów jest większe lub równe ΔV , tj. $V_{ds} \gg V_{gs} - V_t$.

W symetrycznym obwodzie lustra prądowego (jak np. na rys. 3a), złożonym z łańcucha n tranzystorów pMOS od strony szyny zasilającej V_{DD} oraz łańcucha n tranzystorów nMOS od strony szyny zasilającej V_{SS} , liniowy obszar napięcia wyjściowego ograniczony jest do wartości:

$$V_{linear} = V_{DD} + |V_{SS}| - (n-1)V_{Tn} - (n-1)|V_{Tp}| - n\Delta V_N - n|\Delta V_P|. \quad (9)$$

O ile napięcia ΔV mogą być zmniejszone przez zwiększenie szerokości W kanałów tranzystorów, to napięcia progowe V_t stanowią istotne ograniczenie liniowego obszaru napięcia wyjściowego. Obszar ten może być znacznie zwiększony w kaskadowym lustrze prądowym przedstawionym na rys. 5b [10].

W układzie tym bramki każdej pary tranzystorów $T_i \div M_i$ kaskady polaryzowane są napięciem o wartości $(V_t + i\Delta V)$, wytworzonym jako spadek napięcia na tranzystorze Q_i , przez który płynie prąd I_{bias} . Zakładając, że tranzystor $T_1 \div T_n$, $M_1 \div M_n$ mają jednakowe wymiary W/L oraz $I_{bias} = I_{in} = I_{out}$, dla uzyskania spadku napięcia o wartości $(V_t + i\Delta V)$, tranzystor Q_i powinien posiadać wymiary $(1/i^2)(W/L)$.

Jak sugerują autorzy w pracy [11], aby zapewnić pracę tranzystorów $M_1 \div M_n$ przy napięciach $V_{ds} = \Delta V$ nieco większych niż na granicy obszaru nasycenia,

każdy z tranzystorów Q_i powinien posiadać wymiary $[1/(i^2 + i - 1)](W/L)$. Napięcia polaryzujące tranzystory łańcucha $M_1 \div M_n$ wynoszą: $V_{gs} = V_t + \Delta V$, $V_{ds} = \Delta V$. W tym przypadku liniowy obszar napięcia wyjściowego w obwodzie zawierającym łańcuch n tranzystorów pMOS i n tranzystorów nMOS wynosi

$$V_{linear} = V_{DD} + |V_{SS}| - n\Delta V_N - n|\Delta V_P| \quad (10)$$

Schemat ideowy konweyora CCII⁺ z opisanym lustrem prądowym przedstawiono na rys. 6a. Zastosowanie takiego lustra pozwoliło na obniżenie napięć zasilających do wartości ± 3.3 V, przy zachowaniu szerokiego zakresu dynamicznego liniowej pracy układu konweyora. Dodatkowy stopień wzmacniający w układzie wtórnika napięciowego pozwolił, zgodnie z zależnością (4), zmniejszyć rezystancję r_x do wartości 9 Ω . Oznacza to, że przy rezystancji obciążenia $R_x = 100$ Ω napięcie U_x jest mniejsze o 8% od napięcia wejściowego U_y . Na rys. 6b przedstawiono charakterystyki transmisji prądowej z wejścia X na wyjście Z: $I_x = f(U_y)$, $I_z = f(U_y)$ dla trzech różnych wartości rezystancji $R_x = R_z = R$, zaś na rys. 6c przedstawiono zależności $(I_x - I_z) = f(U_y)$ dla $R_x = 1$ k Ω i trzech różnych wartości rezystancji R_z . Z przedstawionej analizy wynika, że prezentowany układ na rys. 6a posiada lepsze parametry, niż układy opisane w punkcie 2. Posiada on dużo mniejszą rezystancję wejściową r_x , małe napięcia zasilające ± 3.3 V oraz charakteryzuje się równie małymi błędami transmisji prądowej z wejścia X na wyjście Z jak układ z rys. 3a, zawierający klasyczne lustra prądowe.

Jeszcze lepsze parametry posiada konweyor CCII⁺, którego schemat ideowy przedstawiono na rys. 7a. Układ ten stanowi rozszerzoną wersję konweyora z rys. 4a o lustra prądowe ze zwiększonym zakresem dynamicznym napięcia wyjściowego.

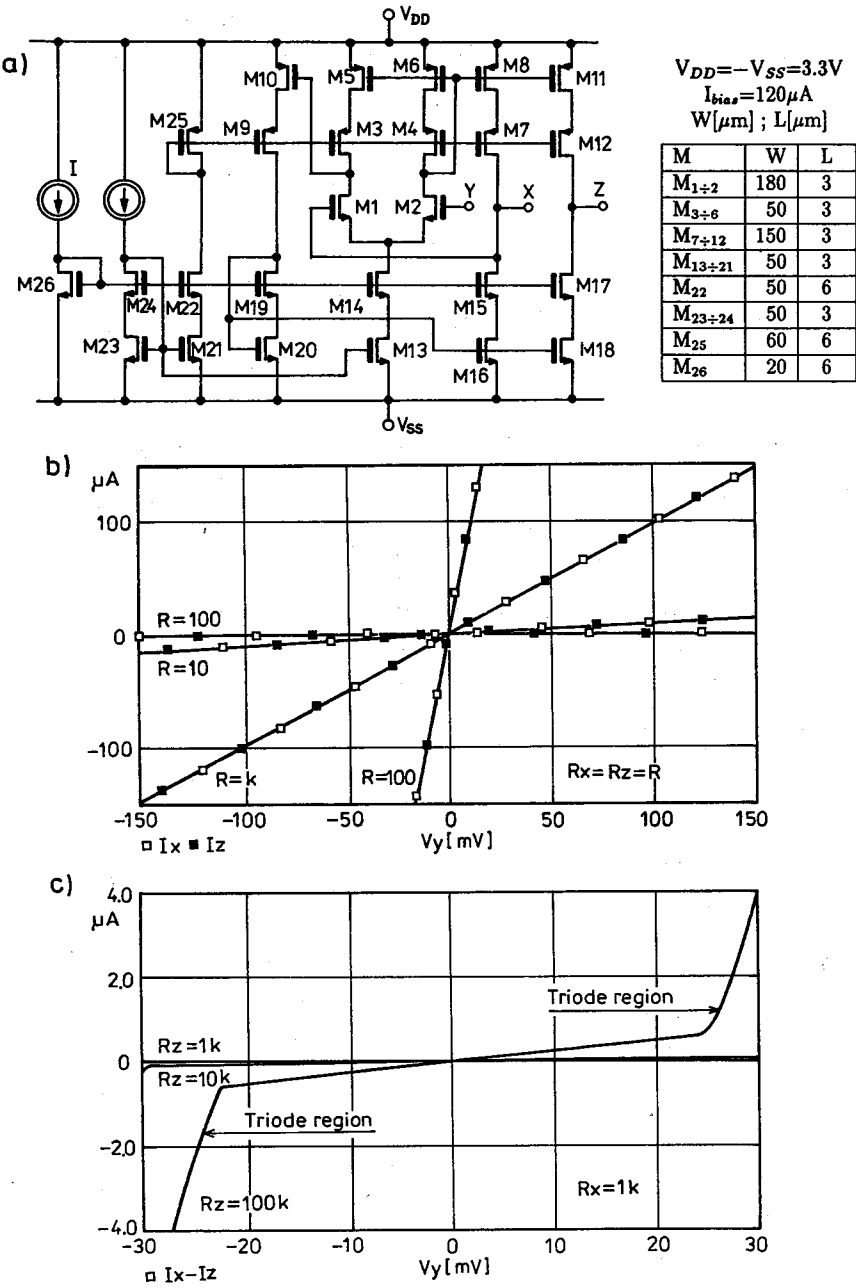
Wymiary tranzystorów i parametry układu są następujące:

$$V_{DD} = -V_{SS} = 3.3 \text{ V}; I_{bias} = 120 \text{ } \mu\text{A}$$

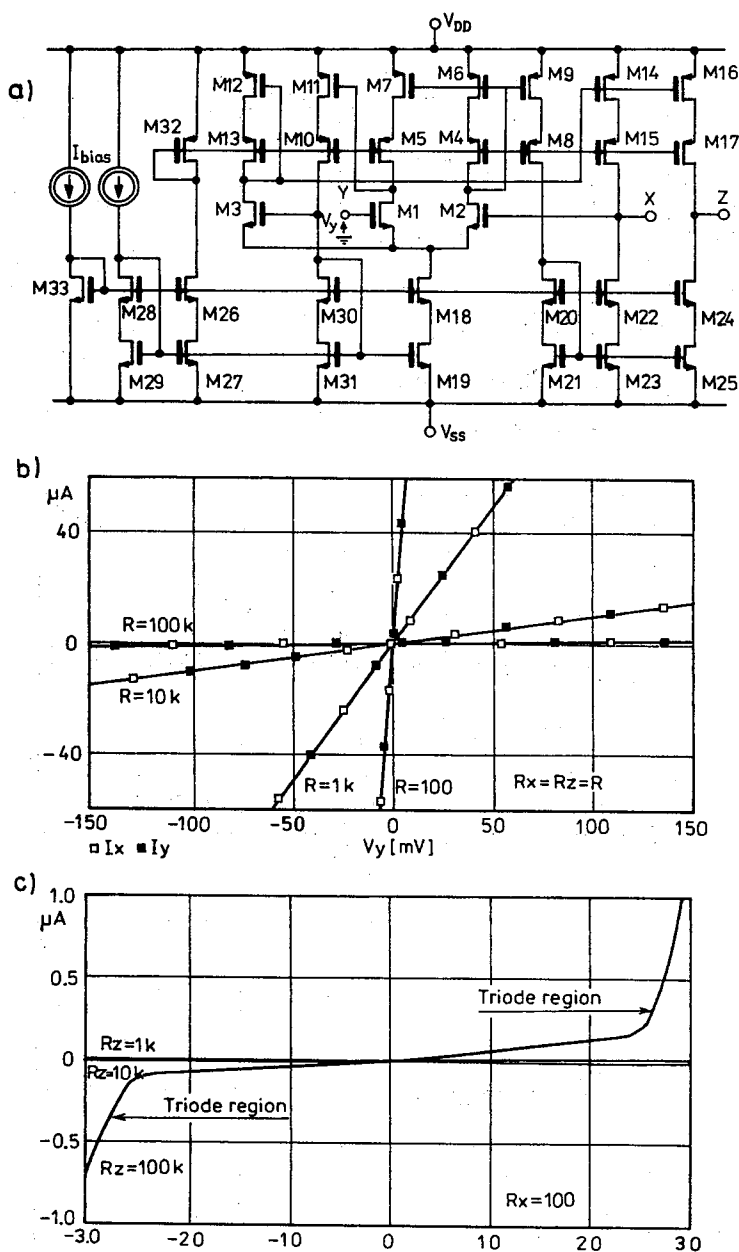
M	$M_{1\pm 3}$	$M_{4\pm 17}$	$M_{18\pm 29}$	$M_{30\pm 31}$	M_{32}	M_{33}
W [μm]	180	150	50	40	60	30
L [μm]	3	3	3	6	6	6

Rezystancja wejściowa r_x układu wynosi 4.1 Ω co przy rezystancji obciążenia $R_x = 100$ Ω daje błąd powtarzania napięcia U_y około 4%.

Układ cechuje się dużą liniowością charakterystyk $U_x = f(U_y)$, $I_z = f(U_y)$ przy różnych wartościach $R_x = R_y = R$ (rys. 7b) oraz małymi błędami transmisji prądowej $(I_x - I_z) = f(U_y)$ przy $R_x = \text{const}$ i różnych wartościach R_z (rys. 7c). Świadczy to o wyeliminowaniu zjawiska modulacji długości kanałów tranzystorów w układzie.



Rys. 6. Konweijor prądowy CCII⁺ ze zwiększonym obszarem liniowej pracy: a) schemat ideowy; b) charakterystyki transmisji prądowej: $I_x = f(U_y)$, $I_z = f(U_y)$; c) błąd bezwzględny transmisji prądowej: $(I_x - I_z) = f(U_y)$ dla $R_x \neq R_z$



Rys. 7. Konweyor prądowy CCH⁺ ze zwiększonym obszarem liniowej pracy i małą rezystancją wejściową na wejściu X: a) schemat ideowy; b) charakterystyki transmisji prądowej: $I_x = f(U_y)$, $I_z = f(U_y)$; c) błąd bezwzględny transmisji prądowej: $(I_x - I_z) = f(U_y)$ dla $R_x \neq R_z$

WNIOSKI KOŃCOWE

Ocenę różnych rozwiązań układowych konweytorów CCII⁺, wybranych z literatury (opisanych w punkcie 2) oraz zaprojektowanych układów oryginalnych (opisanych w punkcie 3), przeprowadzono w oparciu o ich parametry i charakterystyki statyczne. W tej ocenie zaprojektowane układy charakteryzują się bardzo dobrymi parametrami, jakich oczekuje się od układów profesjonalnych. Należy również podkreślić, że zaprojektowanie tych układów nie nastęrcza większych trudności. Większość tranzystorów posiada te same wymiary, przy czym założono, że szerokość kanału W_p tranzystorów pMOS jest trzykrotnie większa niż W_n tranzystorów nMOS.

Zróżnicowane wymiary W/L posiadają głównie tranzystory wykorzystywane do polaryzacji bramek tranzystorów w kaskadach.

Ocena właściwości częstotliwościowych rozpatrywanych układów wymaga oddzielnego i nieco szerszego omówienia. Należy jednak zauważyć, że układy o strukturze kaskadowej charakteryzują się również dobrymi właściwościami częstotliwościowymi.

BIBLIOGRAFIA

1. K.C. Smith, A. Sedra: *The current conveyor — a new circuit building block*. IEEE Proc., 1968, 56, pp. 1368–1369
2. A. Sedra, K.C. Smith: *A second generation current conveyor and its applications*. IEEE Trans., 1970, CT-17, pp. 132–134
3. W. Surakampontorn, V. Riewruja, K. Kumwachara, K. Dehjan: *Accurate CMOS based Current Conveyors*. IEEE Trans. on Instrumentation and Measurement, VOL., 40, No. 4, August 1991, pp. 699–701
4. S.I. Liu, H.W. Tsao, J. Wu: *CCII — based continuous — time filters with reduced gain — bandwidth sensitivity*. IEE Proc.-G, Vol. 138, No. 2, April 1991, pp. 210–215
5. Th. Laopoulos, S. Siskos, M. Bafleur, Ph. Givelin: *CMOS Current Conveyor*. Electronics Letters, 19th November 1992, Vol. 28, No. 24, pp. 2261–2262
6. B. Wilson: *Recent developments in current conveyors and current — mode circuits*. IEE Proc., Vol. 137, Pt.G, No. 2, April 1990, pp. 63–77
7. M.C.H. Cheng, C. Toumazou: *3V MOS Current Conveyor Cell for VLSI Technology*. Electronics Letters, 4th February 1993, Vol. 29, No. 3, pp. 317–318
8. A. Guziński, A. Kulej: *CMOS Second Generation Current Conveyor*. Proc. of the XVI-th National Conference — Circuit Theory and Electronic Circuits, Kołobrzeg, Pl., Oct.26–28, 1993, pp. 72–77
9. W. Surakampontorn, K. Kumwachara: *CMOS — based Electronically Tunable Current Conveyor*. Electronics Letters, 2nd July 1992, Vol. 28, Nr 14, pp. 1316–1317
10. P.J. Crawley, G.W. Roberts: *High — Swing MOS Current mirror with Arbitrarily High Output Resistance*. Electronics Letters 13th February 1992, Vol. 28, No. 4, pp. 361–363
11. P.J. Crawley, G.W. Roberts: *Designing Operational Transconductance Amplifiers for Low Voltage Operation*. Proc. of IEEE International Symposium on Circuits and Systems, Chicago, May 3–6 1993, Vol. 2, pp. 1455–1458
12. HSPICE Version H92 User's Manual Meta — Software, Inc. 1992

S. KUTA, W. MACHOWSKI, J. JASIELSKI

IMPROVED CMOS IMPLEMENTATION OF CURRENT CONVEYORS

Summary

The scope of presented paper is the comparison of basic CMOS implementation of current conveyors (CCII). The review was done on the basis of SPICE simulations.

The improved CMOS conveyor's designs with inner closed loop and high-swing cascoded current mirrors are also presented. Proposed circuits can work at reduced (± 3.3 V) voltages and exhibit good accuracy and linearity of signal processing.

Analiza i projektowanie falowodowego dzielnika mocy wytwarzającego sygnały wyjściowe o zadanych proporcjach amplitud

JERZY GORZKOWSKI

Przemysłowy Instytut Telekomunikacji, Warszawa

Otrzymano 1994.02.24

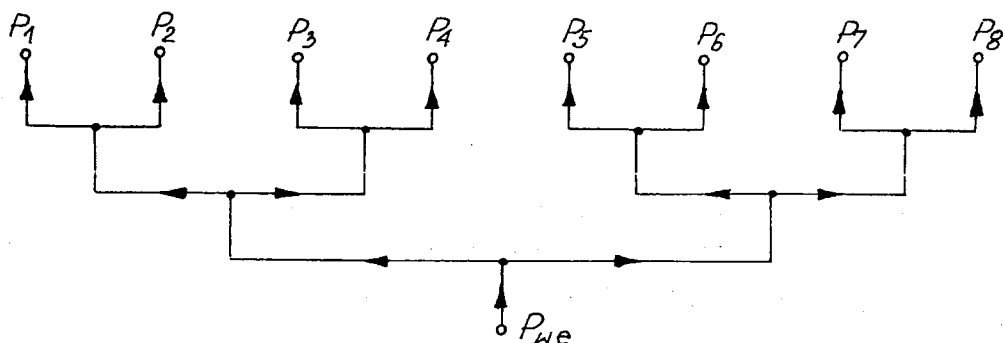
Autoryzowano do druku 1994.04.26

Rozważany jest dzielnik mocy utworzony z niesymetrycznych rozgałęzień falowodowych typu „T” w płaszczyźnie E. W oparciu o uproszczony układ zastępczy wyprowadzono zależności, z których można obliczyć przybliżone wymiary dzielnika, gdy są zadane współczynniki podziału mocy. Dokładne wymiary obliczane są za pomocą programu optymalizującego uwzględniającego wpływ niejednorodności pola elektromagnetycznego w rozgałęzieniach falowodowych. Funkcja celu użyta w programie optymalizującym osiąga minimum, gdy dzielnik jest dopasowany, a amplitudy sygnałów wyjściowych mają zadany rozkład. Uzyskuje się bardzo dobrą zgodność obliczonych i zmierzonych parametrów dzielnika.

1. WSTĘP

Mikrofalowe dzielniki mocy, z punktu widzenia ich struktury, podzielić można na dwie podstawowe grupy. Pierwszą z nich tworzą dzielniki o strukturze równoległej, których budowę ilustruje rys. 1.1. Sygnał doprowadzony na wejście dzielnika dzielony jest wielokrotnie na dwie części przez kolejne jego ogniwa aż do wytworzenia odpowiedniej liczby sygnałów wyjściowych, których amplitudy mają pożądane proporcje.

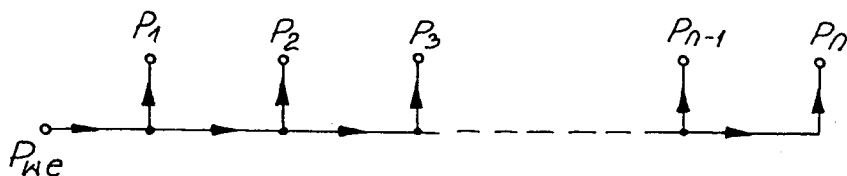
Podział mocy w każdym ogniwie dzielnika na dwie części nie jest obligatoryjny. Równie dobrze można wyobrazić sobie np. dzielnik utworzony z ogniw dzielących sygnał wejściowy na trzy części. Strukturę taką jest jednak zwykle znacznie trudniej zrealizować w praktyce. Zastosowanie dzielników o strukturze równoległej jest szczególnie wygodne gdy pożądane jest wytworzenie synfazowych sygnałów wyjściowych o równych amplitudach. Każde ogniwo dzielnika realizuje



Rys. 1.1. Dzielnik o strukturze równoległej

wtedy podział sygnału wejściowego na dwie równe części, a synfazowość sygnałów wyjściowych zapewniona jest przez symetrię struktury dzielnika.

Do budowy dzielników mogą być wykorzystane różne typy przewodnic falowych poczynając od linii paskowych i linii współosiowej, a kończąc na falowodzie prostokątnym. W przypadku linii paskowych i linii współosiowej, ogniwa dzielnika utworzyć można z równoległych rozgałęzień typu „T” lub „Y” albo z dzielników Wilkinsona [3,4] zapewniających separację między ramionami wyjściowymi, co powiększa walory dzielnika dzięki mniejszej wrażliwości jego parametrów na wewnętrzne i zewnętrzne niedopasowania. W przypadku dzielnika falowodowego, ogniwa najwygodniej jest utworzyć z rozgałęzień typu „T” lub „Y” w płaszczyźnie E. Stosunek mocy w ramionach wyjściowych rozgałęzienia, równy stosunkowi impedancji charakterystycznych tworzących je falowodów, równy jest wtedy stosunkowi ich wysokości, co zapewnia stałość podziału mocy w funkcji częstotliwości.



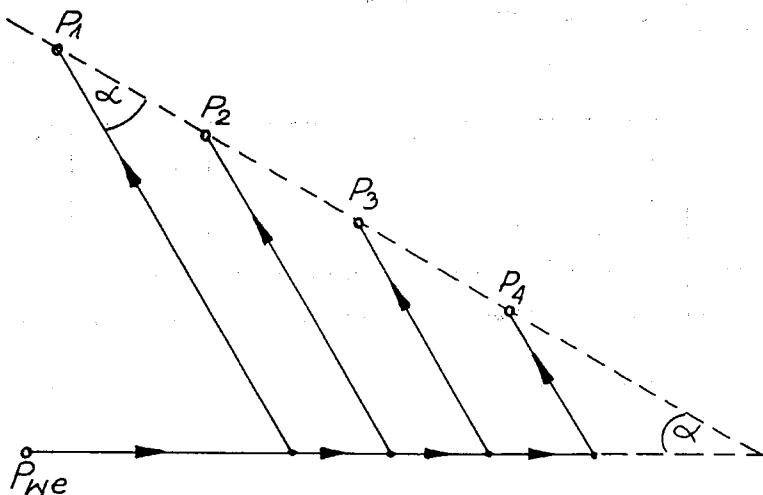
Rys. 1.2. Dzielnik o strukturze szeregowej.

Drugą grupę tworzą dzielniki o strukturze szeregowej, których budowę ilustruje rys. 1.2. Wszystkie ramiona wyjściowe sprzęgnięte są teraz bezpośrednio z linią, wzdłuż której przesyłana jest energia mikrofalowa doprowadzona na wejście dzielnika. Linie tę tworzyć mogą np. połączone łańcuchowo kolinearne ramiona rozgałęzień typu „T” lub tory główne sprzęgaczy kierunkowych. Niezależnie od typu przewodnicy falowej użytej do jego budowy, dzielnik utworzony ze sprzęgaczy kierunkowych góruje nad dzielnikiem zbudowanym ze zwykłych roz-

gałęzię dzięki mniejszej wrażliwości jego parametrów na niedopasowanie obciążeń w ramionach wyjściowych. Użycie sprzęgaczy kierunkowych powoduje jednak powiększenie gabarytów dzielnika zwłaszcza gdy jest on realizowany w technice falowodowej.

Warto zwrócić uwagę, że typ przewodnicy falowej użytej do budowy dzielnika decyduje o rodzaju rozgałęzień, z których utworzony jest dzielnik. W przypadku dzielnika na linii paskowej lub współosiowej będą to rozgałęzienia równoległe. Dzielnik falowodowy najkorzystniej jest natomiast utworzyć z rozgałęzień w płaszczyźnie E, które mają cechy rozgałęzień szeregowych.

Często wymaganą synfazowość sygnałów wyjściowych, w przypadku dzielników o strukturze szeregowej, łatwo jest uzyskać stosując odpowiednio dobrane długości ramion wyjściowych. Najwygodniej jest wykorzystać w tym celu „zasadę trójkąta równoramiennego”, którą ilustruje rys. 1.3. Z prostych zależności geometrycznych wynika, że w przypadku ukształtowania ramion wyjściowych zgodnie ze schematem pokazanym na rys. 1.3., długości dróg od wejścia do wszystkich wyjść dzielnika są takie same, co zapewnia synfazowość sygnałów wyjściowych.



Rys. 1.3. Ilustracja zasady trójkąta równoramiennego dla uzyskania synfazowości sygnałów wyjściowych.

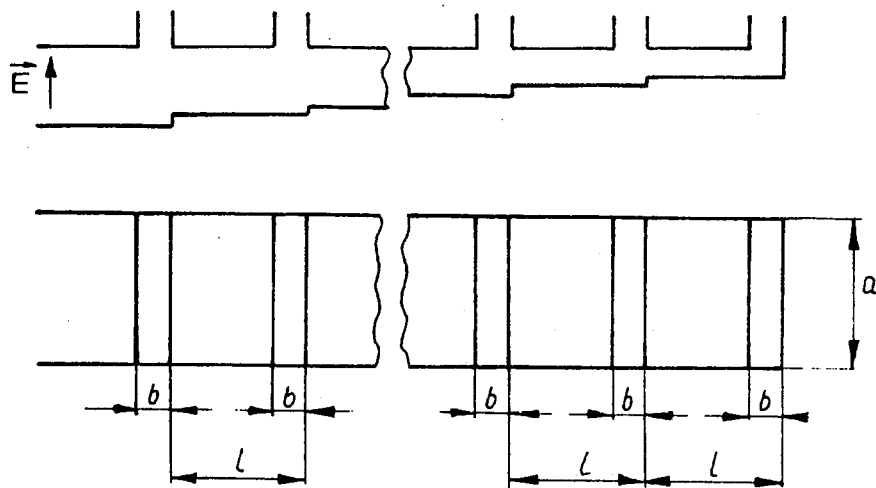
Dzielniki o strukturze szeregowej są chętnie stosowane gdy amplitudy sygnałów wyjściowych tworzą malejący ciąg o dużej liczbie elementów. Wymagania takie spełniać muszą zwykle dzielniki mocy służące do centralnego zasilania wierszy lub kolumn anten ścianowych [1,2]. O wyborze typu przewodnicy falowej użytej do budowy dzielnika decyduje najczęściej wymagana wytrzymałość mocowa dzielnika oraz jego gabaryty. W przypadku dzielników falowodowych, ogniwa budowane są

najczęściej z rozgałęzień typu „T” w płaszczyźnie E. Rozwiązania wykorzystujące sprzęgacze kierunkowe, pomimo niepodważalnych zalet, są rzadko używane ze względu na duże gabaryty.

Artykuł niniejszy poświęcony jest omówieniu właściwości i metody projektowania falowodowego dzielnika o strukturze szeregowej, którego ogniwa utworzone są z niesymetrycznych rozgałęzień typu „T” w płaszczyźnie E.

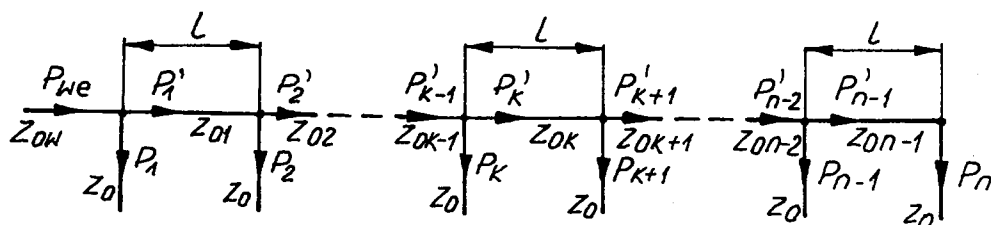
2. ANALIZA UPROSZCZONEGO UKŁADU ZASTĘPCZEGO DZIELNIKA

Szkic rozważanego dzielnika przedstawiono na rys. 2.1. Dzielnik zbudowany jest z falowodów prostokątnych o tej samej szerokości i składa się z ogniw utworzonych z niesymetrycznych rozgałęzień typu „T” w płaszczyźnie E.



Rys. 2.1. Szkic analizowanego dzielnika

Odległość l między osiami falowodów tworzących sąsiednie ramiona wyjściowe dzielnika jest zawsze jednakowa, a same falowody mają tę samą wysokość b , co oznacza, że wyjścia dzielnika zbudowane są na falowodach o tym samym przekroju. Najprostszy układ zastępczy dzielnika utworzyć można zaniedbując zaburzenia pola elektromagnetycznego powstające w obszarze rozgałęzień falowodowych. Układ zastępczy każdego rozgałęzienia uważać można wtedy za szeregowe połączenie linii przesyłowych o impedancjach charakterystycznych równych impedancjom łączonych falowodów. Ponadto, stosunki impedancji charakterystycznych falowodów, ze względu na jednakową ich szerokość, zastępować można stosunkami ich wysokości.



Rys. 2.2. Impedancyjny układ zastępczy dzielnika.

Utworzony w opisany wyżej sposób, impedancyjny układ zastępczy dzielnika umożliwia wyprowadzenie prostych zależności łączących jego wymiary z parametrami elektrycznymi. Z zależności tych wyznaczyć można przybliżone wartości wymiarów dzielnika, przy których realizowany jest zadany rozkład amplitud sygnałów wyjściowych. Otrzymane tą drogą wartości wymiarów dzielnika mogą być następnie użyte jako pierwsze przybliżenie podczas obliczeń dokładniejszych, w których każde rozgałęzienie reprezentowane jest przez układ zastępczy otrzymany przy uwzględnieniu niejednorodności pola elektromagnetycznego w obszarze połączenia falowodów.

Na rys. 2.2 pokazano impedancyjny układ zastępczy dzielnika, w którym moc wejściowa P_{we} zostaje podzielona na n części.

Ciąg mocy odprowadzanych przez kolejne wyjścia dzielnika P_k ($k = 1, 2, \dots, \dots, n$) jest nierosnący, a suma jego wyrazów jest równa mocy wejściowej ($\sum_{k=1}^n P_k = P_{we}$). Wszystkie wyjścia dzielnika wykonane są z linii przesyłowych o tej samej impedancji charakterystycznej Z_0 . Odcinki linii zawarte między kolejnymi odgałęzieniami, przez które przesyłane są moce P'_k , ($k = 1, 2, \dots, n-1$), mają tę samą długość l , a ich impedancje charakterystyczne równe są Z_{0k} ($k = 1, 2, \dots, n-1$).

Łatwo zauważyć, że dzielnik o konfiguracji pokazanej na rys. 2.2. nie może być utworzony z rozgałęzień dopasowanych od strony wejścia. Rozpływ mocy w każdym węźle, a tym samym rozkład amplitud sygnałów wyjściowych dzielnika, byłby wtedy zdeterminowany przez warunki dopasowania rozgałęzień oraz poziom mocy odprowadzanej przez pierwsze wyjście. Aby możliwe było uzyskanie dowolnego rozkładu amplitud sygnałów wyjściowych odcinki linii zawarte między kolejnymi wyjściami muszą spełniać rolę transformatorów ćwierćfalowych. Trzeba zatem przyjąć, że $l = \lambda_{f_0}/4$, gdzie λ_{f_0} jest długością fali w falowodzie dla środkowej częstotliwości pasma pracy f_0 .

Niech Z_{wk} , ($k = 1, 2, \dots, n-1$) oznacza impedancję wejściową części dzielnika leżącej na prawo od k -tego wyjścia w miejscu jego przyłączenia. Dla środkowej częstotliwości pasma pracy, gdy wszystkie Z_{wk} są rzeczywiste, rozpływ mocy w k -tym i $k+1$ węźle opisać można przy pomocy prostych zależności:

$$\frac{P'_k}{P_k} = \frac{P_{we} - \sum_{i=1}^k P_i}{P_k} = \frac{Z_{wk}}{Z_0} \quad (2.1)$$

$$\frac{P'_{k+1}}{P_{k+1}} = \frac{P_{we} - \sum_{i=1}^{k+1} P_i}{P_{k+1}} = \frac{Z_{wk+1}}{Z_0} \quad (2.2)$$

Równocześnie, ze znanych własności transformatora ćwierćfalowego wynika, że:

$$Z_{wk}(Z_{wk+1} + Z_0) = Z_{0k}^2$$

skąd, przez podzielenie obu stron równości przez Z_0^2 , dostaje się

$$\left(\frac{Z_{0k}}{Z_0}\right)^2 = \frac{Z_{wk}}{Z_0} \left(1 + \frac{Z_{wk+1}}{Z_0}\right). \quad (2.3)$$

Podstawienie (2.1) i (2.2) do wzoru (2.3) prowadzi do zależności:

$$\left(\frac{Z_{0k}}{Z_0}\right)^2 = \frac{P'_k}{P_k} \left(1 + \frac{P'_{k+1}}{P_{k+1}}\right). \quad (2.4)$$

pokazującej, że impedancja charakterystyczna transformatora ćwierćfalowego położonego między k -tym i $k+1$ węzłem wyznaczona jest przez rozływ mocy w tych węzłach. Wzór (2.4) słuszny jest dla wszystkich transformatorów z wyjątkiem ostatniego tzn. dla $k = 1, 2, \dots, n-2$. Aby wyznaczyć impedancję charakterystyczną Z_{0n-1} ostatniego transformatora wystarczy zauważyć, że:

$$Z_{wn-1} \cdot Z_0 = Z_{0n-1}^2$$

tzn.

$$\left(\frac{Z_{0n-1}}{Z_0}\right)^2 = \frac{Z_{wn-1}}{Z_0} \quad (2.5)$$

oraz:

$$\frac{P'_{n-1}}{P_{n-1}} = \frac{P_n}{P_{n-1}} = \frac{Z_{wn-1}}{Z_0} \quad (2.6)$$

Ze wzorów (2.5) i (2.6) otrzymuje się poszukiwaną zależność:

$$\left(\frac{Z_{0n-1}}{Z_0}\right)^2 = \frac{P_n}{P_{n-1}}. \quad (2.6)$$

Warunek dopasowania na wejściu dzielnika zapisać można w postaci:

$$Z_{0w} = Z_0 + Z_{w1}$$

gdzie Z_{0w} oznacza impedancję charakterystyczną linii na wejściu dzielnika. Stąd po podzieleniu obu stron przez Z_0 dostaje się:

$$\frac{Z_{0w}}{Z_0} = 1 + \frac{Z_{w1}}{Z_0},$$

a po wykorzystaniu zależności (2.1) dla $k = 1$, ostatecznie:

$$\frac{Z_{0w}}{Z_0} = 1 + \frac{P'_1}{P_1} = \frac{P_{we}}{P_1}. \quad (2.8)$$

Ze wzoru (2.8) wynika, że o stosunku impedancji charakterystycznej linii na wejściu i wyjściach dzielnika decyduje współczynnik sprzężenia pierwszego wyjścia.

Ze względu na jednakową szerokość falowodów użytych do budowy dzielnika, stosunki ich impedancji charakterystycznych, występujące w zależnościach (2.4), (2.7) i (2.8), zastąpić można stosunkami ich wysokości. Otrzymane w ten sposób wzory pozwalają na wyznaczenie względnych wymiarów dzielnika (wysokości wszystkich falowodów zredukowane są względem wysokości falowodów wyjściowych) gdy zadane są współczynniki podziału mocy.

Podczas ustalania wymiarów bezwzględnych należy pamiętać, że wysokości wszystkich falowodów muszą być dostatecznie małe, aby nie dopuścić do propagacji wyższych rodzajów. Powyższa uwaga dotyczy w pierwszym rzędzie falowodu wejściowego, którego wysokość jest największa. Szczególnie groźne są tu, pobudzane przez każde rozgałęzienie, rodzaje H_{11} i E_{11} .

Warto zauważyć, że układ zastępczy podobny do pokazanego na rys. 2.2 przypisać można także dzielnikom o strukturze szeregowej, zbudowanym na linii paskowej lub współosiowej, jeżeli ich ogniwa utworzone są ze zwykłych rozgałęzień typu „T”. Adekwatny układ zastępczy uzyskuje się przez zamianę szeregowych rozgałęzień linii na rozgałęzienia równoległe i zastąpienie impedancji charakterystycznych linii odpowiednimi wartościami admitancji. Aby otrzymać wzory wiążące wymiary dzielnika ze współczynnikami podziału wystarczy w zależnościach (2.1)–(2.8) zastąpić stosunki impedancji charakterystycznych linii stosunkami ich admitancji.

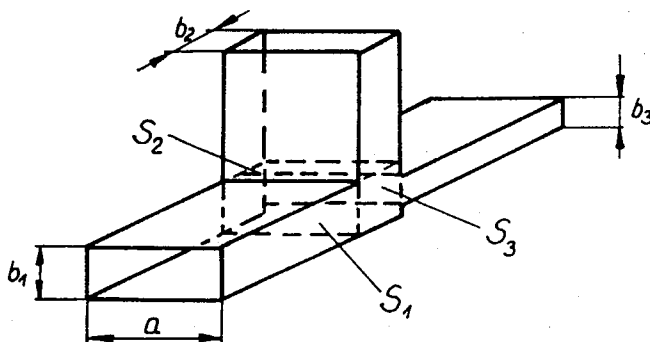
3. CZĘSTOTLIWOŚCIOWE CHARAKTERYSTYKI DZIELNIKA

Parametry dzielnika tzn. współczynniki podziału mocy oraz współczynnik fali stojącej na jego wejściu zależą od częstotliwości. Przebieg charakterystyk częstotliwościowych dzielnika, tzn. przebieg zależności wymienionych wyżej parametrów w funkcji częstotliwości, wyznaczyć można zgrubnie posługując się opisanym, impedancyjnym układem zastępczym. Otrzymane w ten sposób charakterystyki odbiegają jednak znacznie od rzeczywistych na skutek zaniedbania wpływu zaburzenia pola elektromagnetycznego w obszarze rozgałęzień falowodowych dzielnika. Ich przebieg uwzględnia tylko zmianę własności transformatorów ćwierćfalowych w funkcji częstotliwości. Niedokładność obliczeń współczynników podziału wzrasta w miarę oddalania się od wejścia dzielnika ponieważ przy łańcuchowym łączeniu jego ogniw błędy się sumują. Błędy rosną też, oczywiście, wraz ze wzrostem liczby wyjść dzielnika.

Zmniejszenie błędów teoretycznego wyznaczenia charakterystyk dzielnika wymaga wprowadzenia dokładniejszego układu zastępczego, w którym każde roz-

gałęzienie reprezentowane jest przez układ zastępczy uwzględniający niejednorodność pola elektromagnetycznego w obszarze połączenia falowodów. Dotyczy to również ostatniego ogniwa dzielnika, w którym dwa falowody o różnych wysokościach połączone są pod kątem prostym.

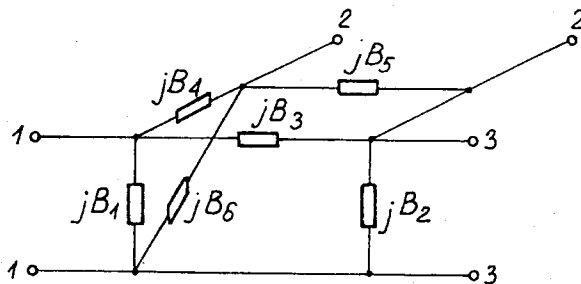
Szkic rozgałęzienia użytego do budowy ogniw rozważanego dzielnika, tzn. niesymetrycznego rozgałęzienia falowodowego typu „T” w płaszczyźnie E, przedstawiono na rys. 3.1.



Rys. 3.1. Szkic niesymetrycznego rozgałęzienia falowodowego typu „T” w płaszczyźnie E

Rozgałęzienie jest bezstratnym trójwrotnikiem odwracalnym, którego macierz admitancji jest symetryczna i składa się z elementów czysto urojonych. Właściwości rozgałęzienia charakteryzuje zatem sześć rzeczywistych parametrów zależnych od częstotliwości. Tyle samo niezależnych parametrów rzeczywistych występować musi w każdym układzie zastępczym, który można przyporządkować rozgałęzieniu. Jeden z takich układów, w którym występuje sześć niezależnych susceptancji pokazano na rys. 3.2. [5, 7, 8].

Zaciski wejściowe układu zastępczego odpowiadają dowolnie wybranym, ale ustalonym płaszczyznom odniesienia określającym wrota trójwrotnika. Mogą być nimi np. płaszczyzny S_1 , S_2 i S_3 zaznaczone linią przerywaną na rys. 3.1.

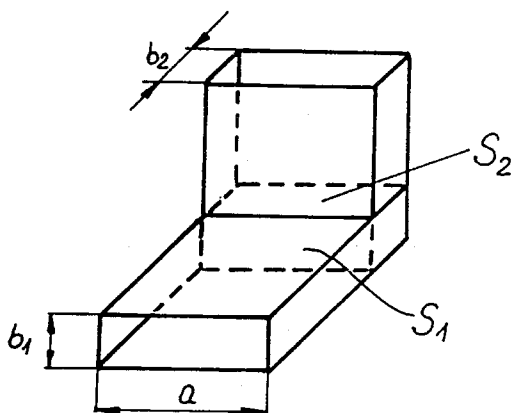


Rys. 3.2. Układ zastępczy niesymetrycznego rozgałęzienia falowodowego w płaszczyźnie E

Wzory określające wartości elementów układu zastępczego, w przypadku rozgałęzienia symetrycznego, w którym wysokości falowodów tworzących ramiona kolinearne są jednakowe ($b_1 = b_3$ wg oznaczeń wprowadzonych na rys. 3.1.), podał Marcuvitz [5]. Układ zastępczy takiego rozgałęzienia upraszcza się, ponieważ rozgałęzienie symetryczne charakteryzują w pełni cztery rzeczywiste parametry zależne od częstotliwości. Jeden z możliwych układów zastępczych takiego rozgałęzienia otrzymać można z układu pokazanego na rys. 3.2. przyjmując $B_6 \equiv 0$ i $B_1 \equiv B_2$.

W dostępnej literaturze poświęconej technice falowodowej brak jest zależności pozwalających obliczyć wartości elementów układu zastępczego rozgałęzienia niesymetrycznego. Podjętej przez G.Coxa próby obejścia powstającego stąd problemu, przez przypisanie rozgałęzieniu niesymetrycznemu układu zastępczego, który powstaje z połączenia układów zastępczych skoku wysokości falowodu prostokątnego oraz rozgałęzienia symetrycznego [2], nie można uznać za udaną. Otrzymany przez Coxa układ zastępczy jest niepełny, zawiera bowiem tylko pięć niezależnych parametrów rzeczywistych. Brak w nim gałęzi zawierającej susceptancję B_6 , tzn. $B_6 \equiv 0$ (wg oznaczeń wprowadzonych na rys.3.2.), co oczywiście nie jest prawdą.

Podobny problem występuje w przypadku ostatniego ogniwa dzielnika gdzie łączone są pod kątem prostym w płaszczyźnie E dwa falowody prostokątne o różnej wysokości, jak pokazano na rys. 3.3.



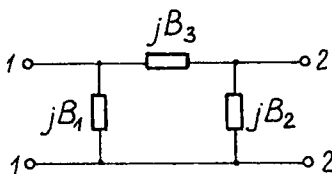
Rys. 3.3. Połączenie pod kątem prostym w płaszczyźnie E dwóch falowodów prostokątnych o różnej wysokości

Utworzony w ten sposób, odwracalny i bezstratny dwuwrotnik charakteryzują w pełni trzy rzeczywiste parametry zależne od częstotliwości. (Macierz admittancji jest symetryczna i czysto urojona). Jeden z możliwych układów zastępczych dwuwrotnika, w postaci niesymetrycznego układu typu II, pokazano na rys. 3.4.

[5, 7, 8]. Zaciski wejściowe układu zastępczego odpowiadają dowolnie wybranym płaszczyznom określającym położenie wrót dwuwrotnika, np. płaszczyznom S_1 i S_2 zaznaczonym linią przerywaną na rys. 3.3.

W dostępnej literaturze podawane są zależności pozwalające obliczyć wartości elementów układu zastępczego tylko w przypadku gdy łączone pod kątem prostym falowody mają jednakową wysokość [5] ($b_1 = b_2$ wg oznaczeń na rys. 3.3.), tzn. tworzą załamanie falowodu prostokątnego w płaszczyźnie E. Z warunków symetrii dwuwrotnika wynika wtedy, że $B_1 \equiv B_2$ tzn., że charakteryzują go w pełni dwa parametry rzeczywiste.

Do wyznaczenia wartości susceptancji występujących w układach zastępczych przedstawionych na rys. 3.2. i 3.4. użyta została „metoda uogólnionej macierzy admitancji” wielowrotnika. Koncepcja wprowadzenia takiej macierzy opiera się na znanej zasadzie [6], że składowe poprzeczne dowolnego pola elektromagnetycznego w falowodzie cylindrycznym o bezstratnych ściankach są superpozycją składowych poprzecznych pola pochodzących od rodzajów E i H, które tworzą razem układ zupełny.



Rys. 3.4. Jeden z możliwych układów zastępczych dwuwrotnika z rys. 3.3

Pole elektromagnetyczne w falowodach tworzących ramiona wielowrotników pokazanych na rys. 3.1. i 3.3. charakteryzuje się takim samym rozkładem wzdłuż szerszej ścianki jak pobudzające pole rodzaju H_{10} . Wynika to z jednorodności warunków brzegowych wzdłuż osi normalnej do węższych ścianek tych falowodów. W przypadku takich pól układ zupełny i ortogonalny tworzą rodzaje, które są kombinacją liniową rodzajów E_{1n} oraz H_{1n} i mają tę własność, że ich pole elektryczne pozbawione jest składowej normalnej do węższych ścianek falowodu. Rodzaje te, oznaczane dalej symbolem EH_{1n} , mogą być wyprowadzone z magnetycznego wektora Hertza posiadającego tylko jedną składową, która jest normalna do pary węższych ścianek falowodu [6]. Wskaźnik n może przyjmować wartości naturalne $n = 0, 1, 2, \dots$, przy czym rodzaj EH_{10} należy utożsamiać z rodzajem podstawowym H_{10} . Z ortogonalności rodzajów EH_{1n} wynika, że funkcje rozkładu składowych poprzecznych pola elektrycznego $\vec{e}_{S_n}$ i magnetycznego $\vec{h}_{S_n}$ dla tych rodzajów spełniają zależności:

$$\int_S \vec{e}_{S_n} \cdot \vec{e}_{S_m}^* dS = \begin{cases} Z_n & \text{dla } m = n \\ 0 & \text{dla } m \neq n \end{cases} \quad (3.1)$$

$$\int_S \vec{h}_{S_n} \cdot \vec{h}_{S_m}^* dS = \begin{cases} \tau_n^2 / Z_n & \text{dla } m = n \\ 0 & \text{dla } m \neq n \end{cases} \quad (3.2)$$

gdzie Z_n jest impedancją falową falowodu dla danego rodzaju gdy rodzaj ten jest propagowany lub wartością bezwzględną impedancji falowej gdy rodzaj ten jest tłumiony, a τ_n^2 wyrazić można równością:

$$\tau_n^2 = 1 + \frac{n^2}{\left(\frac{b}{a}\right)^2 \left\{ \left(\frac{2a}{\lambda}\right)^2 - 1 \right\}^2}, \quad (3.3)$$

w której a oznacza szerokość falowodu, b — jego wysokość, a λ długość fali w wolnej przestrzeni.

Należy tu przypomnieć, że wymiary falowodów tworzących ramiona wyjściowe wielowrotników z rys. 3.1. i 3.3. dopuszczają propagację tylko rodzaju podstawowego $H_{10} \equiv EH_{10}$.

Poprzeczne składowe pola elektromagnetycznego w płaszczyznach S_p tworzących wrota rozważanych wielowrotników ($p = 1, 2, 3$ dla trójwrotnika i $p = 1, 2$ dla dwuwrotnika) można przedstawić w postaci rozwinięć w szeregi Fouriera względem ortogonalnych układów funkcji rozkładu pola $\vec{e}_{S_k}$ i $\vec{h}_{S_k}$ dla odpowiednich falowodów. Zadowolając się skończoną liczbą wyrazów tych szeregów, poprzeczne składowe pola elektromagnetycznego $\vec{E}_s$ i $\vec{H}_s$ w płaszczyźnie S_p można wyrazić w przybliżeniu za pomocą wzorów:

$$\vec{E}_s = \sum_{n=1}^{N_p} (a_n + b_n) \vec{e}_{S_n} + \sum_{n=1}^{N_p} u_n \vec{e}_{S_n} \quad (3.4)$$

$$\vec{H}_s = \sum_{n=1}^{N_p} (a_n - b_n) \vec{h}_{S_n} + \sum_{n=1}^{N_p} i_n \vec{h}_{S_n} \quad (3.5)$$

gdzie N_p oznacza liczbę rodzajów na p -tych wrotach, uwzględnianych w obliczeniach. Współczynniki a_n oznaczają przy tym zredukowane amplitudy fal padających poszczególnych rodzajów, b_n zredukowane amplitudy fal odbitych, a u_n i i_n są odpowiednio zredukowanymi wartościami napięcia i prądu dla tych rodzajów.

Poprzeczne składowe pola elektromagnetycznego w płaszczyźnie odniesienia S_p wielowrotnika reprezentowane są zatem przez jednokolumnowe macierze (wektory) napięcia:

$$U_p = \begin{bmatrix} u_1^p \\ u_2^p \\ \vdots \\ u_{N_p}^p \end{bmatrix} \quad (3.6)$$

Nie wszystkie elementy uogólnionej macierzy admitancji wielowrotnika są urojone nawet gdy jest on bezstratny i odwracalny. Elementy Y_{nk}^{qp} dla wskaźników n odpowiadających rodzajom tłumionym są rzeczywiste. Urojone są tylko elementy Y_{nk}^{qp} odpowiadające rodzajom podstawowym.

W przypadku dwuwrotnika pokazanego na rys. 3.3, wyznaczenie pola elektromagnetycznego w obszarze V, a tym samym i elementów uogólnionej macierzy admitancji, nie stwarza żadnych trudności. Obszar V stanowi bowiem naturalne przedłużenie jednego lub drugiego falowodu, za pośrednictwem którego pole jest pobudzane. Podobna sytuacja występuje także w przypadku trójwrotnika z rys. 3.1, gdy pole pobudzane jest od strony wrót 1 i 3. Do wyznaczenia pola w obszarze V od strony wrót 2 użyto metody wariacyjnej. Metoda ta pozwala łatwo wyznaczyć współczynniki w szeregach aproksymujących pole w obszarze V przez rozwiązanie układu algebraicznych równań liniowych względem tych współczynników. W omawianym przypadku, dzięki ortogonalności rodzajów EH_{1n} użytych do aproksymacji pola w obszarze V, otrzymuje się szczególnie prostą postać tych równań, gdyż macierz współczynników przy niewiadomych jest diagonalna.

Konwencjonalną macierz admitancji wielowrotnika, wiążącą ze sobą napięcia i prądy dla fal rodzaju podstawowego, wyznacza się z macierzy uogólnionej wykorzystując warunek, że źródłem fal wyższych rodzajów w falowodach dołączonych do obszaru V jest tylko wielowrotnik. Oznacza to, że w falowodach dołączonych do obszaru V fale wyższych rodzajów „poruszają się” tylko w kierunku - od wielowrotnika, a między odpowiadającymi im prądami i napięciami zachodzi związek $i_k = -u_k$. Znak minus wynika z faktu, że prąd i_k skierowany jest przeciwnie niż poruszająca się fala tzn. w kierunku - do wielowrotnika. Problem wyznaczenia konwencjonalnej macierzy admitancji sprowadza się w rezultacie do problemu rozwiązania układu rzeczywistych, algebraicznych równań liniowych, aby napięcia odpowiadające pobudzonym, wyższym rodzajom wyrazić jako kombinacje liniowe napięć odpowiadających rodzajom podstawowym.

Warto przypomnieć, że elementy konwencjonalnej macierzy bezstratnego wielowrotnika odwracalnego, a więc także trójwrotnika z rys. 3.1 i dwuwrotnika z rys. 3.3, są czysto urojone.

Do numerycznego obliczania charakterystyk częstotliwościowych rozważanego dzielnika użyto, napisanego w języku FORTRAN, programu dla komputera PC/AT. Najważniejszymi segmentami tego programu były podprogramy wyznaczające wartości susceptancji występujących w układach zastępczych z rys. 3.2 i 3.4. Algorytmy obliczeń tych susceptancji oparto na naszkicowanej wyżej metodzie „uogólnionej macierzy admitancji”. Obydwa podprogramy przetestowano porównując obliczone przy ich pomocy wartości susceptancji z wartościami obliczonymi wg zależności podanych przez Marcuvitza [6] dla przypadku symetrycznego trójkąta falowodowego w płaszczyźnie E oraz załamania falowodu w płaszczyźnie E. Różnice obliczonych oboma metodami wartości nie przekraczały 1.5%.

Dodatkowym sprawdzianem poprawności działania podprogramów była dokładność spełnienia warunków wynikających z symetrii wielowrotników. W przypadku trójkąta warunki te sprowadzają się do spełnienia równości $B_1 \equiv B_2$ i $B_6 \equiv 0$ (dla oznaczeń jak na rys. 3.2), a dla połączenia falowodów pod kątem prostym - równości $B_1 \equiv B_2$ (dla oznaczeń jak na rys. 3.4). Stwierdzono, że powyższe równości były spełnione z dokładnością do 6-7 cyfr znaczących.

Obliczone przy pomocy opracowanego programu charakterystyki częstotliwościowe dzielnika z 14 wyjściami wykazywały różnice dochodzące do 2 dB w stosunku do parametrów dzielnika wyznaczonych na podstawie najprostszego, impedancyjnego układu zastępczego pokazanego na rys. 2.2. Maksymalne różnice występowały dla ostatniego wyjścia, najbardziej oddalonego od wejścia dzielnika.

4. PROJEKTOWANIE DZIELNIKA

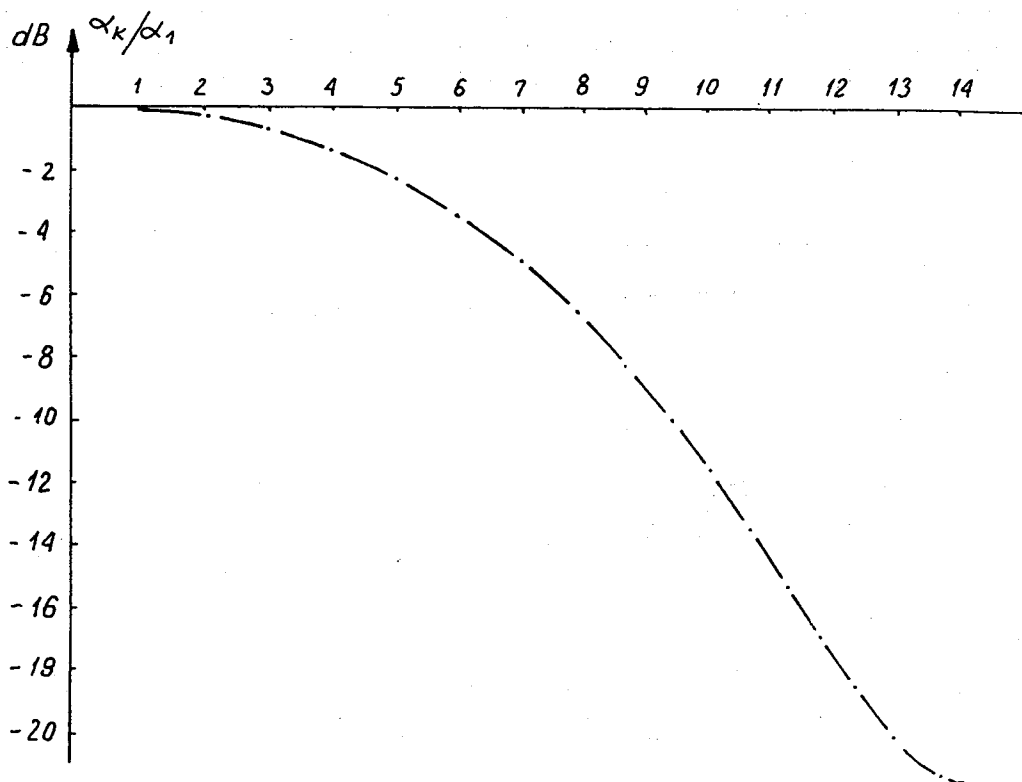
W oparciu o program obliczający charakterystyki częstotliwościowe dzielnika opracowano program wyznaczający jego wymiary dla zadanych wartości współczynników podziału mocy. Program ten, działający na zasadzie optymalizacji wymiarów dzielnika w celu uzyskania parametrów mikrofalowych maksymalnie zbliżonych do wartości zadanych, pozwala na projektowanie dzielnika wytwarzającego zadany rozkład sygnałów wyjściowych. użytą w programie optymalizującym funkcję celu utworzono w taki sposób, aby osiągała minimum równe zero w przypadku dzielnika idealnego tzn. dzielnika, którego WFS (współczynnik fali stojącej) wejściowy jest równy jedności, a współczynniki podziału mocy są stałe i równe zadanim w całym pasmie pracy dzielnika. Zmiennymi niezależnymi funkcji celu, względem których w procesie optymalizacji poszukiwane jest jej minimum, są wysokości i długość odcinków falowodów zawartych między sąsiednimi wyjściami dzielnika oraz wysokość falowodu wejściowego, zredukowane względem wysokości falowodów wyjściowych.

Opracowany program, oparty na procedurze optymalizującej Hooke'a-Jeevesa, umożliwił zaprojektowanie 14 - to wyjściowego dzielnika, w którym rozkład amplitud sygnałów wyjściowych określony wzorem:

$$a_k = a + b \cos^2 \frac{\pi}{2} \frac{2k-1}{2n-1} \quad k = 1, 2, \dots, n \quad (4.1)$$

gdzie a i b oznaczają stałe (cosinus kwadrat na piedestale), pokazany jest na rys. 4.1.

Poprawność przeprowadzonych obliczeń potwierdziły badania dwóch wykonanych egzemplarzy zaprojektowanego dzielnika, pracujących w pasmie L. Pomiarzy charakterystyk dzielnika przeprowadzono za pomocą automatycznego ana-



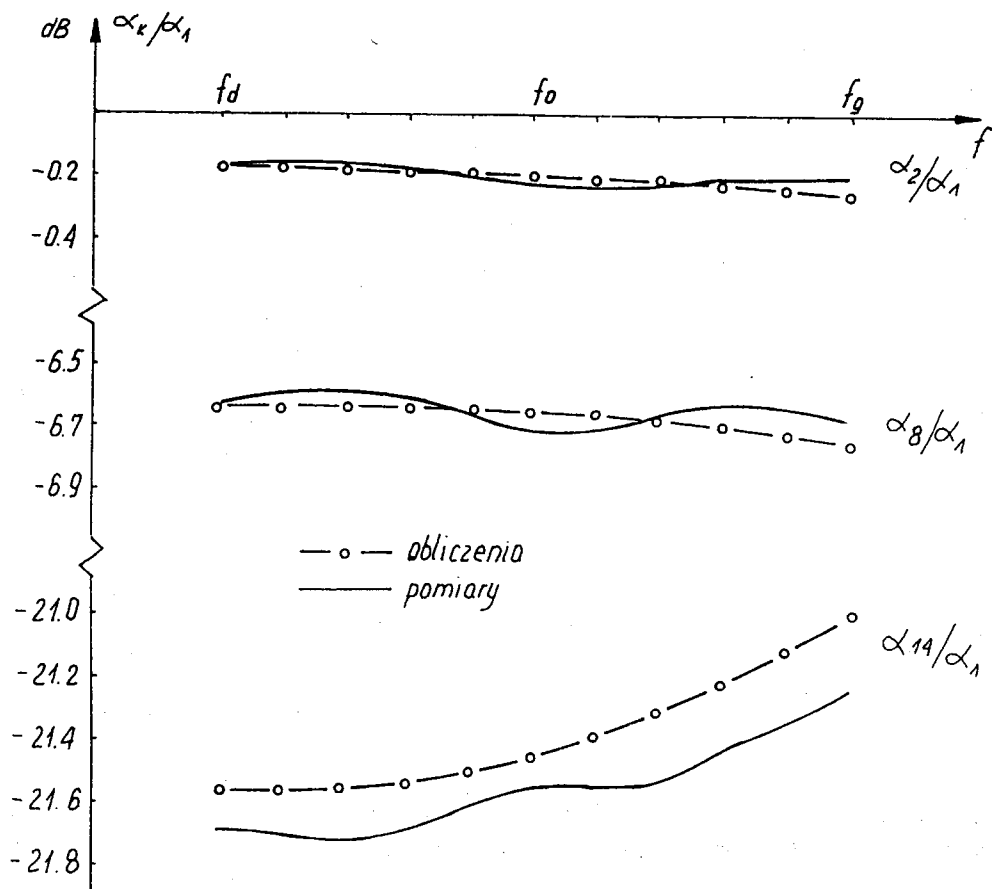
Rys. 4.1. Rozkład amplitud sygnałów wyjściowych zaprojektowanego, 14-to wyjściowego dzielnika

lizatora obwodów mikrofalowych firmy Hewlett-Packard. Ze względu na dużą wrażliwość współczynników podziału mocy na niedopasowanie ramion wyjściowych dzielnika, w pomiarach używane były specjalnie w tym celu opracowane obciążenia oraz sprzęgacze kierunkowe. WFS obciążeń nie przekraczał 1.02, a tor główny sprzęgaczy kierunkowych miał WFS tak mały, że nie udało się go zmierzyć za pomocą dostępnych układów pomiarowych.

Rozbieżności między policzonymi i zmierzonymi charakterystykami współczynników podziału mocy nie przekraczają wartości 0.25 dB. Można przy tym zaobserwować że różnice te rosną w miarę oddalania się od wejścia dzielnika tzn. na ogół są tym większe im wyjście dzielnika położone jest dalej od jego wejścia. Wynika to z szeregowej struktury dzielnika, która ma tę wadę, że na charakterystykę sprzężenia k -tego wyjścia mają wpływ przebiegi charakterystyk sprzężenia wszystkich wyjść poprzedzających.

Na rys. 4.2. pokazano przykładowo przebieg zmierzonych i obliczonych charakterystyk sprzężenia drugiego, ósmego i czternastego (ostatniego) wyjścia dziel-

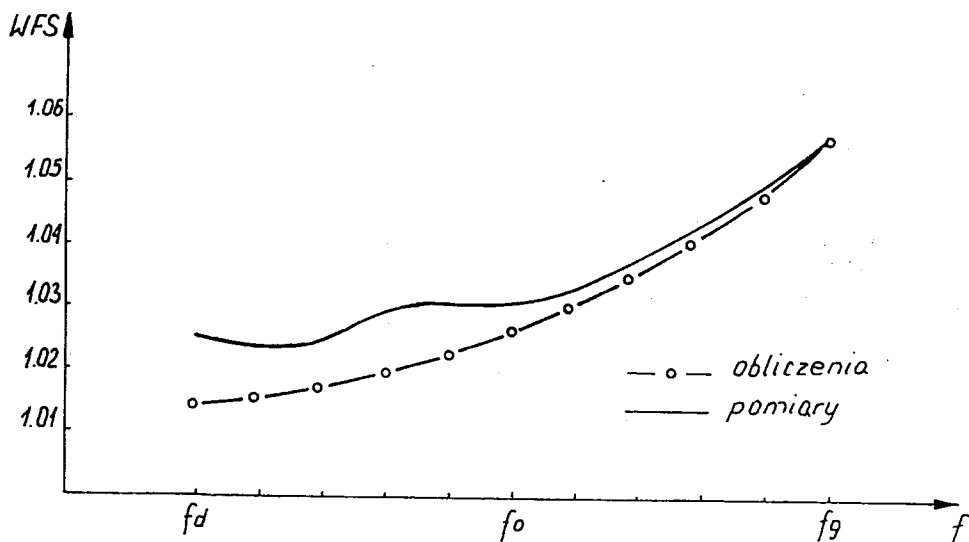
nika odniesionych do sprzężenia pierwszego wyjścia, w pasmie częstotliwości o szerokości względnej ok. 10%. Jak widać, największe różnice między zmierzoną i obliczoną charakterystyką sprzężenia występują dla ostatniego wyjścia.



Rys. 4.2. Przebieg zmierzonych i obliczonych charakterystyk sprzężenia drugiego, ósmego i czternastego wyjścia dzielnika odniesionych do sprzężenia pierwszego wyjścia, w pasmie częstotliwości o szerokości względnej ok 10%

Charakterystyki są dość płaskie w funkcji częstotliwości (skala decybelowa na rys. 4.2 jest celowo rozciągnięta), przy czym stałość sprzężenia w pasmie pracy jest tym lepsza im wyjście położone jest bliżej wejścia dzielnika. Na rys. 4.3 pokazano obliczoną i zmierzoną charakterystykę WFS-u na wejściu dzielnika.

Rozbieżność między obliczonymi i zmierzonymi wartościami WFS-u jest największa dla dolnej części pasma pracy gdzie dopasowanie dzielnika jest bardzo dobre. Spowodowane jest to prawdopodobnie dużymi błędami pomiaru małych



Rys. 4.3. Zmierzona i obliczona charakterystyka WFS-u na wejściu dzielnika

wartości WFS-u. Zmierzona i obliczona, maksymalna wartość WFS-u wejściowego dzielnika w pasmie pracy nie przekracza 1.06.

LITERATURA

1. A. Rogers: *Wideband Squintless Linear Arrays*. The Marconi Review, Nr 187, Fourth Quarter 1972
2. G. Cox Graham: *An Analysis of the Waveguide Squintless Feed*. The Marconi Review, Third Quarter 1981
3. L.I. Parad, R.L. Moynihan: *Split-tee Power Divider*. IEEE Trans. on MTT 13, Jan. 1965.
4. S. Cohn: *A Class of Broadband, Three-Port TEM Mode Hybrides*. IEEE Trans. on MTT, vol.16, Nr 2, Feb. 1968.
5. N. Marcuvitz: *Waveguide Handbook*. MIT Rad. Lab. Series, vol. 10. McGraw-Hill Book Co.,Inc, New York, 1951
6. R.E. Collin: *Prowadzenie Fal Elektromagnetycznych*. Wydawnictwo Naukowo-Techniczne, NT Warszawa 1966
7. C. G. Montgomery: *Principles of Microwave Circuits*. MIT Rad. Lab. Series, vol. 8, McGraw-Hill Book Co.,Inc., New York 1948
8. J.L. Altman: *Microwave Circuits*. D.Van Nostrand Company, Inc. New Jersey, Toronto, New York, London 1964
9. *Metody Optymalizacji w Języku FORTRAN*. Praca zbiorowa pod redakcją Jacka Szymanowskiego, PIW, Warszawa 1984

J. GORZKOWSKI

ANALYSIS AND DESIGN OF WAVEGUIDE DIVIDER FORMING OUTPUT SIGNALS
WITH AMPLITUDES WHICH PROPORTION ARE GIVEN

Summary

The waveguide divider formed by nonsymmetric E-plane Tee junctions is considered. On base of the simplified equivalent circuit of the divider formulas for calculation approximate dimensions of divider for given coefficients of power division were derived. Exact dimensions are calculated by aid of the optimization program where the influence of the electromagnetic field's inhomogeneities in waveguides' junctions is taken into account. The objective function use in program have minimum when the divider is matched and amplitudes of output signals have desirable distribution. The numerical results agree very well with the experimental values.

621.319.1
621.314.2.042

Ceramiczne ferroelektryki perowskitowe o półprzewodnikowym i metalicznym przewodnictwie elektrycznym

ZYGMUNT SUROWIAK, DARIUSZ KRAWCZYK

Instytut Problemów Techniki, Uniwersytet Śląski

EVGENIJ G. FESENKO

Instytut Fizyki, Uniwersytet Rostowski, Rostów nad Donem

Otrzymano 1994.03.15

Autoryzowano do druku 1994

Przedstawiono jakościowy model przejścia ceramicznych ferroelektryków perowskitowych od stanu dielektrycznego do stanu półprzewodnikowego oraz do stanu przewodnictwa elektrycznego typowego dla metali. Na przykładzie tworzyw ceramicznych BaTiO_3 i SrTiO_3 przedstawiono metody zwiększania przewodnictwa elektrycznego dielektrycznych perowskitów oraz przeanalizowano na poziomie mikroskopowym fizyczne mechanizmy, prowadzące do wzrostu koncentracji swobodnych nośników ładunku elektrycznego w tych materiałach. Wnioski wynikające z teoretycznej analizy procesów transportu elektrycznego w modelowych ferroelektrykach, jakimi są BaTiO_3 i SrTiO_3 , zweryfikowano za pomocą badań eksperymentalnych; uzyskano dobrą zgodność teoretycznego modelu z rezultatami doświadczalnymi. Wzrost przewodnictwa elektrycznego BaTiO_3 i SrTiO_3 uzyskiwano poprzez ich domieszkowanie oraz redukcję tlenową w wyniku obróbki termiczno-próżniowej. Badania oporności elektrycznej prowadzono w zakresie temperatur 4.2–300 K.

1. WPROWADZENIE

Materiały o idealnej strukturze typu perowskitu $(\text{CaTi})\text{O}_3$ i o idealnym składzie stechiometrycznym (ABO_3) ze względu na przewodnictwo elektryczne zalicza się do dielektryków. Otrzymanie takich idealnych perowskitów w formie monokryształów, lub polikryształów jest w praktyce bardzo trudne. Najczęściej tego typu związki tlenooktaedryczne charakteryzują się pewnymi odstępstwami

od składu stechiometrycznego (np. wakancjami anionowymi lub kationowymi) i (lub) deformacją idealnej struktury perowskitowej. Charakterystyczny dla perowskitów znaczny izomorfizm, łatwość odchyień od stechiometrii i łatwość naruszenia translacyjnej symetrii pozwalają sterować ich właściwościami fizycznymi (przez dobór warunków technologicznych oraz domieszkowanie). Dlatego też realne perowskity (zarówno monokryształy jak i materiały ceramiczne) mogą wykazywać wiele unikalnych właściwości elektrycznych (piezoelektryki, piroelektryki ferro- lub antyferroelektryki), magnetycznych (ferromagnetyczne ferroelektryki), optycznych, elektrooptycznych, ciepłych itd. [1–2].

Od wielu lat w rodzinie perowskitów szczególne zainteresowanie wywołują związki o właściwościach ferroelektrycznych. Krystalochemiczne uwarunkowania stanu ferroelektrycznego w perowskitach autorzy przedstawili już wcześniej w pracach [3–4].

Ferroelektryki tlenowe o strukturze typu perowskitu (F-TSP) w zależności od składu chemicznego i stopnia doskonałości struktury krystalicznej mogą mieć właściwości typowe dla nieliniowych dielektryków lub właściwości półprzewodnikowe [5]. Większość F-TSP pod względem przewodnictwa elektrycznego zajmuje obszar pośredni pomiędzy typowymi dielektrykami i półprzewodnikami, zachowując się jak „półdielektryki”, bądź półprzewodniki o szerokiej przerwie energetycznej [6].

Badania nad przewodnictwem elektrycznym i innymi zjawiskami transportu nośników ładunku w F-TSP zostały skoncentrowane głównie na modelowym związku, jakim jest tytanian baru (BaTiO_3) [7–10]. Natomiast badania nad możliwością wystąpienia w F-TSP stanu nadprzewodnictwa wysokotemperaturowego (HTS) prowadzone były przede wszystkim na przykładzie tytanianu strontu (SrTiO_3) [np. 11–12] lub związków o warstwowej perowskito-podobnej strukturze typu $\text{A}_{m-1}\text{Bi}_2\text{B}_m\text{O}_{3m+3}$ [np. 13–15].

Pomimo, iż na temat przewodnictwa elektrycznego F-TSP ukazało się dotychczas ponad 1500 prac, a w tym kilka obszernych monografii, to w dalszym ciągu teoria tego zjawiska jest jeszcze słabo zaawansowana. Żaden z opracowanych modeli dotyczących przewodnictwa elektrycznego F-TSP nie wyjaśnia w zadowalającym stopniu całokształtu zaobserwowanych faktów doświadczalnych [6]. Problemem podstawowym jest nadal określenie natury nośników ładunku i mechanizmu przewodnictwa elektrycznego. Dzisiaj można już wprawdzie odpowiedzieć na pytanie czy przewodnictwo elektryczne F-TSP ma charakter jonowy czy elektronowy, samoistny, czy domieszkowy, lecz w dalszym ciągu brak jest pełnego opisu w skali mikroskopowej mechanizmu przejścia dielektrycznego perowskitu w stan półprzewodnikowy.

Autorzy niniejszej pracy w oparciu o badania własne oraz dane literaturowe podjęli próbę opisu mechanizmu przejścia dielektrycznych perowskitów w stan półprzewodnikowego, a następnie w stan metalicznego przewodnictwa na przykładzie ceramicznych materiałów BaTiO_3 i SrTiO_3 . Oddzielnego potraktowania

wymaga problem wytworzenia w perowskito-podobnych związkach stanu HTS (zagadnieniu temu poświęcona będzie oddzielna praca).

Problematyka niniejszej pracy obok aspektów czysto poznawczych ma również duże znaczenie dla praktycznych zastosowań F-TSP w elektronice i elektrotechnice. Wykazano bowiem już dawno [np. 6], że swobodne nośniki ładunku elektrycznego wywierają zasadniczy wpływ na właściwości ferroelektryczne perowskitów. W makroskopowym ujęciu wpływ ten sprowadza się do ekranowania polaryzacji spontanicznej (P_s). W wyniku ekranowania swobodne nośniki ładunku wpływają na efektywną wartość polaryzacji spontanicznej, na parametry procesu przepolaryzowania F-TSP w zewnętrznym polu elektrycznym (pole koercji, czas przepolaryzowania) oraz na zachowanie F-TSP pod wpływem zmian temperatury (punkt Curie, efekt piroelektryczny) i zewnętrznych naprężeń mechanicznych (efekt piezoelektryczny) [2]. Duża koncentracja swobodnych nośników ładunku elektrycznego może w niektórych przypadkach ograniczać możliwość praktycznego wykorzystania F-TSP (np. jako ferroelektrycznych elementów pamięci), lecz może równocześnie sprzyjać innym zastosowaniom tych materiałów w elektronice (np. jako ferroelektrycznych fotoprzekaźników i fotopowielaczy, termistorowych mierników temperatury, elementów ferroelektryczno-półprzewodnikowych typu FEFED, pozystorowych ograniczników natężenia prądu elektrycznego itd.).

2. MODEL TEORETYCZNY

2.1. PEROWSKITY Z WAKANCIAMI

Ferroelektryki tlenkowe o strukturze typu perowskitu (F-TSP) i o małej koncentracji defektów charakteryzują się słabym przewodnictwem elektrycznym ($\sigma \approx 10^{-13} \Omega^{-1} \text{cm}^{-1}$). Obok słabego przewodnictwa jonowego takie F-TSP wykazują również słabe samoistne przewodnictwo elektronowe, polegające na przejściach elektronów z pasma walencyjnego do pasma przewodnictwa [2].

Elektryczne przewodnictwo właściwe σ takich materiałów wyraża się równaniem:

$$\sigma = en\mu, \quad (1)$$

gdzie:

e — ładunek elementarny,

n — koncentracja swobodnych nośników ładunku elektrycznego,

μ — ruchliwość swobodnych nośników ładunku elektrycznego

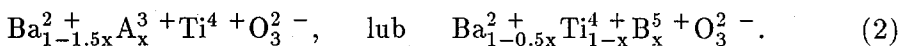
($\mu = \frac{v}{E}$ [$\text{m}^2 \text{V}^{-1} \text{s}^{-1}$], gdzie v prędkość poruszania się ładunku w zewnętrznym polu elektrycznym E).

Szerokość przerwy energetycznej w dielektrycznych F-TSP jest duża ($E_g \approx 3 \text{ eV}$) a w związku z tym koncentracja równowagowych nośników ładunku jest

bardzo niska ($n \approx 10^6 \text{ cm}^{-3}$). Również ruchliwość elektronów jest bardzo mała ($\mu \approx (10^{-4} - 10^{-1}) \text{ cm}^2 \text{ V}^{-1} \text{ s}^{-1}$), co wskazuje na polaronowy mechanizm przewodnictwa elektrycznego [7].

W przypadku F-TSP z defektami punktowymi, a więc praktycznie we wszystkich materiałach polikrystalicznych, w temperaturach niższych od 873 K ma miejsce przewodnictwo elektronowe, uwarunkowane istnieniem w paśmie wzbronionym lokalnych poziomów energetycznych. Poziomy te związane są przede wszystkim z samoistnymi punktowymi defektami struktury, wśród których główną rolę odgrywają wakancje (V) i domieszki.

Nie każdy jednak mechanizm tworzenia się wakancji w położeniach A(V_A) perowskitu ABO_3 prowadzi do powstania poziomów lokalnych w przerwie wzbronionej. I tak np. w przypadku podstawień do BaTiO_3 w podsieć Ba^{2+} jonów A^{3+} lub w podsieć Ti^{4+} jonów B^{5+} o koncentracji większej od 1 % at. powstają obojętne elektrycznie struktury z wakancjami w położeniach A. Ich wzory chemiczne przedstawić można następująco:



Związki takie charakteryzują się właściwościami dielektrycznymi.

Koncentracja wakancji, wywołanych modyfikowaniem BaTiO_3 , może być bardzo wysoka, lecz nie zmienia to stabilności struktury perowskitu. Wzrost wakancji wywołany być może poprzez podgrzewanie F-TSP do temperatur przewyższających 1073–1273 K. W związkach ABO_3 wakancje jako nie obsadzone węzły odpowiednich podsieci, przyjęto oznaczać następująco: V_A , V_B , V_O . Badania eksperymentalne wykazały, że dla F-TSP charakterystyczne są przede wszystkim wakancje V_O , wakancje V_B powstają w ilościach niewielkich, a V_A mają małą koncentrację w niedomieszkowanych F-TSP na bazie Ba, Sr i Ca, zaś bardzo dużą — w F-TSP zawierających Pb lub Cd. Domieszkowe przewodnictwo elektryczne w F-TSP jest dominujące w przypadku podstawień w podsięci kationowe jonów o innej wartościowości.

Sterowanie koncentracją wakancji oraz kontrolowane domieszkowanie są podstawą wszystkich znanych metod zmiany przewodnictwa elektrycznego F-TSP. Szczególnie wysoka koncentracja nośników ładunku elektrycznego charakteryzuje nadprzewodniki (NP). Stąd też w trakcie poszukiwań nowych NP w rodzinie F-TSO otrzymano materiały o półprzewodnikowym, a nawet o metalicznym typie przewodnictwa elektrycznego.

Efektywną metodą wytwarzania wakancji jest obróbka termiczno-próżniowa (TPP). Stosując ją można podwyższyć przewodnictwo elektryczne niektórych F-TSP o 11–12 rzędów wartości. Najczęściej spotykanymi defektami w zredukowanych F-TSP są kompleksy utworzone z wakancji tlenowej (V_O) oraz dwóch sąsiadujących z nią jonów B. W miarę wzrostu koncentracji V_O odpowiadające defektom poziomy lokalne w paśmie wzbronionym rozmywają się i przekrywają z dnem pasma przewodnictwa. W konsekwencji poziom Fermiego (E_F) przemiesz-

cza się z pasma wzbronionego do pasma przewodnictwa, tzn. następuje zwyrodnienie półprzewodnika.

Jeśli przyjąć, że ubytek masy próbki BaTiO_3 poddanej TPO wiąże się z ubytkiem tlenu, to zmianę koncentracji V_O można określić następująco:

$$\Delta V_O = \frac{\Delta m}{m} \cdot \frac{\rho \cdot N_A}{M_0}, \quad (3)$$

gdzie:

$\frac{\Delta m}{m}$ — względna zmiana masy próbki poddanej TPO;

ρ — gęstość próbki;

N_A — stała Avogadra;

M_0 — masa cząsteczkowa tlenu.

Zmianie koncentracji V_O towarzyszy zmiana koncentracji elektronów swobodnych:

$$\Delta n = 2\Delta V_O \quad (4)$$

Wynika z tego, że dla otrzymania zwyrodniałego gazu elektronowego ($n = 5 \times 10^{19} \text{ cm}^{-3}$) w półprzewodnikowym BaTiO_3 należy wytworzyć 0.03 % deficyt tlenu. Natomiast metaliczne przewodnictwo ($n \gg 5 \times 10^{21} \text{ cm}^{-3}$) może być osiągnięte przy 0.5 % ubytku tlenu.

W F-TSP z ułatwiającym się jonem A (np. w perowskitach PbZrO_3 i PbTiO_3 zawierających Pb) powstają kompleksy utworzone z wakancji V_A i otaczających tę wakancje jonów tlenu. W tym przypadku przewodnictwo elektryczne ma charakter dziurowy.

Strukturę perowskitu ma również SrTiO_3 . Zwiększając koncentrację nośników ładunku elektrycznego poprzez redukcję tlenu można wytworzyć w nim nawet stan nadprzewodnictwa w temperaturze 0.05–0.06 K. Podobną metodą starano się otrzymać nadprzewodniki na bazie innych F-TSP [np. 16].

2.2. MECHANIZMY PRZEJŚCIA F-TSP DO STANU PÓLPRZEWODNIKOWEGO I METALICZNEGO

Rozpatrzmy teraz bardziej szczegółowo mechanizm przejścia dielektrycznych F-TSP w stan półprzewodnikowego i metalicznego przewodnictwa elektrycznego poprzez zwiększenie koncentracji swobodnych nośników ładunku elektrycznego.

Do najczęściej stosowanych metod zwiększania koncentracji swobodnych nośników ładunku zaliczyć można:

- procesy redukcji tlenowej;
- podstawianie do sieci krystalicznej obcych atomów (M) o wartościowości przewyższającej wartościowość atomów własnych (sterowanie wartościowością może zwiększyć przewodnictwo elektryczne półprzewodnika o 10–11 rzędów wartości [7]).

Bardzo efektywne są również kombinacje obydwu metod.

Tablica 1

**Typy reakcji podstawienia (1–4) kationów M
w miejsce kationów A i (lub) B oraz reakcja redukcji (5)**

Stosunek wartościowości	Stosunek promieni jonowych	Reakcje podstawienia dla $x \ll 0.005$
$N_M \gg N_A$	$R_A \approx R_M \gg R_B$	$A^{2+}B^{4+}O_3^{2-} + xM^{3+} \rightarrow A_{1-x}^{2+}M_x^{3+}B^{4+}(e)_x^-O_3^{2-} \rightarrow$ $\rightarrow A_{1-x}^{2+}M_x^{3+}B_{1-x}^{4+}B_x^{3+}O_3^{2-}$
$N_M \gg N_B$	$R_A \gg R_M \approx R_B$	$A^{2+}B^{4+}O_3^{2-} + xM^{5+} \rightarrow A^{2+}B_{1-x}^{4+}M_x^{5+}(e)_x^-O_3^{2-} \rightarrow$ $\rightarrow A^{2+}B_{1-2x}^{4+}B_x^{3+}M_x^{5+}O_3^{2-}$
$N_M \gg N_B$	$R_A \gg R_M \approx R_B$	$A^{2+}B^{4+}O_3^{2-} + xM^{5+} \rightarrow A^{2+}B_{1-x}^{4+}M_x^{5+}(e)_x^-O_3^{2-} \rightarrow$ $\rightarrow A^{2+}B_{1-x}^{4+}M_x^{4+}O_3^{2-}$
$N_M \gg N_B$	$R_A \gg R_M \approx R_B$	$A^{2+}B^{4+}O_3^{2-} + xM^{6+} \rightarrow A^{2+}B_{1-x}^{4+}M_x^{6+}(e)_{2x}^-O_3^{2-} \rightarrow$ $\rightarrow A^{2+}B_{1-2x}^{4+}B_{2x}^{3+}M_x^{6+}O_3^{2-}$
—	—	Reakcja redukcji: $A^{2+}B^{4+}O_3^{2-} - xO \rightarrow A^{2+}B^{4+}(e)_{2x}^-O_{3-x}^{2-} \rightarrow$ $\rightarrow A^{2+}B_{1-2x}^{4+}B_{2x}^{3+}V_{ax}O_{3-x}^{2-}$ gdzie V_a — wakancja anionowa

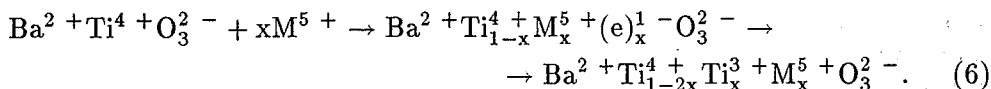
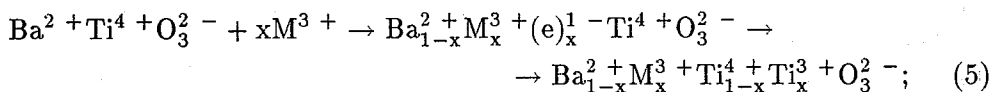
Schemat procesów prowadzących do zwiększenia koncentracji nośników ładunku elektrycznego przedstawiono w tabeli 1. Za podstawę rozpatrzenia tych procesów przyjęto związek chemiczny o ogólnym wzorze $A^{2+}B^{4+}O_3^{2-}$ o strukturze typu perowskitu. W tabeli uwzględniono stosunki wartościowości (N_A, N_B, N_M) i promieni jonowych (R_A, R_B, R_M) kationów zastępowanych (A, B) oraz zastępujących (M). Schematy reakcji dotyczą tylko jednego przypadku zmiennych wartościowości, a mianowicie kationów B^{4+} i M^{5+} . Dzięki zmiennej wartościowości kationów B^{4+} i M^{5+} spełniony jest jeden z głównych warunków stabilności struktur typu perowskitu, a mianowicie warunek obojętności elektrycznej [1, 4].

W charakterze materiału badań również w tym przypadku wybraliśmy $BaTiO_3$ i $SrTiO_3$. Podwyższenie wartości przewodnictwa elektrycznego $BaTiO_3$ o 9–11 rzędów wielkości metodą wprowadzenia odpowiednich modyfikatorów było po raz pierwszy uzyskane w pracy [17]. Przeprowadzone przez nas badania dotyczące wpływu różnorodnych modyfikatorów na przewodnictwo elektryczne $BaTiO_3$ wykazały, że w przypadku zastępowania jonów Ba^{2+} najbardziej efektywny wpływ na σ mają jony Y^{3+} , La^{3+} , Sm^{3+} i inne lantanowce, skłonne do koordynacji o liczbie 12, zaś przy zastosowaniu jonów Ti^{4+} najbardziej efektywnym jest jon Sb^{5+} , skłonny do koordynacji o liczbie 6.

Naruszenie neutralności elektrycznej struktur $BaTiO_3$ lub $SrTiO_3$, w czasie wprowadzenia do nich małej ilości kationów o wyższej wartościowości przy nieobecności jakichkolwiek wakancji, jest kompensowane w wyniku powstawania tej samej liczby jonów Ti^{3+} . Duże przewodnictwo elektryczne w tym przypadku jest

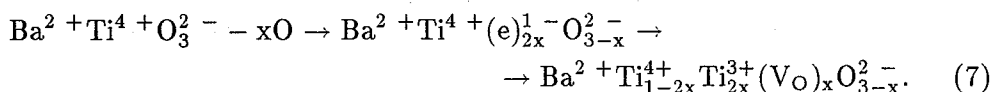
uwarunkowane wymianą elektronów pomiędzy Ti^{4+} i Ti^{3+} . Analogiczny mechanizm przywrócenia neutralności elektrycznej ma miejsce również w przypadku redukcji tlenowej.

Zgodnie z tabelą 1 dla domieszkowanego $BaTiO_3$ można zapisać następujące reakcje przejść do stanu obojętnego elektrycznie:

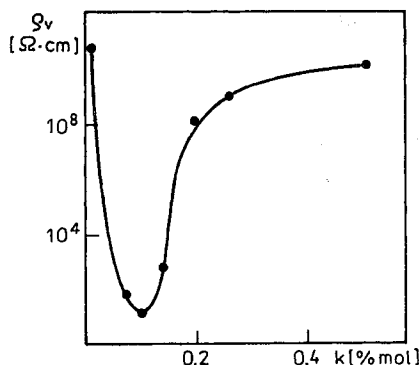


Reakcje takie mogą mieć miejsce jeśli promienie jonowe domieszek M^{3+} i M^{5+} są odpowiednio bliskie promieniom jonowym Ba^{2+} i Ti^{4+} .

W przypadku redukcji $BaTiO_3$ proces neutralizacji elektrycznej przebiega następująco:



Analogiczne równania (5)–(7) można zapisać dla $SrTiO_3$.



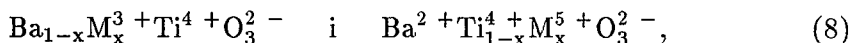
Rys. 1. Zależność elektrycznej oporności właściwej (ρ_v) ceramiek $BaTiO_3$ od koncentracji (k) modyfikatora Sm_2O_3

Przewodnictwo elektryczne modyfikowanego $BaTiO_3$ w sposób istotny zależy od koncentracji domieszki i osiąga maksimum dla 0.1–0.3 % at. Ze wzrostem koncentracji jonów domieszki (od zera do wyżej podanej wartości) w strukturze $BaTiO_3$ nie powstają wakanecje w żadnej podsioci, a liczba nośników ładunku elektrycznego zwiększa się. Widoczne to jest na przykładzie ceramicznego

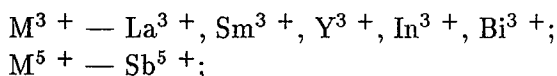
BaTiO_3 domieszkowanego Sm (rys. 1). Przy dalszym wzroście koncentracji modyfikatorów zwiększa się prawdopodobieństwo formowania roztworu stałego, w którym atomy M podstawione są w miejsce Ba. W takim przypadku obojętność elektryczna struktury zachowana jest bez zmiany wartościowości przemiennych jonów Ti; stąd też nie obserwuje się wyraźnego wzrostu przewodnictwa elektrycznego, które zachowuje wartość charakterystyczną dla nie modyfikowanego BaTiO_3 . Przy wzroście koncentracji modyfikatorów sfery oddziaływania centrów domieszkowych przekrywają się i następuje przejście od kompensacji elektronowej do jonowej samokompensacji. I tak np. w BaTiO_3 jony lantanu wywołują powstawanie wakancji Ba bez podwyższenia koncentracji swobodnych nośników ładunku elektrycznego. Przejście do kompensacji jonowej zależy od warunków powstawania wakancji w konkretnych F-TSP. Na przykład w F-TSP zawierających ołów przy modyfikowaniu ich lantanem lub niobem ułatwiony jest proces powstawania wakancji V_{Pb} , a tym samym wzrasta możliwość kompensacji jonowej. Oznacza to, że domieszki takie nie wywołują wzrostu przewodnictwa elektrycznego PbTiO_3 . W BaTiO_3 i SrTiO_3 , gdzie proces powstawania wakancji V_A jest utrudniony, wzrasta możliwość kompensacji elektronowej. Stąd też domieszkowanie tych F-TSP niobem (do 0.5 % at.) pozwala zwiększyć pierwotne przewodnictwo elektryczne o 9–10 rzędów wielkości.

3. EKSPERYMENTALNA WERYFIKACJA MODELU TEORETYCZNEGO

W niniejszej pracy dla otrzymania półprzewodnikowego BaTiO_3 stosowano domieszki o koncentracji niższej od 0.5 % at. Do badań przygotowano ceramiczne materiały o następujących składach chemicznych:



gdzie:



x — określa koncentrację kationów M (wartości podano w tabeli 2 i 3).

Przygotowano dwie serie materiałów: o składzie stechiometrycznym i z wakancjami V_{Ba} . Analogiczne próbki były otrzymane na bazie SrTiO_3 modyfikowanym Y^{3+} i Sm^{3+} . Synteza mieszaniny składników wyjściowych z modyfikatorami prowadzona była metodą reakcji w fazie stałej w temperaturze 1423 K, w czasie 5 h, w atmosferze powietrza. Badania struktury metodą mikroanalizy rentgenowskiej wykazały we wszystkich przypadkach, że otrzymane materiały są monofazowe i należą do rodziny perowskitów. Do syntezy użyto materiałów o następującym stopniu czystości: BaCO_3 — chemicznie czysty, SrTiO_3 — czysty do analizy, TiO_2 — spektralnie czysty, pozostałe — czyste.

Tablica 2

**Elektryczny opór właściwy modyfikowanej ceramiki BaTiO₃
w $T = 300$ K (w zakresie koncentracji nie naruszających
elektrycznej obojętności struktury)**

Modyfikator		ρ_v [Ωcm] przed TPO	Warunki TPO			ρ_v [Ωcm] po TPO
Pierwiastek	k [% at.]		T [K]	t [h]	p [Pa]	
—	0	4×10^{10}	1200	4	1.73	1.09
La	0.067	1.09×10^9	1473	4	1.73	1.20
	0.13	2.14×10^{10}	1473	4	1.73	2.03
	0.20	7.72×10^3	1473	4	1.73	1.20
	0.27	5.09×10^9	1473	4	1.73	0.97
	0.33	1.02×10^{10}	1473	4	1.73	2.30
Sm	0.067	2.9×10^5	1473	1	1.47	2.60
	0.13	1.5×10^4	1473	1	1.47	2.70
	0.20	3.0×10^{10}	1473	1	1.47	5.80
	0.33	7.12×10^9	1473	1	1.47	2.90
Y	0.067	1.46×10^{10}	1548	4	1.87	1.00
	0.13	2.41×10^{11}	1548	4	1.87	0.74
	0.20	1.90×10^4	1548	4	1.87	1.00
	0.27	7.34×10^3	1548	4	1.87	0.44
	0.33	5.05×10^3	1548	4	1.87	0.55
In	0.067	5.22×10^9	1473	4	1.73	1.10
	0.13	3.72×10^9	1473	4	1.73	3.70
	0.20	8.67×10^9	1473	4	1.73	3.10
	0.27	1.11×10^{10}	1473	4	1.73	2.30
	0.33	4.13×10^9	1473	4	1.73	5.40
Bi	0.13	1.01×10^9	1473	4	1.60	2.30
	0.20	4.23×10^{10}	1473	4	1.60	0.91
	0.27	6.17×10^{10}	1473	4	1.60	1.50
	0.33	1.12×10^{11}	1473	4	1.60	0.91
Sb	0.08	1.65×10^{10}	1473	4	1.60	2.70
	0.16	1.87×10^4	1473	4	1.60	1.10
	0.24	6.15×10^5	1473	4	1.60	2.70
	0.32	1.55×10^{10}	1473	4	1.60	0.96

Mielenie zsyntetyzowanego materiału prowadzono w kwarcowym moździerzu w środowisku spirytusu etylowego. Formowanie wyprasek w formie dysków ($2r = 10$ mm; $h = 2$ mm) odbywało się pod wpływem dwustronnego, jednoosiowego prasowania (do proszku wprowadzono nieznaczne ilości spirytusu poliwinilowego). Spiekanie prowadzono w temperaturze 1673 K, w atmosferze powietrza, utrzymując wsad w maksymalnej temperaturze przez 2 godziny. Badanie mikroskopowe powierzchni i przełomów wykazały, że otrzymane próbki mają gęsto upakowaną mikrostrukturę z rozmiarami ziaren do 10 μm .

Tablica 3

**Elektryczny opór właściwy ceramiki BaTiO modyfikowanej itrem
w $T = 300$ K (w zakresie koncentracji nie naruszających
elektrycznej obojętności struktury)**

Modyfikator		ϱ_v [Ωcm] przed TPO	Warunki TPO			ϱ_v [Ωcm] po TPO
Pierwiastek	k [% at.]		T [K]	t [h]	p [Pa]	
Y	0.05	5.53×10^8	1508	4	1.60	2.30
	0.10	2.94×10^{10}	1508	4	1.60	0.27
	0.15	1.17×10^5	1508	4	1.60	1.30
	0.20	8.74×10^3	1508	4	1.60	3.00
	0.25	4.34×10^4	1508	4	1.60	1.30
	0.30	3.27×10^4	1508	4	1.60	1.10
	0.35	3.20×10^5	1508	4	1.60	3.50
	0.40	1.94×10^{12}	1508	4	1.60	2.80
	0.45	3.46×10^{11}	1508	4	1.60	2.20
	0.50	3.64×10^{11}	1508	4	1.60	3.60
Y	0.05	5.53×10^8	1473	4	1.87	1.70
	0.10	2.94×10^{10}	1473	4	1.87	0.76
	0.15	1.17×10^5	1473	4	1.87	2.70
	0.20	8.74×10^3	1473	4	1.87	2.30
	0.25	4.34×10^4	1473	4	1.87	1.40
	0.30	3.27×10^4	1473	4	1.87	3.60
	0.35	3.20×10^8	1473	4	1.87	1.05
	0.40	1.94×10^{12}	1473	4	1.87	1.40
	0.45	3.46×10^{11}	1473	4	1.87	1.60
	0.50	3.64×10^{11}	1473	4	1.87	2.40

Celem zwiększenia ilości swobodnych nośników ładunku elektrycznego, a tym samym przewodnictwa elektrycznego otrzymanych próbek ceramicznych, poddawano je procesowi redukcji tlenowej poprzez wygrzewanie w temperaturze 1173 K, w próżni (ciśnienie parcjale tlenu $p_{O_2} = 0.2666$ Pa); w różnych okresach czasu. Ubytki tlenu z próbek określano na podstawie zmiany ich masy po termiczno-próżniowej obróbce. Zmiana masy kontrolowana była metodą termograwimetryczną [18]. Ubytek tlenu w komórce w ułamkach jego masy atomowej można wyrazić następująco:

$$\delta = \frac{\Delta m}{m_1} \cdot \frac{M}{16}, \quad (9)$$

gdzie:

m_1 — wyjściowa masa próbki (przed TPO),

Δm — ubytek masy tlenu w wyniku TPO,

M — masa cząsteczkowa związku F-TSP;

Stopień redukcji tlenu (K) w komórce określano w oparciu o wzór:

$$K = \frac{\Delta m}{m_1} \cdot \frac{M}{48} \cdot 100 \%. \quad (10)$$

Miarą K jest procent termodesorbowanego tlenu obliczony względem jego całkowitej koncentracji w komórce o składzie stechiometrycznym.

Dla domieszkowanego BaTiO_3 otrzymano $\delta = 0.01$ i $K = 0.3 \%$, zaś dla domieszkowanego SrTiO_3 — $\delta = 0.022$ i $K = 0.7 \%$.

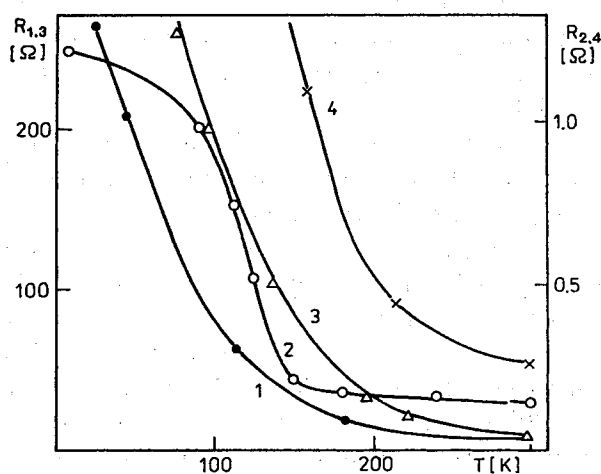
Wprowadzenie donorowych modyfikatorów Sm, Y, La i In nie wpływa w sposób istotny na straty masy BaTiO_3 w wyniku TPO, natomiast modyfikatory Sb i Bi w sposób wyraźny zwiększają ubytek masy (nawet dwa lub więcej razy).

Modyfikowanie SrTiO_3 pierwiastkami donorowymi Y, Tb, Sm zwiększa ubytek masy nieznacznie — maksimum o 30 % (tj. $\delta < 0.013$). Przy zmianie koncentracji (k) donorowych modyfikatorów La^{3+} , Sm^{3+} , Y^{3+} , Sb^{3+} wprowadzonych do BaTiO_3 stwierdzono występowanie minimum funkcji $\rho_v(k)$, gdzie ρ_0 — elektryczny opór właściwy. Rezultat ten potwierdza wyniki badań przedstawionych w pracach [19–20]. Obniżenie elektrycznego oporu właściwego BaTiO_3 dla pewnych koncentracji domieszek wynosi 6–7 rzędów wielkości. Natomiast w przypadku SrTiO_3 modyfikowanego donorami Y, Tb lub Sm, ρ_v obniża się zaledwie o 1 rząd wielkości.

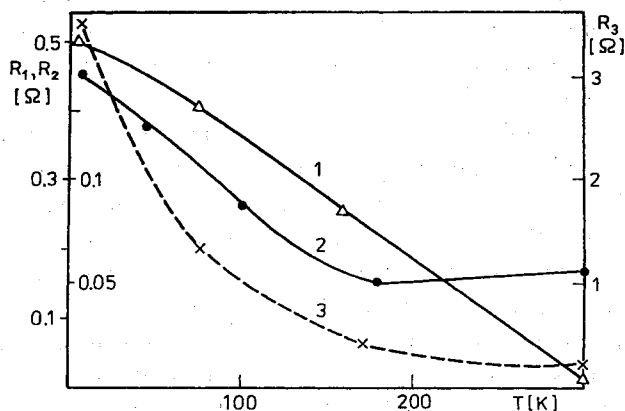
Elektryczny opór właściwy (ρ_v) jest bardziej czułym parametrem przy określaniu deficytu tlenu niż ubytek masy próbki (Δm). Eksperyment wykazał, że TPO przy ciśnieniu 1.333 Pa ($p_{\text{O}_2} = 0.2666$ Pa) i w odpowiednio dobranej temperaturze i czasie wyprężania, opór elektryczny czystych i domieszkowanych BaTiO_3 i SrTiO_3 można zmniejszyć o 10–11 rzędów wielkości. Oznacza to, że wytworzenie deficytu tlenu jest efektywnym sposobem obniżenia oporu elektrycznego ceramicznych materiałów typu F-TSP.

Przygotowując próbki ceramiczne do pomiaru elektrycznego oporu właściwego nanoszono na nie metodą rozpylania katodowego elektrody srebrne lub platynowe. W tabeli 2 przedstawiono wartości elektrycznego oporu właściwego (ρ_v) ceramicznych próbek BaTiO_3 modyfikowanych różnymi domieszkami o koncentracji nie naruszającej elektrycznej obojętności struktury oraz poddanych TPO. Natomiast w tabeli 3 podane są rezultaty badań ρ_v próbek BaTiO_3 modyfikowanych itrem o koncentracji naruszającej bilans wartościowości. Badania nad zjawiskiem Seebecka, polegające na pomiarze siły termoelektrycznej przy określonym gradiencie temperatury między elektrodami próbki, wykazały, że w temperaturze pokojowej nośnikami ładunku w polikrystalicznym BaTiO_3 są elektrony (tak jak w półprzewodnikach donorowych typu n , lub w metalach).

W lewej części tabel 2 i 3 jest widoczne, że dla każdego modyfikatora przy pewnej określonej jego koncentracji (k) ma miejsce obniżenie wartości ρ_v o 6–7 rzędów wielkości w porównaniu z elektryczną opornością właściwą próbek wyjściowych ($k = 0$). W rezultacie poddanie modyfikowanych próbek obróbce termopróżniowej (PTO) obserwuje się dalsze obniżenie wartości elektrycznego oporu właściwego do wartości $\rho_v = 1\text{--}10 \Omega\text{cm}$. Tak więc w wyniku procesu modyfikowania, a następnie procesu PTO można obniżyć opór elektryczny polikrystalicznych F-TSP nawet o 10–11 rzędów wielkości.

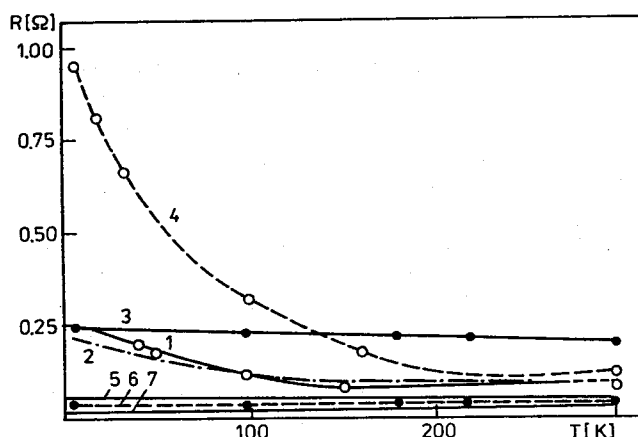


Rys. 2. Zależność oporu elektrycznego (R) ceramiek BaTiO_3 domieszkowanych indem (1,3) i antymonem (2,4) po termiczno-próżniowej obróbce ($T = 1473 \text{ K}$, $t = 4 \text{ h}$, $p = 1.33 \text{ Pa}$ (dla In) i $p = 1.60 \text{ Pa}$ (dla Sb)) od temperatury (T) dla różnych koncentracji In i Sb: 1 — 0.067; 2 — 0.16; 3 — 0.13; 4 — 0.24; [% at.]



Rys. 3. Zależność oporu elektrycznego (R) od temperatury (T) ceramiek BaTiO_3 domieszkowanych lantanem po termiczno-próżniowej obróbce ($T = 1473 \text{ K}$, $t = 4 \text{ h}$, $p = 1.33 \text{ Pa}$) dla różnych koncentracji La: 1 — 0.067; 2 — 0.2; 3 — 0.27 [% at.]

Na rys. 2, 3 i 4 przedstawiono zależności oporu elektrycznego modyfikowanej ceramiki BaTiO_3 od temperatury w zakresie od 300 do 4.2 K. Dla próbek modyfikowanych indem i antymonem obserwuje się przebieg $R(T)$ typowy dla półprzewodników, tzn. w miarę zmniejszania temperatury opór elektryczny ceramiek wzrasta od kilku do 10^3 – $10^4 \Omega$ (rys. 2). W przypadku modyfikowania lantanem (rys. 3) w miarę zmniejszania temperatury opór elektryczny ceramiki wzrasta nie



Rys. 4. Zależność oporu elektrycznego (R) od temperatury (T) ceramiek BaTiO_3 domieszkowanych itrem po termiczno-próżniowej obróbce ($T = 1508 \text{ K}$, $t = 4 \text{ h}$, $p = 1.33 \text{ Pa}$) dla różnych koncentracji Y: 1 — 0.05; 2 — 0.01; 3 — 0.15; 4 — 0.20; 5 — 0.25; 6 — 0.30; 7 — 0.35 [% at.]

więcej jak o rząd wielkości, ma miejsce również niemonotoniczna zmiana $R(T)$. Opór elektryczny ceramiek BaTiO_3 modyfikowanych itrem praktycznie nie wzrasta wraz z obniżeniem temperatury próbek, tzn., że zależność $R(T)$ dla ceramiek staje się zbliżona do charakterystyk typowych dla przewodnictwa metalicznego (rys. 4).

4. PODSUMOWANIE

Heterowalencyjne podstawienia w podsioci A i (lub) B perowskitowych ferroelektryków (ABO_3) oraz redukcja koncentracji tlenu prowadzi do znacznego wzrostu przewodnictwa elektrycznego tych dielektrycznych materiałów. Odnosi się to szczególnie do BaTiO_3 i SrTiO_3 . Wykazano, że najefektywniej można wpływać na przewodnictwo elektryczne BaTiO_3 poprzez podstawienia w miejsce Ba^{2+} jonów Y^{3+} , La^{3+} i Sm^{3+} , a w miejsce Ti^{4+} jonów Sb^{5+} oraz w wyniku redukcji tlenowej wywołanej obróbką termiczno-próżniową. W przypadku SrTiO_3 podstawienia w miejsce Sr^{2+} i (lub) Ti^{4+} innych jonów nie są tak efektywne jak w przypadku BaTiO_3 , lecz w połączeniu z obróbką termiczno-próżniową mogą zwiększyć koncentrację swobodnych nośników ładunku elektrycznego nawet do wartości 10^{19} – 10^{20} cm^{-3} w temperaturze pokojowej. W odpowiednio niskiej temperaturze ($T_c \approx 0.7 \text{ K}$) w takim materiale może mieć miejsce przejście w stan nadprzewodzący. Analogiczna sytuacja ma miejsce również w innych perowskitach: $\text{BaPb}_{1-x}\text{Bi}_x\text{O}_3$ ($T_c \approx 10 \text{ K}$); $\text{Ba}_{1-x}\text{K}_x\text{BiO}_3$ ($T_c \approx 30 \text{ K}$).

BIBLIOGRAFIA

1. E.G. Fesenko: *Semeistvo perovskita i segnetoelektrichestvo*. Atomizdat, Moskva 1972, s. 74.
2. Z. Surowiak: *Dielektryczne i półprzewodnikowe właściwości cienkowarstwowych ferroelektryków*. Uniwersytet Śląski, Katowice 1989, s. 42.
3. Z. Surowiak: *Prace Naukowe Uniwersytetu Śląskiego w Katowicach*, ser. Prace Wydziału Techniki, 1987, t. 24
4. E.G. Fesenko, Z. Surowiak: *Wiadomości Chemiczne*, 1988, t. 42, s. 409
5. Z. Surowiak, E.G. Fesenko: *Chemia Stosowana*. 1986, t. 30, 4, s. 475
6. J. Handerek, Z. Ujma, K. Wójcik: *Zjawiska kinetyczne w ferroelektrycznych i antyferroelektrycznych kryształach o strukturze perowskitu*. Uniwersytet Śląski, Katowice 1979, s. 6
7. E.V. Bursyan: *Nieliniejnij kristall — titanat bariya*. Izd. Nauka, Moskva 1974, s. 48
8. V.M. Fridkin: *Segnetoelektriki — poluprovodniki*. Izd. Nauka, Moskva 1976, s. 18
9. M.E. Lines, A.M. Glass: *Principles and applications of ferroelectrics and related materials*. Clarendon Press Oxford 1977, s. 52
10. Y. Xu: *Ferroelectric Materials and Their Applications*. North-Holland, Amsterdam 1991, s. 140
11. M.A. Saifi, L.E. Cross: *Phys. Rev. B.*, 1970, vol. 2, 3, pp. 677
12. E.R. Pfeifer, J.F. Schocley: *Physics Letters A*, (1969) vol. 29, 10, pp. 589
13. Z. Surowiak, A. Lisińska, G.A. Geguzina, E.G. Fesenko, E.T. Skurayeva: *Ceramika*, 1992, t. 41, s. 249
14. Z. Surowiak, G.A. Geguzina, A. Lisińska, E.G. Fesenko, D. Czekaj, E.T. Skurayeva: *Proc. of "Third Euro-Ceramics"*, Madrid (Spain), 13–17.09.1993. Edit. P. Duran and J.F. Fernandez. Faenza Editrice Iberica, Madrid 1993, p. 1171
15. Z. Surowiak, G.A. Geguzina, A. Lisińska, E.G. Fesenko, D. Czekaj, E.T. Skurayeva: *Polish J. Appl. Chem.*, 1, 12 (1994)
16. J.F. Schocley, W.R. Hosler, M.L. Cohen, *Phys. Rev. Lett.*, 1964, vol. 12, p. 474
17. P.W. Haayman, R. Dam, H.A. Klassens: *Patent RFG No 929350* (23.06.55)
18. H.J. Hagemann, D. Hennings: *J. Amer. Ceram. Soc.*, vol. 64, 10, 590 (1981)
19. B.A. Rotenberg: *Fiz. Tverd. Tela*, t. 7, 10, s. 3048 (1965)
20. O.I. Prokopalo: *Fiz. Tverd. Tela*, t. 21, 10, s. 3073 (1979)

Z. SUROWIAK, E.G. FESENKO, D. KRAWCZYK

THE PEROVSKITE-TYPE CERAMIC
FERROELECTRIC MATERIALS CHARACTERIZED BY SEMICONDUCTIVE
AND METALLIC ELECTRIC CONDUCTION

Summary

The qualitative model of transition of the perovskite-type ceramic ferroelectrics from dielectric state to the semiconductive state and to the state of the electric conduction typical for metals is described. On the basis of the ceramic materials BaTiO_3 and SrTiO_3 the methods of increasing the electric conduction of dielectric perovskites is presented. The physical mechanisms which lead to increasing in free electric charge carriers in those materials were analysed

thoroughly on the microscopic level. The conclusions issuing from the theoretical analysis of the electric transport processes in the model ferroelectric materials (BaTiO_3 , SrTiO_3) have been examined by experimental investigations; a good consistence between the theoretical model and the experimental results has been achieved. The increase in the electric conduction of BaTiO_3 and SrTiO_3 was reached by doping and oxygen reduction in consequence of the thermal-vacuum treatment. The investigations of the electric resistance were performed within the range of temperatures from 4.2 to 300 K.

Układowo-zorientowany model statyczny IGBT

WITOLD PAWELSKI

Instytut Elektroniki, Politechnika Łódzka

Otrzymano 1994.03.10

Autoryzowano do druku 1994.05.30

W artykule przedstawiono nowy model statyczny IGBT o dużej dokładności odwzorowania charakterystyk rzeczywistych, przeznaczony do komputerowej symulacji i projektowania układów elektronicznych. Podano założenia modelu, jego opis analityczny dla stałej i zmiennej temperatury oraz zdefiniowano parametry. Omówiono również właściwości oraz niektóre wyniki badań i weryfikacji doświadczalnej tego modelu.

1. WSTĘP

Aktualnie w dziedzinie modelowania IGBT można dostrzec dwie drogi rozwoju:

- jedno lub wielowymiarowe modelowanie w oparciu o parametry fizyczne i technologiczne, ukierunkowane głównie na modyfikację lub optymalizację struktury wewnętrznej i parametrów tego przyrządu [1], [2], [3],
- modelowanie lub makromodelowanie, w oparciu o parametry identyfikowane metrologicznie na zaciskach przyrządu, przeznaczone głównie do wykorzystania układowego w CAD [4], [5].

Próby adaptacji modeli z pierwszej grupy do celów symulacji i projektowania układów mogą być efektywne tylko wtedy gdy użytkownik dysponuje wymiarami struktury i danymi procesu technologicznego. Dane te nie są jednak najczęściej dostępne dla szerokiego ogółu użytkowników modeli, projektantów układów. W modelowaniu zorientowanym na symulację i projektowanie komputerowe znaczenie ma prostota modelu i łatwość ekstrakcji jego parametrów a także zgodność jego opisu matematycznego z rzeczywistymi relacjami fizycznymi.

Z dotychczasowych doświadczeń w modelowaniu wielowarstwowych przyrządów mocy wynika, że akceptację szerokiego grona użytkowników uzyskały jedynie

proste modele układowo-zorientowane. Najważniejszymi cechami tych modeli powinna być zdolność dobrego odwzorowania rzeczywistych charakterystyk w modelu statycznym oraz przejściowych przebiegów łączeniowych w modelu dynamicznym. Unikalne właściwości IGBT w stanie załączenia są jednymi z najważniejszych atrybutów tego przyrządu, odróżniającymi go od innych łączników mocy [3], [6]. Charakterystyki statyczne IGBT właściwe dla stanu załączenia określają straty mocy w stanach statycznych i quasistatycznych jak też zdolność samoograniczania prądów zwarciovych i przeciążeniowych w obwodzie kolektora [7].

W niniejszej pracy przedstawione zostały założenia, opis analityczny oraz właściwości nowego układowo zorientowanego modelu statycznego IGBT o dużej dokładności odwzorowania charakterystyk rzeczywistych.

2. ZAŁOŻENIA I KRYTERIA KOMPOZYCJI MODELU

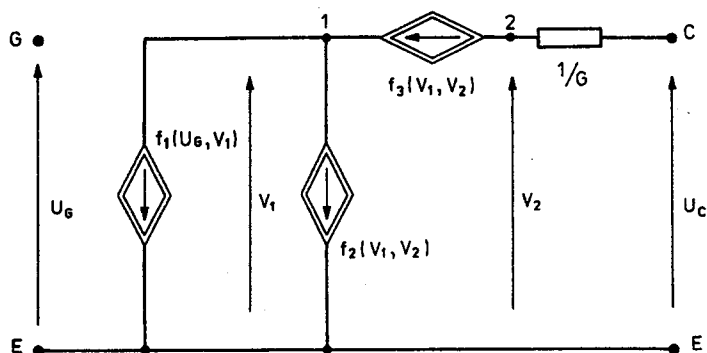
Proponowany model jest hierarchicznie globalnym modelem statycznym IGBT. Pozwala na modelowanie charakterystyk w obszarze bezpiecznej pracy (FSOA) zarówno tranzystorów rodzaju NPT jak i PT. Nie zakłada się zatem modelowania zjawiska latch-up oraz przebicia lawinowego lub skrośnego w strukturze wewnętrznej. Model nie obejmuje również zjawisk występujących w zakresie ujemnych napięć kolektor-emiter (RSOA), które są zróżnicowane dla struktur NPT i PT oraz ich konfiguracji z diodą wstecznie przewodzącą najczęściej zintegrowaną w jednym module z IGBT.

Podstawowymi założeniami proponowanego modelu są:

- 1 — wyrażenie zjawisk statycznych w reprezentacji układem o parametrach skupionych
- 2 — przyjęcie układu zastępczego oraz jego opisu analitycznego w formie maksymalnie uproszczonej przy zachowaniu warunku maksimum dokładności odwzorowania rzeczywistych elektrycznych relacji zaciskowych modelowanego przyrządu
- 3 — przyjęcie minimalnej liczby parametrów modelu przy zachowaniu warunku z p. 2
- 4 — przyjęcie jednolitej procedury identyfikacji wartości parametrów modelu drogą optymalizacji z wykorzystaniem minimalnej liczby pomiarów elektrycznych wielkości zaciskowych modelowanego przyrządu w określonych punktach pracy
- 5 — uwzględniając złożoność zjawisk wielowymiarowych w cztero- lub pięciowarstwowej strukturze IGBT, dla spełnienia warunku z p. 2 dopuszczalne są parametry korekcyjne, nie mające bezpośrednich relacji do parametrów fizycznych
- 6 — przyjęcie zerowej konduktancji wejściowej na wrotach bramka-emiter.

Przyjęciu założenia w/g p. 1 sprzyja struktura wewnętrzna IGBT, zawierająca kilkaset tysięcy do kilku milionów heksagonalnych lub oktagonalnych komórek połączonych równolegle. Powoduje to między innymi odpowiednie unormowanie rozplywu wewnętrznych prądów elektronowych i dziurowych we wspólnej słabo domieszkowanej warstwie n^- .

Generalna koncepcja kompozycji układowej modelu wykorzystuje „naturalną” konfigurację układu pseudo-Darlingtona połączenia tranzystora MOSFET mocy i tranzystora bipolarnego (BJT) mocy. Jako kontakty połączeń wirtualnych baz BJT i drenów komórek MOSFET przyjmuje się za Hefnerem [1] punkty zlokalizowane na granicy obszaru zubożonego i neutralnego w warstwie n^- w pobliżu złączy kolektorowych $n^- - p^+$.



Rys. 1. Nieliniowy model IGBT

Na rys. 1 przedstawiony został układ zastępczy modelu, którego submodele BJT i MOSFET mogą być opisane analitycznymi relacjami prądowych źródeł sterowanych do odpowiednich napięć węzłowych. Układ zawiera ponadto konduktancję G , będącą optymalizowanym parametrem modelu. Reprezentuje ona rozproszone konduktancje obu stronnych obszarów przy złączu emiterowym BJT oraz częściowo obszarów quasineutralnych warstwy n^- . Parametr ten, podobnie jak w modelach diod mocy, pozwala na uwzględnienie niektórych efektów wynikających z przepływu przez to złącze znacznych prądów kolektora.

Układ z rys. 1 jest pobudzany dwoma niezależnymi napięciami wrotowymi definiowanymi przez wektor

$$U = [U_G \cdot U_C]^T, \quad (1)$$

gdzie U_G — napięcie bramka-emiter określane dalej jako napięcie bramki

U_C — napięcie kolektor-emiter (napięcie kolektora)

Układ może być opisany nieliniowymi równaniami węzłowymi o postaci

$$\begin{aligned} f_1(U_G, V_1) + f_2(V_1, V_2) - f_3(V_1, V_2) &= 0 \\ f_3(V_1, V_2) + GV_2 - GU_C &= 0, \end{aligned} \quad (2)$$

gdzie funkcja $f_1(U_G, V_1)$ opisuje zależność między prądem drenu, napięciem bramki i napięciem dren-źródło tranzystora MOSFET

funkcje $f_2(V_1, V_2)$ i $f_3(V_1, V_2)$ opisują zależności napięciowe prądu kolektora i emitera BJT

G — parametr modelu.

Rozwiązanie równań (2) może być dokonane za pomocą wielowymiarowego algorytmu Newtona-Raphsona. Dla każdej pary pobudzeń określonych wartością składowych wektora (1) można w wyniku otrzymać wartość prądu kolektora $I_c(U_G, U_C)$, których zbiór wyznacza charakterystyki statyczne IGBT. Warto w tym miejscu zwrócić uwagę na znane [9] trudności w iteracji N-R równań nieliniowych typu diodowego, co wiąże się z koniecznością zastosowania specjalnych strategii w doborze punktów startowych tych iteracji. Argumentami funkcji w zależnościach (2) oprócz wymuszeń i potencjałów węzłowych są parametry modelu, których liczba i wartości zależą od analitycznego opisu submodeli. Podstawą do ekstrakcji głównych parametrów modelu jest metrologiczne określenie prądów kolektora $I_{cm}(U_{Gm}, U_{Cm})$ w kilku wybranych punktach pracy na płaszczyźnie charakterystyk rzeczywistego modelowanego przyrządu. Liczba pomiarów dla jednej temperatury nie musi przy tym przekraczać liczby głównych parametrów. Ze względu na zmiany temperatury struktury wewnętrznej mierzonego obiektu pod wpływem średnich i chwilowych strat mocy przy przepływie dużych prądów, pomiary muszą wykorzystywać metody impulsowe.

Wartości większości parametrów modelu są identyfikowane w procesie optymalizacji globalnej. Minimalizowana jest przy tym funkcja celu określona przez różnicę między wartością pomierzoną I_{cmi} i obliczoną I_{ci} prądu kolektora, dla identycznych wartości składowych wektora (1) przy uzmiennieniu m -wymiarowego wektora parametrów

$$p = [p_1 \dots p_m]^T. \quad (3)$$

W niniejszej pracy przyjęto minimalizację funkcji celu w postaci

$$\min e(p_1 \dots p_m) = \min \sum_{i=1}^m \left(\frac{I_{cmi} - I_{ci}}{I_{cmi}} \right)^2. \quad (4)$$

Zbiór optymalnych parametrów może być uzyskany przy zastosowaniu jednej ze znanych procedur optymalizacyjnych.

Wybór punktów pracy (U_{Gi}, U_{Ci}) dla pomiarów i optymalizacji dokonywany jest arbitralnie. Ma on dość istotny wpływ na końcową dokładność odwzorowania. Zagadnienie to jest przedmiotem oddzielnej publikacji [8].

3. SUBMODELE I PARAMETRY

Jednym z podstawowych warunków uzyskania dużej dokładności modelu IGBT w/g koncepcji omówionej w p. 2. jest trafność przyjęcia analitycznego opisu submodeli tranzystorów BJT i MOSFET, wyrażonego funkcjami $f_{1..3}$ w (2). Zagadnienie to było przedmiotem poprzednich badań autora, a ich wyniki opisane w [5] dają się streścić następującymi konkluzjami:

1 — w odniesieniu do submodelu MOSFET

— popularny model Shichmana-Hodgesa (wbudowany w programie PSPICE) umożliwia relatywnie dobre odwzorowanie charakterystyk IGBT w obszarze pentodowym⁽¹⁾, natomiast w ważnym aplikacyjnie obszarze triodowym jest nieadekwatny. Powoduje to że średnie błędy odwzorowania przekraczają 20%, co potwierdziły badania innych autorów [4]. Uwzględniając dodatkowe błędy w metrologicznej ekstrakcji parametrów, wybór tego modelu należy uznać za nieodpowiedni

2 — w odniesieniu do submodelu BJT

— zastosowanie modelu Gummela-Poona nie zapewnia satysfakcjonujących rezultatów ze względu na konieczność decydowania o wyborze i wartościach dużej liczby parametrów. Jednocześnie możliwości tego modelu w zakresie opisu nieliniowego wzmocnienia prądowego i efektu Earlyego nie mogą być racjonalnie wykorzystane w modelowaniu IGBT — wykorzystanie modeli Beaufoya-Sparkesa lub Ebersa-Molla z linowymi współczynnikami wzmocnienia prądowego jest dogodnym rozwiązaniem z punktu widzenia dokładności i liczby parametrów modelu IGBT.

Uwzględniając powyższe konkluzje BJT może być reprezentowany przez zmodyfikowany model E-M ze stałym wzmocnieniem prądowym. W świetle określonego w p.1. zakresu stosowności modelu, w konfiguracji dwutranzystorowej z Rys. 1. może on pracować wyłącznie w obszarze aktywnym, dlatego współczynniki inwersyjne mogą być pominięte. Funkcje f_2 i f_3 w zależnościach (2) przyjmują zatem postać

$$f_2(V_1, V_2) = I_s \left[\exp \left(\frac{V_2 - V_1}{n\varphi_T} \right) - 1 \right] + I_{CEO} \quad (5)$$

$$f_3(V_1, V_2) = I_s (1 + \beta^{-1}) \left[\exp \left(\frac{V_2 - V_1}{n\varphi_T} \right) - 1 \right] + I_{CEO}, \quad (6)$$

gdzie I_s — prąd nasycenia złącza emiterowego BJT

φ_T — potencjał termiczny

⁽¹⁾ Terminy „diodowy”, „triodowy” i „pentodowy” określają obszary charakterystyk IGBT (Rys. 2.) i wykluczają konflikty terminologiczne z pojęciami „aktywny”, „liniowy”, „nasycenia” odnoszone do BJT i MOSFET.

β — współczynnik wzmocnienia prądowego

n — współczynnik emisji

I_{CEO} — wynikowy prąd zerowy BJT z otwartą bazą.

Wielkości I_s, β, n , i I_{CEO} są parametrami modelu a ich podstawowe właściwości zostały zestawione w tabeli 1.

Tablica 1

Podstawowe właściwości parametrów modelu

parametr	obszar oddziaływania		sposób ekstrakcji				zależność od temperatury
	triiod.	pentod.	pomiar I_c i optymal.	oddzielny pomiar	ekstrakcja w modelu dynamicznym	stała wartość	
I_s	+	+	—	—	—	+	+
β	+	+	—	+	+	—	—
n	+	+	+	—	—	—	—
G	+	+	+	—	—	—	+
k_p	+	+	+	—	—	—	+
U_T	+	+	+	—	—	—	+
θ	+	—	—	—	—	+	—
θ_i	+	+	+	—	—	—	—
k_a	+	—	+	—	—	—	+
k_k	+	—	+	—	—	—	—
λ	—	+	—	+	—	—	—
I_{CEO}	+	+	—	+	—	—	—

Wzmocnienie prądowe β ma związek z fizycznymi zjawiskami w IGBT, określając przybliżoną relację między prądem dziurowym wstrzykiwanym przez emiter BJT do szerokiej bazy n^- i prądem elektronowym kanału indukowanego pod wpływem napięcia bramki U_G . Wartość tego parametru jest istotna w modelowaniu zjawiska „tail current” przy tworzeniu modelu dynamicznego i zwykle mieści się w zakresie od 0.1 do 10. Może być określona metrologicznie na podstawie analizy wielkości uskoju prądu kolektora w przebiegu przejściowym przy wyłączeniu [1].

Prąd nasycenia I_s wpływa decydująco na właściwości i charakterystyki IGBT w zakresie diodowym przy małych wartościach I_c i U_c . Parametr ten ma odniesienie do fizycznych właściwości złącza emiterowego BJT, w postaci przybliżonej relacji, słusznej dla asymetrycznych złączy i niskiego poziomu wstrzykiwania nośników

$$I_s \approx q \frac{n_i^2 A}{N_D} \sqrt{\frac{D_p}{\tau_p}}, \quad (7)$$

gdzie q — ładunek elektronu w C

N_D — gęstość domieszkowania bazy w cm^{-3}

A — aktywna powierzchnia emitera w cm^2

n_i — samoistna gęstość nośników w cm^{-3}

D_p — stała dyfuzji dziur w cm^2s^{-1}

τ_p — czas życia dziur w s.

W stosowanych współcześnie technologiach gęstości prądu nasycenia złącza emiterowego są praktycznie utrzymywane na stałym poziomie [10] o typowej wartości $3\text{e-}13 \text{ Acm}^{-3}$, niezależnie od innych parametrów fizycznych i strukturalnych. Stąd przyjmuje się stałą wartość I_s jako parametru modelu.

Parametr I_{CEO} umożliwia określenie strat mocy w IGBT w stanie wyłączenia, w warunkach $U_C > 0$ i gdy napięcie bramki U_G nie przekracza wartości progowej U_T . W tym stanie prąd elektronowy kanału, równy prądowi bazy BJT, może być uznany za równy zeru a I_{CEO} za maksymalny prąd zerowy tranzystora z otwartą bazą. Parametr ten odzwierciedla wówczas łączny prąd upływu wstecznie spolaryzowanych złączy kolektorowych wielu tysięcy komórek połączonych równolegle, a jego metrologiczna ekstrakcja jest łatwa. Współczynnik emisji n pozwala na częściowe uwzględnienie efektów bipolarnego przepływu prądu emiterowego BJT, i oddziałującej znacząco na kształt charakterystyk w zakresie diodowym i przy dużych wartościach U_G .

Źródło prądowe z Rys. 1., reprezentujące właściwości submodelu MOSFET, jest opisywane funkcją $f_1(U_G, V_1)$ w zależności (2) o postaci

$$f_1(U_G, V_1) = \frac{k_p V_1 (k_k U_{GT} - 0.5 V_1) (1 + U_{GT} \theta_t)}{k_k^2 [1 + \theta (k_k U_{GT} - V_1) k_1 (1 + U_{GT} \theta_t)]} \quad \text{dla } V_1 < k_k U_{GT} \quad (8a)$$

$$k_1 = \exp(-k_a V_1) - \exp(-k_a k_k U_{GT}) \quad (8b)$$

$$f_1(U_G, V_1) = 0.5 k_p U_{GT}^2 [1 + \lambda (V_1 - k_k U_{GT})] (1 + U_{GT} \theta_t) \quad \text{dla } V_1 \geq k_k U_{GT} \quad (8c)$$

$$f_1(U_G, V_1) = 0 \quad \text{dla } U_G \leq U_T \quad \text{lub dla } V_1 \leq 0 \quad (8d)$$

$$U_{GT} = U_G - U_T \quad (8e)$$

$$k_k > 0 \quad (8f)$$

gdzie k_k jest bezwymiarowym parametrem modelu a parametry modelu $\lambda, \theta, \theta_t, k_a$ wyrażane są w V^{-1} , U_T w V i k_p w AV^{-2} .

W przyjętym opisie analitycznym zapewniona została ciągłość funkcji (8a) i (8c), a tym samym identyczne wartości prądów kolektora I_c obustronnie na granicy obszaru triodowego i pentodowego charakterystyk IGBT.

Podstawowe cechy parametrów modelu, wykorzystywanych w zależnościach (8) zostały zestawione w tabeli 1. Niektóre z nich jak np. napięcie progowe bramki U_T , mają jednoznaczną i znaną interpretację fizyczną. Jednym z najistotniejszych parametrów jest k_p , który posiada odniesienie do współczynnika przewodzenia kanału, powiązanego z parametrami fizycznymi i wymiarami struktury w/g relacji

$$k_p \approx \frac{\mu_0}{L} C_{ox} \sum_{i=1}^{k_0} W_i, \quad (9)$$

gdzie μ_0 — ruchliwość powierzchniowa elektronów w kanale w $\text{cm}^2\text{V}^{-1}\text{s}^{-1}$
 C_{0x} — pojemność na jednostkę powierzchni warstwy tlenku w Fcm^{-2}
 L — długość kanału w komórce w cm
 W_i — szerokość czynna kanału w komórce w cm.
 k_0 — liczba komórek

Parametr λ uwzględnia modulację długości kanału i pozwala na modelowania różniczkowej konduktancji wyjściowej IGBT w obszarze pentodowym. Natomiast θ_t jest parametrem rezerwowym, który w razie potrzeby umożliwia korekcję kwadratowej funkcji (8c), opisującej charakterystykę przejściową w obszarze pentodowym.

Zastosowane dodatkowo parametry korekcyjne θ , k_a i k_k mają istotny wpływ na prawidłowe modelowanie właściwości IGBT w obszarze triodowym. W tym zakresie uwzględniają głównie oddziaływanie podłużnego i poprzecznego pole elektrycznego od U_G i U_C na ruchliwość nośników.

4. ZALEŻNOŚCI TEMPERATUROWE MODELU

Unikalną właściwością IGBT, w porównaniu z unipolarnym lub bipolarnym tranzystorem mocy, jest jego relatywnie mała wrażliwość na zmiany temperatury. Jest to wynikiem częściowego i strefowego kompensowania się wpływu zmian termicznych określonych parametrów. Do najważniejszych z nich należy zaliczyć ruchliwość elektronów w warstwie inwersyjnej kanału i w warstwie n^- , czas życia nośników, szerokość pasma zabronionego w krzemie, potencjał Fermiego [11], [12]. Spektakularny efekt tej strefowej autokompensacji termicznej widoczny jest w zmianach napięcia $U_C(T)$ w stanie załączenia IGBT przy $U_G = \text{const.}$ i $I_c = \text{const.}$ W zakresie małych wartości I_c współczynnik termiczny zmian tego napięcia jest ujemny, dla dużych natomiast jest dodatni. Powodem jest zmniejszanie się napięcia progowego U_T , oraz spadku napięcia na złączu emiterowym a także zmniejszanie się ruchliwości nośników ze wzrostem temperatury.

W omawianym modelu, od temperatury uzależniono pięć parametrów, wymienionych w tabeli 1. Zmiany temperaturowe napięcia progowego IGBT są z dobrym przybliżeniem liniowe podobnie jak w przypadku tranzystorów MOSFET mocy, i mogą być opisane jako

$$U_T(T) = U_T(T_0)[1 + \gamma(T - T_0)], \quad (10)$$

gdzie T, T_0 — temperatura aktualna i odniesienia w K
 γ — współczynnik zmian termicznych U_T w K^{-1} .

Współczynnik γ jako parametr modelu może być identyfikowany w drodze dwukrotnego pomiaru napięcia progowego $U_{Tm}(T_1, Q)$, $U_{Tm}(T_2, Q)$ w znacznie różniących się temperaturach T_1, T_2 oraz w identycznym punkcie pracy

$Q(I_{cm}, U_{Gm}) = \text{const.}$ zakresu podprogowego i obliczony wg zależności

$$\gamma = \frac{U_{Tm}(T_2) - U_{Tm}(T_1)}{T_1 - T_2} \quad (11)$$

Zależność parametru k_p od temperatury dana jest w postaci

$$k_p(T) = k_p(T_0) \left(\frac{T}{T_0} \right)^\alpha, \quad (12)$$

gdzie α — współczynnik zmian termicznych parametru k_p .

Parametr termiczny α powinien być identyfikowany indywidualnie w oparciu o procedurę jednowymiarowej optymalizacji z funkcją celu (4) przy wykorzystaniu przynajmniej jednego dodatkowego pomiaru $I_{cm}(T, U_{Cm}, U_{Gm})$ w temperaturze T ekstremalnie różnej względem T_0 i określonym punkcie pracy $Q(U_{Cm}, U_{Gm})$. Ze względu na oddziaływanie obu powyższych współczynników termicznych we wszystkich trzech obszarach charakterystyk, korzystne jest identyfikowanie ich w pierwszej kolejności.

Parametr $G(T)$ oddziałuje na charakterystyki głównie w obszarze triodowym, a dla $\lambda = 0$ wyłącznie w tym obszarze. Zakres oddziaływania parametru $k_a(T)$ jest natomiast zawsze ograniczony do obszaru triodowego. Zależności od temperatury obu tych parametrów mają postać

$$k_a(T) = k_a(T_0)[1 + \alpha_1(T - T_0)] \quad (13)$$

$$G(T) = G(T_0)[1 + \alpha_2(T - T_0)]^{-1}. \quad (14)$$

gdzie α_1, α_2 — współczynniki zmian temperaturowych parametrów k_a i G w K^{-1} .

W zakresie diodowym i triodowym temperatura wpływa również na parametr I_s , który może być modelowany według zależności stosowanej między innymi w programie SPICE [13] o postaci

$$I_s(T) = I_s(T_0) \exp \left(\frac{E_g(T)q(T - T_0)}{nkT} \right) \left(\frac{T}{T_0} \right)^{XTI} \quad (15)$$

$$E_g(T) = 1.16 - 0.000702 \frac{T^2}{T + 1108} \quad (16)$$

gdzie k — stała Boltzmanna w JK^{-1}

XTI — współczynnik temperaturowy prądu I_s

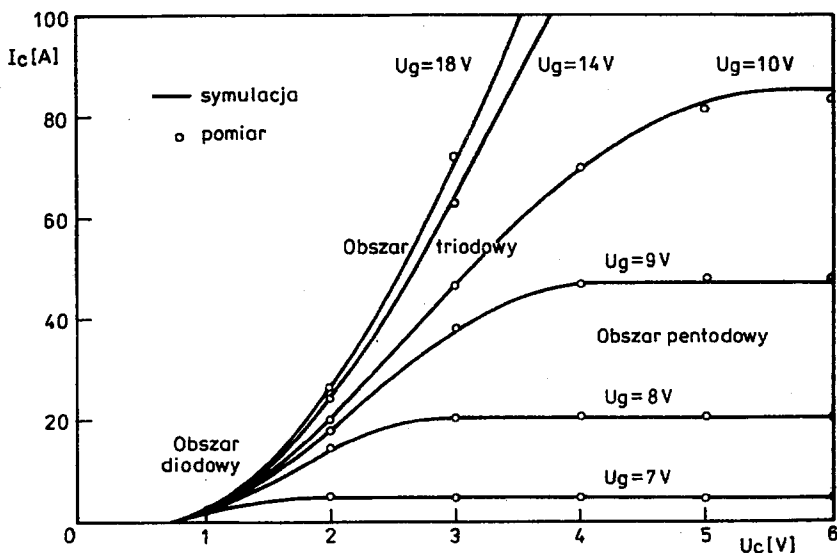
q — ładunek elektronu w C

$E_g(T)$ — wartość przerwy energetycznej w krzemie w eV .

Parametry $k_a(T)$, $G(T)$ i $I_s(T)$ mogą być identyfikowane z pomocą trójwymiarowej procedury identyfikacyjnej. Wielkościami odniesienia funkcji celu (4) powinny być conajmniej trzy wartości prądów kolektora $I_{cm}(T, Q_i)$ modelowanego IGBT, dodatkowo pomierzone w temperaturze T ekstremalnie różnej od T_0 i w określonych punktach pracy $Q_i(U_{Ci}, U_{Gi})$.

5. BADANIA I WERYFIKACJA DOŚWIADCZALNA MODELU

Badania modelu zostały przeprowadzone na przykładzie typowego IGBT rodzaju NPT o nominalnym prądzie kolektora 50 A i napięcia 1000 V. Ekstrakcja parametrów modelu dokonywana była metrologicznie z zastosowaniem impulsowego testera tego przyrządu [14]. Tester ten umożliwiał również weryfikację doświadczalną i ocenę dokładności modelu w zakresie prądów kolektora od 0 do 100 A, napięć bramki od 0 do 20 V, napięć kolektora od 0 do 60 V i temperatur od 25 do 100°C.

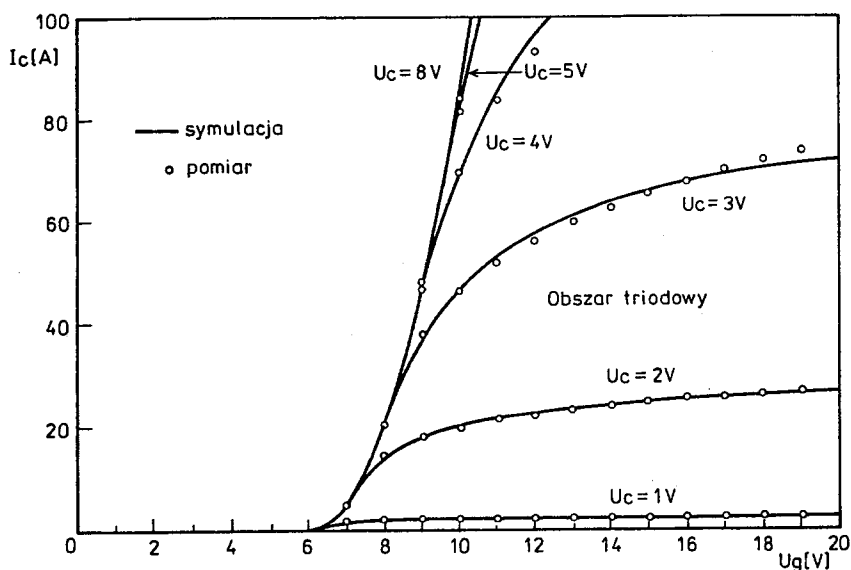


Rys. 2. Porównanie wyników modelowania i pomiarów dla $T = 25^\circ\text{C}$. Typowe charakterystyki wyjściowe IGBT w stanie załączenia z wyróżnionymi obszarami pracy

Wyniki porównania wyjściowych charakterystyk pomierzonych i obliczonych⁽²⁾ dla temperatury 25°C i kilka napięć bramki przedstawiono na rys. 2. Widoczna duża zgodność tych przebiegów, zwłaszcza w ważnym użytkowo obszarze triodowym, jeszcze lepiej może być obserwowana na charakterystyce statycznej przejściowej⁽²⁾ pokazanej na rys. 3. Wartości parametrów modelu, określone według zasad podanych w p. 3., zostały wyszczególnione w tabeli 2. Przyjęcie zerowej wartości pomocniczego parametru θ_t wynika z wystarczającej dokładności odwzorowania charakterystyki przejściowej w obszarze pentodowym badanego IGBT

(2) Przebieg charakterystyki przejściowej, zwłaszcza w obszarze triodowym, ostro uwypukla niedokładności modelowania IGBT w stanie załączenia. Autorzy poprzednich prac z dziedziny modelowania tego przyrządu prezentowali jej przebieg wyłącznie w obszarze pentodowym (np. dla $U_C = 8 \div 12\text{ V}$).

zależnością kwadratową (8c). Zerowa wartość parametru λ jest następstwem ograniczonego zakresu pomiarowego napięcia U_C i dokładności pomiarów, a jej konsekwencją jest przyjęcie zerowej dynamicznej konduktancji wyjściowej przyrządu w obszarze pentodowym.



Rys. 3. Weryfikacja doświadczalna wyników modelowania na tle statycznych charakterystyk przejściowych IGBT dla $T = 25^\circ\text{C}$

Dla oceny dokładności odwzorowania przez model pomierzonych charakterystyk przyjęto kryterium wartości trzech rodzajów błędu średniego, które są definiowane zależnościami

$$\delta = \frac{1}{k} \sum_{i=1}^k \delta_i \quad (17)$$

$$\delta_1 = \frac{1}{k} \sum_{i=1}^k |\delta_i| \quad (18)$$

$$\delta^2 = \frac{1}{k} \sqrt{\sum_{i=1}^k (\delta_i)^2} \quad (19)$$

gdzie

$$\delta_i = \frac{I_{cmi} - I_{ci}}{I_{cmi}} 100$$

k — liczba weryfikowanych punktów pomiarowych.

Tabela 2
Wartości parametrów modelu badanego IGBT

parametr	wartość	jednostka
U_T	6.068	V
k_p	5.499	AV^{-2}
n	1.008	—
k_a	1.350	V^{-1}
k_k	1	—
G	101.0	S
I_s	1E-13	A
β	1	—
θ	1	V^{-1}
θ_i	0	V^{-1}
λ	0	V^{-1}
I_{CEO}	0.8	mA
γ	1.3E-3	K^{-1}
α	-1.65	—
α_1	-4.7E-3	K^{-1}
α_2	2E-3	K^{-1}
XTI	2.4	—

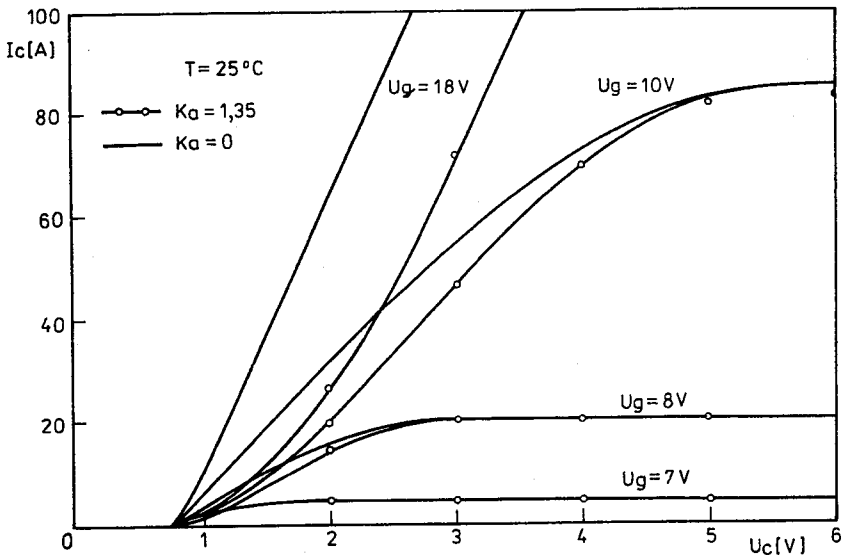
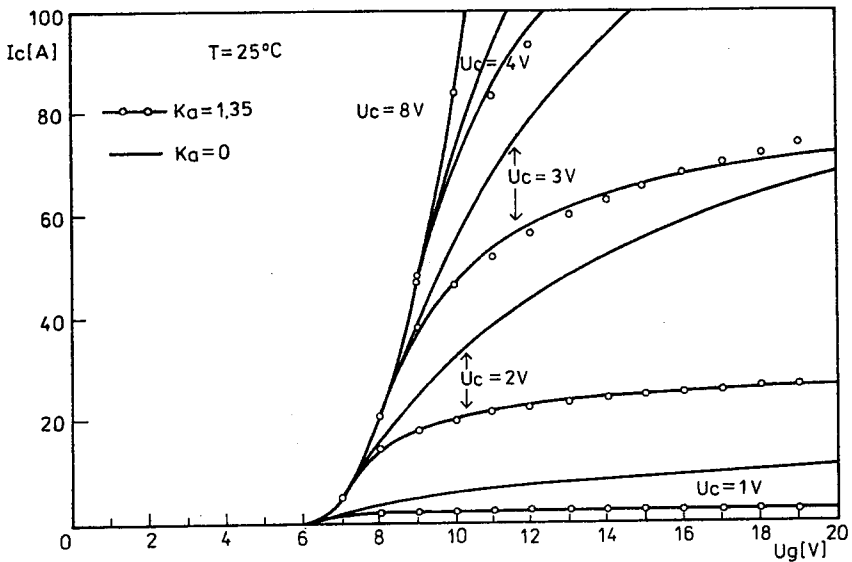
Średni błąd odwzorowania δ , przypadający na jeden punkt pomiarowy, określa stopień asymetrii rozkładu punktów pomiarowych wokół charakterystyki i może przybierać zarówno wartości dodatnie jak i ujemne. Jak wynika z zestawienia zamieszczonego w tabeli 3., maksymalna wartość błędu

Tabela 3
Typowe wartości błędów modelowania

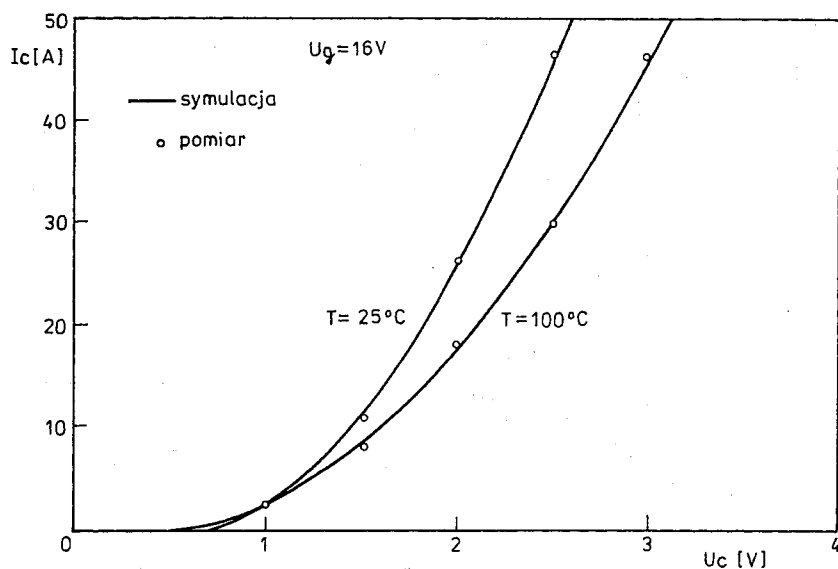
$T[^\circ\text{C}]$	25	50	75	100
δ [%]	-0.3	1.0	1.2	0.7
δ_1 [%]	2.0	2.6	2.5	2.8
δ^2 [%]	0.4	0.4	0.4	0.4
k	61	54	60	63

średniego dla szerokiego zakresu zmian temperatury nie przekracza 3%. Ze względu na ograniczoną dokładność pomiarów $I_{cm}(U_{Gm}, U_{Cm})$, wynik ten jest w pełni satysfakcjonujący, czyniąc nieracjonalnymi dalsze dążenia do poprawy dokładności modelu kosztem zwiększonej liczby parametrów i stopnia złożoności opisu analitycznego. Rezultaty weryfikacji doświadczalnej modelu w oparciu o inne egzemplarze i rodzaje IGBT były podobne.

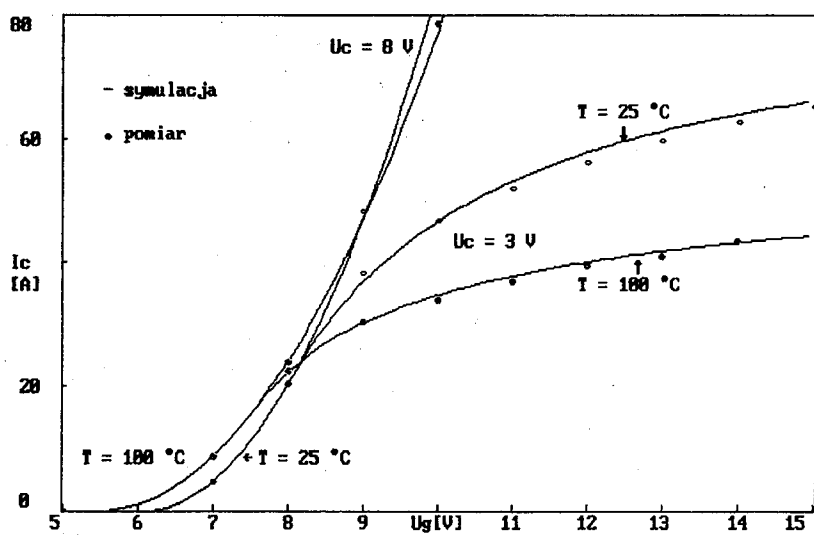
Parametrem modelu, który w największym stopniu przyczynia się do tak dużej dokładności odwzorowania, jest k_a , występujący w zależności (8b). Na rys.4. i na rys. 5. przedstawiono wpływ tego parametru odpowiednio na kształt charakterystyki wyjściowej i przejściowej, w stałej temperaturze. Warto zauważyć,

Rys. 4. Wpływ parametru k_a na kształt charakterystyki wyjściowej IGBTRys. 5. Wpływ parametru k_a na kształt charakterystyki przejściowej IGBT

że założenie zerowej wartości parametru k_a lub θ praktycznie sprowadza submodel MOSFET do postaci modelu S-H, a cały model IGBT do modelu opisanego między innymi w pracy [4]. Na rys. 4. i 5 można zatem obserwować różnice w wynikach symulacji z udziałem obu tych modeli.



Rys. 6. Wpływ temperatury na charakterystykę wyjściową IGBT w obszarze diodowym i triodowym



Rys. 7. Wpływ temperatury na kształt charakterystyki przejściowej IGBT w obszarze triodowym i pentodowym

Niektóre rezultaty badań relacji temperaturowych omawianego modelu zostały przedstawione w tabeli 3. oraz na rys.6. i rys.7., a także w pracy [8]. Potwierdzają one z jednej strony unikalne właściwości termiczne IGBT i zdolność autokompensacji tych oddziaływań w określonych obszarach pracy, z drugiej zaś strony dużą dokładność modelu również w odwzorowywaniu wpływu temperatury na przebieg charakterystyk statycznych.

PODSUMOWANIE

W ostatnich latach IGBT, dzięki swym unikalnym właściwościom i osiąganym parametrom stał się jednym z najatrakcyjniejszych łączników energoelektroniki, stąd wzrost zainteresowania jego modelami do komputerowej symulacji i projektowania układowego. Przy modelowaniu do celów CAD tak złożonych wielowarstwowych przyrządów mocy jak IGBT, istotny jest kompromis między liczbą parametrów i dokładnością odwzorowania rzeczywistych właściwości elektrycznych a wartościami użytkowymi modelu. W niniejszej pracy przedstawiony został po raz pierwszy pełny model statyczny IGBT, który w ekstremalnym stopniu spełnia warunki takiego kompromisu. Do głównych zalet modelu należy duża dokładność odwzorowania z błędami na poziomie kilku procent oraz korzystne do implementacji komputerowej struktura układowa i opis analityczny. Zastosowana metoda ekstrakcji parametrów modelu wykorzystuje, współcześnie popularne, metody optymalizacyjne z minimalną liczbą danych metrologicznych rzeczywistego modelowanego przyrządu. Dodatkową zaletą jest również możliwość modelowania zarówno przyrządów o strukturze PT jak i NPT bez konieczności zmiany opisu analitycznego.

Istotną cechą prezentowanego modelu IGBT, warunkującą jego dokładność i powyższe zalety, jest jego oryginalny opis analityczny z zastosowaniem parametrów korygujących relacje między ruchliwością i gęstością wstrzykiwanych nośników a obu napięciami zaciskowymi. Wprowadzenie tych parametrów pozwoliło na radykalną poprawę odwzorowania charakterystyk statycznych IGBT, zwłaszcza w ważnym użytkowo obszarze trydowym. Model stanowi też dobrą podstawę wyjściową do opracowania modelu dynamicznego tego przyrządu.

BIBLIOGRAFIA

1. A.R. Hefner, D.L. Blackburn: *An Analytical Model for the steady-state and transient characteristics of the Power Insulated-Gate Bipolar Transistor*. Solid-State Electronics, vol. 31, pp. 1513-1532, nr. 10, 1988
2. D.S. Kuo, C. Hu, S.P. Sapp: *An Analytical Model for the Power Bipolar — MOS Transistor*. Solid-State Electronics, vol. 29, pp. 1229-1237, nr. 12, 1986
3. B.J. Baliga: *Modern Power Devices*. J. Wiley & Sons Inc. New York, 1987

4. F.F. Protiva, P. Apeldoorn, N. Groos: *New IGBT Model for PSPICE*. EPE'93 Proceedings, pp. 226–231
5. W. Pawelski: *A New Approach to the Modelling of On-state Characteristics of the IGBT*. Circuit Theory and Electronics Circuits Conference Proceedings, Kolobrzeg'93, pp. 266–271
6. W. Pawelski: *The analysis of the IGBT properties and parameters*. Proceedings of the XII Symposium on Electromagnetic Phenomena in Nonlinear Circuits, October 1991, Poznań
7. H.G. Eckel, L. Sack: *Optimization of the Turn-off Performance of IGBT at Overcurrent and Short-Circuit Current*. EPE'93 Proceedings pp. 317–322
8. W. Pawelski: *Analiza wrażliwości nowego modelu statycznego IGBT*, Kwart. Elektron. i Telekom. 1994, t. 40, z. 2, ss. 265–276
9. M. Tadeusiewicz: *Metody komputerowej analizy stałoprądowej nieliniowych układów elektronicznych*. WNT Warszawa 1991
10. A. Napieralski: *Computer Aided Design of Power Semiconductor Devices with Special Emphasis of Thermal Problems*. DSC Thesis, W.N.P.L., Łódź 1988
11. T.P. Chow, K.C. So, D. Lau: *Performance of 600-V n-Channel IGBT's at Low Temperatures*. IEEE Electron Device Letters, vol. 12, nr. 9, 1991, pp. 498–499
12. B.J. Baliga: *Temperature Behaviour of Insulated Gate Bipolar Transistor Characteristics*. Solid-State Electronics, vol. 28, nr.3, 1985, pp. 289–297
13. J. Porębski, P. Korohoda: *SPICE — program analizy nieliniowej układów elektronicznych*. WNT Warszawa 1992
14. A. Gabryelczyk, M. Grecki, Z. Kuśmierek, A. Napieralski, W. Pawelski: *Tester charakterystyk statycznych IGBT*. Kwart. Elektron. i Telekom. 1994, t. 40, z. 2, ss. 277–290

W. PAWELSKI

CIRCUIT-ORIENTED STATIC MODEL OF THE IGBT

Summary

In the paper a new circuit-oriented model of the IGBT has been presented. Extraction of the model parameters uses an optimization method and requires of a few two-port measurements in a simple pulse tester. Model gives the facilities for excellent fitting of a real IGBT static characteristics in all SOA range.

Analiza wrażliwości nowego modelu statycznego IGBT

WITOLD PAWELSKI

Instytut Elektroniki, Politechnika Łódzka

Otrzymano 1994.03.10

Autoryzowano do druku 1994.05.30

W artykule analizowane są właściwości nowego modelu statycznego IGBT w oparciu o przebiegi wrażliwości funkcji układowej na zmiany podstawowych parametrów. Wyniki analizy zostały wykorzystane do usprawnienia procedury identyfikacji parametrów oraz zwiększenia dokładności odwzorowania charakterystyk rzeczywistych modelowanych tranzystorów.

1. WSTĘP

W artykule [1] został przedstawiony nowy model statyczny IGBT, zorientowany na komputerowe projektowanie i symulację działania układów elektronicznych. Model odznacza się bardzo dużą wiernością odtwarzania rzeczywistych charakterystyk statycznych obu rodzajów (PT i NPT) wytwarzanych obecnie tranzystorów bipolarnych z izolowaną bramką. Umożliwia on użytkownikowi, nie dysponującemu parametrami fizycznymi ani danymi technologiczno-wymiarowymi, na identyfikację parametrów modelu w oparciu o szereg pomiarów zaciskowych wielkości elektrycznych rzeczywistego modelowanego IGBT.

Kilka parametrów modelu oraz kilka współczynników termicznych podlega ekstrakcji w drodze optymalizacji z odniesieniem metrologicznym w funkcji celu (4)[1]. Jednakże dla zachowania dokładności odwzorowania charakterystyk w szerokim zakresie użytkowym, arbitralny wybór punktów pomiarowych tego odniesienie nie może być zupełnie dowolny. Jednym z korzystnych kryteriów wyboru tych punktów mogą być wyniki analizy wrażliwości modelu na zmiany parametrów. Analiza wrażliwościowa, która jest przedmiotem niniejszej pracy pozwala również na lepsze poznanie właściwości modelu.

2. IDENTYFIKACJA PARAMETRÓW A WRAŻLIWOŚĆ MODELU

Jednym z podstawowych założeń niniejszego modelu jest przyjęcie metody skalowania, odniesionego do metrologicznych danych rzeczywistego modelowanego przyrządu. Jak wynika z [1] omawiany model IGBT wymaga określenia dwunastu parametrów, z których pięć zależnych od temperatury, korzysta z pięciu dodatkowych współczynników termicznych. Sposoby ich identyfikacji są jednak zróżnicowane. Niektóre z parametrów jak np. I_s , β i θ mogą być przyjmowane jako stałe przy tworzeniu modelu statycznego. Parametr I_{CE0} najłatwiej określić metrologicznie w warunkach maksymalnej temperatury, bez korzystania z impulsowych metod pomiaru. W przypadku uwzględniania w modelu IGBT wyjściowej różniczkowej konduktancji w obszarze pentodowym [1], koniecznym staje się określenie parametru λ . W tym celu należy jednak korzystać z impulsowych metod pomiarowych, stosując tester podobny do omówionego w pracy [2].

Pozostałe siedem parametrów tj. U_T , k_p , n , k_a , k_k , θ_t , G i pięć współczynników termicznych a mianowicie γ , α , α_1 , α_2 , XTI powinny być określone metodą optymalizacji z funkcją celu, wykorzystującą zgodnie z założeniem odniesienie metrologiczne do modelowanego IGBT. Efektywność procedury optymalizacyjnej, warunki jej zbieżności i czas trwania są tym korzystniejsze im mniejszy jest wymiar przestrzeni optymalizacji czyli liczba jednocześnie poszukiwanych parametrów.

Jednym z celów niniejszej pracy jest zatem wyselekcjonowanie grup parametrów do sekwencyjnej optymalizacji globalnej oraz ustalenie kolejności uruchamiania tych procedur.

Drugim celem jest określenie współrzędnych punktów pracy dla metrologicznego odniesienia funkcji celu w ramach poszczególnych grup optymalizowanych parametrów. W przyjętym bowiem sposobie ekstrakcji właściwe wartości parametrów, a tym samym dokładność modelu w szerokim zakresie charakterystyk, zależą w większym stopniu od miarodajności wyboru obszarów lokalizacji niż od liczby punktów pomiarowych. Powyższe zamierzenia można zrealizować w oparciu o analizę wrażliwości modelu na zmiany parametrów.

Dla uogólnień wyników warto wprowadzić jednostki względne do opisu elektrycznych wielkości zaciskowych modelu. Powinny one być odniesione do wartości w miarę jednoznacznych, charakterystycznych dla modelu i dostępnych dla użytkownika przed etapem skalowania. Dla temperatur proponuje się wartość odniesienia $T_0 = 25^\circ\text{C}$ i jednostkę względną t_w oraz dla prądów kolektora I_c i jego jednostki względnej y_w wartość nominalną I_{cn} przepływu prądu ciągłego, podawaną w danych przyrządu zwykle dla określonej temperatury. Napięcia kolektora U_C i bramki U_G można wyrazić odpowiednio przez jednostki względne x i z , odniesione do pomierzonej umownej wartości napięcia progowego U_{Tm0} . Wartość tą można łatwo określić drogą nieimpulsowego pomiaru przy umownej wartości prądu kolektora np. dla $y_w = 0.001$, w temperaturze odniesienia według definicji

$$U_{Tm_0} = \min(U_G) \Big|_{\substack{y_w=0.001 \\ t_w=1 \\ x>0}} \quad (1)$$

Ponadto celowe wydaje się wykorzystanie, zawsze dostępnej wartości maksymalnej dopuszczalnej napięcia bramki i odpowiadającej jej wartości względnej $z_{\max}$, do wprowadzenia krotności napięcia różnicowego bramki, zdefiniowanego zależnością

$$z_k = \frac{z - 1}{z_{\max} - 1} \quad (2)$$

Jeżeli przez $\mathbf{p}_0$ określi się wektor parametrów modelu podlegających optymalizacji o wartościach optymalnych, to funkcją układową właściwą dla badań wrażliwościowych jest $y_w(\mathbf{p}_0)$. Dla zachowania możliwości analiz porównawczych szczególnie przydatna jest wielokoprzyrostowa względna wrażliwość [3], [4] funkcji układowej $y_w(\mathbf{p})$ względem i -tego parametru modelu p_i w punkcie $\mathbf{p}_0$, zdefiniowana jako

$$S_{p_i}^{y_w(\mathbf{p})} = S(p_i) = \frac{p_i}{y_{w0}(\mathbf{p}_0)} \frac{y_w(\mathbf{p}) - y_w(\mathbf{p}_0)}{p_i - p_{i0}} \quad (3)$$

$$\mathbf{p} = \mathbf{p}_0 + a(p_i - p_{i0}) \quad (4)$$

gdzie a — znormalizowany wskaźnik zmiany parametrów.

Wrażliwość (3) funkcji układowej y_w jest nieodłączną cechą modelu IGBT. Wyniki badań wrażliwości względnej przy różnych wartościach x , z i t_w mogą dostarczać pożytecznych informacji dla prawidłowej ekstrakcji parametrów na etapie skalowania modelu. Główną trudnością wówczas jest określenie zbioru punktów pracy $Q(x, z, t_w)$, dla których powinny być pomierzone wartości odniesienia $y_{wm}(x, z, t_w)$, wchodzące w skład funkcji celu optymalizacji.

Jako kryterium główne wyboru tych punktów może być uznane ich położenie w obszarze występowania maksimum modułu wrażliwości względnej (3) na zmiany i -tego parametru p_i . Minimalną liczbą tych punktów, zapewniającą dostateczną liczbę stopni swobody w procesie optymalizacji globalnej jest liczba optymalizowanych parametrów

$$m = \dim \mathbf{p}_0 \quad (5)$$

Jak widać z zależności analitycznych (8) [1], opisujących omawiany model statyczny IGBT, niektóre parametry oddziałują celowo na funkcję układową wyłącznie w obszarze pentodowym, triodowym lub też w obu tych obszarach, warunkowo bądź bezwarunkowo. Jest to okoliczność sprzyjająca możliwości wyselekcjonowania pewnych podzbiorów tych parametrów do sekwencyjnej optymalizacji globalnej. Może ona być korzystnie przeprowadzona dla temperatury odniesienia ($t_w = 1$) w dwóch etapach ekstrakcji wartości parametrów w/g kolejności

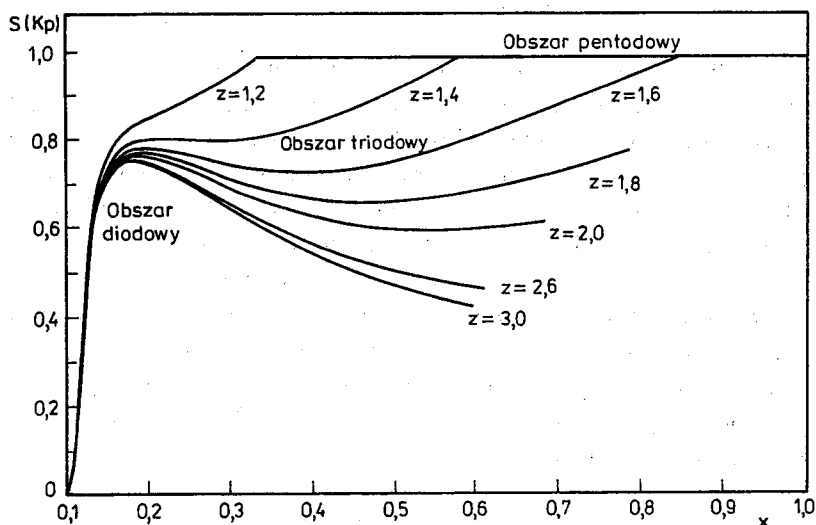
1 — U_T, k_p i θ_t , oddziaływujących we wszystkich trzech obszarach charakterystyk,

2 — k_t, k_{an} i G , oddziaływujących tylko w obszarach diodowym i triodowym dla $\lambda = 0$.

Zadanie identyfikacji zbioru siedmiu optymalizowanych parametrów modelu IGBT może być zatem przeprowadzone w maksymalnie czterowymiarowej podprzestrzeni tych parametrów w dwóch kolejnych procedurach. W dalszej kolejności powinny być identyfikowane wartości współczynników termicznych.

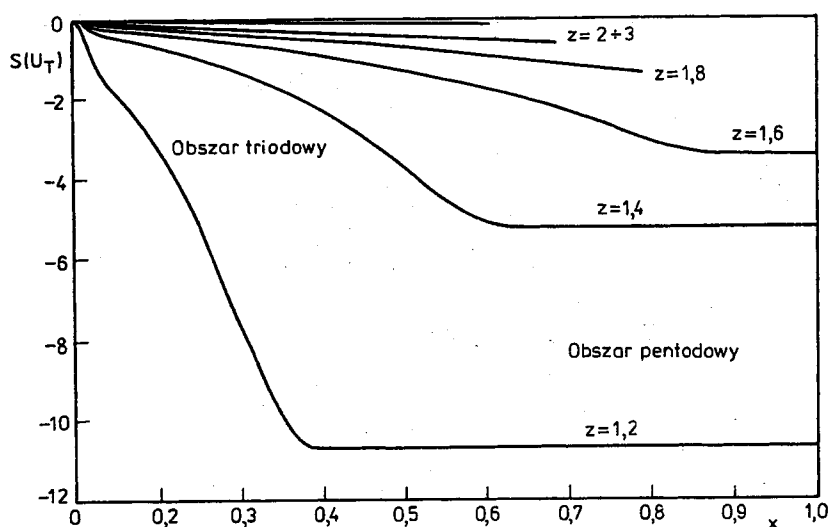
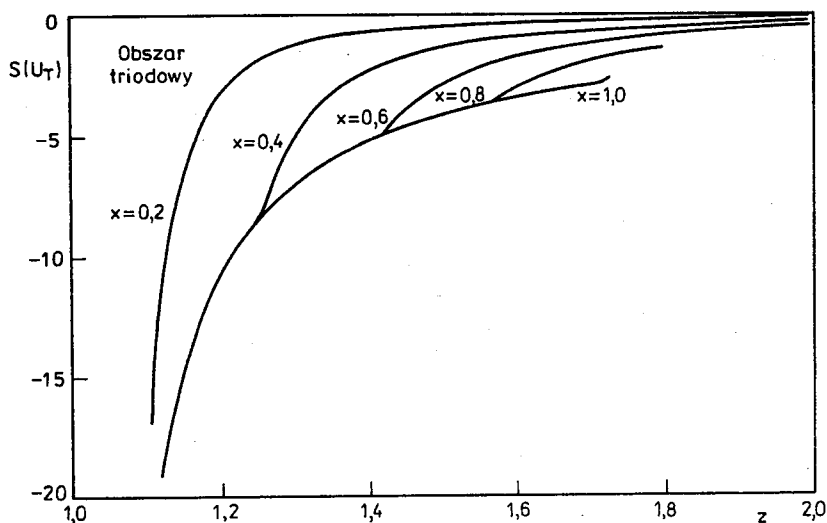
3. WRAŻLIWOŚĆ MODELU W OBSZARZE PENTODOWYM CHARAKTERYSTYK

Analiza wrażliwości modelu zostanie przeprowadzona dla $a = 0.1$, $t_w = 1$ oraz dla zakresu zmienności funkcji układowej y_w ograniczonego w przedziale $< 0, 2 >$. Wartości elementów wektora p_0 są zgodne z przykładowymi danymi zawartymi w tabeli 2 [1].



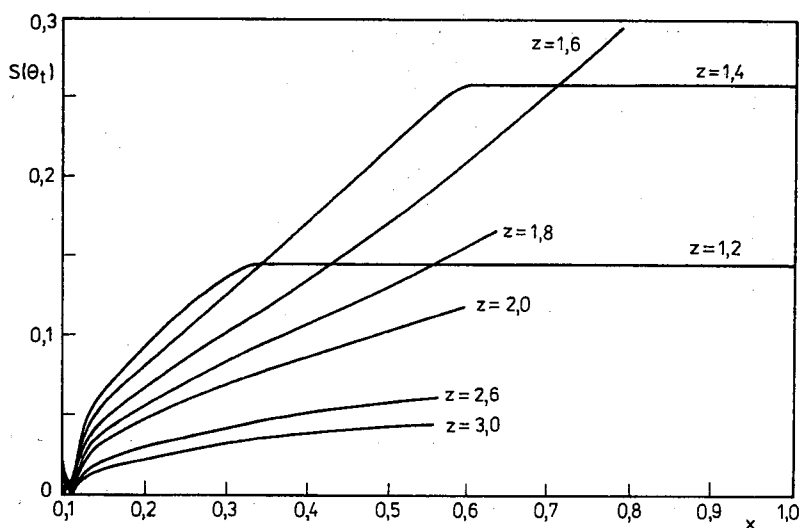
Rys. 1. Przebiegi wrażliwości względnej $S(k_p, x)$, $z = \text{const}$.

Na rys. 1 przedstawiono przebiegi wrażliwości $S(k_p)$ wielkości y_w na zmianę parametru k_p , wyrażone w funkcji x dla szeregu stałych wartości z , której dziedziną odpowiada charakterystyce wyjściowej IGBT. Z przebiegów wynika że moduł wrażliwości modelu osiąga w tych warunkach maksimum-maksimum w obszarze pentodowym dla $x > 0.32$ i dla z ograniczonych od góry maksymalną

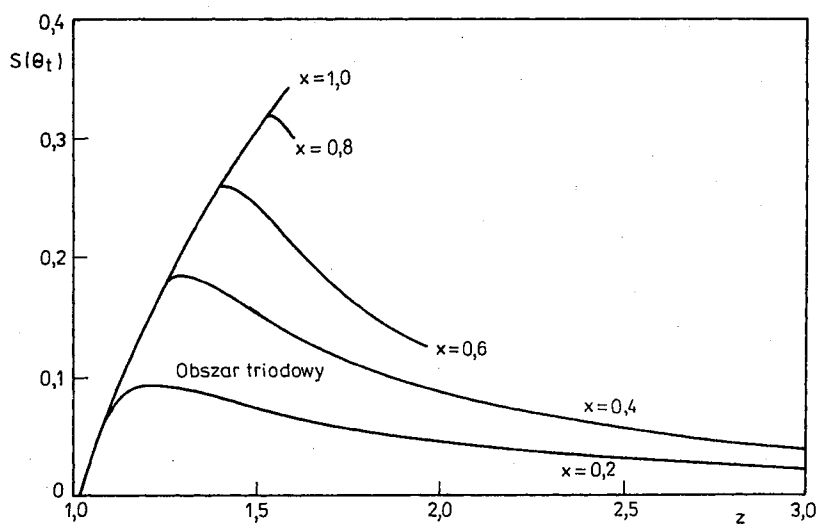
Rys. 2. Przebiegi wrażliwości względnej $S(U_T, x)$, $z = \text{const.}$ Rys. 3. Przebiegi wrażliwości względnej $S(U_T, z)$, $x = \text{const.}$

wartością $y_w = 2$. W obszarze diodowym i triodowym wrażliwość modelu na parametr k_p jest wyraźnie mniejsza.

Przebiegi wrażliwości $S(U_T)$ na zmiany napięcia progowego wyrażone w dziedzinie charakterystyk wyjściowych i przejściowych pokazane zostały odpowiednio na rysunkach 2 i 3. Można na nich zaobserwować bardzo dużą wrażliwość modelu w zakresie pentodowym dla $z = z_{\min} \approx 1$ i $x > 0,1$.



Rys. 4. Przebiegi wrażliwości względnej $S(\theta_t, x)$, $z = \text{const.}$



Rys. 5. Przebiegi wrażliwości względnej $S(\theta_t, z)$, $x = \text{const.}$

Przebiegi wrażliwości $S(\theta_t)$ na parametr korekcyjny θ_t zostały pokazane na rys. 4 i 5 w dziedzinie charakterystyk wyjściowych i przejściowych. Można na nich zaobserwować, relatywnie największe wartości wrażliwości w zakresie pentodowym dla $z_{\max} = z(y_w = 2)$ i $x > z_{\max}$.

Tabela 1

Orientacyjna lokalizacja węzłów optymalizacji parametrów
 $U_T, k_p, \theta_t, \alpha, \gamma$

Typ	$Q(x, z_k, t_w)$		
	x	z_k	t_w
U_T, k_p, θ_t	> 1.3	0.07	1
	> 1.3	0.2	
	> 1.3	$0.28 \div z_k \ (y_w = 2)$	
α, γ	> 1.3	0.07	t_{wmax}
	> 1.3	0.2	

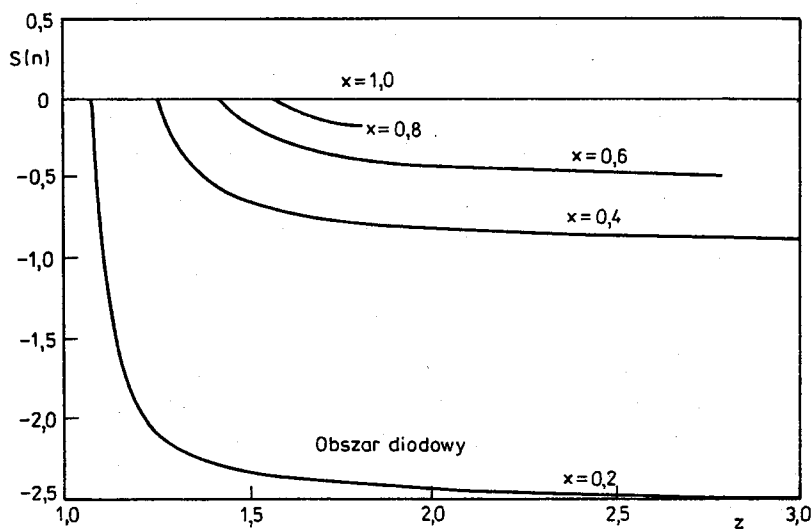
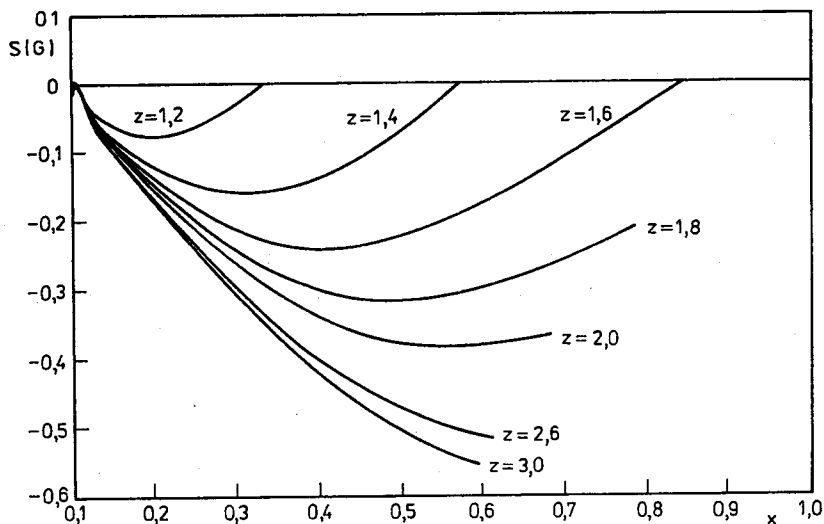
Na podstawie przedstawionych przebiegów wrażliwości modelu na powyższe trzy parametry można określić właściwą lokalizację punktów pracy dla pomiarów i jednocześnie węzłów optymalizacji globalnej. Została ona przedstawiona w tabeli 1. W tabeli tej podano również orientacyjne współrzędne lokalizacji punktów pomiarowych odniesienia w przypadku zastosowania dwuwymiarowej optymalizacji globalnej dla identyfikacji parametrów termicznych α i γ .

4. WRAŻLIWOŚĆ MODELU W OBSZARZE DIODOWYM I TRIODOWYM

Na rys. 6 pokazane zostały przebiegi wrażliwości względnej $S(n)$ funkcji układowej modelu, wyrażonej w dziedzinie charakterystyki przejściowej, na współczynnik emisji n . Jak widać, największe wartości modułu tej wrażliwości występują w obszarze diodowym charakterystyk IGBT już przy stosunkowo małych wartościach względnych x . W obszarze triodowym występuje natomiast maksimum modułu wrażliwości $S(G)$ na parameter G w przebiegu z rys. 7, przedstawionym w funkcji x . Zwraca uwagę niewrażliwość modelu na ten parametr w obszarze pentodowym w warunkach zerowej wartości parametru λ .

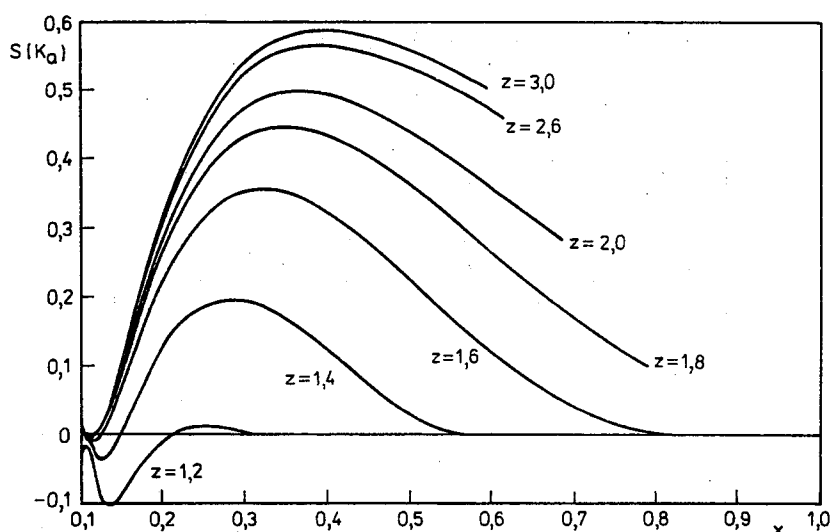
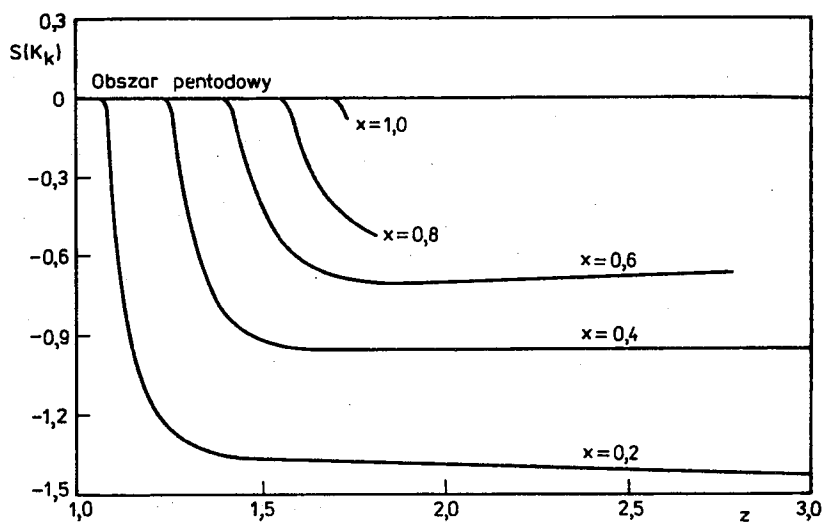
Parametry k_a i k_k decydująco wpływają na prawidłową reprezentację właściwości IGBT w zakresie triodowym charakterystyk. Na rys. 8 pokazane zostały przebiegi wrażliwości $S(k_a)$ w funkcji x przy $z = const$. Widoczne jest zmniejszanie się gradientu modułu tej wrażliwości w miarę wzrostu z przy $z_k \approx (z = z_{max})$ oraz niewrażliwość modelu na ten parametr w obszarze pentodowym. Na rys. 9 pokazane zostały przebiegi wrażliwości $S(k_k)$ modelu na parameter k_k w funkcji z przy $x = const$. Widoczne są na tym rysunku znaczne zmiany wrażliwości przy małych wartościach z i x w obszarze triodowym oraz niewrażliwość modelu na zmiany tego parametru w obszarze pentodowym.

Jak wynika z przedstawionych przebiegów; obszary występowania dużych wartości lub maksimum modułu wrażliwości modelu w zakresie diodowym i triodowym są różne dla różnych parametrów, co stwarza korzystne warunki dla wyboru węzłów optymalizacji n, G, k_a i k_k . W tabeli 2 zestawiono orientacyjne dane dla wyboru punktów pracy optymalizacyjnej identyfikacji tych parametrów oraz

Rys. 6. Przebiegi wrażliwości względnej $S(n, z)$, $x = \text{const.}$ Rys. 7. Przebiegi wrażliwości względnej $S(G, x)$, $z = \text{const.}$

dla ewentualnej trójwymiarowej optymalizacji współczynników termicznych α_1 , α_2 i XTI, którą można przeprowadzić w dalszej kolejności.

Spośród parametrów modelu IGBT przyjmowanych jako stałe istotny wpływ na przebiegi charakterystyk statycznych w obszarze diodowym wywiera wielkość I_s . Na rys. 10 pokazane zostały przebiegi wrażliwości względnej $S(I_s)$ prądu y_w na zmiany tego parametru w dziedzinie charakterystyk przejściowych. W

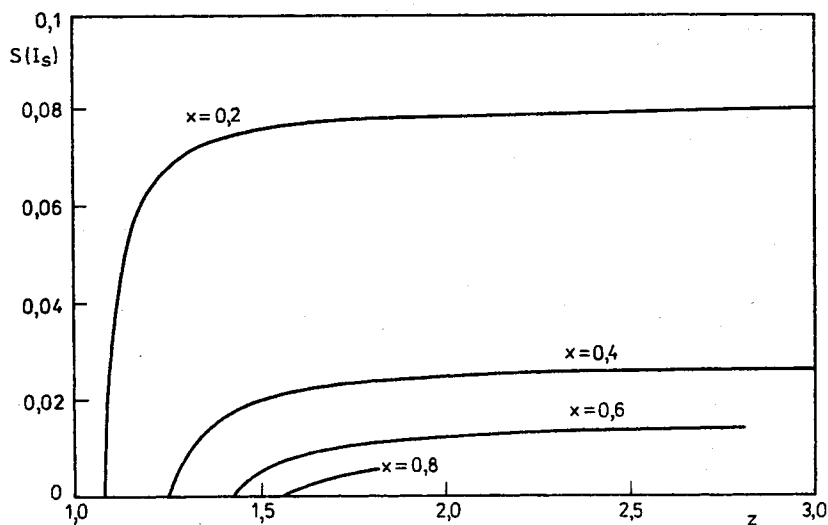
Rys. 8. Przebiegi wrażliwości względnej $S(k_a, x)$, $z = \text{const.}$ Rys. 9. Przebiegi wrażliwości względnej $S(k_k, z)$, $x = \text{const.}$

przebiegach tych znamienne są małe wartości modułu $S(I_s)$ nawet w bardzo wąskim zakresie diodowym charakterystyk, gdzie wykazuje on maksimum. Potwierdza to trafność zaproponowanego w [1] sposobu ustalania wartości tego parametru.

Tabela 2

Orientacyjna lokalizacja węzłów optymalizacji parametrów
 $n, G, k_a, k_k, \alpha_1, \alpha_2, XTI$

Typ	$Q(x, z_k, t_w)$		
	x	z_k	t_w
n, G, k_a, k_k	$0.8 \div 0.9$	$0.7 \div 0.8$	1
	$0.7 \div 0.8$	$0.9 \div 1$	
	$0.1 \div 0.2$	$0.2 \div 1$	
	0.5	$0.4 \div 0.8$	
α_1, α_2, XTI	$0.8 \div 0.9$	$0.7 \div 0.8$	t_{wmax}
	$0.7 \div 0.8$	$0.9 \div 1$	
	$0.1 \div 0.2$	$0.2 \div 1$	



Rys. 10. Przebiegi wrażliwości względnej $S(I_s, z)$, $x = \text{const.}$

PODSUMOWANIE

Nowy model statyczny IGBT, będący przedmiotem niniejszej pracy, posiada szereg zalet, do których należy zaliczyć w pierwszej kolejności dużą dokładność odwzorowania rzeczywistych charakterystyk, małą liczbę parametrów i minimalną liczbę pomiarów koniecznych przy ich ekstrakcji.

Analiza wrażliwości modelu na zmiany głównych parametrów wykazała, że model posiada w określonych warunkach obszary całkowitej niewrażliwości na określone parametry w zakresie pentodowym. Ponadto przebiegi wrażliwości funkcji układowej modelu dla różnych x i z są nader zróżnicowane. Wykazano, że obie te cechy modelu mogą być efektywnie wykorzystane przy identyfikacji jego parametrów.

Identyfikację głównych parametrów można ponadto usprawnić i znacznie przyspieszyć przez jej segmentację na kilka sekwencyjnie dokonywanych procedur optymalizacyjnych o małych wymiarach przestrzeni parametrów. Dla zachowania dużej dokładności modelu przy wyborze węzłów identyfikacji poszczególnych parametrów powinno się uwzględniać obszary maksymalnej wrażliwości modelu.

BIBLIOGRAFIA

1. W. Pawelski: *Układowo zorientowany model statyczny IGBT*. Kwart. Elektron. i Telekom., 1994, t. 40, z. 2, ss. 249-264
2. A. Gabryelczyk, M. Grecki, A. Napieralski, W. Pawelski, Z. Kuśmerek: *Tester charakterystyk statycznych IGBT*. Kwart. Elektron. i Telekom. 1994, t. 40, z. 2, ss. 277-290
3. R.K. Brayton, R. Spence: *Sensitivity and Optimization*. Elsevier Scientific Publishing Co. Amsterdam-Oxford-New York, 1980
4. J. Ogrodzki: *Komputerowa analiza układów elektronicznych*. W.N.T., Warszawa 1994

W. PAWELSKI

THE SENSITIVITY ANALYSIS OF NEW IGBT STATIC MODEL

Summary

In the paper some properties of new IGBT static model has been presented. The sensitivity of the model circuit function on main parameter change is analyzed. On this ground the improvement ways of the extraction parameter modes are found.

Tester charakterystyk statycznych IGBT

ANTONI GABRYELCZYK, MARIUSZ GRECKI,
ANDRZEJ NAPIERAŁSKI, WITOLD PAWELSKI

Institut Elektroniki, Politechnika Łódzka

ZYGMUNT KUŚMIEREK

Institut Podstaw Elektrotechniki, Politechnika Łódzka

Otrzymano 1994.03.11

Autoryzowano do druku 1994

W pracy omówiono założenia, budowę, właściwości i zakresy pomiarowe testera charakterystyk IGBT, przeznaczonego do identyfikacji parametrów i weryfikacji doświadczalnej modeli tego przyrządu. Dokonano oceny metrologicznej dokładności testera oraz oszacowano obciążeniowe zmiany temperatury badanych obiektów.

1. WPROWADZENIE

IGBT (Insulated Gate Bipolar Transistor) jest chronologicznie piątym podstawowym półprzewodnikowym przyrządem łączeniowym energoelektroniki po tyrystorze GTO, tranzystorze bipolarnym mocy i tranzystorze MOSFET mocy. Łącząc w sobie zalety tranzystora bipolarnego i polowego mocy, IGBT w ostatnich latach stał się obiektem szczególnego zainteresowania zarówno technologów i producentów jak też projektantów układów energoelektronicznych. Olbrzymi postęp technologii w dziedzinie wytwarzania struktur komórkowych VDMOS z krótkim kanałem oraz konstrukcji modułów elektroizolowanych sprawił, że w chwili obecnej kolektorowe napięcia i prądy dopuszczalne IGBT osiągają wartości 1600 V i 600 A. Szczególnie korzystne warunki autokompensacji termicznej rozprywu prądów przy połączeniach równoległych tych przyrządów stwarzają dodatkowe możliwości poszerzenia zakresu mocy przełączalnych.

W dziedzinie scalonych układów sterowania polowych przyrządów mocy postęp jest nie mniejszy, a prace nad ich zintegrowaniem w jednej półprzewodnikowej

strukturze (smart-power) z układami IGBT, są bardzo zaawansowane. Wszystko to pozwala na prognozę [1] że w najbliższym czasie tranzystory bipolarne mocy będą w aplikacjach łącznikowych, a zwłaszcza falownikowych, całkowicie zastąpione przez IGBT.

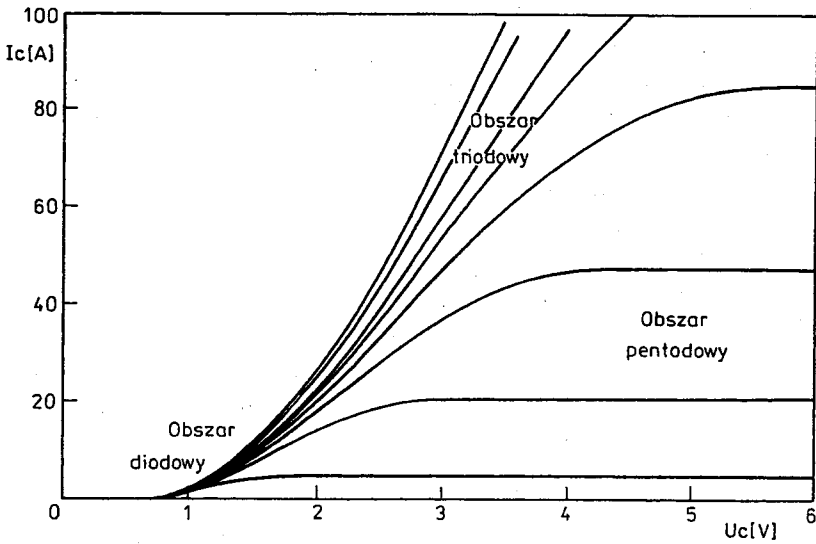
Wzrost zainteresowania IGBT z jednej strony i rozwój komputerowych metod analizy i projektowania układów elektronicznych z drugiej, stymulują zapotrzebowanie na układowo zorientowany model tego przyrządu [2]. Jedynym odniesieniem dla rezultatów takiego modelowania jest pomiar rzeczywistych charakterystyk statycznych IGBT w miarodajnie szerokim zakresie zmian elektrycznych wielkości zaciskowych i temperatury. W ramach prac nad modelowaniem IGBT, podjętych na wydziale Elektrotechniki i Elektroniki Politechniki Łódzkiej, opracowano i wykonano impulsowy tester charakterystyk statycznych tego przyrządu. Tester ten jest przeznaczony zarówno do weryfikacji doświadczalnych wyników teoretycznych modelowania jak również do ekstrakcji parametrów określonych rodzajów modeli statycznych IGBT [2].

Przyjęta metoda pomiarów impulsowych, ze względu na nagrzewanie się badanego obiektu pod wpływem chwilowych (rzędu kW) i średnich mocy strat, jest najbardziej uzasadnioną w odniesieniu do charakterografów przyrządów mocy. W niniejszej pracy przedstawione zostaną założenia i podstawowe cechy metrologiczne testera, estymacja warunków termicznych testowanego tranzystora oraz ocena dokładności pomiaru charakterystyk.

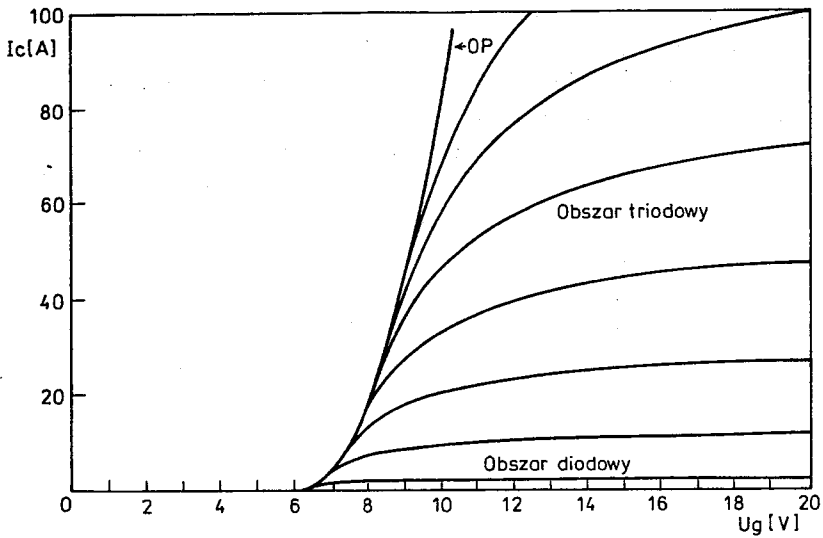
2. CHARAKTERYSTYKI STATYCZNE IGBT

Na rys. 1 przedstawiona została rodzina statycznych charakterystyk wyjściowych typowego IGBT. Na charakterystykach tych zostały wyróżnione trzy obszary określone jako: diodowy, triodowy i pentodowy. Są one znamienne dla tego przyrządu zarówno z punktu widzenia fizycznych właściwości, modelowania jak też cech i parametrów jego układowych aplikacji.

Unikalny kształt charakterystyk IGBT w obszarze triodowym jest jeszcze lepiej widoczny na statycznej charakterystyce przejściowej pokazanej na rys. 2. Przy założeniu zerowej dynamicznej konduktancji wyjściowej IGBT obszar pentodowy zawęży się do postaci charakterystyki OP (rys. 2), w przeciwnym przypadku obszar ten zlokalizowany jest w bezpośrednim sąsiedztwie tej charakterystyki. W zakresie diodowym przebieg charakterystyk warunkowany jest decydująco właściwościami asymetrycznego złącza emiterowego, utworzonego przez warstwę p kolektora IGBT i szeroką warstwę n^- (lub warstwę buforową n^+ w przyrządach o strukturze asymetrycznej rodzaju Punch-Through). Obszar obejmuje zarówno zakres małych napięć i prądów kolektora, o niskim poziomie wstrzykiwania nośników, jak też zakres dużych prądów kolektora I_c o dużym poziomie wstrzykiwania, występujących w następstwie dużych napięć bramki U_G . Charakterystyki



Rys. 1. Typowy przebieg charakterystyk wyjściowych IGBT



Rys. 2. Przykładowe charakterystyki przejściowe IGBT

statyczne w obszarze diodowym, w zakresie małych prądów kolektora I_c mogą być zatem aproksymowane zależnością

$$I_c \approx q \frac{n_i^2 A}{N_D} \sqrt{\frac{D_p}{\tau_p}} \left[\exp \left(\frac{U_C}{n\varphi} \right) - 1 \right], \quad (1)$$

gdzie q — ładunek jednostkowy w C

N_D — gęstość domieszkowania bazy w cm^{-3}

A — czynna powierzchnia przekroju struktury półprzewodnikowej w cm^2

n_i — gęstość nośników w półprzewodniku samoistnym w cm^{-3}

D_p — stała dyfuzji dziur w cm^2s^{-1}

τ_p — czas życia dziur w s

U_C — napięcie kolektor-emiter w V

n — współczynnik emisji złącza emiterowego

φ — potencjał termiczny złącza emiterowego w V .

Dla określania strat mocy w IGBT oraz symulacji komputerowej łączeniowych stanów przejściowych istotne znaczenie posiada dokładne określenie przebiegu charakterystyk tego przyrządu w obszarze triodowym. Przebieg ten może być opisany [2] złożoną funkcją nieliniową

$$I_c = F(U_C, U_G, U_T, T_j, k_p, n, G) \quad (2)$$

której głównymi argumentami są napięcia kolektora i bramki U_C, U_G , napięcie progowe U_T , temperatura w strukturze wewnętrznej T_j , parametr transkonduktancyjny k_p , rozproszona konduktancja G obszarów przyłączowych $p-n^+-n^-$ i współczynnik emisji n tego złącza.

Obszary triodowe na charakterystykach przedstawionych na rys. 1 i 2 odpowiadają stanowi załączenia lub quasi-załączenia IGBT. Dokładność określenia tych charakterystyk ma bezpośredni wpływ nie tylko na dokładność określenia statycznych strat mocy w stanie załączenia. Poprzez zależność (2) i silny wpływ efektu Millera dokładność ta rzutuje na wierność odwzorowania przebiegów czasowych a tym samym dokładność obliczenia strat komutacyjnych w stanach dynamicznych przełączenia.

W obszarze pentodowym przebieg charakterystyk statycznych IGBT zależy głównie od napięcia bramki U_G . W przypadku przyjęcia zerowej wyjściowej konduktancji dynamicznej może on być z dobrym przybliżeniem aproksymowany zależnością Shichmana i Hodgesa w postaci

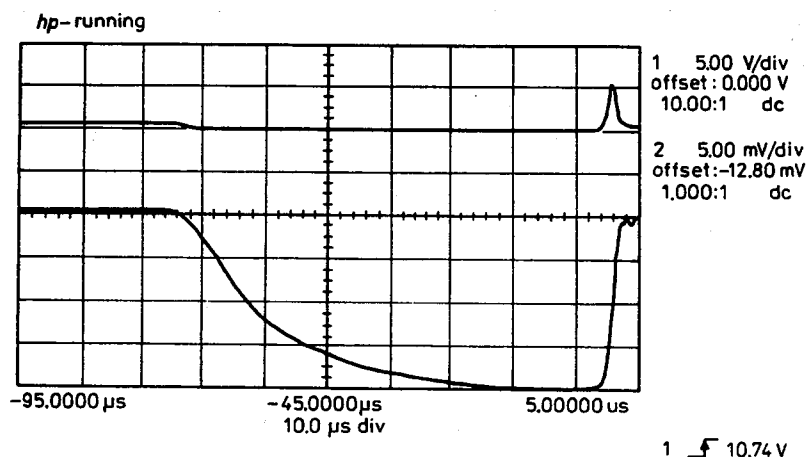
$$I_c \cong 0.5k_p(U_G - U_T)^2, \quad (3)$$

której wyrazem graficznym jest charakterystyka OP pokazana na rys. 2. Możliwie dokładna identyfikacja przebiegu charakterystyk w obszarze pentodowym jest przydatna podczas symulacji łączeniowych przebiegów przejściowych zwłaszcza przy obciążeniach nierezystancyjnych. Jest niezbędna przy projektowaniu systemów zabezpieczenia zwarciovego lub przetężeniowego metodą S-L (self-limiting), stosowaną powszechnie w układach z IGBT.

W zależności od obszaru lokalizacji mierzonego punktu pracy, w badanym IGBT może wydzielać się znaczna moc strat. Straty te są powodem wzrostu chwilowej temperatury w strukturze warstw półprzewodnikowych IGBT, i interakcji

przekraczające poziom około 300 mV. Rozszerzenie zakresu badań charakterystyk w ograniczonym obszarze pentodowym umożliwia alternatywne zasilanie tego obwodu z zewnętrznego źródła napięcia E_{C2} o wartości do 60 V. Pomiar napięcia U_C jest dokonywany na wyjściu układu pomiarowego zawierającego pamięć analogową, której czas próbkowania jest zsynchronizowany z impulsami prądu kolektora. Możliwość szybkiego obniżenia napięcia na baterii kondensatorów C o znacznej pojemności a tym samym napięcia U_C , zapewnia impulsowy układ rozładowujący z wyzwaniem ręcznym.

Badany IGBT jest włączany i wyłączany impulsami napięcia bramki U_G o amplitudzie regulowanej w zakresie od 4 do 20 V. Generator tych impulsów posiada możliwość zarówno pracy astabilnej ze współczynnikiem wypełnienia impulsów równym 0.001, jak też monostabilnej z wyzwaniem ręcznym. W obu rodzajach pracy czas trwania impulsu prądowego wynosi około 70 μ s. Na wyjściu generatora, za pośrednictwem układu pamięci analogowej dokonywany jest pomiar amplitudy impulsów napięcia bramki. Generator wyposażony został w układ ograniczania amplitudy impulsów prądu kolektora na zadanym poziomie w granicach od 0 do 100 A przez regulację napięcia bramki i wykorzystanie sprzężenia zwrotnego od sondy pomiaru prądu I_C . Układ ten zabezpiecza testowane tranzystory o mniejszych maksymalnych dopuszczalnych prądach kolektora przed przypadkowym uszkodzeniem.



Rys. 4. Oscylogramy przebiegu zmian napięcia kolektor-emiter i prądu kolektora badanego IGBT

Na rys. 4 przedstawiony został oscylogram impulsu prądu kolektora i napięcia kolektora. Pomiar amplitudy impulsu I_C odbywa się w układzie pomiarowym zawierającym sondę prądową w postaci przekładnika prądowego ze wzmacniaczem oraz pamięć analogową. Zakres pomiarowy prądów kolektora wynosi 100 A.

4. OSZACOWANIE ZMIAN TEMPERATURY BADANEGO IGBT

Dokładne wyznaczenie zmian temperatury wewnątrz badanego IGBT nie jest łatwe ani na drodze pomiarowej ani obliczeniowej. Ze względu na relatywnie duże rozmiary struktury półprzewodnikowej występuje na ogół przestrzenne zróżnicowanie rozkładu temperatury. Wyznaczenie tego rozkładu na drodze pomiarowej jest praktycznie niemożliwe, ze względu na konieczność zastosowania wielu czujników, które dodatkowo nie powinny wpływać na przebieg mierzonego procesu cieplnego. Zastosowanie metod opartych na detekcji promieniowania podczerwonego nie ma wady zakłócania przebiegu pomiaru, ale pozwala jedynie na wyznaczenie rozkładu temperatury na powierzchni i wnioskowanie na tej podstawie o temperaturze wewnątrz struktury. Dodatkowym utrudnieniem jest konieczność przynajmniej częściowego usunięcia obudowy przyrządu, tak aby powierzchnia struktury półprzewodnikowej była widziana przez termograf. W przypadku przyrządów mocy w obudowach ceramicznych i epoksydowych nie jest to możliwe bez zniszczenia przyrządu.

Wyznaczenie rozkładu temperatury na drodze obliczeniowej również wiąże się z poważnymi problemami. Często można jednak skoncentrować się jedynie na najważniejszych zjawiskach i dokonać pewnych uproszczeń. Jeżeli ograniczyć się do symulacji przepływu ciepła tylko z elementu mocy do radiatora, zakładając że temperatura radiatora jest ustalona, to można pominąć radiację i konwekcję. W tych warunkach transfer ciepła opisany jest dobrze znanym równaniem przewodnictwa ciepła

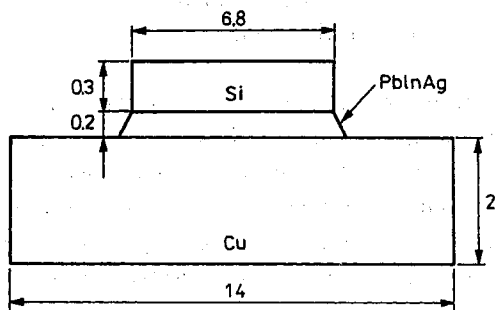
$$\frac{\partial^2 T}{\partial x^2} + \frac{\partial^2 T}{\partial y^2} + \frac{\partial^2 T}{\partial z^2} = \frac{c\rho\partial T}{\lambda\partial t} - \frac{1}{\lambda}g(x, y, z, t) \quad (4)$$

gdzie T — temperatura, x, y, z — współrzędne kartezjańskie, t — czas, λ — przewodność ciepła, ρ — masa właściwa, c — ciepło właściwe, g — gęstość generowanej mocy.

Niestety, rozwiązanie tego równania nie jest w ogólnym przypadku łatwe, a analityczne rozwiązanie może być otrzymane tylko dla bardzo prostych przypadków. Numeryczne rozwiązanie równania przewodnictwa ciepła również wiąże się z wieloma poważnymi problemami. Zazwyczaj stosowane metody oparte na dyskretyzacji metodą elementów skończonych są uciążliwe i wymagają użycia komputerów o dużej mocy obliczeniowej.

Wyznaczenie rozkładu temperatury dla stanu ustalonego wymaga rozwiązania równania (4) dla średniej mocy strat w przyrządzie i uwzględnienia otoczenia termicznego struktury półprzewodnikowej, tzn. rodzaju materiału i wymiarów warstw napotykaných przez strumień cieplny w drodze do radiatora. Oprócz metod ogólnych można tutaj zastosować metodę współczynników sprzężenia cieplnego zastosowaną w programach NAPOM i PYRATHERM [3]. Obliczenie przyrostu temperatury w stanach dynamicznych jest na ogół łatwiejsze niż dla stanu ustalonego. Jeżeli czas trwania procesu przejściowego jest dostatecznie krótki,

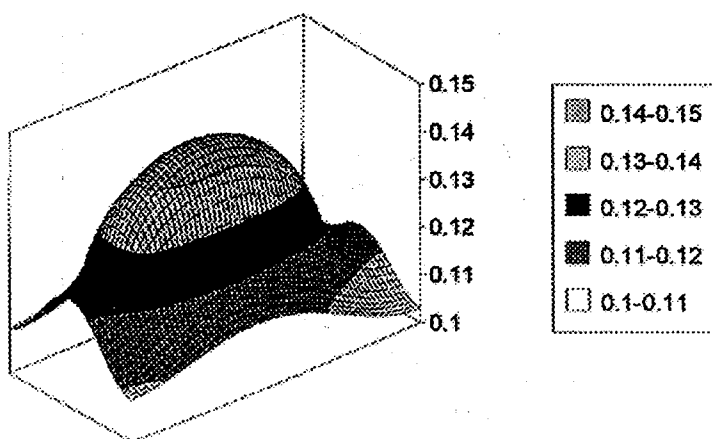
to wystarczy przeprowadzić analizę zjawisk termicznych zachodzących w pobliżu obszaru wydzielania ciepła, czyli w samej strukturze półprzewodnikowej IGBT. Uproszczenie to można uzasadnić niewielką szybkością rozchodzenia się ciepła, wskutek czego dla krótkich czasów strumień ciepły „nie widzi” zewnętrznego otoczenia struktury [3]. Jeżeli dodatkowo generacja ciepła charakteryzuje się symetrią przestrzenną to można zastąpić model trójwymiarowy modelem dwu- a nawet jednowymiarowym. W takich przypadkach wygodnie jest testować model termiczny w postaci odpowiedniej sieci elektrycznej, wykorzystując analogie pomiędzy wielkościami termicznymi i elektrycznymi. Rozwiązanie równania (4) osiąga się wówczas przez rozwiązanie równań Kirchhoffa, co najłatwiej jest wykonać stosując do tego celu uniwersalne programy analizy obwodów elektrycznych jak NAP2, SPICE, ESACAP lub inne [4].



Rys. 5. Przekrój poprzeczny struktury badanego tranzystora IGBT (wymiar w mm)

W celu oszacowania wzrostu temperatury wewnątrz struktury półprzewodnikowej podczas trwania pomiaru charakterystyk statycznych przeprowadzono przykładową analizę termiczną, wykorzystując dane typowego IGBT o prądzie nominalnym 25 A i napięciu 900 V. Dla stanu ustalonego, przy opcji okresowych impulsów mocy testera, obliczenia zostały przeprowadzone przy pomocy programu PYRATHERM dla impulsów mocy o amplitudzie 0,25 kW. Temperatura dla dolnej płaszczyzny miedzi (rys. 5) została przyjęta za stałą, ze względu na umocowanie badanego tranzystora na dużym radiatorze termostatu (rys. 3), którego temperatura była stabilizowana. Na skutek bardzo małych średnich strat mocy przy współczynniku wypełnienia równym 0,001, ustalony przyrost temperatury wewnątrz struktury jest niewielki i wynosi 0,15 K (rys. 6). Przy wykorzystywaniu opcji testowania z pojedynczymi impulsami mocy przyrost ten praktycznie jest równy zero.

Przyrost temperatury w stanie przejściowym wyznaczono w oparciu o jednowymiarowy model termiczny w postaci sieci elektrycznej RC [4]. Komórkowa struktura tranzystora IGBT powoduje równomierny rozptył prądu i w konsekwencji równomierny rozkład generowanej mocy w płaszczyźnie przekroju po-



Rys. 6. Przyrost temperatury w tranzystorze IGBT (stan ustalony)

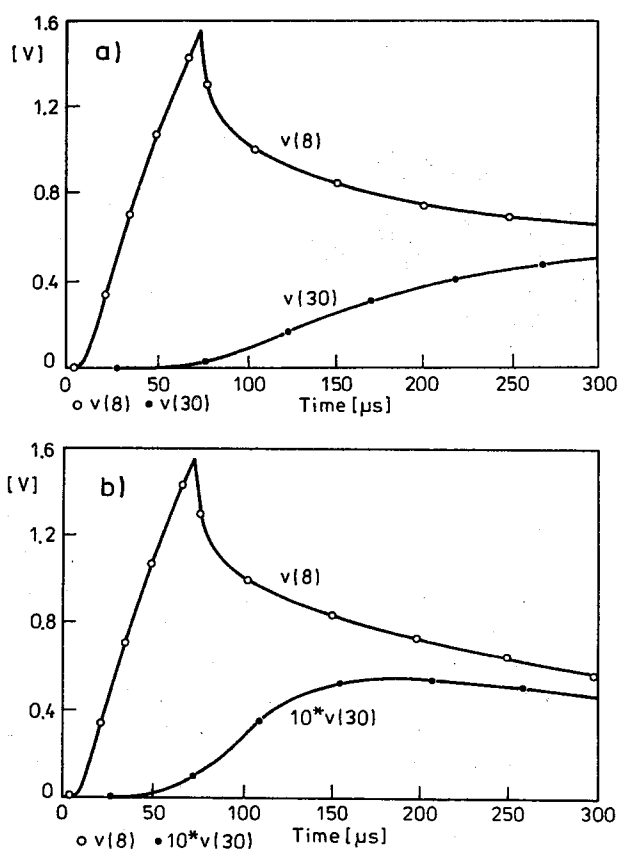
przecznego. Wybór zatem jednowymiarowego modelu termicznego jest w tym przypadku wystarczająco uzasadniony. Rozpatrzono najgorszy przypadek mocy generowanej w jednej płaszczyźnie. Model uwzględnia jedynie warstwę krzemu ponieważ podczas trwania impulsu mocy (rys. 4) ciepło nie może zdążyć rozprzestrzenić się poza nią. Na górnej krawędzi krzemu przyjęto adiabatyczne warunki brzegowe. Na dolnej krawędzi warunki brzegowe nie są istotne dla analizowanego czasu i niezależnie od przyjętych tam warunków adiabatycznych czy też izotermicznych, wewnątrz struktury temperatura jest identyczna. Sytuację tę ilustruje rys. 7, z którego wynika że ciepło dociera w pobliże dolnej krawędzi krzemu dopiero po czasie trwania impulsu mocy wynoszącym ok. 70 μs .

**Przyrosty temperatury struktury półprzewodnikowej
IGBT przy ustalonej temperaturze radiatora**

Temperatura radiatora [$^{\circ}\text{C}$]	27	50	75	100
Δ [$^{\circ}\text{C}$]	1.55	1.61	1.67	1.73

Wyniki symulacji przeprowadzonych dla różnych temperatur radiatora, zestawione w tabeli 1 pokazują, że przyrost temperatury w strukturze podczas trwania impulsu mocy nie przekracza w najgorszym przypadku 1.73 $^{\circ}\text{C}$. Z tego względu można przyjąć, że temperatura wewnątrz struktury półprzewodnikowej podczas trwania pomiarów jest stała i praktycznie równa temperaturze radiatora (ewentualny przyrost temperatury dla stanu ustalonego rzędu 0.15 K jest do pominięcia).

Chociaż symulacja termiczna została przeprowadzona tylko dla jednego typu tranzystora IGBT, to można jednak na jej podstawie wnioskować w przypadku innych tranzystorów. Tranzystory o większych od badanego prądach znamionowych



Rys. 7. Przebieg temperatury wewnątrz struktury półprzewodnikowej w miejscu wstrzykiwania mocy — $V(8)$ oraz w pobliżu dolnej krawędzi krzemu — $V(30)$, temperatura radiatora 27°C , (skala $1\text{V}=1\text{K}$): a) — warunki adiabatyczne, b) — warunki izotermiczne.

mają struktury o większych powierzchniach. Czasami są to po prostu połączone równolegle mniejsze struktury bez dodatkowych elementów. Jeżeli porówna się tranzystory o różnych prądach znamionowych to okaże się, że gęstość prądu a zatem i gęstość wydzielonej mocy jest stała i nie zależy od wartości prądu znamionowego. W takim przypadku przyrost temperatury będzie zbliżony do otrzymanego dla badanego tranzystora.

5. OCENA DOKŁADNOŚCI TESTERA

Dokładność wyznaczania charakterystyk statycznych IGBT za pomocą testera można oszacować przez estymację błędów, z jakim są określone współrzędne punktów tworzących charakterystykę.

W układzie pomiarowym testera z rys. 3 można wyodrębnić cztery tory pomiarowe:

- napięcia kolektora U_C ,
- napięcia bramki U_G ,
- prądu kolektora I_c ,
- ustalonej temperatury radiatora badanego IGBT.

Pierwsze trzy wielkości zmieniają się impulsowo i do ich pomiaru zastosowano szybkie układy próbkująco-pamiętające z układami scalonymi LF 198. W układach pomiarowych napięć U_C i U_G oraz w członie wyjściowym układu pomiarowego prądu I_c znajdują się po dwa takie układy. Według danych katalogowych ich błąd dla sygnału wyjściowego nie przekracza wartości 0.01%.

Ocenę dokładności poszczególnych torów pomiarowych przeprowadzono zgodnie z zaleceniami Międzynarodowego Komitetu Metrologii [5] przez wyznaczenie wariancji poszczególnych składowych błędów.

Dla oceny dokładności układu do pomiaru napięcia kolektora porównano jego wskazania ze wskazaniem wzorcowego woltomierza. W wyniku otrzymano graniczną wartość błędu $\delta_{g1} = 0.4\%$. Dla równomiernego rozkładu tych błędów kwadrat wariancji jest określony wzorem

$$S_{c1}^2 = \frac{\delta_{g1}^2}{3}. \quad (5)$$

Według danych katalogowych błąd graniczny dla sygnału wyjściowego układu próbkująco-pamiętającego wynosi $\delta_{g2} = 0.01\%$. Przy równomiernym rozkładzie błędów, kwadrat wariancji jest równy

$$S_{c2}^2 = \frac{\delta_{g2}^2}{3}. \quad (6)$$

W celu określenia wpływu czynników zewnętrznych, jak również wyznaczenia cech testera dotyczących powtarzalności wyników pomiaru wykonano serię 10 pomiarów napięcia U_C . Pomiarów powtarzano następnie dla różnych wartości tego napięcia, otrzymując w wyniku kwadraty wariancji S_{c3} dla różnych wartości napięcia U_C . Wariancję wypadkową wyniku pomiaru wartości napięcia U_C można wyznaczyć ze wzoru

$$S_{UC} = \sqrt{S_{c1}^2 + S_{c2}^2 + S_{c3}^2}. \quad (7)$$

Niepewność wyniku pomiaru oblicza się z zależności

$$\Delta U_C = \pm k S_{UC}, \quad (8)$$

gdzie k — współczynnik równy 2, co odpowiada prawdopodobieństwu 0.95.

Dla podanych wartości błędów pomiarowych i napięcia $U_C = 10V$ otrzymano niepewność równą 0.5%. W podobny sposób została dokonana analiza dokładności

układu pomiarowego napięcia bramki U_G . Ze względu na zastosowanie w układzie pomiarowym tego toru elementów o identycznych danych jak przy pomiarze napięcia U_C , w wyniku otrzymano również niepewność 0.5%.

Wyznaczenie błędu pomiaru amplitudy impulsu prądowego I_c jest nieco trudniejsze. Pomiar prądu dokonywany jest w układzie z przekładnikiem prądowym o rdzeniu ferrytowym. Sygnał wyjściowy przekładnika jest wzmacniany i wprowadzany do układu pomiarowego z pamięcią analogową. Do oceny dokładności toru pomiarowego jak też wzorcowania wykorzystano oscyloskop typu HP 54503A z sondą pomiarową prądu A6303 (Tektronix) i własnym wzmacniaczem. Mając na uwadze błędy wprowadzane przez sondę, wzmacniacz sygnałów jak i sam oscyloskop, w oparciu o dane tych przyrządów można oszacować niepewność pomiaru prądu na poziomie nie przekraczającym 5%. Oznacza to, że głównie tor prądowy rzutuje na dokładność wyznaczania charakterystyk IGBT przy pomocy omawianego testera w warunkach stałej temperatury pomiaru.

Dodatkową wielkością mierzoną jest temperatura radiatora badanego IGBT, która ma wpływ na kształt charakterystyk prądowo-napięciowych. Tor pomiarowy temperatury jest oddzielony od torów pomiarów wielkości elektrycznych testera. Pomiar dokonywany jest za pomocą oddzielnego cyfrowego miernika typu MCT-401 z czujnikiem półprzewodnikowym. Błąd pomiaru wynosi $\pm 1^\circ\text{C}$, przy rozdzielczości 0.1°C .

Punkt na charakterystyce prądowo-napięciowej IGBT przy stałej temperaturze jest określony współrzędnymi I_c, U_C, U_G . Każda z tych wielkości, jak wykazano, wyznaczona jest z określoną dokładnością.

Jeżeli wektor łączący początek układu współrzędnych z punktem leżącym na charakterystyce oznaczyć przez $\mathbf{K}$, to moduł tego wektora wynosi

$$|\mathbf{K}| = \sqrt{U_C^2 + U_G^2 + I_c^2}. \quad (9)$$

Dla ułatwienia obliczeń poszczególne wielkości określone są w jednostkach względnych. Wariancja niepewności wyznaczenia modułu wektora $\mathbf{K}$ jest określona zależnością

$$S_K = \sqrt{\left(\frac{\partial K}{\partial U_C}\right)^2 S_{U_C}^2 + \left(\frac{\partial K}{\partial U_G}\right)^2 S_{U_G}^2 + \left(\frac{\partial K}{\partial I_c}\right)^2 S_{I_c}^2}, \quad (10)$$

gdzie $\frac{\partial K}{\partial U_C}, \frac{\partial K}{\partial U_G}, \frac{\partial K}{\partial I_c}$ — pochodne cząstkowe odpowiednio względem U_C, U_G, I_c , $S_{U_C}, S_{U_G}, S_{I_c}$ — wariancje niepewności, z jakimi są wyznaczone U_C, U_G, I_c . Niepewność wyznaczenia wektora $\mathbf{K}$ może być określona jako

$$\Delta K = \pm k S_K \quad \text{gdzie } k = 2. \quad (11)$$

Po podstawieniu wartości liczbowych, wynikających z niepewności wyznaczenia U_C, U_G , i I_c , okazuje się że wektor określający punkt na charakterystyce prądowo-

napięciowej w warunkach stałej temperatury wyznaczony jest, dla jej środkowego położenia z niepewnością 3.4%.

Błędy wynikające z pomiaru ustalonej temperatury radiatora i z oszacowanego przyrostu obciążeniowego temperatury struktury półprzewodnikowej badanego IGBT wpływają na dokładność wyznaczania jego charakterystyk statycznych. Parametry badanego przyrządu zależą nieliniowo od temperatury zatem błąd wyznaczenia odpowiednich charakterystyk, wynikający z niedokładności określenia temperatury będzie różny dla różnych obszarów tych charakterystyk. Błąd ten, pomimo bardzo małych wartości, powinien być uwzględniony przy analizie porównawczej wyników modelowania IGBT.

PODSUMOWANIE

W artykule przedstawiono właściwości i możliwości pomiarowe impulsowego testera charakterystyk statycznych IGBT. Jego wykorzystanie przy modelowaniu zapewnia odniesienie do rzeczywistych charakterystyk IGBT. Tester ten może być więc stosowany zarówno podczas identyfikacji parametrów modeli układowo-zorientowanych jak też weryfikacji doświadczalnych wyników symulacji. W ten sposób umożliwia on ocenę wierności odwzorowania przez model kształtu unikalnych charakterystyk IGBT oraz ocenę dokładności modelu.

Dokonano oceny błędów pomiarowych testera oraz estymacji chwilowych przyrostów temperatury w strukturze wewnętrznej badanego przyrządu. Wykazano dużą dokładność pomiarów w całym zakresie pracy przyrządu, prognozując symulacyjnie niewielkie dynamiczne przyrosty temperatury badanego obiektu. Dzięki temu tester może być oceniony jako wiarygodne narzędzie metrologiczne, niezmiennie pomocne przy modelowaniu i badaniach IGBT.

BIBLIOGRAFIA

1. B.J. Baliga: *Power Semiconductor Devices for the 1990's*. EPE'91 Proceedings—pp. 0-001-0-002
2. W. Pawelski: *Układowo-zorientowany model statyczny IGBT*. Kwart. Elekt. i Telekom., 1994 t. 40, z. 2, ss. 249–264
3. A. Napieralski: *Komputerowe projektowanie układów półprzewodnikowych dużej mocy ze szczególnym uwzględnieniem ich właściwości termicznych*. Zeszyty Naukowe PŁ, Nr 562, Łódź 1988
4. A. Napieralski, M. Grecki: *2D simulation of power bipolar transistor with temperature computation during switching process*. XVI KKTOiUE Kołobrzeg 1993
5. Guide to Expression of Uncertainty in Measurement. ISO, 1992

A. GABRYELCZYK, M. GRECKI, Z. KUŚMIEREK, A. NAPIERALSKI, W. PAWELSKI

STATIC CHARACTERISTICS TESTER OF THE IGBT

Summary

In this paper the static characteristics tester of the IGBT is presented. Tester can be used to experimental verification of the IGBT models and to its parameters extraction. The construction assumptions, properties and measurement ranges of the tester are described. The accuracy of the tester and the temperature variation in investigated IGBT are also put to estimate.

Wykorzystanie komory TEM i GTEM w pomiarach emisji zakłóceń radiowych

JAN SROKA

*Swiss Federal Institute of Technology, Electromagnetic Group
Zurich*

Otrzymano 1993.11.26

Autoryzowano do druku 1994.03.15

Pomiary emisji zakłóceń radiowych dostarczają istotnych informacji o kompatybilności elektromagnetycznej (EMC) urządzeń elektronicznych. Europejskie normy dotyczące emisji promieniowania elektromagnetycznego nakazują przeprowadzanie testów w otwartej przestrzeni (OATS — Open Area Test Site) lub w komorach bezdechowych. Pierwszy wariant jest kłopotliwy, drugi bardzo kosztowny. Pomiary w komorze TEM lub GTEM są alternatywnym, tańszym sposobem mierzenia emisji. Metody alternatywne muszą dawać możliwość przetransponowania pomierzonych sygnałów na wielkości określone w normach, w tym wypadku na charakterystykę promieniowania w otwartej przestrzeni.

Komora TEM jest wycinkiem przewodnicy falowej o powiększonej przestrzeni między przewodem wewnętrznym, zwanym septum, a zewnętrznym, stanowiącym płaszczyznę metalową. Ma ona dwa przyłącza kablowe, jest więc dwuwrotnikiem. Komora GTEM jest jak gdyby wycinkiem komory TEM z jednym przyłączem kablowym i częścią lejkową ograniczoną ścianą wyłożoną materiałem absorbcyjnym. Zaletą komory GTEM jest większe pasmo częstotliwości, w którym można dokonywać pomiarów bez obaw istnienia w przyłączy kablowym innych rodzajów pól oprócz TEM. Dla małych komór GTEM pasmo to przekracza 1 GHz, podczas gdy w komorze TEM sięga kilkuset MHz [3]. Köpke [6] podał sposób identyfikowania testowanego urządzenia (EUT — Equipment Under Test) w komorze TEM, jako układ trzech dipoli elektrycznych i magnetycznych z różnymi fazami początkowymi. Rozważania teoretyczne w [6] rozpoczynają się podaniem zależności między sygnałami mierzonymi na przyłączach kablowych, a momentami dipoli elektrycznego i magnetycznego, będących źródłem promieniowania w komorze. W niniejszej pracy przedstawiono wyprowadzenie wzorów podanych w [6], modelując EUT jako lokalne źródło prądu. Dalej opisano procedurę pomiaru emisji zakłóceń w komorze TEM, tak jak to uczyniono w [6]. W artykule zwrócono uwagę na różnice uzyskane w przypadku komory TEM i GTEM.

W założeniach przyjmowanych w trakcie wyprowadzenia wzorów, od których rozpoczyna się artykuł [6] tkwią ograniczenia stosowania tej metody pomiarowej.

Ograniczenia te znikają z pola widzenia wielu autorów, co prowadzi wielokrotnie do mylnych stwierdzeń w publikacjach [3], [7]. Jest o tym mowa we wnioskach. Identyfikowanie EUT jako układu trzech dipoli elektrycznych o różnych fazach początkowych i trzech magnetycznych o różnych fazach początkowych wymaga równoczesnego mierzenia dwu sygnałów w sześciu ortogonalnych ustawieniach EUT. W komorze TEM, jako dwuwrotniku jest to możliwe. Komora GTEM jest jednowrotnikiem. Dlatego nie daje ona możliwości uzyskania informacji o przesunięciach fazowych poszczególnych dipoli modelujących EUT.

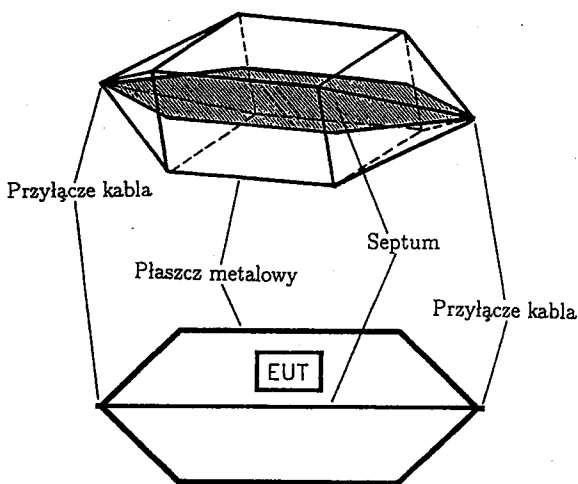
W publikacjach dotyczących pomiarów emisji w GTEM [4], [3], [12] mówi się o pomiarze tylko w trzech ortogonalnych ustawieniach utrzymując, że otrzymany w ten sposób model odpowiada układowi trzech dipoli elektrycznych i trzech magnetycznych o identycznych fazach początkowych. W istocie taki pomiar umożliwia zidentyfikowanie EUT jako jednego dipola elektrycznego albo magnetycznego.

Z rozważań przedstawionych w niniejszym artykule wynika, że komora TEM daje możliwość dokładniejszego zidentyfikowania EUT.

1. WSTĘP

Obecny etap rozwoju cywilizacji cechuje się nasyceniem życia codziennego wszelkiego rodzaju urządzeniami sterowanymi mikroprocesorami (PC, telefony komórkowe, FAX, drukarki, kopiarki, kamery video, aparaty fotograficzne, TV, kuchenki mikrofalowe itp.). Rozwój mikroelektroniki sprawia, że częstotliwość pracy mikroprocesorów jest coraz większa. Dlatego urządzenia posiadające mikroprocesory są potencjalnym źródłem emisji zakłóceń częstotliwości radiowej oraz są potencjalnie narażone na takie zakłócenia z zewnątrz. Koniecznością stało się konstruowanie i dopuszczanie do sprzedaży urządzeń, które tych potencjalnych zagrożeń nie niosą, czyli kompatybilnych elektromagnetycznie. Sposoby sprawdzania kompatybilności elektromagnetycznej (EMC — ElectroMagnetic Compatibility) określone są w normach. Europejskie normy EMC przewidują przeprowadzanie testów w otwartej przestrzeni, zwanych OATS (Open Area Test Site) lub w komorze bezechowej (anechoic chamber). Testowanie w otwartej przestrzeni jest kłopotliwe, komory bezechowe są bardzo kosztowne. Stąd próby opracowania alternatywnych, prostszych i tańszych metod przeprowadzania testów EMC.

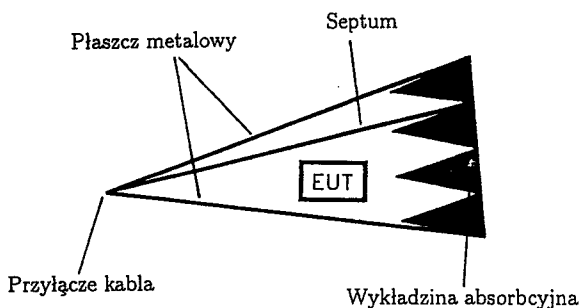
Jedną z alternatyw jest komora TEM przedstawiona na rys. 1. Jej nazwa wywodzi się stąd, że w pustej komorze TEM, w zakresie częstotliwości roboczej mamy do czynienia tylko z polem rodzaju TEM, pozbawionym wpływów pól zewnętrznych. Komora jest wycinkiem przewodnicy falowej o powiększonej przestrzeni między przewodem wewnętrznym, zwanym w komorze septum, a zewnętrznym, stanowiącym płaszczyznę metalową. Septum stanowi szeroką powierzchnię metalową, dzięki czemu w dużej części przestrzeni roboczej, w której umieszcza się testowane urządzenie (EUT — Equipment Under Test) uzyskuje się jednorodne pole. Komora ma dwa przyłącza kablowe. Jest więc dwuwrotnikiem. Komory projektuje się tak, aby impedancja charakterystyczna wzdłuż całej komory



Rys. 1. Widok perspektywiczny i przekrój podłużny komory TEM

była równa impedancja przylączy kablowych. Dlatego w strefach między częścią roboczą, a przylączami kablowymi, septum zwęża się, a płaszcz ma kształt lejka.

Crawford w 1974 roku opublikował artykuł o zastosowaniu komory TEM w pomiarach EMC [2]. Pierwszym zastosowaniem komory była kalibracja sond do mierzenia natężenia pola elektrycznego. Drugim obszarem zastosowań komory TEM był pomiar wrażliwości urządzeń na zakłócenia pól zewnętrznych (immunity tests). Najpóźniej zaczęto stosować komory TEM w pomiarach emisji promieniowania elektromagnetycznego (radiated emission tests). Na ten temat pojawiło się parę artykułów na przełomie lat 70/80 [9, 10, 6].



Rys. 2. Przekrój podłużny komory GTM

Pod koniec lat 80-tych pojawiła się modyfikacja komory TEM tzw. komora GTM (TRIGATEM [7]), przedstawiona na rys. 2. Komora ta jest jak gdyby wycinkiem komory TEM z jednym przylącem kablowym i częścią lejkową ograniczoną ścianą wyłożoną materiałem absorbcyjnym. Konstruktorzy tej komory

utrzymują, że dzięki małemu kątowni rozwarcia lejka, pole w pustej komorze jest praktycznie polem rodzaju TEM [3]. Zaletą komory GTEM jest większe pasmo częstotliwości, w którym można dokonywać pomiarów bez obaw istnienia w przyłączach kablowych innych rodzajów pól, oprócz TEM. Dla małych komór GTEM pasmo to przekracza 1 GHz, podczas gdy w komorze TEM sięga kilkuset MHz [3].

Trudność przy pomiarze emisji w komorze TEM i GTEM polega na tym, że EUT musi być zidentyfikowane, po czym dla modelu uzyskanego w identyfikacji należy określić charakterystykę promieniowania w otwartej przestrzeni, gdyż takie warunki pomiaru określone są w normach. Kopke [6] podał sposób identyfikowania EUT w komorze TEM, jako układu trzech dipoli elektrycznych i magnetycznych z różnymi fazami początkowymi. Rozważania teoretyczne w [6] rozpoczynają się podaniem zależności między sygnałami mierzonymi na przyłączach kablowych, a momentami dipoli elektrycznego i magnetycznego, będących źródłem promieniowania w komorze. W niniejszej pracy przedstawiono wyprowadzenie wzorów podanych w [6], modelując EUT jako lokalne źródło prądu. Dalej opisano procedurę pomiaru emisji zakłóceń w komorze TEM, tak jak to uczyniono w [6]. W artykule zwrócono uwagę na różnice uzyskane w przypadku komory TEM i GTEM.

W założeniach przyjmowanych w trakcie wyprowadzenia wzorów tkwią ograniczenia stosowania tej metody pomiarowej. Ograniczenia te znikają z pola widzenia wielu autorów, co prowadzi wielokrotnie do mylnych stwierdzeń w publikacjach [3], [7]. Jest o tym mowa we wnioskach.

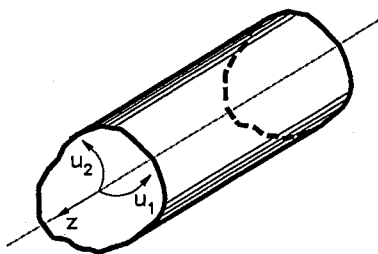
Aby wyprowadzeniom wzorów nadać zwartość, dołączono do artykułu dodatkowe. Są w nim zawarte zagadnienia ogólne wykorzystywane w artykule, a więc prezentacja charakterystyki promieniowania dipoli elektrycznego i magnetycznego w strefie dalekiej oraz wyprowadzenie charakterystyki promieniowania układu dipoli.

2. PODSTAWY TEORETYCZNE IDENTYFIKACJI EUT W KOMORZE TEM I GTEM

2.1. ORTONORMALNE RODZAJE PÓL W PROWADNICY FAŁOWEJ

Do opisu pól w przewodnicy fałowej użyte będą ortogonalne, krzywoliniowe współrzędne o nazwie u_1, u_2 w przekroju poprzecznym oraz z wzdłuż przewodnicy. Wersory tych osi nazwane będą odpowiednio l_1, l_2, l_z .

Rozwiązaniem równania fałowego w przewodnicy, której ściany i przewód wewnętrzny są idealnymi przewodnikami elektrycznymi jest kombinacja liniowa ortogonalnych rodzajów pól. Każdemu rodzajowi przyporządkowana jest para liczb naturalnych l, m . W dalszej części artykułu dla uproszczenia zapisu, symbol l, m zastąpiony będzie symbolem jednoliterowym n . Jednoznaczne rozwiązanie jednorodnego równania fałowego jest możliwe po przyjęciu dodatkowego warunku,



Rys. 3. Układ współrzędnych w przewodnicy falowej

zwanego normalizacją. Pole w przewodnicy o bezstratnych ścianach, uzyskane po narzuceniu na rozwiązanie warunku normalizacji, nazywa się polem ortonormalnym. Pola ortonormalne oznaczone będą symbolami odpowiednio: wektor pola elektrycznego $\vec{E}_n$, wektor pola magnetycznego $\vec{H}_n$.

Ortonormalne natężenia pola elektrycznego fali rodzaju n , rozchodzącej się w dodatnim kierunku osi z można wyrazić wzorem

$$\vec{E}_n^{(+)} = [\vec{E}_{Tn}(u_1, u_2) + \vec{E}_{zn}(u_1, u_2)]e^{ik_n z}. \quad (1)$$

Dla fali rozchodzącej się w ujemnym kierunku osi z natężenie to wyraża się wzorem

$$\vec{E}_n^{(-)} = [\vec{E}_{Tn}(u_1, u_2) - \vec{E}_{zn}(u_1, u_2)]e^{-ik_n z}. \quad (2)$$

Aby pole elektryczne było polem bezźródłowym ($\text{div } \vec{E}_n = 0$), składowa wzdłużna $\vec{E}_{zn}$, w równaniu (2) musi mieć znak przeciwny do składowej poprzecznej $\vec{E}_{Tn}^{(1)}$.

(1) Rozpatrzmy spełnienie warunku $\text{div } \vec{E}_n = 0$ dla pola opisanego równaniem (1), (2)

$$\vec{E}_n = \vec{E}_n^{(+)} + \vec{E}_n^{(-)} \quad (i)$$

$$\text{div } \vec{E}_n = (\text{div } \vec{E}_{Tn} + ik_n \vec{E}_{zn})(e^{ik_n z} + e^{-ik_n z}) \quad (ii)$$

Pierwszy czynnik w tym równaniu jest niczym innym jak operacją div rozbitą na składową poprzeczną i podłużną.

Przyjmijmy zamiast (2) równanie

$$\vec{E}_n^{(-)} = [\vec{E}_{Tn}(u_1, u_2) + \vec{E}_{zn}(u_1, u_2)]e^{-ik_n z} \quad (iii)$$

Wówczas

$$\text{div } \vec{E}_n = \text{div } \vec{E}_{Tn}(e^{ik_n z} + e^{-ik_n z}) + ik_n \vec{E}_{zn}(e^{ik_n z} - e^{-ik_n z}) \quad (iv)$$

Równanie to jest spełnione dla dowolnego z tylko wtedy gdy $\text{div } \vec{E}_{Tn} = 0$ i $\vec{E}_{zn} = 0$. Warunek ten jest możliwy do spełnienia tylko dla pola rodzaju TEM, dla którego problem znaku składowej wzdłużnej nie istnieje.

Warunek ten dopuszcza inną postać równania (2)

$$\vec{E}_n^{(+)} = [-\vec{E}_{Tn}(u_1, u_2) + \vec{E}_{zn}(u_1, u_2)]e^{-ik_n z}. \quad (3)$$

Ortonormalne natężenia pola magnetycznego fali rodzaju n , rozchodzącej się w dodatnim kierunku osi z można wyrazić wzorem

$$\vec{H}_n^{(+)} = [\vec{H}_{Tn}(u_1, u_2) + \vec{H}_{zn}(u_1, u_2)]e^{ik_n z}. \quad (4)$$

Dla fali rozchodzącej się w ujemnym kierunku osi z natężenie to wyraża się wzorem

$$\vec{H}_n^{(-)} = [-\vec{H}_{Tn}(u_1, u_2) + \vec{H}_{zn}(u_1, u_2)]e^{-k_n z}. \quad (5)$$

Składowa poprzeczna w równaniu (5) musi mieć znak „-”, aby moc transportowana falą rozchodzącą się w kierunku ujemnym była przekazywana rzeczywiście w tym kierunku. Jeśli jako $\vec{E}_n^{(-)}$ wykorzystać wzór (3), to $\vec{H}_{Tn}$ we wzorze (5) będzie poprzedzone znakiem „+”. Składowa wzdłużna w równaniu (5) jest dodatnia, gdyż pole $\vec{H}_n$, podobnie jak pole $\vec{E}_n$, musi spełniać warunek bezźródłowości $\text{div } \vec{H}_n = 0$.

Żądamy, aby pola $\vec{E}_{Tn}$, $\vec{H}_{Tn}$ były funkcjami rzeczywistymi ($\vec{E}_{Tn} = E_{Tn}$), ($\vec{H}_{Tn} = H_{Tn}$). Ponadto narzucamy warunek [1]

$$\iint_S \vec{E}_{Tn} \times \vec{H}_{Tm} \cdot \vec{1}_z dS = \delta_{nm}, \quad (6)$$

przy czym S jest przekrojem poprzecznym falowodu, a δ_{nm} jest deltą Kroneckera

$$\delta_{mn} = \begin{cases} 0 [W] & \text{jeśli } m \neq n \\ 1 [W] & \text{jeśli } m = n. \end{cases} \quad (7)$$

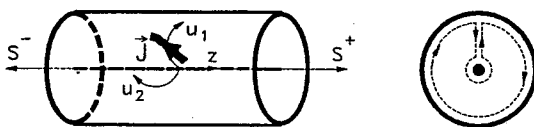
Wzór (6) wyraża ortogonalność pól dla $n \neq m$ oraz normalizację dla $n = m$.

2.2. POLE W PROWADNICY FALOWEJ WZBUDZONE LOKALNYM ŹRÓDŁEM PRĄDU

Rozpatrzmy wycinek przewodnicy falowej ograniczony powierzchniami S^+ , S^- , prostopadłymi do osi z jak pokazano na rys. 4. W wycinku tym znajduje się lokalne źródło prądu o gęstości $\vec{J}$, rozłożonej dowolnie w objętości V .

Od źródła tego, w obu kierunkach przewodnicy rozchodzą się fale elektromagnetyczne. Falę rozchodzącą się w dodatnim kierunku osi z opisujemy następująco

$$\vec{E}^{(+)} = \sum_n a_n \vec{E}_n^{(+)} \quad (8.a)$$



Rys. 4. Lokalne źródło prądu w przewodnicy falowej

$$\vec{H}^{(+)} = \sum_n a_n \vec{H}_n^{(+)} \quad (8.b)$$

oraz w ujemnym kierunku osi z

$$\vec{E}^{(-)} = \sum_n b_n \vec{E}_n^{(-)} \quad (9.a)$$

$$\vec{H}^{(-)} = \sum_n b_n \vec{H}_n^{(-)} \quad (9.b)$$

przy czym a_n, b_n są nieznanymi amplitudami pól, zależnymi od źródła prądu $\vec{J}$. Są one bezwymiarowe.

Wzory (8.a), (8.b) określają pola $\vec{E}, \vec{H}$ na ścianie S^+ , podczas gdy wzory (9.a), (9.b) na ścianie S^- .

Do wyznaczenia amplitud a_n, b_n wykorzystamy zasadę wzajemności Lorenza, której postać wygodna do dalszych rozważań jest następująca

$$\oint_S (\vec{E}_n^{(\pm)} \times \vec{H} - \vec{E} \times \vec{H}_n^{(\pm)}) \cdot \vec{1}_n dS = \iiint_V \vec{J} \cdot \vec{E}_n^{(\pm)} dV, \quad (10)$$

przy czym jako $\vec{1}_n$ wybrano kierunek normalny, wewnętrzny do objętości V . Przez powierzchnię zamkniętą we wzorze (10) należy rozumieć powierzchnię z wyłączeniem przewodu wewnętrznego, jak zaznaczono linią przerywaną na rys. 4.

Zajmijmy się lewą stroną równ. (10). Jeżeli ściany przewodnicy oraz przewód wewnętrzny są idealnie przewodzące, to znika na nich całka powierzchniowa gdyż $(\vec{E}_n^{(\pm)} \times \vec{H}) \cdot \vec{1}_n = \vec{H} \cdot (\vec{1}_n \times \vec{E}_n^{(\pm)})$, przy czym zarówno $\vec{E}_n^{(\pm)}$, jak i $\vec{E}$ mają na powierzchni idealnie przewodzących ścian tylko składową normalną. To samo dotyczy składnika $\vec{E} \times \vec{H}_n^{(\pm)} \cdot \vec{1}_n$. Tylko całki na powierzchniach S^+, S^- są niezerowe.

Rozpatrzmy tę całkę na ścianie S^+ dla górnego indeksu „+”

$$\iint_{S^+} [\vec{E}_n^{(+)} \times \vec{H}^{(+)} - \vec{E}^{(+)} \times \vec{H}_n^{(+)}] \cdot \vec{1}_n dS =$$

$$\iint_{S^+} \left\{ (\vec{E}_{Tn} + \vec{E}_{zn}) e^{ik_n z} \times \left[\sum_n \mathbf{a}_n (\vec{H}_{Tn} + \vec{H}_{zn}) e^{ik_n z} \right] - \left[\sum_n \mathbf{a}_n (\vec{E}_{Tn} + \vec{E}_{zn}) e^{ik_n z} \right] \times (\vec{H}_{Tn} + \vec{H}_{zn}) e^{ik_n z} \right\} \cdot \vec{1}_n dS. \quad (11)$$

Wartość tej całki jest zerem, bowiem:

- iloczyn wektorowy składowej transwersalnej przez wzdłużną daje wektor styczny do powierzchni całkowania, strumień tego wektora jest zerem,
- iloczyn wektorowy obu składowych wzdłużnych jest zerem, gdyż są to wektory równoległe do siebie,
- jeśli wykorzystać warunek normalizacji do iloczynów wektorowych składowych transwersalnych to otrzymamy

$$\iint_{S^+} \left[\vec{E}_{Tn} e^{ik_n z} \times \left(\sum_n \mathbf{a}_n \vec{H}_{Tn} e^{ik_n z} \right) - \left(\sum_n \mathbf{a}_n \vec{E}_{Tn} e^{ik_n z} \right) \times \vec{H}_{Tn} e^{ik_n z} \right] \cdot \vec{1}_n dS = \mathbf{a}_n e^{2ik_n z} - \mathbf{a}_n e^{2ik_n z} = 0. \quad (12)$$

Rozpatrzmy tę samą całkę (dla górnego indeksu „+”) na ścianie S^-

$$\iint_{S^-} [\vec{E}_n^{(+)} \times \vec{H}_n^{(-)} - \vec{E}_n^{(-)} \times \vec{H}_n^{(+)}] \cdot \vec{1}_n dS = \iint_{S^-} \left\{ (\vec{E}_{Tn} + \vec{E}_{zn}) e^{ik_n z} \times \left[\sum_n \mathbf{b}_n (-\vec{H}_{Tn} + \vec{H}_{zn}) e^{-ik_n z} \right] - \left[\sum_n \mathbf{b}_n (\vec{E}_{Tn} - \vec{E}_{zn}) e^{-ik_n z} \right] \times (\vec{H}_{Tn} + \vec{H}_{zn}) e^{ik_n z} \right\} \cdot \vec{1}_n dS. \quad (13)$$

Z tego samego względu co w równ (11) wszystkie całki zerują się za wyjątkiem

$$\iint_{S^-} \left[-\vec{E}_{Tn} e^{ik_n z} \times \left(\sum_n \mathbf{b}_n \vec{H}_{Tn} e^{-ik_n z} \right) - \left(\sum_n \mathbf{b}_n \vec{E}_{Tn} e^{-ik_n z} \right) \times \vec{H}_{Tn} e^{ik_n z} \right] \cdot \vec{1}_n dS = -2\mathbf{b}_n [W]. \quad (14)$$

Wymiarem tego wyrażenia są W , co wynika z przyjętej normalizacji (wzór (6)). Amplituda $\mathbf{b}_n$ jest bezwymiarowa, ale jest pomnożona przez całkę (wzór (6)), której wartość wynosi 1 $[W]$.

Równanie (10) można więc wyrazić następująco:

$$\iiint_V \vec{J} \cdot \vec{E}_n^{(+)} dV = -2\mathbf{b}_n [W]. \quad (15)$$

Rozpatrując analogicznie równ. (10) dla dolnego indeksu „-” otrzymujemy

$$\iiint_V \vec{J} \cdot \vec{E}_n^{(-)} dV = -2a_n [W]. \quad (16)$$

2.3 ELEKTRYCZNIE MAŁE ŹRÓDŁO PRĄDU W KOMORZE TEM LUB GTM

Ponieważ zarówno w komorze TEM oraz GTM pożądanym jest, aby w zakresie częstotliwości roboczej do miejsca pomiaru docierało tylko pole rodzaju TEM ($n = 0$), w dalszych rozważaniach ograniczymy się do tego rodzaju. Ograniczenie to nie wyklucza istnienia wyższych rodzajów pól w pobliżu źródła zakłócenia. Stawia tylko wymóg, aby górna granica zakresu częstotliwości roboczej leżała poniżej częstotliwości odcięcia wszystkich wyższych rodzajów pól.

Powtórzmy wzory na amplitudy fal rozchodzących się od lokalnego zaburzenia w przewodnicy falowej (15), (16).

$$\iiint_V \vec{J} \cdot \vec{E}_0^{(\pm)} dV = - \left(\frac{2b_0}{2a_0} \right) [W]. \quad (17)$$

W przypadku komory TEM ściany S^+ , S^- na rys. 4 odpowiadają przyłączom kablowym. W komorze TEM, dzięki dwóm przyłączom kablowym jest możliwe zmierzenie obu amplitud a_0 , b_0 .

Przyjmijmy, że w przypadku komory GTM ściana S^+ na rys. 4 jest wyłożona materiałem absorbcyjnym, natomiast ściana S^- jest przyłączem kablowym. Uzyskuje się wówczas wzory identyczne do (17), ale istnieje możliwość pomierzenia tylko jednej amplitudy (b_0).

W dalszej części, dla uproszczenia zapisu opuścimy górny indeks przy $\vec{E}_0^{(\pm)}$. Zostawimy natomiast dolny indeks. Przypomnijmy, że $\vec{E}_0$ oznacza ortonormalne pole rodzaju TEM, natomiast $\vec{E}$ całkowite pole rodzaju TEM (por. wzory 8.a, 8.b, 9.a, 9.b).

Jeśli wektor $\vec{E}_0$ w otoczeniu źródła nie zmienia się znacznie, można całkę $\iiint_V \vec{J} \cdot \vec{E}_0 dV$ rozwinąć w szereg Taylora

$$\iiint_V \vec{J} \times \vec{E}_0 dV = \sum_{i=1}^3 \iiint_V \mathbf{J}_i \left[E_{0i}(0) + \sum_{j=1}^3 x_j \frac{\partial E_{0i}(0)}{\partial x_j} + \dots \right] dV. \quad (18)$$

Pokażemy, że pierwszy szereg wyraża pole dipola elektrycznego. Korzystając dla wektora $\vec{J}$ z tej samej tożsamości co w przypadku wzoru (A.3) otrzymujemy

$$\iiint_V \mathbf{J}_i dV = - \iiint_V x_i \operatorname{div} \vec{J} dV + \oint_S x_i \vec{J} \cdot \vec{1}_n dS. \quad (19)$$

Ponieważ wzdłuż dipola Hertza, z założenia prąd nie zmienia się ($I(x_i) = \text{const.}$), więc $\text{div } \vec{J} = 0$. Dlatego

$$\iiint_V \vec{J} dV = \iint_S r(\vec{J} \cdot \vec{1}_n) dS. \quad (20)$$

Na końcach dipola zgromadzony jest całkowity ładunek Q odpowiedzialny za prąd dipola I . Jeśli zdefiniować prąd jako ruch dodatnich ładunków, to zasada zachowania ładunku $\text{div } \vec{J} + d\rho/dt = 0$ przyjmuje następującą postać całkową $I = dQ/dt$. Dla amplitud zespolonych zależność ta wygląda następująco $I = -i\omega Q$.

Ostatecznie, dla amplitud zespolonych

$$\iint_S \vec{r}(\vec{J} \cdot \vec{1}_n) dS = I\Delta r \vec{1}_r = -i\omega\Delta r Q \vec{1}_r. \quad (21)$$

Korzystając z definicji momentu dipolowego wzór (A.6) możemy pierwszy składnik rozwinięcia w szereg Taylora, z wzoru (18) zapisać następująco

$$\vec{E}_0(0) \iiint_V \vec{J} dV = -i\omega \vec{p} \vec{E}_0(0). \quad (22)$$

Drugi składnik szeregu Taylora przegrupujemy następująco

$$\begin{aligned} \sum_{i=1}^3 \sum_{j=1}^3 \mathbf{J}_i x_j \frac{\partial E_{0i}(0)}{\partial x_j} &= \frac{1}{4} \sum_{i=1}^3 \sum_{j=1}^3 (\mathbf{J}_i x_j - \mathbf{J}_j x_i) \left[\frac{\partial E_{0i}(0)}{\partial x_j} - \frac{\partial E_{0j}(0)}{\partial x_i} \right] + \\ &\quad \frac{1}{2} \sum_{i=1}^3 \sum_{j=1}^3 (\mathbf{J}_i x_j + \mathbf{J}_j x_i) \frac{\partial E_{0i}(0)}{\partial x_j}. \end{aligned} \quad (23)$$

Pierwsze sumy zostały zapisane tak, aby dostrzec w nich iloczyn skalarny $2(\vec{r} \times \vec{J}) \cdot \text{rot } \vec{E}_0(0)$. Pokażemy, że sumy te wyrażają pole dipola magnetycznego. Pozostałe sumy we wzorze (23) dają pole kwadropola. Tę część rozwinięcia w szereg pominiemy w dalszych rozważaniach.

Zgodnie z równaniem Maxwella dla amplitud zespolonych

$$\text{rot } E_0(0) = i\omega\mu H_0(0),$$

natomiast $\vec{r} \times \vec{J}$ jest podwojoną gęstością objętościową dipolowego momentu magnetycznego, która występuje we wzorze (A.9). Ostatecznie możemy „część dipolową” drugiego wyrazu rozwinięcia w szereg Taylora przedstawić następująco

$$\iiint_V \sum_{i=1}^3 \sum_{j=1}^3 x_j \mathbf{J}_i \frac{\partial E_{0i}(0)}{\partial x_j} dV = i\omega\mu \vec{m} \vec{H}_0(0) + \dots \quad (24)$$

Uwzględniając zatem tylko pole dipoli, można amplitudy $\mathbf{b}_0, \mathbf{a}_0$ ze wzoru (17) przedstawić jak następuje

$$\begin{pmatrix} \mathbf{b}_0 \\ \mathbf{a}_0 \end{pmatrix} [W] = \frac{1}{2} i \omega [\vec{\mathbf{p}} \cdot \vec{E}_0^{(\pm)}(0) - \mu \vec{\mathbf{m}} \cdot H_0^{(\pm)}(0)]. \quad (25)$$

Jeśli przyjąć układ współrzędnych w komorze tak, aby oś z leżała wzdłuż komory i wskazywała dodatni kierunek rozchodzenia się pola (S^+), to pole rodzaju TEM można wyrazić następująco

$$\vec{E}_0^{(\pm)} = \vec{E}_0 e^{\pm i k_0 z} \quad (26.a)$$

$$\vec{H}_0^{(\pm)} = \vec{H}_0 e^{\pm i k_0 z}, \quad (26.b)$$

przy czym k_0 jest stałą propagacji pola rodzaju TEM ($n = 0$).

Miedzy $\vec{E}_0^{(\pm)}$, a $\vec{H}_0^{(\pm)}$ istnieje związek

$$\vec{H}_0^{(\pm)} = \pm \vec{1}_z \times \vec{E}_0^{(\pm)} \frac{1}{Z_{f0}}, \quad (27)$$

przy czym Z_{f0} jest impedancją falową rodzaju TEM. W komorze TEM jest ona równa impedancji właściwej próżni ζ_0 i wynosi $120\pi \approx 377 [\Omega]$.

Ponieważ $k_0 = \omega \sqrt{\mu_0 \epsilon_0}$ oraz $\zeta_0 = \sqrt{\mu_0 / \epsilon_0}$, można równanie (25) przedstawić następująco

$$\begin{pmatrix} \mathbf{b}_0 \\ \mathbf{a}_0 \end{pmatrix} [W] = \frac{1}{2} i (\omega \vec{\mathbf{p}} \mp k_0 \vec{\mathbf{m}} \times \vec{1}_z) \cdot \vec{E}_0. \quad (28)$$

W dalszej części opuścimy symbol $[W]$ pamiętając, że korzystamy z zasady wzajemności Lorenza wzór (17), w której bezwymiarowe amplitudy $\mathbf{b}_0, \mathbf{a}_0$ są współczynnikami wyrażającymi moc w $[W]$. Wzór (28) można zapisać jako układ dwu równań

$$\mathbf{a}_0 + \mathbf{b}_0 = i\omega \vec{\mathbf{p}} \cdot \vec{E}_0 \quad (29.a)$$

$$\mathbf{a}_0 - \mathbf{b}_0 = i k_0 (\vec{\mathbf{m}} \times \vec{1}_z) \cdot \vec{E}_0 \quad (29.b)$$

Równania (29.a), (29.b) dają się przekształcić do następującej postaci

$$P_s = (\mathbf{a}_0 + \mathbf{b}_0)(\mathbf{a}_0 + \mathbf{b}_0)^* = \omega^2 |\vec{\mathbf{p}} \cdot \vec{E}_0|^2 \quad (30.a)$$

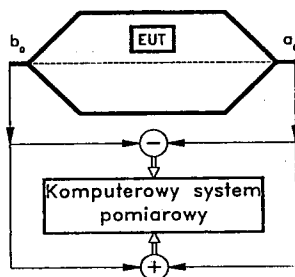
$$P_d = (\mathbf{a}_0 - \mathbf{b}_0)(\mathbf{a}_0 - \mathbf{b}_0)^* = k_0^2 |(\vec{\mathbf{m}} \times \vec{1}_z) \cdot \vec{E}_0|^2. \quad (30.b)$$

Wzory (30.a, 30.b) mają wymiar kwadratu mocy. Równ. (30.a) zależy tylko od dipola elektrycznego, a równ. (30.b) tylko od dipola magnetycznego. Ponadto należy zauważyć, że na wartości pomierzonych sygnałów mają wpływ tylko rzuty

wektorów momentów $\vec{p}$, $\vec{m}$ na kierunek wektora $\vec{E}_0$. Oznaczmy je symbolami $\vec{p}_T$, $\vec{m}_T$.

3. PROCEDURA POMIAROWA

Wzory (30.a), (30.b) pokazują, że w przypadku komory TEM znając sumę amplitud $a_0 + b_0$ można określić dipol elektryczny. Znając różnicę $a_0 - b_0$ można określić dipol magnetyczny. Z tych rozważań wynika procedura pomiarowa. Schemat ideowy układu pomiarowego przedstawia rys. 5.



Rys. 5. Schemat systemu mierzącego promieniowanie EUT w komorze TEM

Zorientujmy globalny układ współrzędnych x, y, z unieruchomiony w komorze tak, aby jego oś z pokrywała się z wzdłużną osią symetrii komory, a płaszczyzna xy leżała w przekroju poprzecznym. Wprowadźmy również lokalny, ruchomy układ współrzędnych ξ, η, ζ , sprzęgnięty z EUT.

Rozpoczynając cykl pomiarowy wstawiamy EUT do komory tak, aby osie ξ, η pokrywały się odpowiednio z osiami x, y oraz aby kierunki osi ζ i z były takie same. Poszukiwać będziemy składowych wektorów

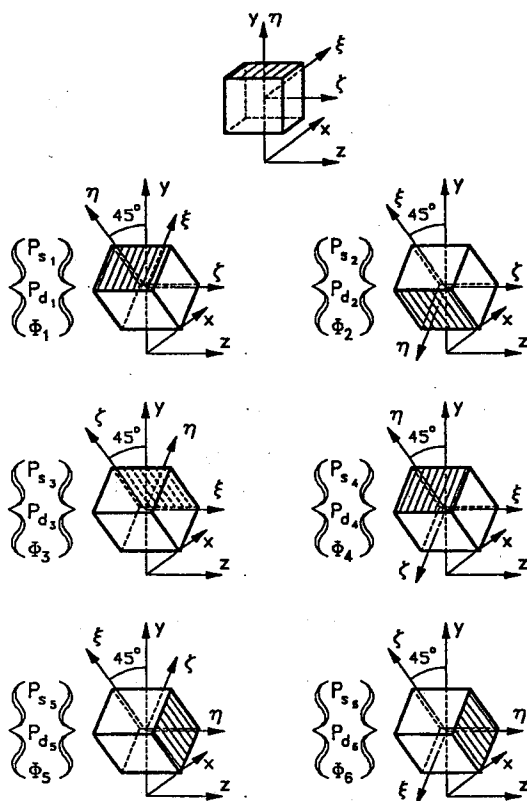
$$\vec{p} = p_x e^{i\psi_{px}} \vec{1}_x + p_y e^{i\psi_{py}} \vec{1}_y + p_z e^{i\psi_{pz}} \vec{1}_z \quad (31.a)$$

$$\vec{m} = m_x e^{i\psi_{mx}} \vec{1}_x + m_y e^{i\psi_{my}} \vec{1}_y + m_z e^{i\psi_{mz}} \vec{1}_z \quad (31.b)$$

w tym ustawieniu EUT. Jest to ustawienie w normalnej pracy.

Należy przypomnieć, że E_0 jest znormalizowanym natężeniem pola, rodzaju TEM w punkcie umieszczenia dipoli, czyli w punkcie $(0, y_0, 0)$. Ważną cechą dobrze skonstruowanej komory TEM jest to, że $\vec{E}_0 = q \vec{1}_y$, tzn. ma tylko składową y , przy czym q jest współczynnikiem wynikającym z normalizacji.

Pierwszego zestawu pomiarów P_{s1}, P_{d1}, ϕ_1 dokonujemy po obróceniu EUT o kąt $\pi/4$ względem osi ζ , tak aby oś ξ zbliżyła się do osi y , (rys. 6). W tym ustawieniu $\vec{p}_T = p_x \vec{1}_\xi + p_y \vec{1}_\eta$, natomiast $\vec{E}_0 = q(\vec{1}_\xi + \vec{1}_\eta)/\sqrt{2}$. Wobec



Rys. 6. Ustawienie EUT w czasie pomiarów w komorze TEM

tego

$$\mathbf{a}_0 + \mathbf{b}_0 = i\omega \vec{\mathbf{p}}_T \cdot \vec{\mathbf{E}}_0 = \frac{i\omega q}{\sqrt{2}}(p_x e^{i\psi_{p_x}} + p_y e^{i\psi_{p_y}}) \quad (32)$$

jest wskazem (wektorem na płaszczyźnie zespolonej), którego kwadrat długości, zgodnie z twierdzeniem cosinusów wynosi

$$P_{s_1} = |\mathbf{a}_0 + \mathbf{b}_0|^2 = \frac{\omega^2 q^2}{2} [p_x^2 + p_y^2 + 2p_x p_y \cos(\psi_{p_x} - \psi_{p_y})]. \quad (33)$$

Zanim wyprowadzimy wyrażenie na P_{d_1} , przypomnijmy, że $\vec{\mathbf{1}}_z = \vec{\mathbf{1}}_\zeta$. Wobec tego

$$\mathbf{a}_0 - \mathbf{b}_0 = ik_0(\vec{\mathbf{m}}_T \times \vec{\mathbf{1}}_\zeta) \cdot \vec{\mathbf{E}}_0 = ik_0(-m_x e^{i\psi_{m_x}} \vec{\mathbf{1}}_\eta + m_y e^{i\psi_{m_y}} \vec{\mathbf{1}}_\xi) \cdot \vec{\mathbf{E}}_0. \quad (34)$$

Postępując analogicznie jak w przypadku wyznaczania P_{s_1} , otrzymujemy

$$\mathbf{a}_0 - \mathbf{b}_0 = \frac{iqk_0}{\sqrt{2}}(-m_x e^{i\psi_{m_x}} + m_y e^{i\psi_{m_y}}) \quad (35)$$

$$P_{d_1} = |\mathbf{a}_0 - \mathbf{b}_0|^2 = \frac{k_0^2 q^2}{2} [m_x^2 + m_y^2 - 2m_x^2 m_y^2 \cos(\psi_{m_x} - \psi_{m_y})]. \quad (36)$$

Teraz zajmijmy się wyznaczeniem fazy początkowej ϕ_{a+b} . Równanie (32) można przekształcić następująco

$$\mathbf{a}_0 + \mathbf{b}_0 = \frac{\omega q}{\sqrt{2}} e^{i(\psi_{p_x} - \frac{\pi}{2})} (p_x + p_y e^{i(\psi_{p_y} - \psi_{p_x})}). \quad (37)$$

Faza początkowa wyrażenia (37) ϕ_{a+b} ma następującą postać

$$\phi_{a+b} = \psi_{p_x} + \alpha_{x_1} - \frac{\pi}{2}, \quad (38)$$

przy czym

$$\alpha_{x_1} = \arctg \frac{-p_y \sin(\psi_{p_x} - \psi_{p_y})}{p_x + p_y \cos(\psi_{p_x} - \psi_{p_y})}. \quad (39)$$

Równie dobrze można równanie (32) przedstawić jako

$$\mathbf{a}_0 + \mathbf{b}_0 = \frac{\omega q}{\sqrt{2}} e^{i(\psi_{p_y} - \frac{\pi}{2})} (p_x e^{i(\psi_{p_x} - \psi_{p_y})} + p_y) \quad (40)$$

co doprowadzi do następującej postaci na fazę początkową ϕ_{a+b} wyrażenia (40)

$$\phi_{a+b} = \psi_{p_y} + \alpha_{y_1} - \frac{\pi}{2}, \quad (41)$$

przy czym

$$\alpha_{y_1} = \arctg \frac{p_x \sin(\psi_{p_x} - \psi_{p_y})}{p_x \cos(\psi_{p_x} - \psi_{p_y}) + p_y}. \quad (42)$$

Przechodzimy do wyznaczenia fazy początkowej ϕ_{a-b} . Równanie (35) można przekształcić następująco

$$\mathbf{a}_0 - \mathbf{b}_0 = \frac{k_0 q}{\sqrt{2}} e^{i(\psi_{m_x} - \frac{\pi}{2})} (-m_x + m_y e^{i(\psi_{m_y} - \psi_{m_x})}) \quad (43)$$

Faza początkowa ϕ_{a-b} wyrażenia (43) ma następującą postać

$$\phi_{a-b} = \psi_{m_x} + \beta_{x_1} - \frac{\pi}{2}, \quad (44)$$

przy czym

$$\beta_{x_1} = \arctg \frac{m_y \sin(\psi_{m_x} - \psi_{m_y})}{m_x - m_y \cos(\psi_{m_x} - \psi_{m_y})}. \quad (45)$$

Równie dobrze można równanie (35) przedstawić jako

$$\mathbf{a}_0 - \mathbf{b}_0 = \frac{k_0 q}{\sqrt{2}} e^{i(\psi_{m_y} - \frac{\pi}{2})} (-m_x e^{i(\psi_{m_x} - \psi_{m_y})} + m_y) \quad (46)$$

co doprowadzi do następującej postaci na fazę początkową ϕ_{a+b} wyrażenia (46)

$$\phi_{a-b} = \psi_{m_y} + \beta_{y_1} - \frac{\pi}{2}, \quad (47)$$

przy czym

$$\beta_{y_1} = \arctg \frac{m_x \sin(\psi_{m_x} - \psi_{m_y})}{m_x \cos(\psi_{m_x} - \psi_{m_y}) - m_y}. \quad (48)$$

Ostatecznie kąt fazowy ϕ_1 między fazą początkową sumy sygnałów ϕ_{a+b} , a fazą początkową ich różnicy ϕ_{a-b} wyraża się wzorem

$$\phi_1 = \phi_{a+b} - \phi_{a-b} = \left\{ \psi_{p_x} + \alpha_{x_1} \right\} - \left\{ \psi_{m_x} + \beta_{x_1} \right\} \quad (49)$$

Drugiego zestawu pomiarów P_{s_2}, P_{d_2}, ϕ_2 dokonujemy po obroceniu EUT o dalsze $\pi/2$ radianów względem osi ζ (rys. 6). W tym ustawieniu $\vec{p} = p_x \vec{1}_\xi + p_y \vec{1}_\eta$, natomiast $\vec{E}_0 = q(\vec{1}_\xi - \vec{1}_\eta)/\sqrt{2}$. Wobec tego

$$\mathbf{a}_0 + \mathbf{b}_0 = i\omega \vec{p} \cdot \vec{E}_0 = \frac{i\omega q}{\sqrt{2}}(p_x e^{i\psi_{p_x}} - p_y e^{i\psi_{p_y}}) \quad (50)$$

Zgodnie z twierdzeniem cosinusów

$$P_{s_2} = |\mathbf{a}_0 + \mathbf{b}_0|^2 = \frac{\omega^2 q^2}{2}[p_x^2 + p_y^2 - 2p_x^2 p_y^2 \cos(\psi_{p_x} - \psi_{p_y})] \quad (51)$$

Zauważmy, że wzór (34) jest słuszny również i przy tym ustawieniu układu współrzędnych ξ, η, ζ względem x, y, z

$$\mathbf{a}_0 - \mathbf{b}_0 = ik_0(\vec{m}_T \times \vec{1}_\zeta) \cdot \vec{E}_0 = ik_0(-m_x e^{i\psi_{m_x}} \vec{1}_\eta + m_y e^{i\psi_{m_y}} \vec{1}_\xi) \cdot \vec{E}_0. \quad (52)$$

Dlatego otrzymujemy

$$\mathbf{a}_0 - \mathbf{b}_0 = \frac{iqk_0}{\sqrt{2}}(m_x e^{i\psi_{m_x}} + m_y e^{i\psi_{m_y}}) \quad (53)$$

oraz

$$P_{d_2} = |\mathbf{a}_0 - \mathbf{b}_0|^2 = \frac{k_0^2 q^2}{2}[m_x^2 + m_y^2 + 2m_x^2 m_y^2 \cos(\psi_{m_x} - \psi_{m_y})]. \quad (54)$$

Równanie (50) można przekształcić następująco

$$\mathbf{a}_0 + \mathbf{b}_0 = \frac{\omega q}{\sqrt{2}} e^{i(\psi_{p_x} - \frac{\pi}{2})} (p_x - p_y e^{i(\psi_{p_y} - \psi_{p_x})}). \quad (55)$$

Faza początkowa ϕ_{a+b} wyrażenia (55) ma następującą postać

$$\phi_{a+b} = \psi_{p_x} + \alpha_{x_2} - \frac{\pi}{2} \quad (56)$$

przy czym

$$\alpha_{x_2} = \arctg \frac{p_y \sin(\psi_{p_x} - \psi_{p_y})}{p_x - p_y \cos(\psi_{p_x} - \psi_{p_y})}. \quad (57)$$

Równie dobrze można równanie (50) przedstawić jako

$$\mathbf{a}_0 + \mathbf{b}_0 = \frac{\omega q}{\sqrt{2}} e^{i(\psi_{py} - \frac{\pi}{2})} (p_x e^{i(\psi_{px} - \psi_{py})} - p_y) \quad (58)$$

co doprowadzi do następującej postaci na fazę początkową ϕ_{a+b} wyrażenia (58)

$$\phi_{a+b} = \psi_{py} + \alpha_{y2} - \frac{\pi}{2}, \quad (59)$$

przy czym

$$\alpha_{y2} = \arctg \frac{p_x \sin(\psi_{px} - \psi_{py})}{p_x \cos(\psi_{px} - \psi_{py}) - p_y}. \quad (60)$$

Przechodzimy do wyznaczenia fazy początkowej ϕ_{a-b} . Równanie (53) można przekształcić następująco

$$\mathbf{a}_0 - \mathbf{b}_0 = \frac{k_0 q}{\sqrt{2}} e^{i(\psi_{mx} - \frac{\pi}{2})} (m_x + m_y e^{i(\psi_{my} - \psi_{mx})}). \quad (61)$$

Faza początkowa ϕ_{a-b} wyrażenia (61) ma następującą postać

$$\phi_{a-b} = \psi_{mx} + \beta_{x2} - \frac{\pi}{2} \quad (62)$$

przy czym

$$\beta_{x2} = \arctg \frac{m_y \sin(\psi_{mx} - \psi_{my})}{m_x + m_y \cos(\psi_{mx} - \psi_{my})}. \quad (63)$$

Równie dobrze można równanie (53) przedstawić jako

$$\mathbf{a}_0 - \mathbf{b}_0 = \frac{k_0 q}{\sqrt{2}} e^{i(\psi_{my} - \frac{\pi}{2})} (m_x e^{i(\psi_{mx} - \psi_{my})} + m_y) \quad (64)$$

co doprowadzi do następującej postaci na fazę początkową ϕ_{a-b} wyrażenia (64)

$$\phi_{a-b} = \psi_{my} + \beta_{y2} - \frac{\pi}{2}, \quad (65)$$

przy czym

$$\beta_{y2} = \arctg \frac{m_x \sin(\psi_{mx} - \psi_{my})}{m_x \cos(\psi_{mx} - \psi_{my}) + m_y}. \quad (66)$$

Ostatecznie kąt fazowy ϕ_2 między fazą początkową sumy sygnałów ϕ_{a+b} , a fazą początkową ich różnicy ϕ_{a-b} wyraża się wzorem

$$\phi_2 = \phi_{a+b} - \phi_{a-b} = \left\{ \psi_{px} + \alpha_{x2} \right\} - \left\{ \psi_{mx} + \beta_{x2} \right\}. \quad (67)$$

Obrócenie EUT o następne $\pi/2$ radianów nie wnosi niczego, gdyż doprowadzi do powtórzenia pomiarów z pozycji 1.

Przed dokonaniem trzeciego i czwartego zestawu pomiarów orientujemy EUT tak, aby $\vec{1}_\xi = \vec{1}_z$, $\vec{1}_\eta = \vec{1}_x$, $\vec{1}_\zeta = \vec{1}_y$. Następnie obracamy EUT o kąt

$\pi/4$ względem osi ξ , tak aby oś η zbliżyła się do osi y i dokonujemy pomiarów P_{s3}, P_{d3}, ϕ_3 . Z kolei obracamy EUT o następne $\pi/2$ radianów i mierzymy P_{s4}, P_{d4}, ϕ_4 .

Przed ostatnim zestawem pomiarów trzeba zorientować EUT następująco: $\vec{1}_\eta = \vec{1}_z, \vec{1}_\zeta = \vec{1}_x, \vec{1}_\xi = \vec{1}_y$. Następnie dla uzyskania pomiarów P_{s5}, P_{d5}, ϕ_5 obracamy EUT o kąt $\pi/4$ względem osi η po czym obracamy o następne $\pi/2$ radianów dla wyznaczenia P_{s6}, P_{d6}, ϕ_6 .

Zestawmy wszystkie zależności potrzebne do wyznaczenia dipoli. Z sumy sygnałów otrzymujemy

$$P_{s1} = \frac{\omega^2 q^2}{2} [p_x^2 + p_y^2 + 2p_x p_y \cos(\psi_{p_x} - \psi_{p_y})] \quad (68.a)$$

$$P_{s2} = \frac{\omega^2 q^2}{2} [p_x^2 + p_y^2 - 2p_x p_y \cos(\psi_{p_x} - \psi_{p_y})] \quad (68.b)$$

$$P_{s3} = \frac{\omega^2 q^2}{2} [p_y^2 + p_z^2 + 2p_y p_z \cos(\psi_{p_y} - \psi_{p_z})] \quad (68.c)$$

$$P_{s4} = \frac{\omega^2 q^2}{2} [p_y^2 + p_z^2 - 2p_y p_z \cos(\psi_{p_y} - \psi_{p_z})] \quad (68.d)$$

$$P_{s5} = \frac{\omega^2 q^2}{2} [p_z^2 + p_x^2 + 2p_z p_x \cos(\psi_{p_z} - \psi_{p_x})] \quad (68.e)$$

$$P_{s6} = \frac{\omega^2 q^2}{2} [p_z^2 + p_x^2 - 2p_z p_x \cos(\psi_{p_z} - \psi_{p_x})]. \quad (68.f)$$

Z różnicy natomiast

$$P_{d1} = \frac{k^2 q^2}{2} [m_x^2 + m_y^2 - 2m_x m_y \cos(\psi_{m_x} - \psi_{m_y})] \quad (69.a)$$

$$P_{d2} = \frac{k^2 q^2}{2} [m_x^2 + m_y^2 + 2m_x m_y \cos(\psi_{m_x} - \psi_{m_y})] \quad (69.b)$$

$$P_{d3} = \frac{k^2 q^2}{2} [m_y^2 + m_z^2 - 2m_y m_z \cos(\psi_{m_y} - \psi_{m_z})] \quad (69.c)$$

$$P_{d4} = \frac{K^2 q^2}{2} [m_y^2 + m_z^2 + 2m_y m_z \cos(\psi_{m_y} - \psi_{m_z})] \quad (69.d)$$

$$P_{d5} = \frac{k^2 q^2}{2} [m_z^2 + m_x^2 - 2m_z m_x \cos(\psi_{m_z} - \psi_{m_x})] \quad (69.e)$$

$$P_{d6} = \frac{k^2 q^2}{2} [m_z^2 + m_x^2 + 2m_z m_x \cos(\psi_{m_z} - \psi_{m_x})]. \quad (69.f)$$

Przesunięcia fazowe

$$\phi_1 = \left\{ \begin{matrix} \psi_{p_x} + \alpha_{x1} \\ \psi_{p_y} + \alpha_{y1} \end{matrix} \right\} - \left\{ \begin{matrix} \psi_{m_x} + \beta_{x1} \\ \psi_{m_y} + \beta_{y1} \end{matrix} \right\} \quad (70.a)$$

$$\phi_2 = \left\{ \begin{matrix} \psi_{p_x} + \alpha_{x_2} \\ \psi_{p_y} + \alpha_{y_2} \end{matrix} \right\} - \left\{ \begin{matrix} \psi_{m_x} + \beta_{x_a} \\ \psi_{m_y} + \beta_{y_2} \end{matrix} \right\} \quad (70.b)$$

$$\phi_3 = \left\{ \begin{matrix} \psi_{p_y} + \alpha_{y_3} \\ \psi_{p_x} + \alpha_{z_3} \end{matrix} \right\} - \left\{ \begin{matrix} \psi_{m_y} + \beta_{y_3} \\ \psi_{m_z} + \beta_{z_3} \end{matrix} \right\} \quad (70.c)$$

$$\phi_4 = \left\{ \begin{matrix} \psi_{p_y} + \alpha_{y_4} \\ \psi_{p_z} + \alpha_{z_4} \end{matrix} \right\} - \left\{ \begin{matrix} \psi_{m_y} + \beta_{y_4} \\ \psi_{m_z} + \beta_{z_4} \end{matrix} \right\} \quad (70.d)$$

$$\phi_5 = \left\{ \begin{matrix} \psi_{p_z} + \alpha_{z_5} \\ \psi_{p_x} + \alpha_{x_5} \end{matrix} \right\} - \left\{ \begin{matrix} \psi_{m_z} + \beta_{z_5} \\ \psi_{m_x} + \beta_{x_5} \end{matrix} \right\} \quad (70.e)$$

$$\phi_6 = \left\{ \begin{matrix} \psi_{p_x} + \alpha_{z_6} \\ \psi_{p_x} + \alpha_{x_6} \end{matrix} \right\} - \left\{ \begin{matrix} \psi_{m_x} + \beta_{z_6} \\ \psi_{m_x} + \beta_{x_6} \end{matrix} \right\}, \quad (70.f)$$

przy czym wyrażenia na $\alpha_{y_3}, \alpha_{z_3}, \beta_{y_3}, \beta_{z_3} \dots$ otrzymuje się z $\alpha_{x_1}, \alpha_{y_1}, \beta_{x_1}, \beta_{y_1}$ — wzory (39), (42), (45), (48) po cyklicznym przestawieniu x, y, z . Podobnie z dal-
szymi kątami.

Rozwiązując układ równań (68.a)–(68.f) otrzymujemy

$$p_x^2 = \frac{1}{2\omega^2 q^2} (P_{s_1} + P_{s_2} - P_{s_3} - P_{s_4} + P_{s_5} + P_{s_6}) \quad (71.a)$$

$$p_y^2 = \frac{1}{2\omega^2 q^2} (P_{s_1} + P_{s_2} + P_{s_3} + P_{s_4} - P_{s_5} - P_{s_6}) \quad (71.b)$$

$$p_z^2 = \frac{1}{2\omega^2 q^2} (-P_{s_1} - P_{s_2} + P_{s_3} + P_{s_4} + P_{s_5} + P_{s_6}) \quad (71.c)$$

$$\cos(\psi_{p_x} - \psi_{p_y}) = \frac{1}{2\omega^2 q^2 p_x p_y} (P_{s_1} - P_{s_2}) \quad (72.a)$$

$$\cos(\psi_{p_y} - \psi_{p_x}) = \frac{1}{2\omega^2 q^2 p_y p_z} (P_{s_3} - P_{s_4}) \quad (72.b)$$

$$\cos(\psi_{p_y} - \psi_{p_z}) = \frac{1}{2\omega^2 q^2 p_z p_x} (P_{s_5} - P_{s_6}). \quad (72.c)$$

Rozwiązując układ równań (69.a)–(69.f) otrzymujemy

$$m_x^2 = \frac{1}{2k^2 q^2} (P_{d_1} + P_{d_2} - P_{d_3} - P_{d_4} + P_{d_5} + P_{d_6}) \quad (73.a)$$

$$m_y^2 = \frac{1}{2k^2 q^2} (P_{d_1} + P_{d_2} + P_{d_3} + P_{d_4} - P_{d_5} - P_{d_6}) \quad (73.b)$$

$$m_z^2 = \frac{1}{2k^2 q^2} (-P_{d_1} - P_{d_2} + P_{d_3} + P_{d_4} + P_{d_5} + P_{d_6}) \quad (73.c)$$

$$\cos(\psi_{m_x} - \psi_{m_y}) = \frac{1}{2k^2 q^2 m_x m_y} (P_{d_2} - P_{d_1}) \quad (74.a)$$

$$\cos(\psi_{m_y} - \psi_{m_x}) = \frac{1}{2k^2 q^2 m_x m_y} (P_{d_4} - P_{d_5}) \quad (74.b)$$

$$\cos(\psi_{m_z} - \psi_{m_x}) = \frac{1}{2k^2 q^2 m_x m_y} (P_{d_6} - P_{d_5}). \quad (74.c)$$

Uzyskane moduły i przesunięcia fazowe nie dają pełnej informacji o dipolach. Do wyznaczenia charakterystyki promieniowania układu dipoli w strefie dalekiej, w otwartej przestrzeni potrzebne są jeszcze przesunięcia fazowe $\psi_{p_i} - \psi_{m_j}$ między składowymi dipoli elektrycznych i magnetycznych, przy czym $i, j \in \{x, y, z\}$ ($i \neq j$) (patrz dodatek B). W celu wyznaczenia tych przesunięć należy, korzystając z wyników uzyskanych z wzorów (74)–(77) obliczyć kąty $\alpha_{x_1}, \alpha_{y_1}, \beta_{x_1}, \beta_{y_1}, \dots$ — wzory (39), (42), (45), (48),... Wstawiając te kąty oraz pomierzone fazy ϕ_1, ϕ_2, ϕ_3 do wzorów (70.a)–(70.f) możemy wyznaczyć brakujące przesunięcia fazowe $\psi_{p_i} - \psi_{m_j}$.

Ustalając procedurę pomiarową w komorze GTM trzeba posłużyć się wzorami (28). Wzory (29.a), (29.b) są w tym względzie nieprzydatne. Z analizy wzorów (28) wynika, że model identyfikowanego na ich podstawie EUT może składać się co najwyżej z trzech dipoli elektrycznych i trzech magnetycznych, przy czym wszystkie muszą mieć te same fazy początkowe. W takim przypadku należy przeprowadzić pomiar w sześciu ortogonalnych ustawieniach EUT. Jeśli EUT jest bryłą prostopadłościenną, należy wykonać pomiary ustawiając EUT każdą ze ścian w kierunku przyłącza kablowego. Jeśli z innych przesłanek wynika, że zadowalający jest model EUT składający się tylko z trzech dipoli elektrycznych o jednakowych fazach początkowych lub tylko trzech dipoli magnetycznych o jednakowych fazach początkowych, wystarczy dokonać pomiarów w trzech ortogonalnych ustawieniach.

4. WNIOSKI

Niniejszy artykuł jest rozszerzeniem pracy [6]. Artykuł [6] rozpoczyna się podaniem zależności (28). Dalej przedstawiona jest w nim procedura pomiarowa w komorze TEM. W centrum uwagi niniejszej pracy jest również komora GTM oraz różnice w procedurze pomiarowej dla obu komór. Zaprezentowanie wyprowadzenia wzorów (28) jest celowe, gdyż pozwala zrozumieć istotę procedury pomiarowej oraz dostrzec jej ograniczenia wynikające z założeń upraszczających takich jak: elektrycznie małe EUT, pole rodzaju TEM jako jedyne występujące w pustej komorze nieskończenie duża przewodność elektryczna ścian komory. W wyprowadzeniach tych tkwi też sposób określenia natężenia pola elektrycznego $\vec{E}_0 = q \vec{1}_y$ w punkcie umieszczenia EUT.

Testowane urządzenie w komorze TEM lub GTM musi być elektrycznie małe. Oznacza to, że jego największy wymiar musi być mniejszy od najkrótszej generowanej przez niego lub rozproszonej na nim fali oraz równocześnie wymiar ten musi być dużo mniejszy od poprzecznych i podłużnych wymiarów komory.

Jednym z etapów wyprowadzenia wzorów (28) jest rozwinięcie wyrażenia podcałkowego w całce po objętości we wzorze (17) w szereg Taylora i wzięcie pierwszych wyrazów rozwinięcia, odpowiadających dipolom elektrycznemu i magnetycznemu. Dlatego EUT musi być małe zarówno w stosunku do wymiarów poprzecznych jak i podłużnych komory. EUT sprzęga się nie tylko z rodzajem TEM, ale również z rodzajami wyższymi, które stanowią dla niego obciążenie reaktancyjne. Z drugiej strony zakłada się, że sygnał na przyłączach kablowych jest tylko rodzaju TEM. Założenie to jest słuszne, jeśli pole TEM jest jedynym rodzajem występującym w pustej komorze oraz, jeśli przyłącza kablowe są na tyle daleko od EUT, aby wyższe rodzaje pól zostały wytłumione zanim dotrą do przyłączy. Z założenia pola TEM jako jedynego rodzaju występującego w pustej komorze wynika zakres częstotliwości roboczej komory.

Wymóg nieskończonej przewodności płaszcza komory nie jest oczywisty. Niech świadczy o tym artykuł [7]. Autor porusza w nim możliwość wykorzystania komory GTEM do pomiaru emisji zakłóceń. W celu uniknięcia odbić od ścian komory i związanych z tym błędów pomiaru, proponuje wyłożyć ściany komory materiałem absorbcyjnym. Nie można tak zrobić. Wtedy bowiem stosując wzór (10) do obiektu przedstawionego na rys. 4 nie można założyć, że całka po powierzchni płaszcza komory jest zerem. Nie uzyska się wtedy wzoru (28). W komorze GTEM, w której tylko ściana naprzeciw przyłącza kablowego jest wyłożona materiałem absorbcyjnym uzyskuje się wzór (28) dzięki spełnieniu na niej warunku normalizacji (wzór (6)).

Zgodnie z warunkiem normalizacji, natężenie pola elektrycznego $\vec{E}_0 = q \vec{1}_y$ w punkcie umieszczenia EUT, występujące we wzorach (28), należy określić zakładając jednostkową moc (1 [W]) przepływającą przez komorę. $\vec{E}_0$ można uzyskać na drodze symulacji numerycznej lub z pomiaru.

Wszystkie wyszczególnione tu założenia upraszczające wpływają na wynik pomiaru. Problem dokładności pomiaru wymaga osobnej prezentacji gdyż jest bardzo obszerny. Raport NBS na ten temat [11] obejmuje 43 strony.

Identyfikowanie EUT jako układu trzech dipoli elektrycznych o różnych fazach początkowych i trzech magnetycznych o różnych fazach początkowych wymaga równoczesnego mierzenia dwu sygnałów w sześciu ortogonalnych ustawieniach EUT. W komorze TEM, jako dwuwrotniku jest to możliwe. Komora GTEM jest jednowrotnikiem. Dlatego nie daje ona możliwości uzyskania informacji o przesunięciach fazowych poszczególnych dipoli modelujących EUT. W publikacjach dotyczących pomiarów emisji w GTEM [4], [3], [12] mówi się o pomiarze tylko w trzech ortogonalnych ustawieniach utrzymując, że otrzymany w ten sposób model odpowiada układowi trzech dipoli elektrycznych i trzech magnetycznych o identycznych fazach początkowych.

Można zgodzić się z uzasadnieniem że emisja promieniowania EUT jest zjawiskiem ubocznym i dlatego nie należy spodziewać się dużych różnic faz po-

czątkowych poszczególnych dipoli modelujących EUT. Jeśli nawet one są, to ich pominięcie jest założeniem najbardziej niekorzystnego przypadku (worst case). Natomiast z całą pewnością identyfikacja taka nie doprowadzi do modelu składającego się z trzech dipoli elektrycznych i magnetycznych. W istocie, w komorze GTEM na podstawie pomiarów w trzech ortogonalnych ustawieniach identyfikuje się EUT jako jeden dipol elektryczny albo magnetyczny. Pomiarów te wystarczają tylko do określenia jego orientacji w przestrzeni i pozwalają wyznaczyć charakterystykę promieniowania, zakładając „worst case” tzn., że zysk kierunkowy EUT jako radiatora nie może być większy niż dipola, którym EUT jest modelowane.

A. MOMENT DIPOLA ELEKTRYCZNEGO I MAGNETYCZNEGO

Do opisu pól w dielektrykach używa się wielkości zwanej elektrycznym momentem dipolowym. Wyraża się on wzorem

$$\vec{p} = \iiint_V \vec{P} dV, \quad (\text{A.1})$$

przy czym $\vec{P}$ jest wektorem polaryzacji występującym we wzorze $\vec{D} = \epsilon_0 \vec{E} + \vec{P}$. Wektor ten jest związany z gęstością ładunku spolaryzowanego (dipoli), tak jak wektor $\vec{D}$ z gęstością ładunków swobodnych.

$$\text{div } \vec{P} = -\rho_p(\vec{r}) \quad P_n = \sigma_p. \quad (\text{A.2})$$

Gęstość objętościowa ładunku spolaryzowanego $\rho_p(\vec{r})$ istnieje tylko w niejednorodnym dielektryku, lub mówiąc inaczej gdy $\text{div } \vec{P} = -\rho_p(\vec{r}) \neq 0$. Jest to wynikiem niecałkowitego znoszenia się pól dipoli sąsiadujących ze sobą. Ponadto źródła pola wektora $\vec{P}$ występują wszędzie tam, gdzie składowa normalna wektora $\vec{P}$ zmienia się skokowo. Ma to miejsce na powierzchni zewnętrznej dielektryka oraz na powierzchni styku z innymi dielektrykami. Za tę nieciągłość wektora odpowiedzialny jest spolaryzowany ładunek powierzchniowy o gęstości σ_p . W dielektryku spolaryzowanym jednorodnie źródło wektora $\vec{P}$ znajduje się jedynie na jego powierzchni. Tam składowa normalna wektora $\vec{P}$ zmienia się skokowo o wartość σ_p .

Jeśli x_1, x_2, x_3 są współrzędnymi punktu w przestrzeni trójwymiarowej to składową x_i momentu dipolowego można wyznaczyć posługując się tożsamością różniczkową

$$\text{div}(x_i \vec{P}) = x_i \text{div } \vec{P} + P_{x_i}. \quad (\text{A.3})$$

Dzięki niej

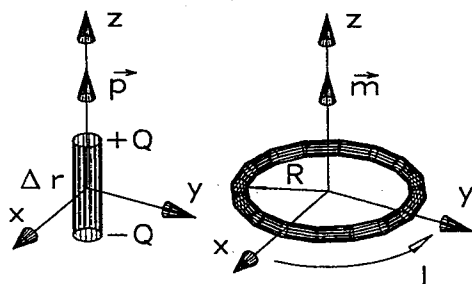
$$p_{x_i} = \iiint_V P_{x_i} dV = - \iiint_V x_i \operatorname{div} \vec{P} dV + \iiint_V \operatorname{div}(x_i \vec{P}) dV. \quad (\text{A.4})$$

Stosując prawo Gaussa do równ. (A.4) otrzymujemy

$$p_{x_i} = - \iiint_V x_i \operatorname{div} \vec{P} dV + \iint_S x_i \vec{P} \cdot \vec{1}_n dS \quad (\text{A.5})$$

Zamiast wektora $\vec{P}$ można w równ. (A.5) użyć gęstości ładunku spolaryzowanego, zgodnie ze wzorem (A.2). Dla pary ładunków $\pm Q$ oddalonych od siebie od Δr wzór (A.5) przyjmuje postać

$$\vec{p} = \lim_{\substack{Q \rightarrow \infty \\ \Delta r \rightarrow 0}} Q \Delta r \vec{1}_r \quad (\text{A.6})$$



Rys. A.1 Definicja momentu elektrycznego, magnetycznego

Do opisu pól w magnetykach używa się wielkości zwanej momentem magnetycznym. Moment magnetyczny wyraża się wzorem

$$m = \iiint_V \vec{M} dV \quad (\text{A.7})$$

przy czym $\vec{M}$ jest wektorem magnetyzacji, występującym we wzorze $\vec{B} = \mu_0(\vec{H} + \vec{M})$. Wektor ten ma bezpośredni związek z prądem magnetyzacji $\vec{j}_m$

$$\vec{j}_m = \operatorname{rot} \vec{M} \quad \vec{\tau}_m = \vec{M} \times \vec{1}_n. \quad (\text{A.8})$$

Prąd magnetyzacji $\vec{j}_m$, zwany również prądem atomowym jest to prąd płynący wewnątrz materiału na skutek tego, że pola spinów elektronów na orbitach atomowych nie kompensują się całkowicie, dając przyczyną lokalnym prądom.

Na granicy między dwoma różnie namagnesowanymi materiałami, jak również na granicy materiału namagnesowanego z powietrzem wektor $\vec{M}$ zmienia się skokowo o wartość ($\vec{\tau}_m = \vec{M} \times \vec{1}_n$) tzn. o gęstość powierzchniowego prądu magnetyzacji. Moment magnetyczny można zapisać następująco przy pomocy prądu magnetyzacji

$$\vec{m} = \frac{1}{2} \iiint_V (\vec{r} \times \vec{j}_m) dV = \frac{1}{2} \iiint_V (\vec{r} \times \text{rot } \vec{M}) dV + \frac{1}{2} \oint_S (\vec{r} \times \vec{M} \times \vec{1}_n) dS. \quad (\text{A.9})$$

Jeśli środowisko jest namagnesowane jednorodnie, wszystkie prądy atomowe w jego wnętrzu kompensują się ($\vec{j}_m = \text{rot } \vec{M} = \vec{0}$). Istnieją wówczas tylko prądy powierzchniowe ($\vec{\tau}_m$). Często powierzchnię zewnętrzną takiego materiału nazywa się skorupą magnetyczną (magnetic shell). Dla pętli w kształcie okręgu o promieniu r z prądem I moment magnetyczny sprowadza się do postaci

$$\vec{m} = \lim_{\substack{I \rightarrow \infty \\ \pi r^2 \rightarrow 0}} \pi r^2 I \vec{1}_r \times \vec{1}_j. \quad (\text{A.10})$$

B. CHARAKTERYSTYKA PROMIENIOWANIA UKŁADU DIPOLI W OTWARTEJ PRZESTRZENI

Zanim przedstawimy charakterystykę promieniowania mocy składu dipoli przypomnijmy wzory na natężenie pola elektrycznego i magnetycznego w strefie dalekiej dipola elektrycznego

$$\vec{E} = -\frac{\omega k_0 \zeta_0}{4\pi} \mathbf{p}_z \sin \theta \frac{e^{ik_0 r}}{r} \vec{1}_\theta \quad (\text{B.1.a})$$

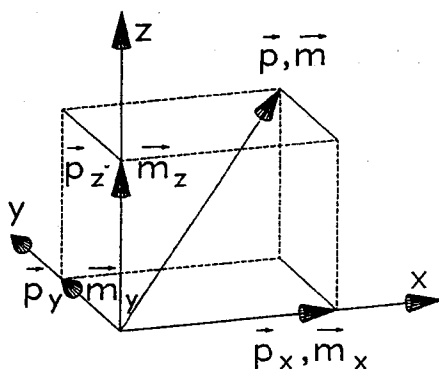
$$\vec{H} = -\frac{\omega k_0}{4\pi} \mathbf{p}_x \sin \theta \frac{e^{ik_0 r}}{r} \vec{1}_\phi \quad (\text{B.1.b})$$

i magnetycznego

$$\vec{E} = \frac{\omega \mu_0 k_0}{4\pi} \mathbf{m}_z \sin \theta \frac{e^{ik_0 r}}{r} \vec{1}_\phi \quad (\text{B.2.a})$$

$$\vec{H} = -\frac{\omega \mu_0 k_0}{4\pi \zeta_0} \mathbf{m}_z \sin \theta \frac{e^{ik_0 r}}{r} \vec{1}_\theta. \quad (\text{B.2.b})$$

Wzory te opisują pole dipoli, których momenty są zorientowane w kierunku z tzn. mają tylko składową z . Stąd indeksy $\mathbf{p}_z, \mathbf{m}_z$.



Rys. B.1 Orientacja dipoli w trzech ortogonalnych kierunkach

Dla dipoli zorientowanych w kierunku x wprowadzamy współrzędne r', θ', ϕ' oraz r'', θ'', ϕ'' dla zorientowanych w kierunku y . Wykorzystując aparat trygonometryczny otrzymujemy

$$\vec{1}_{r'} = \vec{1}_r \quad (\text{B.3.a})$$

$$\vec{1}_{\theta'} = -\frac{\cos \phi \cos \theta}{\sqrt{1 - \sin^2 \theta \cos^2 \phi}} \vec{1}_{\theta} + \frac{\sin \phi}{\sqrt{1 - \sin^2 \theta \cos^2 \phi}} \vec{1}_{\phi} \quad (\text{B.3.b})$$

$$\vec{1}_{\phi'} = -\frac{\sin \phi}{\sqrt{1 - \sin^2 \theta \cos^2 \phi}} \vec{1}_{\theta} - \frac{\cos \phi \cos \theta}{\sqrt{1 - \sin^2 \theta \cos^2 \phi}} \vec{1}_{\phi} \quad (\text{B.3.c})$$

$$\vec{1}_{r''} = \vec{1}_r \quad (\text{B.3.d})$$

$$\vec{1}_{\theta''} = -\frac{\sin \phi \cos \theta}{\sqrt{1 - \sin^2 \theta \sin^2 \phi}} \vec{1}_{\theta} - \frac{\cos \phi}{\sqrt{1 - \sin^2 \theta \sin^2 \phi}} \vec{1}_{\phi} \quad (\text{B.3.e})$$

$$\vec{1}_{\phi''} = \frac{\cos \phi}{\sqrt{1 - \sin^2 \theta \sin^2 \phi}} \vec{1}_{\theta} - \frac{\sin \phi \cos \theta}{\sqrt{1 - \sin^2 \theta \sin^2 \phi}} \vec{1}_{\phi} \quad (\text{B.3.f})$$

oraz

$$\sin \theta' = \sqrt{1 - \sin^2 \theta \cos^2 \phi}, \quad (\text{B.4.a})$$

$$\sin \theta'' = \sqrt{1 - \sin^2 \theta \sin^2 \phi} \quad (\text{B.4.b})$$

Wypadkowe pole elektryczne wyraża się więc następującym wzorem

$$\begin{aligned} \vec{E}(r, \theta, \phi) = & E_p \sin \theta \vec{1}_{\theta} + E'_p \sin \theta' \vec{1}_{\theta'} + E''_p \sin \theta'' \vec{1}_{\theta''} + E_m \sin \theta \vec{1}_{\phi} + \\ & + E'_m \sin \theta' \vec{1}_{\phi'} + E''_m \sin \theta'' \vec{1}_{\phi''} \quad (\text{B.5}) \end{aligned}$$

co po przekształceniach daje

$$\begin{aligned}\vec{E}(r, \theta, \phi) = & (\mathbf{E}_p \sin \theta - \mathbf{E}'_p \cos \phi \cos \theta - \mathbf{E}''_p \sin \phi \cos \theta + \\ & - \mathbf{E}'_m \sin \phi + \mathbf{E}''_m \cos \phi) \vec{1}_\theta + (\mathbf{E}'_p \sin \phi - \mathbf{E}''_p \cos \phi + \\ & + \mathbf{E}_m \sin \theta - \mathbf{E}'_m \cos \theta \cos \phi - \mathbf{E}''_m \cos \theta \sin \phi) \vec{1}_\phi.\end{aligned}\quad (\text{B.6})$$

Podobnie wypadkowe pole magnetyczne

$$\begin{aligned}\vec{H}(r, \theta, \phi) = & \mathbf{H}_p \sin \theta \vec{1}_\phi + \mathbf{H}'_p \sin \theta' \vec{1}_{\phi'} + \mathbf{H}''_p \sin \theta'' \vec{1}_{\phi''} + \\ & + \mathbf{H}_m \sin \theta \vec{1}_\theta + \mathbf{H}'_m \sin \theta' \vec{1}_{\theta'} + \mathbf{H}''_m \sin \theta'' \vec{1}_{\theta''}\end{aligned}\quad (\text{B.7})$$

po przekształceniach wyraża się następująco

$$\begin{aligned}\vec{H}(r, \theta, \phi) = & (-\mathbf{H}'_p \sin \phi + \mathbf{H}''_p \cos \phi + \mathbf{H}_m \sin \theta - \mathbf{H}'_m \cos \phi \cos \theta - \mathbf{H}''_m \sin \phi \cos \theta) \vec{1}_\theta + \\ & + (\mathbf{H}_p \sin \theta - \mathbf{H}'_p \cos \theta \cos \phi - \mathbf{H}''_p \cos \theta \sin \phi + \mathbf{H}'_m \sin \phi - \mathbf{H}''_m \cos \phi) \vec{1}_\phi.\end{aligned}\quad (\text{B.8})$$

Jeśli przyjąć, że przez $\vec{E}$ i $\vec{H}$ rozumiemy amplitudy zespolone, a nie wartości skuteczne, to gęstość mocy wypromieniowanej wyraża się wzorem

$$P(\theta, \phi) = \frac{1}{2} (\text{Re} \{ \vec{E} \times \vec{H}^* \}) \cdot \vec{1}_r \quad (\text{B.9})$$

W przypadku układu dipoli, z jakim mamy do czynienia wzór ten przyjmuje postać

$$P(\theta, \phi) = \frac{1}{2} \text{Re} \{ \mathbf{E}_\theta \mathbf{H}_\phi^* - \mathbf{E}_\phi \mathbf{H}_\theta^* \}. \quad (\text{B.10})$$

Wykorzystując zależności $k_0^2 = \omega^2 \mu_0 \varepsilon_0$, $\lambda k_0 = 2\pi$ oraz $\zeta_0 = 120\pi$ otrzymujemy

$$\begin{aligned}P(r, \theta, \phi) = & \frac{15\pi}{\lambda^2 r^2} \{ (\omega^2 p_x^2 + k_0^2 m_x^2) (\cos^2 \theta \cos^2 \phi + \sin^2 \phi) + \\ & + (\omega^2 p_y^2 + k_0^2 m_y^2) (\cos^2 \theta \sin^2 \phi + \cos^2 \phi) + (\omega^2 p_z^2 + k_0^2 m_z^2) \sin^2 \theta - \\ & + 2[\omega^2 p_x p_y \cos(\psi_{p_x} - \psi_{p_y}) + k_0^2 m_x m_y \cos(\psi_{m_x} - \psi_{m_y})] \sin^2 \theta \sin \phi \cos \phi + \\ & - 2[\omega^2 p_y p_z \cos(\psi_{p_y} - \psi_{p_z}) + k_0^2 m_y m_z \cos(\psi_{m_y} - \psi_{m_z})] \sin \theta \cos \theta \sin \phi + \\ & - 2[\omega^2 p_z p_x \cos(\psi_{p_z} - \psi_{p_x}) + k_0^2 m_z m_x \cos(\psi_{m_z} - \psi_{m_x})] \sin \theta \cos \theta \cos \phi + \\ & + 2\omega k_0 [p_x m_y \cos(\psi_{p_x} - \psi_{m_y}) - p_y m_x \cos(\psi_{p_y} - \psi_{m_x})] \cos \theta + \\ & + 2\omega k_0 [p_y m_z \cos(\psi_{p_y} - \psi_{m_z}) - p_z m_y \cos(\psi_{p_z} - \psi_{m_y})] \sin \theta \cos \phi + \\ & + 2\omega k_0 [p_z m_x \cos(\psi_{p_z} - \psi_{m_x}) - p_x m_z \cos(\psi_{p_x} - \psi_{m_z})] \sin \theta \sin \phi \} \end{aligned}\quad (\text{B.11})$$

Z zasady zachowania ładunku wynika związek $I = -\omega Q$. Oznacza to, że między fazą początkową ładunku ψ^Q , a fazą początkową prądu ψ^I zachodzi związek

$$\psi^Q = \psi^I - 90^\circ \quad (\text{B.12})$$

gdyż $i = -j$ jest operatorem obrotu o kąt -90° . Ten sam związek (B.12) istnieje między momentem dipola zdefiniowanym z użyciem ładunku p^Q lub prądu p^I . Aby zatem uzyskać charakterystykę promieniowania z użyciem p^I należy do wzoru (B.11) wstawić zależność (B.12). Zmieni to charakterystykę tylko w tej części, w której występują czynniki mieszane p z m i to tak, że występujące przy nich funkcje cos zostaną zastąpione funkcjami sin.

BIBLIOGRAFIA

1. R.E. Collin: FIELD THEORY OF GUIDED WAVES, IEEE Press, 1991
2. M.L. Crawford: *Generation of standard EM fields using TEM transmission cells*. IEEE Trans. on EMC, 1974, Vol. EMC-16, No. 4
3. H. Garbe: TEM-Zelle in der EMV-Messtechnik — Ein Überblick. ELEKTROMAGNETISCHE VERTRÄGLICHKEIT, 3. Int. Fachmesse und Kongress für Elektromagnetische Verträglichkeit EMV'92, VDE-Verl. 1992, pp. 223–242
4. D. Hansen, P. Wilson, D. Königstein, H. Scharer: *A broadband alternative EMC test chamber based on a TEM-cell anechoic-chamber hybrid concept*. Int. Symp. on Electromagnetic compatibility, Nagoya 1989, pp. 133–137
5. J.D. Jackson: *Classical electrodynamics*, John Wiley & Sons, 1975
6. G.H. Köpke: *A new method for determining the emission characteristics of an unknown interference source*. EMC-Zürich'83, 1983, pp. 35–40
7. J. Nedtwig: *Vermeidung stoerender Vielfachreflexionen in TEM-Wellenleitern: GTEM → TRIGATEM*. 3. Int. Fachmesse und Kongress fuer Elektromagnetische Vertraeglichkeit EMV'92, VDE-Verlag GMBH Berlin, Offenbach, 1992, pp. 243–250
8. W.K. Panofsky, M. Philips: *Classical electricity and magnetism*, Addison-Wesley Publishing Company, Reading, Massachusetts, 1962
9. I. Sreenivasiah, D.C. Chang, M.T. Ma: *Emission characteristic of electrically small radiating sources from tests inside a TEM cell*. IEEE Trans. on EMC, 1981, Vol. EMC-23, No. 3
10. J.C. Tipper, D.C. Chang: *Radiation characteristic of electrically small devices in a TEM transmission cell*. IEEE Trans. on EMC, 1976, Vol. EMC-18, No. 4
11. *Uncertainties in extracting radiation parameters for an unknown interference source based on power and phase measurements*. U.S. Department of Commerce/National Bureau of Standards, NBS Technical Note 1964
12. P. Wilson, D. Hansen, D. Königstein: *Stimulating open area test site emission measurements based on data obtained in a novel broadband TEM cell*. International IEEE. Symposium on EMC, Denver, USA, 1989, pp. 171–177

USING TEM AND GTEM CELL IN RADIATION EMISSION TESTS

J. SROKA

Summary

Radiation emission tests deliver basic data about electromagnetic compatibility (EMC) of electronic devices. According to the european standards concerning radiation emission, measurements have to be performed in open area (OATS — Open Area Test Site). The area must be free of any electromagnetic disturbances. The same conditions can be guaranteed in the anechoic chamber. The first measurement method is troublesome, the second expensive. TEM and GTEM cells convey alternative, cheaper method of measuring the radiation emission. In the alternative measurements the equipment under test (EUT) has to be identified and for the identified model the radiation pattern in free space must be derived.

The TEM cell is the segment of a waveguide with enlarged space between inner wire, called septum and the metal coat. It has two cable connections, so it is the twoport. The GTEM cell looks like a piece of the TEM cell with one cable connection and the taper ended with the absorber wall. Advantage of the GTEM cell is the frequency range. For the small GTEM cells it lays over 1 GHz versus some hundred MHz in the TEM cell of the same size. Köpke [6] presented the identification of the EUT as the set of three electric and three magnetic dipoles, each with different phase. His paper begins with the relations between signals measured at the cable connections and the electric dipole moments that radiate in the TEM cell. In this paper the derivation of the formulae, Köpke begins with is performed. Then follows description of the measurement procedure as in [6]. Attention is paid on deviation of the GTEM from the TEM cell. In simplifications that are assumed by derivation of the above mentioned formulae the limitations of the procedure are hidden. Some authors forget about it broadening mistaken statements [3], [7]. It is discussed in the conclusions. In order to identify the EUT as the set of three electric the three magnetic dipoles each with different phase, simultaneous measurement of two signals in six orthogonal orientations are required. The TEM cell as a twoport unlike the GREM (oneport) makes it possible.

It is maintained in [4], [3], [12] that in order to identify the EUT in the GTEM as the set of three electric and three magnetic dipoles with the same phases, measurements in three orthogonal orientation suffice. Actually such measurement build up the model consisted of one electric or magnetic dipole. The paper shows that TEM cell guarantees better precision in the EUT identification.

SPIS TREŚCI

W. Niemiec: Studium porównawcze konstytutywnej i Maxwell'owskiej teorii pola magnetycznego. Część I: Porównanie konstytutywnego rozłożonego parametrycznie modelu do Maxwell'owskiego modelu pola elektrycznego. Część II: Porównanie zagadnień analitycznych rozwiązań konstytutywnego rozłożonego parametrycznie modelu do Maxwell'owskiego modelu pola elektrycznego	149
F. Wysocka-Schillak: Projektowanie cyfrowego filtra SOI o liniowej fazie z uwzględnieniem wymagań dotyczących jego odpowiedzi jednostkowej	185
S. Kuta, M. Machowski, J. Jasielski: Ulepszone rozwiązanie układowe konweyżorów prądowych w technice CMOS	201
J. Gorzkowski: Analiza i projektowanie falowodowego dzielnika mocy wytwarzającego sygnały wyjściowe o zadanych proporcjach amplitud	215
Z. Surowiak: Ceramiczne ferroelektryki perowskitowe o półprzewodnikowym i metalicznym przewodnictwie elektrycznym	233
W. Pawelski: Układowo-zorientowany model statyczny IGBT	249
W. Pawelski: Analiza wrażliwości nowego modelu statycznego IGBT	265
A. Gabryelczyk, M. Grecki, A. Napieralski, W. Pawelski: Tester charakterystyk statycznych IGBT	277
J. Sroka: Wykorzystanie komory TEM i GTM w pomiarach emisji zakłóceń radiowych	291

CONTENTS

W. Niemiec: The comparative study of the constitutive and Maxwell's electromagnetic field theories. Part I: The comparison of the constitutive distributed parameter model to the Maxwell's model of the electromagnetic field. Part II: The comparison of the analytical solution problems of the constitutive distributed parameter model to the Maxwell's model of the electromagnetic field	149
F. Wysocka-Schillak: Design of a FIR linear phase digital filter including requirements due to the unit step response	185
S. Kuta, W. Machowski, J. Jasielski: Improved CMOS implementation of current conveyors	201
J. Gorzkowski: Analysis and design of waveguide divider forming output signals with amplitudes which proportion are given	215
Z. Surowiak, E.G. Fesenko, D. Krawczyk: The perovskite-type ceramic ferroelectric materials characterized by semiconductive and metallic electric conduction	233
W. Pawelski: Circuit-oriented static model of the IGBT	249
W. Pawelski: The sensitivity analysis of new IGBT static model	265
A. Gabryelczyk, M. Grecki, Z. Kuśmierek, A. Napieralski, W. Pawelski: Static characteristics tester of the IGBT	277
J. Sroka: Using TEM and GTM cell in radiation emission tests	291

Kwartalnik Elektroniki i Telekomunikacji

WARUNKI PRENUMERATY

Wpłaty na prenumeratę przyjmowane są na okresy roczne

Na terenie kraju prenumeratę przyjmują jednostki kolportażowe przedsiębiorstwa RUCH S.A., urzędy pocztowe oddawcze właściwe dla miejsca zamieszkania lub siedziby prenumeratora oraz doręczyciele w miejscowościach, gdzie dostęp do urzędu jest utrudniony, a także Zakład Kolportażu Wydawnictwa SIGMA - NOT.

Termin przyjmowania prenumeraty rocznej krajowej i na zagranicę przez RUCH S.A. i Zakład Kolportażu Wydawnictwa SIGMA - NOT
do 20 XI roku poprzedzającego

Termin przyjmowania prenumeraty rocznej krajowej przez Pocztcę Polską
do 25 XI roku poprzedzającego

Dostawa zamówionej prasy następuje w sposób uzgodniony z zamawiającym, pod wskazanym adresem, w ramach opłaconej prenumeraty.

Adresy przedsiębiorstw kolportażowych i ich konta bankowe:

Zakład Kolportażu Wydawnictwa SIGMA-NOT, 00-716 Warszawa, ul. Bartycka 20, skr.poczt. 1004

Konto PBK III Oddział Warszawa, nr 370015-1573-139-11

Prenumerata ze zleceniem wysyłki za granicę jest dwukrotnie wyższa od ceny krajowej. Zlecający powinien podać dokładny adres odbiorcy. Dodatkowe informacje można otrzymać pod nr telefonu 40-30-86 lub 40-00-21 wew. 249, 295, 299.

RUCH S.A. Oddział Warszawa, 00-958 Warszawa, ul. Towarowa 28

Konto PBK XIII Oddział Warszawa nr 370044-1195-139-11

Dostawa odbywa się przesyłką zwykłą w ramach opłaconej prenumeraty, z wyjątkiem zlecenia dostawy pocztą lotniczą, której koszt w pełni pokrywa zleceniodawca.

Prenumerata ze zleceniem dostawy za granicę jest o 100% wyższa od krajowej. Od instytucji lub osób zamieszkałych w miejscowościach, w których nie ma jednostek kolportażowych RUCH, prenumeratę Kwartalnika można zamówić w Oddziale Warszawskim RUCH.

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

W roku 1995

przewidywana cena 1 egz. wyniesie 50.000 zł
prenumerata roczna w kraju 200.000 zł
prenumerata roczna za granicę 90 \$ USA

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

Indeks 36318
PL ISSN 0867-674

**KWARTALNIK
ELEKTRONIKI I TELEKOMUNIKACJI**

**ELECTRONICS AND
TELECOMMUNICATIONS
QUARTERLY**

TOM XL — ZESZYT 3

WYDAWNICTWO NAUKOWE PWN
WARSZAWA 1994

Subscription for external subscribers:

The promotional subscription price in 1995 is \$90 including postage for institutions. Subscriptions should be sent to the publisher, Polish Scientific Publishers PWN Ltd, Journal Division, Miodowa 10, 00-251 Warsaw, POLAND, fax (48) (22) 26 09 50, (48) (22) 26 71 63, with a copy by fast to Electronics and Telecommunications Quarterly 00-665 Warsaw, ul. Nowowiejska 15/19, p. 470. Subscription is occupied after showing a cheque or transfer documents. Our bank account is as follows: Bank account: PBK VIII O/Warszawa nr 370028-1052. At subscriber's request this journal will be air mailed at additional postage to 50% dross price to European countries and 65% overseas.

Subscriptions orders available through the local press distribution or through the Foreign Trade Enterprise ARS Polona, 00-068 Warszawa, Krakowskie Przedmieście 7, Poland.

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

KWARTALNIK
ELEKTRONIKI I TELEKOMUNIKACJI
ELECTRONICS AND
TELECOMMUNICATIONS
QUARTERLY

TOM XL — ZESZYT 3

WYDAWNICTWO NAUKOWE PWN

WARSZAWA 1994

RADA REDAKCYJNA

Przewodniczący

prof. dr inż. ADAM SMOLIŃSKI
członek rzeczywisty PAN

Członkowie

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

REDAKCJA

Redaktor Naczelny

prof. dr hab. inż. WIESŁAW WOLIŃSKI

Zastępca Redaktora Naczelnego

doc. dr inż. KRYSTYN PLEWKO

Sekretarz Odpowiedzialny

mgr KRYSTYNA LELAKOWSKA

Wydanie kwartalnika sponsoruje „Spółka Telekomunikacja Polska S.A.”

ADRES REDAKCJI

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

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

tel. 628 89 81 21007-737

Telefony domowe: Redaktora Naczelnego: 12 17 65

Zast. Red. Naczelnego: 26 83 41

Sekretarza Odpowiedzialnego: 25 29 18

WYDAWNICTWO NAUKOWE PWN

Warszawa, ul. Miodowa 10

Ark. wyd. 10.25 Ark. druk. 8.25	Podpisano do druku w grudniu 1994 r.
Papier offsetowy kl. III 70 g. B-1	Druk ukończono w styczniu 1995 r.

Skład: Wyd. **GP-BIT** W-wa ul. Marymoncka 34.

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

THRESHOLD PHOTOEMISSION FROM SEMICONDUCTING REAL SURFACES

JÓZEF WOJAS

Profesor emerytowany, Politechnika Wrszawska

Received 1994.05.4

Authorized 1994.08.13

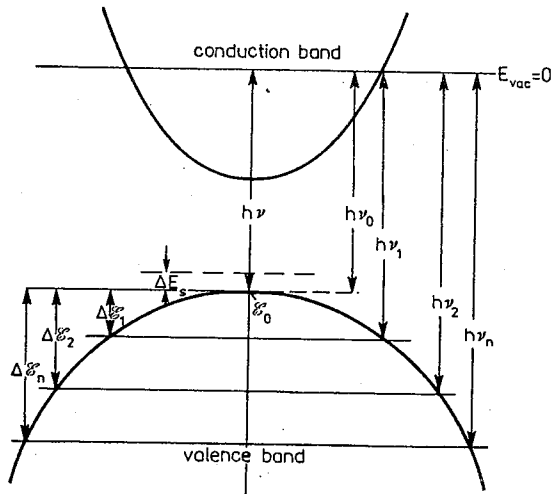
This paper shows the occurrence of over-threshold emission (over-valence band). Particularly it concerns the real surface of the semiconductors. Using the quantum yield (photoemission) measurements data, the quantative calculation of the emission have been carried out and a conclusion has been reached (confirmed by analysis energetic distributions) to occur above threshold photoemissions. The good agreement between measured data and the theory has been obtained.

1. INTRODUCTION

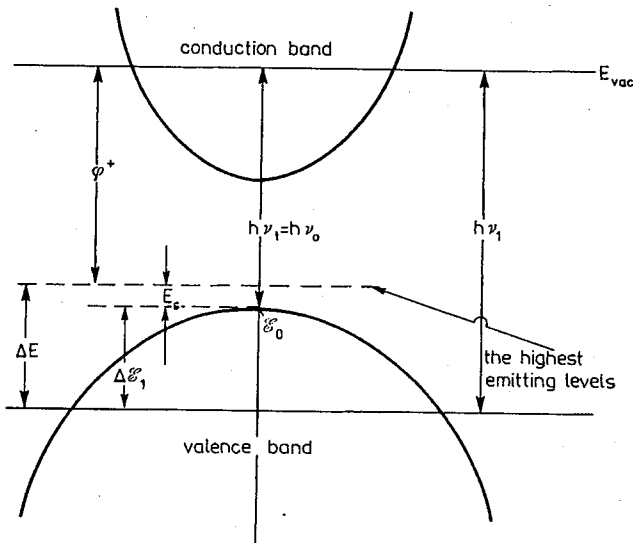
After a semiconductor sample is subjected to mechanical grinding and polishing and chemical etching a real semiconductor surface is created [1, 2]. Since it usually contacts the atmosphere of air, the neutral technical gas or the vacuum which is of the order of 10^{-3} Pa in typical circumstances, it is covered by a thin film of native oxide of the complex physical and chemical structure. This influences the surface electronic effects [3 – 5] including the photoemission. It results the occurrence of over-threshold (surface) emission. Thus, during photoemission process, we met the addind to the emission from the original surface clean in atomic sense the emission from the above mentioned film, changing as a result of it some emissive parameters (the classical model of photoemission). This paper discuss the reason and mechanisms of appearance of the over-valence emission.

2. Theoretical considerations

The energetic diagrams of photoelectron emission mechanism are depicted in Figs 1 – 3. Fig. 1 presents (in a schematic way) the dependence of the penetration depth in an energy sense of radiation beam $h\nu$ on its energy; however Fig. 2 depicts



Rys.1. Schematic of photomission energetic depth against the exciting photon energy



Rys.2. Schematic of energetic mechanism of valence and over-threshold emission

the such schematic of energetic mechanism of the valence band and overthreshold photoemission for the all spectrum

$$h\nu_0 - h\nu_t = \Delta E_0 = 0$$

$$h\nu_1 - h\nu_t = \Delta \mathcal{E}_1$$

$$h\nu_2 - h\nu_t = \Delta \mathcal{E}_2$$

(1)

$$h\nu_n - h\nu_t = \Delta\mathcal{E}_n \quad (\text{cd } 1)$$

$$\text{and } \Delta E_s + \Delta\mathcal{E}_1 = \Delta E$$

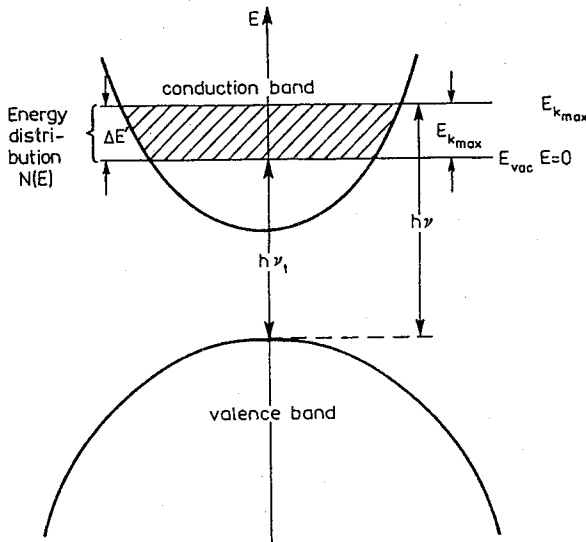
where $\Delta\mathcal{E}_1$ is the penetration depth in an energy sense of the light probe, ΔE_s is the over-valence emission, and ΔE – the width (the spread) of entire energy distribution (including the over-valence part).

The emitting part of the surface state band E_s one can find from the relations shown in Fig. 2 where $\Delta\mathcal{E}_1$ is the emitting part of the valence bands, and E_s is the emitting over-valence band (the surface state bands):

$$\begin{aligned} \varphi^+ + E_s &= h\nu_t \\ \Delta\mathcal{E}_1 + E_s &= \Delta E \\ (\varphi^+ &= h\nu_1 - \Delta E) \\ E_s &= \Delta E - \Delta\mathcal{E}_1 \end{aligned} \quad (2)$$

where φ^+ is the work function of the fastest electrons, measured with respect to the highest filled level. From Fig. 3 it can be found, that the velocity spectrum of electrons

$$E_{k\max} = \Delta E'.$$



Rys.3. Schematic of energy spectrum of the valence-band electron distribution

The „length” $h\nu$ in this model can be “placed” arbitrary,

$$h\nu - h\nu_t = \Delta E' = E_{k\max}$$

(It determines photoelectrons, which have escaped over E_{vac}), it would be so, in the case when only valence emission appears, but in practice also over-threshold emission

appears. This can be seen from Fig. 2. Here ΔE is a real width of the energetic distribution $N(E)$ of electrons. When $\Delta E > \Delta \mathcal{E}_1$, then some electrons flow from the levels over the top of the valence band; they create the over-threshold emission.

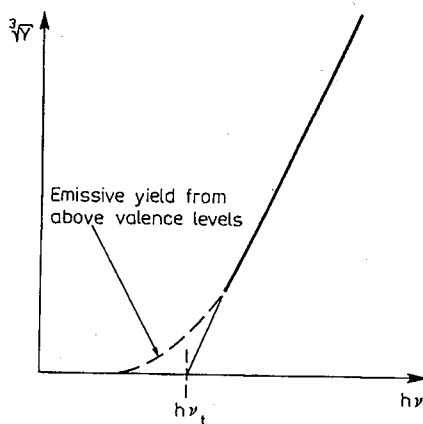
When energy $h\nu$ increases, then distribution width ΔE of electrons has to increase, since $h\nu_i = \text{const}$. Energy $\Delta \mathcal{E}_1$ should be equal to energetic spectrum width of electrons if there were no over-threshold emission.

The energy of photoemission threshold $h\nu_i$ from the ideally clean surface is located at the valence band top. Inside the energy gap E_G the energy levels (called the local levels) are composed not only of the surface states, but also are created by the bulk defects and by some oxides localized on the surface [1]. Different types of surface contaminations result the photemission called the over-threshold or over-valence one. In this case emitted are the electrons from the populated levels localized at the bottom part of the energy gap.

The method of quantitative assessment of the over-valence emission part in the entire spectrum of the emitted electrons has been elaborated by Seroczyńska [6].

The over-threshold photoemission can be determined [7, 8] by the rectilinear extrapolation at high-energy side on the curve of energy distribution of photoelectrons. The extrapolation line divides the area under the curve on two parts: the over-threshold emission and the basic emission from the valence band.

The author's experimental results from the last several years show the existence of small emission from some semiconductors at photon energy less than $h\nu_i$ [9]. The author and Seroczyńska [10, 11] concluded from the experimental points located at the left side of the line $Y = k\sqrt[3]{h\nu}$ (Fig. 4) extrapolating the function that there exist in GaAs (III) crystals the emitting states of exciting energy laess than $h\nu_i$. The author has proposed to call "the over-threshold emission" [8] the emission of photoelectron group the work function of which lies between the values of $h\nu_i$ and φ (measured with respect to the vacuum level). Thew over-threshold emissin can proceed from the surface layer of semiconductor but mainly from different types of surface states (located within the energy gap), resulted mainly by the surface defects.



Rys.4. Yield vs photon energy and determination of its threshold

Different kinds of oxides and defects resulted by treating the samples introduce the additional states giving the above mentioned emission. The method of determining the over-valence emission is supported by the analysis of the $N(E)$ distributions with respect to the photoelectron number as a function of their kinetic energy E_k (Fig. 3). To find if there exists the measurable photoemission in the particular case, from the states within the energy gap, the width of the spectral distribution of photoelectrons have to be determined.

If this width satisfies the equation

$$h\nu - h\nu_t = \Delta E$$

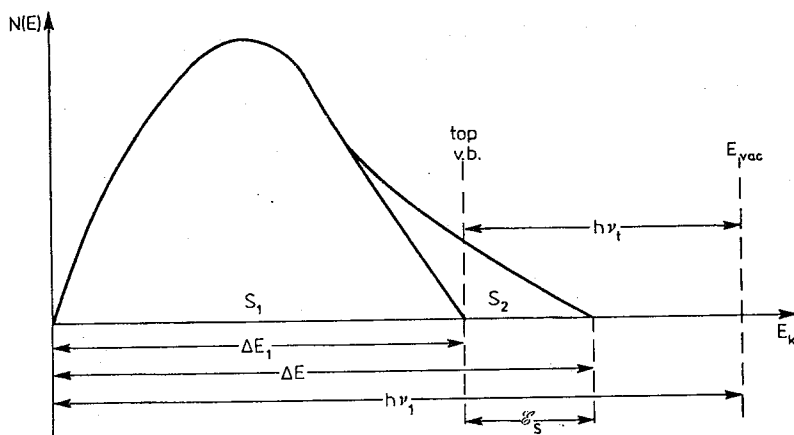
at the assumption, that the total error in determination the position of high-energy ends of $N(E)$ curves does not exceed 0.1 eV, then one can assume, that the quickest photoelectrons originate from the valence band top and then the over-threshold emission does not exist in practice.

When $\Delta E > h\nu - h\nu_t + 0.1$ eV it is possible to estimate the width of the emitting surface band (and its relative yield) localized within the band gap. Informs about it also the shape of $N(E)$ curves on their right-hand-side slopes, on which the high-energy humps occur [12]. But the upper parts of the right-hand-side slopes of $N(E)$ distributions most often appear to be the line segments.

If one performs the linear extrapolation of the line segments up to the intersection with the E_k axis, then after measuring the interval ΔE_1 at the section point one proves if

$$\Delta E_1 \cong h\nu_1 - h\nu_t.$$

Confirmation of it, means that energy ΔE_1 is connected with photoelectrons emitted from the valence band only while ΔE is the energy of photoelectrons emitted from the valence band top. Thus the energetic width $\Delta E - \Delta E_1$ determines the emitting section of the surface band (the over-threshold emission) (Fig. 5). The percentage share of over-valence photoelectrons in the entire spectrum of emission can be estimated by planimetry of fields S_1 and S_2 (Fig. 5 in no scales).



Rys.5. Density of states against the electron energy and determination of photoemission parameters (in no scales)

In the literature it is accepted (Gobeli and Allen [13]), that the surface bands in semiconductors are half filled by electrons. Knowing E_s , i.e. from relation $|\varphi^+ - h\nu_0| = E_s$ one can determine the width of the emitting part of the surface band located within the energy gap (Fig. 2). Expressions with values (Fig. 4):

$$\begin{aligned} \Delta E_1 & \dots \Delta E_n \\ h\nu_t &= h\nu_1 - \Delta E_1 \\ & \dots \dots \dots \\ h\nu_t &= h\nu_n - \Delta E_n \end{aligned} \quad (3)$$

were found to be fulfilled with accuracy of about 0.1 eV, after substitution the respective ΔE_i values found from extrapolation of photoelectron energy distributions. Thus it should be proved justice of assumption in the discussed method concerning extrapolation of the right slopes in energetic distribution curves by line swlections. We have to remember, that photoemission is equivalent to the effect of exciting the valence electrons to the conduction bands up to the levels higher than the vacuum level E_{vac} . The high-energy slopes of curves at the boundary of extrapolated line present the electron groups emitted from the highest levels of the valence bands. In this energetic area a function of emitted number of electrons is approximately proportional to the state density $N_v(\mathcal{E})$, where $|\mathcal{E}|$ is the absolute value of energy of emitting valence levels. It follows from studying the expression in paper [14], where quantum relation of probability for interband electron transitions has the following form:

$$P(\mathcal{E}, E) = \text{const} [f N_c(E) N_v(\mathcal{E})] \quad (4)$$

where f is the oscillator strength, this factor is constant within not particularly large interval of $h\nu$; N_c , N_v are the densities of states in conduction and valence bands respectively.

Influence of density N_c on the given probability function was shown to be negligible, mostly because the exited photoelectrons occupy positions in the high located levels of conduction bands. Particularly it concerns the electrons having the largest energies. These energies after passing the electrons to the conduction bands, with respect to the point Γ_{15} in the centre of Brillouin's zone are equal to $h\nu_t + E_{k_{\max}}$ (Fig. 3).

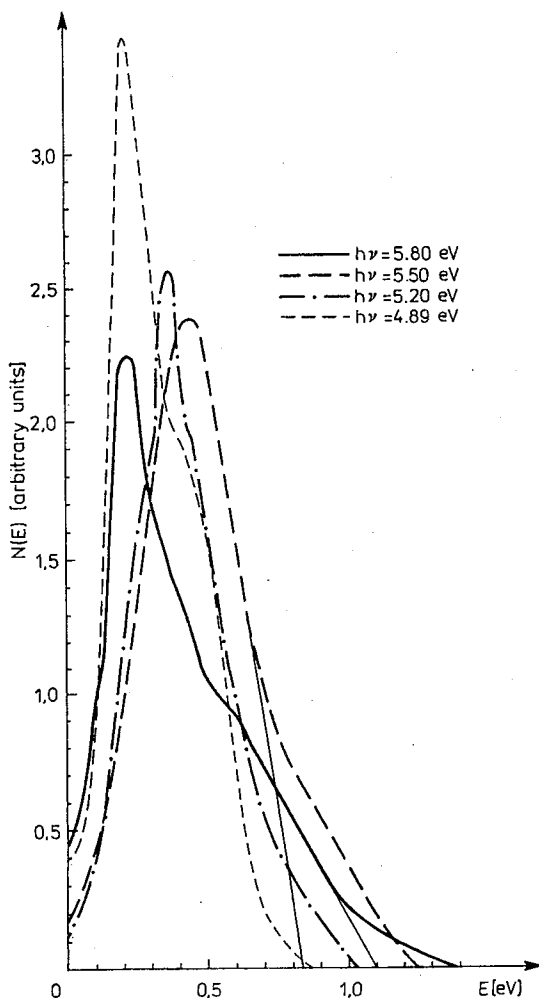
Values of $h\nu_t$ for III – V compounds are between 4.3 and 5.5 eV. At this level of the conduction band of III – V compounds the density of the quantum states was found to be very large due to the overlapping of two or three bands. Moreover, the literature data [15] suggest that at height of few eV with respect to the band edges, the densities of the states were found to be almost constant. Near the bottom of the conduction band the density N_c increases rapidly. Thus the above study leads to the conclusion, that the courses of high energy-parts of $N(E)$ distribution depend exclusively on the state density function N_v in the higher parts of the emitting valence bands.

One can conclude, that the linear extrapolation of the highenergy "valence" ends of $N(E)$ distribution will be right when in the valence band the density of states increases nearly linearly with the energy $\mathcal{E}$. There are known some papers [15]

proving that this condition is satisfied for 1 eV from the top of the valence band for some compounds from the III—V group. $N_v(\mathcal{E})$ distributions presented in [15] for three semiconductors from III—V group, evidently show the linear courses of the state density up to the energy about 1 eV with respect to the point Γ_{15} in the middle of Brillouin's zone. It is reasonable to assume that for electrons emitted from upper parts of the bands the emission probability depends on the state density only and rises linearly with it. Thus described above linear extrapolation of $N(E)$ distributions should give the correct results within the limits of measurement and calculation errors.

3. EXAMPLES OF CALCULATIONS FOR GaAs SAMPLES

The author presents the results of investigations and calculation for GaAs (111) *n*-type samples (Fig. 6).



Rys.6. Energetic distribution of electrons from *n*-type GaAs (111) sample, after annealing at $t = 550^\circ\text{C}$ in ultrahigh vacuum

From graphical differentiation on of current-voltage characteristics measured at the incident photon energy of 5.50 eV the energy distributions of photoelectrons were found to be within the interval of $\Delta E = 1.22$ eV and after extrapolation, cutting out the high-energetic humps, the emission from the valence band was obtained of energy $\Delta E_1 = 0.82$ eV and then the over-threshold emission of $E_s = 0.40$ eV (it is the width of the emitting part of the surface band). The value of $h\nu_t$ determined from the yield Y measurements equals 4.68 eV (fig. 7).

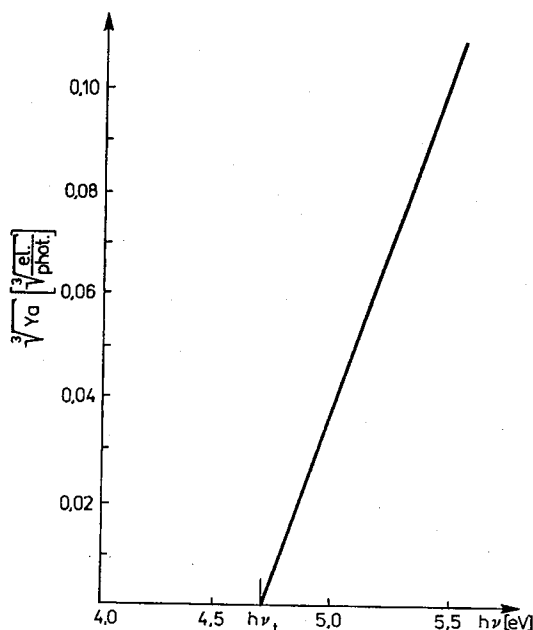
The exact calculations:

$$h\nu_1 = h\nu_t + \Delta E_1$$

$$5.50 = 4.68 + \Delta E_1$$

$$\text{hence } \Delta E_1 = 5.50 - 4.68 = (0.82 \pm 0.1) \text{ eV} \quad (5)$$

$$\text{and } E_s = \Delta E - \Delta E_1 = (1.22 - 0.82) \text{ eV} = (0.4 \pm 0.1) \text{ eV}$$



Rys.7. The average from five samples the cube root of the quantum yield from the (111) *n*-type GaAs surface against the photon energy

The second example:

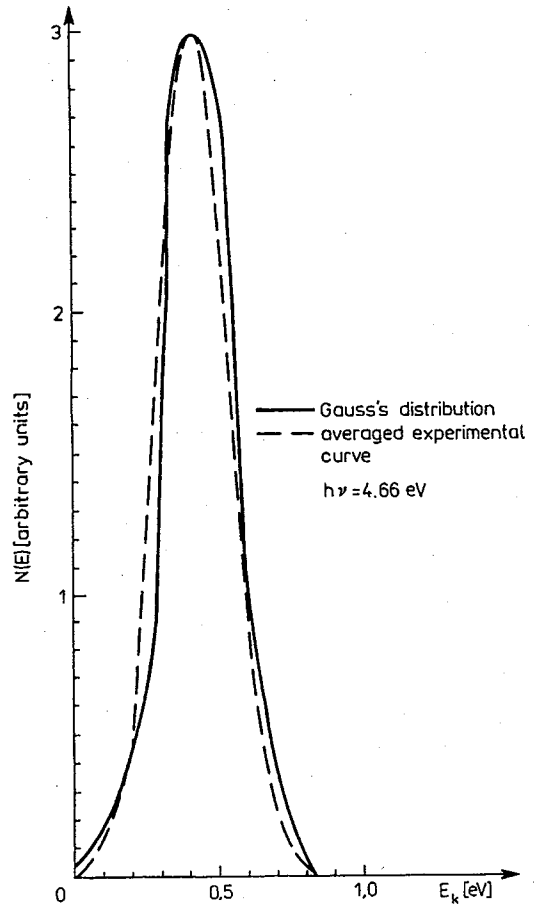
Energy of the exciting radiation of 5.80 eV (no change in $h\nu_t$). The spread of the energy distributions of electrons $\Delta E = 1.38$ eV, emission from the valence band $\Delta E_1 = 1.1$ eV, thus the over-threshold emission $E_s = 0.28$ eV.

$$\begin{aligned}
 5.8 &= 4.48 + \Delta E_1 \\
 \Delta E_1 &= 5.8 - 4.68 = (1.12 \pm 0.1) \text{ eV} \\
 E_s &= \Delta E - \Delta E_1 = (1.38 - 1.1) \text{ eV} = (0.28 \pm 0.1) \text{ eV}.
 \end{aligned}
 \tag{6}$$

In the both cases the calculation results coincides with the values obtained from the spectral yield curves and from the electron energy distribution curves.

4. CONCLUSIONS

Since for $h\nu = h\nu_t$ the photoelectrical yield does not equal to zero (Fig. 7) (and according to Kane's equation $Y = Z(h\nu - \nu_t)^{5/2}$ should be zero) then is apparent that over-valence emission exists with a defined yield. If the surfaces were ideally clean (in the superhigh vacuum), then over-valence emission would be likely absent.



Rys.8. The average over-valence emission intensity from GaAs (111) versus electron energy; The full curve represents the Gauss-function with the parameters adjusted)

In this case, there would be a small amount of surface states, thus emission would be negligible within the limits of error.

So-called real surfaces result that the probability of occurrence of above threshold emission is increased.

Energetic analysis of over-valence emission distribution as well as analysis of ratios of the emission to the total emission leads to the conclusion, that on GaAs (111) surfaces appear the bands of emitting surface states located within the energetic gap. These bands have broadened spectra without sharp thresholds.

It was mathematically verified, that for both types of the surfaces (111) and ($\bar{1}\bar{1}\bar{1}$) the distribution of over-valence photoelectrons are similar to the Gauss-like distributions (Fig. 8).

From measurements of photoemission yield Y it was obtained the photoelectric threshold for surface (111) to be larger by 0.3 eV then $h\nu_i$ ($\bar{1}\bar{1}\bar{1}$) value. The tail of over-valence emission from the surface (111) was found to be also longer, which can be justified by energetic width of distributions and by electron population degree of the bands created by the surface states [8].

The differences in the shape of energetic distributions of electrons from (111) and ($\bar{1}\bar{1}\bar{1}$) surfaces appear mainly as a longer tails of the former surface in the high energy area. This elongation is connected with the broader over-valence level band.

The energetic distributions from (100) p -type samples and from (110) p -type ones were found to be different from distributions for n -type samples by greater interval ΔE and by differing in shape. Near by the main maxima at the right hand side of the distribution appear humps or flat areas (Fig. 5). The author identifies the above mentioned maxima with the valence emission, and the high energetic humps — with the over threshold emission. The large values of absorption coefficient α in the ultraviolet range encourages one to relate the last emission to the populated surface states. From the $\delta_s^{*(1)}$ values and the other data the author concludes, that the surface band appears to be located in the lower part of the energy gap and occupies the width within (0.5–0.6) eV touching the valence band tops.

Author proved that the probability of occurrence of above threshold emission is greater for wideband semiconductors.

REFERENCES

1. V.A. Zuev, V.G. Litovchenko, G.A. Sukach, Ukr. Fiz. Zh. 20, 1147 (1975)
2. O.V. Snitko, Red.: *Problemy fiziki poverkhosti poluprovodnikov*. Naukova Dumka, Kiev, 1981
3. H. Flietner, N. Duong Singh, Phys. Stat. Sol. (a) 37, 533 (1976)
4. S. Gourrier, P. Friedel, P.K. Larsen, Surf.Sci., 152/153, 1147 (1985)
5. L. Szaro, Surf.Sci., 122, 149 (1982)
6. B. Seroczyńska-Wojas, Prace Instytutu Fizyki Politechniki Warszawskiej, (1980) z. 22, 23

⁽¹⁾ δ_s^* is the energy distance between the Fermi level and the top of the valence band at the surface.

7. A. T u l k e, H. L u t h, Surf. Sci., 178, 131 (1986)
8. J. W o j a s, *Badania fotoelektrycznych i powierzchniowych własności monokryształów GaAs o termicznie czyszczonych powierzchniach*. Wyższa Szkoła Pedagogiczna, Olsztyn (1976)
9. J. W o j a s, Acta Phys. Pol., A51, 25 (1977)
10. J. W o j a s, B. S e r o c z y Ń s k a - W o j a s, Annales Universitatis Mariae Curie-Skłodowska, Lublin–Polonia, XLIII/XLIV, 31, 341 (1989)
11. B. S e r o c z y Ń s k a - W o j a s, J. W o j a s, Polish Journal of Chemistry, (1990) 64, 183
12. F.G. A l l e n, W.G. G o b e l i, J. Appl. Phys. 45, 597 (1964)
13. W.G. G o b e l i, F.G. A l l e n, Semiconductors and Semimetals, Vol. 2, Chapter 11, Academic Press, New York and London (1966)
14. N. K i n d i g, W. S p i c e r, Phys. Rev. 138, A561 (1965)
15. J. C h e l i k o w s k y et al., Phys. Rev. B8, 2786 (1973)

J. WOJAS

EMISJA NADPROGOWA Z REALNYCH POWIERZCHNI PÓLPRZEWODNIKOWYCH

Streszczenie

Praca ta przedstawia występowanie fotoemisji nadprogowej (ponad walencyjnej). Szczególnie odnosi się to do powierzchni realnych półprzewodników. Używając danych z pomiarów przeprowadzono obliczenia ilościowe emisji elektronów. Uzyskano dobrą zgodność pomiędzy danymi z pomiarów i teorią.

Speech coding using product code vector quantization, long-term prediction and orthogonal transformation

PRZEMYSŁAW DYMARSKI

Instytut Telekomunikacji, Politechnika Warszawska

SURASIT RATREE

*Faculty of Industrial Education,
King Mongkut's Institute of Technology (KMUTL)
Ladkrabang, Bangkok, Thailand*

Received 1994.05.27

Authorized 1994.08.24

Several product code vector quantizers (PCVQ) have been described: the shape-gain vector quantizers (SGVQ), the multistage vector quantizers using the long-term prediction (LTP) and the transform vector quantizers using the Karhunen-Loeve transform (KLT). New codebook search and codebook design algorithms have been proposed. The SGVQ, LTP and KLT techniques have been combined to obtain the 9.6 kbit/s speech coder.

1. INTRODUCTION

The vector quantization techniques are widely used in the coding of speech, image and video signals. In the standard vector quantizer (VQ), any N -dimensional vector of input speech samples is represented by one of L vectors issued from a codebook $F = \{f_1, f_2, f_L\}$. Shannon [8] has shown that VQ is a theoretically optimal coder: for any bit rate b [bits/sample] the distortion D approaches (as $N \rightarrow \infty$) the rate-distortion bound. This bound is given by the rate-distortion function (RDF) of a signal source [7]. The reduction of distortion is obtained at a cost of an exponential growth of complexity. This is because the codebook contains $L = 2^{bN}$ vectors and the complexity of the search algorithm and storage requirements are proportional to LN .

In order to diminish the distortion and keep the moderate complexity, the product code vector quantizers (PCVQ) are used. The codebook of such a quantizer may be described as a cartesian product of several small codebooks. The most popular PCVQs are the shape-gain vector quantizers (SGVQ), in which any input signal vector is modelled by a shape vector $\mathbf{f}_j$ (issued from the shape codebook) multiplied by a gain g_m (issued from the gain codebook).

In the multistage PCVQ the input vector is modelled as a sum of codebook vectors (or a sum of shape-gain pairs). The most popular multistage PCVQ is the VQ with the long-term prediction (LTP). The model is obtained as a sum of the prediction vector multiplied by its proper gain and the codebook vector multiplied by its gain. The prediction vectors are issued from an adaptive codebook, containing the most recent samples of the reconstructed signal.

The input vector may be also modelled by a concatenation of codebook vectors of different length. This approach, however, makes sense only if the orthogonal transform is used prior to VQ. The very popular CELP coder may be also regarded as a special case of PCVQ, where the signal model is obtained by filtering of codebook vector using the predictive filter issued from the set of available filters (the CELP coders are out of the scope of this work).

In this work we combine the shape-gain, LTP and transform techniques to obtain the PCVQ suitable for speech signal coding at bit rates of about 10 kb/s. In the section 2 several variants of SGVQ are analyzed. New gain encoding algorithms, using multiple codebooks, are presented. In the section 3 the LTP is combined with the SGVQ. The LTP, commonly used in CELP coders [12], has also proved useful in the VQ. In order to improve the speech quality, the codebook orthogonalization algorithms are introduced. These algorithms have been proposed for the CELP coders [4]. We have shown, that they are also suitable for the VQ. Further improvement of speech quality may be obtained by using the modified algorithm of codebook design. Finally open- and close-loop codebook design algorithms are tested. In the section 4 the Karhunen-Loeve transform (KLT) is applied, together with the LTP. The transformed vector is modelled by a concatenation of two codebook vectors of nonequal length. In this way further improvement of speech quality is obtained, without increasing much the complexity of the coder.

2. VECTOR QUANTIZATION AND SHAPE-GAIN STRUCTURED CODEBOOKS

2.1 THE STANDARD VECTOR QUANTIZER

In the basic vector quantizer (fig.1) each N -dimensional vector of speech samples $\mathbf{x}(i)$ is compared with L codebook vectors $\{\mathbf{f}_1, \mathbf{f}_2, \dots, \mathbf{f}_L\}$. The reconstruction vector (model) $\hat{\mathbf{x}}(i)$ is the codebook vector $\mathbf{f}_j$ minimizing the euclidean distance $\|\mathbf{x}(i) - \mathbf{f}_j\|$.

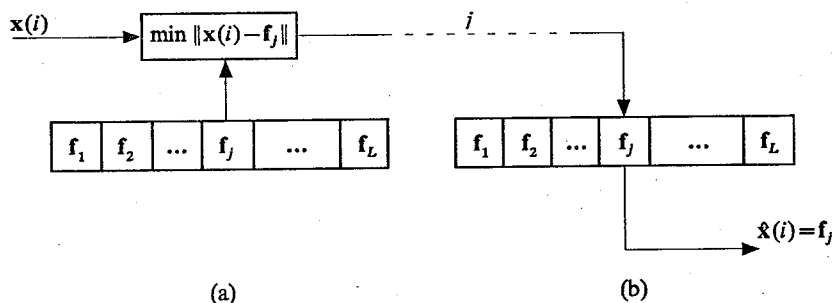


Fig. 1. Standard vector quantizer: (a) encoder, (b) decoder

The euclidean distance is chosen because of its close relation to the energy of quantization error. The output signal quality depends on the codebook $\{f_j\}$. The codebook is designed for minimum distortion D in quantization of a given set of input vectors $\{x(i)\}$ i.e. the training sequence. The generalized Lloyd algorithm [1] is used for codebook design (fig.2).

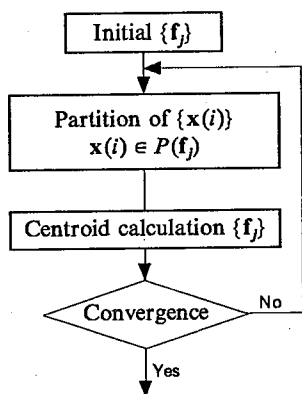


Fig. 2. Flow chart of the generalized Lloyd algorithm

Given the initial codebook $\{f_j\}$ the partition of the training sequence $\{x(i)\}$ may be performed in the following way:

$$\begin{aligned}
 & x(i) \in P(f_j) \text{ if } \hat{x}(i) = f_j \\
 & \text{i.e. } \|x(i) - f_j\| < \|x(i) - f_k\| \\
 & \text{for } k = 1, 2, \dots, L; \quad k \neq j.
 \end{aligned} \tag{1}$$

Thus the N -dimensional space R^N is partitioned in the L Voronoi cells $P(f_1), \dots, P(f_L)$. The input vector $x(i)$ belongs to the Voronoi cell $P(f_j)$ if f_j is the closest codevector to $x(i)$.

The global distortion D may thus be defined as follows:

$$D \stackrel{\text{df}}{=} \sum_i \|x(i) - \hat{x}(i)\|^2 = \sum_{j=1}^L \sum_{\substack{i \\ x(i) \in P(f_j)}} \|x(i) - f_j\|^2. \quad (2)$$

The new codebook vector minimizes the partial distortion $D_j = \sum_{\substack{i \\ x(i) \in P(f_j)}} \|x(i) - f_j\|^2$

for input vectors $x(i)$ closest to the former vector f_j . From the equation $\frac{\partial D_j}{\partial f_j} = 0$ the following formula for the new codebook vector is obtained:

$$f_j^{\text{new}} = \frac{1}{L_j} \sum_{\substack{i \\ x(i) \in P(f_j)}} x(i) \quad (3)$$

where L_j — number of vectors $x(i)$ closest to the former vector f_j .

The new vector f_j is the centroid of the subset of vectors $x(i)$ closest to the former vector f_j . New codebook generates a new partition and the iterations are performed until some convergence criterion is fulfilled. The initial codebook may be obtained by splitting the vectors of a smaller codebook containing $\frac{L}{2}$ vectors as proposed in the LBG algorithm [6].

2.2 THE SHAPE-GAIN VECTOR QUANTIZER

The standard VQ performs well for long vectors and large codebook. In order to reduce complexity the large codebook may be described as a cartesian product of two smaller codebooks: the shape codebook $\{f_j\}$ containing L_s shape vectors and

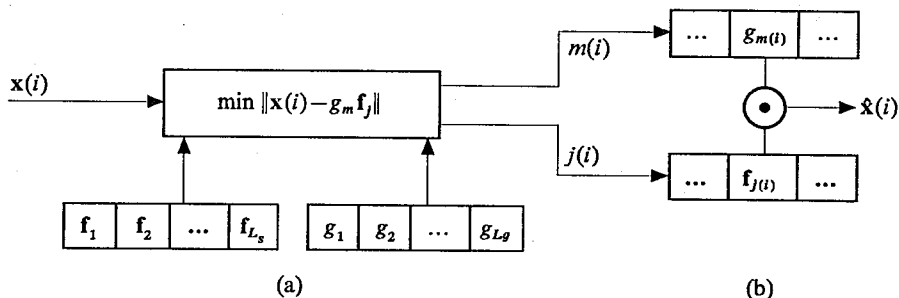


Fig.3. Shape-gain vector quantizer: (a) encoder, (b) decoder

the gain codebook $\{g_m\}$ containing L_g scalars. In the shape-gain vector quantizer (SGVQ—fig.3) each input vector $x(i)$ is modelled as a product

$$\hat{x}(i) = g_{m(i)} \cdot f_{j(i)}. \quad (4)$$

As a whole there are $L = L_g L_s$ available input vector representations (models). The gain $g_{m(i)}$ and shape $\mathbf{f}_{j(i)}$ are selected so as to minimize the local distortion:

$$D(i) = \|\mathbf{x}(i) - \hat{\mathbf{x}}(i)\|^2 = \mathbf{x}^t(i) \mathbf{x}(i) - 2g_{m(i)} \mathbf{x}^t(i) \mathbf{f}_{j(i)} + g_{m(i)}^2 \mathbf{f}_{j(i)}^t \mathbf{f}_{j(i)} \quad (5)$$

where t denotes vector transposition.

The local distortion $D(i)$ may be rewritten using the following formula:

$$D(i) = \mathbf{x}^t(i) \mathbf{x}(i) - \frac{[\mathbf{x}^t(i) \mathbf{f}_{j(i)}]^2}{\mathbf{f}_{j(i)}^t \mathbf{f}_{j(i)}} + \mathbf{f}_{j(i)}^t \mathbf{f}_{j(i)} [g_{m(i)} - g(i)]^2 \quad (6)$$

$$\text{where } g(i) = \frac{\mathbf{x}^t(i) \mathbf{f}_{j(i)}}{\mathbf{f}_{j(i)}^t \mathbf{f}_{j(i)}}.$$

The codebook search algorithm may be deduced from the above formulas. For any shape $\mathbf{f}_{j(i)}$ the best gain $g_{m(i)}$ is the nearest neighbour of $g(i)$. Thus $g(i)$ may be regarded as an optimum, non-quantized gain. It may be noticed that $g(i) \mathbf{f}_{j(i)}$ is the orthogonal projection of $\mathbf{x}(i)$ on $\mathbf{f}_{j(i)}$.

First the best shape $\mathbf{f}_{j(i)}$ must be chosen. If the shape codebook is normalized i.e. $\mathbf{f}_j^t \mathbf{f}_j = 1$ and if the gain codebook contains only positive scalars then, independently on the gain g_m , the best shape maximizes $\mathbf{x}^t(i) \mathbf{f}_j$ (5). This choice is equivalent to minimization of angle $\varphi[\mathbf{x}(i), \mathbf{f}_j]$ between the vector $\mathbf{x}(i)$ and $\mathbf{f}_j$. Notice that $\cos \varphi[\mathbf{x}(i), \mathbf{f}_j] = \mathbf{x}^t(i) \mathbf{f}_j / [\|\mathbf{x}(i)\| \|\mathbf{f}_j\|]$. If the shape codebook is not normalized, the optimal shape depends not only on the angle $\varphi[\mathbf{x}(i), \mathbf{f}_j]$ but also on the gain codebook(6). Nevertheless, the maximization of $\cos^2 \varphi[\mathbf{x}(i), \mathbf{f}_j] = [\mathbf{x}^t(i) \mathbf{f}_j]^2 / [\mathbf{f}_j^t \mathbf{f}_j \cdot \mathbf{x}^t(i) \mathbf{x}(i)]$ i.e. maximization of the second term of (6) yields a good suboptimal solution. Better results may be obtained by calculating the local distortion for several shape candidates forming sufficiently small angle with $\mathbf{x}(i)$ [5].

The codebooks for shapes $\{\mathbf{f}_j\}$ and gains $\{g_m\}$ are designed at the same time, using the modified Lloyd algorithm proposed by Sabin and Gray [2] – fig. 4.

Each codebook yields a different partition of the training sequence: if the optimal model of $\mathbf{x}(i)$ is $\hat{\mathbf{x}}(i) = g_m \mathbf{f}_j$ then $\mathbf{x}(i) \in P(\mathbf{f}_j)$ and $\mathbf{x}(i) \in P(g_m)$. The global distortion for the training sequence may be expressed in two different ways:

$$D = \sum_i D(i) = \sum_{m=1}^{L_g} D_m$$

where

$$D_m = \sum_{\substack{i \\ \mathbf{x}(i) \in P(g_m)}} D(i) = \sum_{\substack{i \\ \mathbf{x}(i) \in P(g_m)}} \mathbf{x}^t(i) \mathbf{x}(i) - 2g_m \sum_{\substack{i \\ \mathbf{x}(i) \in P(g_m)}} \mathbf{x}^t(i) \mathbf{f}_{j(i)} + g_m^2 \sum_{\substack{i \\ \mathbf{x}(i) \in P(g_m)}} \mathbf{f}_{j(i)}^t \mathbf{f}_{j(i)} \quad (7)$$

and

$$D = \sum_{j=1}^{L_s} D_j$$

where

$$D_j = \sum_{\substack{i \\ \mathbf{x}(i) \in P(\mathbf{f}_j)}} D(i) = \sum_{\substack{i \\ \mathbf{x}(i) \in P(\mathbf{f}_j)}} \mathbf{x}^t(i) \mathbf{x}(i) - 2 \left[\sum_{\substack{i \\ \mathbf{x}(i) \in P(\mathbf{f}_j)}} g_{m(i)} \mathbf{x}^t(i) \right] \mathbf{f}_j + \mathbf{f}_j^t \mathbf{f}_j \sum_{\substack{i \\ \mathbf{x}(i) \in P(\mathbf{f}_j)}} g_{m(i)}^2. \quad (8)$$

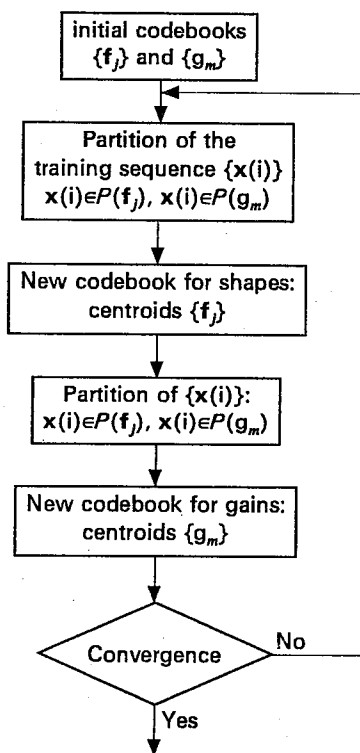


Fig.4 The modified Lloyd algorithm for SGVQ codebooks design

The new gain codebook is obtained from the equations $\frac{\partial D}{\partial g_m} = \frac{\partial D_m}{\partial g_m} = 0$. The solution:

$$g_m^{\text{new}} = \frac{\sum_{\substack{i \\ x(i) \in P(g_m)}} x^T(i) f_{j(i)}}{\sum_{\substack{i \\ x(i) \in P(g_m)}} f_{j(i)}^T f_{j(i)}} \quad (9)$$

may be called the gain centroid for all input vectors $x(i)$ represented by the former gain g_m . The solution of $\frac{\partial D}{\partial f_j} = \frac{\partial D_j}{\partial f_j} = 0$ yields the following shape centroid for all input vectors $x(i)$ represented by the former shape f_j :

$$f_j^{\text{new}} = \frac{\sum_{\substack{i \\ x(i) \in P(f_j)}} g_{m(i)} x(i)}{\sum_{\substack{i \\ x(i) \in P(f_j)}} g_{m(i)}^2} \quad (10)$$

It is always some degree of freedom in choosing the shape and gain codebooks, e.g. the shape vectors may be always multiplied by a scaling factor if the gains are divided by the same scalar. Thus some constraint may be added, e.g. the shape vectors may be normalized as follows

$$\|\mathbf{f}_j\|^2 = \mathbf{f}_j^T \mathbf{f}_j = 1, \quad j = 1, 2, \dots, L_s \quad (11)$$

The condition (11) simplifies the codebook search algorithm, but influences the choice of centroids. Minimization of D with the constraint (11) yields the following solution:

$$\mathbf{f}_j^{\text{new}} = \frac{\sum_{\substack{i \\ \mathbf{x}(i) \in P(\mathbf{f}_j)}} g_{m(i)} \mathbf{x}(i)}{\left\| \sum_{\substack{i \\ \mathbf{x}(i) \in P(\mathbf{f}_j)}} g_{m(i)} \mathbf{x}(i) \right\|}. \quad (12)$$

2.3 MATCHING THE SIGNAL PROPERTIES

The SGVQ with the normalized shape codebook is suitable for quantization of uncorrelated stationary signals. For nonstationary signals better results may be obtained using adaptive codebooks. For correlated signals there is some mismatch between the signal properties and the available reconstruction vectors — fig.5. The available models $g_m \mathbf{f}_j$ lie on L_g N -dimensional spheres of radii $g_1, g_2, \dots, g_{L_g}$ respectively. The *pdf* (probability density function) of a correlated gaussian source is constant on the ellipsoids.

This mismatch results in a bad exploitation of available models: some $g_m \mathbf{f}_j$ combinations are used extremely seldom.

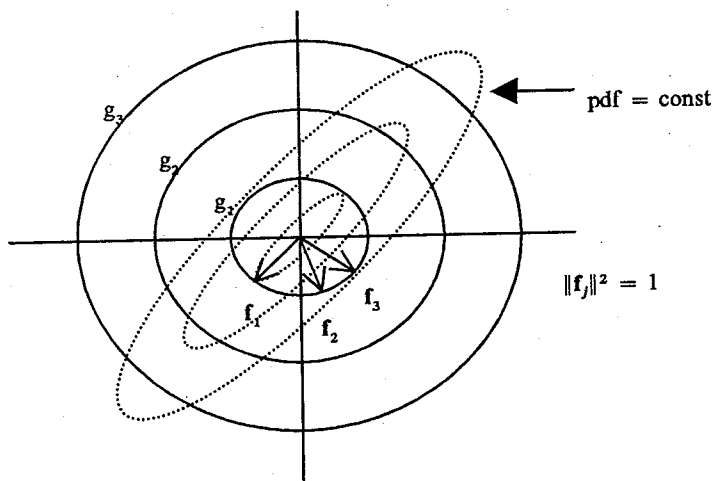


Fig. 5. SGVQ models and *pdf* of a correlated gaussian source

To solve this problem we propose to apply different gain codebook $\{g_m^j\}$ for each shape $\mathbf{f}_j$. The encoding algorithm is straightforward:

For each shape $\mathbf{f}_j$ the gain $g_{m(i)}^j$ is selected from the codebook $\{g_m^j\}$, as the nearest neighbour of $g(i)$ (see formula 6). The pair $\mathbf{f}_{j(i)}, g_{m(i)}^j$ is chosen for minimum of $D(i) = \|\mathbf{x}(i) - g_{m(i)}^j \mathbf{f}_j\|^2$. Using the Sabin and Gray algorithm (fig. 4) for designing L_g gain codebooks and one shape codebook would be quite a prohibitive task. Therefore the shape codebook design should be decoupled from the gain codebooks design. In our shape design algorithm we assume that gain encoding is perfect, i.e. the last term in equation (6) may be omitted. With $\|\mathbf{f}_j\|^2 = 1$ the formula (6) may be rewritten as follows:

$$D(i) = \mathbf{x}^t(i) \mathbf{x}(i) - [\mathbf{x}^t(i) \mathbf{f}_{j(i)}]^2. \quad (13)$$

The generalized Lloyd algorithm (fig. 2) is used, with partition based on the angle minimization: $\mathbf{x}(i) \in P(\mathbf{f}_j)$ if $\varphi[\mathbf{x}(i), \mathbf{f}_j] < \varphi[\mathbf{x}(i), \mathbf{f}_k]$ for each $k \neq j$. The global distortion is to be minimized:

$$D = \sum_i D(i) = \sum_{j=1}^{L_g} D_j$$

where

$$D_j = \sum_{\mathbf{x}(i) \in P(\mathbf{f}_j)} D(i) = \sum_{\mathbf{x}(i) \in P(\mathbf{f}_j)} \mathbf{x}^t(i) \mathbf{x}(i) - \sum_{\mathbf{x}(i) \in P(\mathbf{f}_j)} [\mathbf{x}^t(i) \mathbf{f}_j]^2. \quad (14)$$

Equation (14) implies that the centroid $\mathbf{f}_j$ should maximize

$$Q_j = \sum_{\mathbf{x}(i) \in P(\mathbf{f}_j)} [\mathbf{x}^t(i) \mathbf{f}_j]^2 = \sum_{\mathbf{x}(i) \in P(\mathbf{f}_j)} \mathbf{f}_j^t \mathbf{x}(i) \mathbf{x}^t(i) \mathbf{f}_j = \mathbf{f}_j^t \sum_{\mathbf{x}(i) \in P(\mathbf{f}_j)} \mathbf{x}(i) \mathbf{x}^t(i) \mathbf{f}_j = \mathbf{f}_j^t \mathbf{R}_j \mathbf{f}_j \quad (15)$$

where $\mathbf{R}_j$ is the covariance matrix of vectors $\mathbf{x}(i)$ forming the minimum angle with the former vector $\mathbf{f}_j$. Maximization of (15) with the constraint $\|\mathbf{f}_j\| = 1$ yields the following set of equations

$$\mathbf{R}_j \mathbf{f}_j = Q_j \mathbf{f}_j. \quad (16)$$

The solution of (16) is the eigenvector of $\mathbf{R}_j$ corresponding to the maximum eigenvalue Q_j , i.e. the principal axis of the set of vectors $\mathbf{x}(i)$ represented by the former vector $\mathbf{f}_j$. The principal axis may be found using the following iterations: [17]

$$\mathbf{f}_j^{k+1} = \frac{\mathbf{R}_j \mathbf{f}_j^k}{\|\mathbf{R}_j \mathbf{f}_j^k\|}. \quad (17)$$

The gain codebooks $\{g_m^j\}$ are designed separately for each shape $\mathbf{f}_j$ using the Lloyd algorithm and the training sequence $\mathbf{x}(i) \in P(\mathbf{f}_j)$. The centroids are calculated using the formula (9). It may be noticed that in this case the centroid is the mean value of unquantized gains $g(i)$ for $\mathbf{x}(i) \in P(\mathbf{f}_j)$ and $\mathbf{x}(i) \in P(g_m^j)$. In the case of small

amount of vectors $\mathbf{x}(i) \in P(g_m^j)$, the special algorithm has been used: to the corresponding gains $g(i)$ new values have been appended, equal to $0.01 g(i)$, $0.02 g(i)$, ..., $0.99 g(i)$.

The quantizer may be also adapted to signal properties using one gain codebook and a shape codebook consisting of vectors of unequal length. This approach is called the contour-gain vector quantization (CGVQ[5]). The longest codebook vectors should be the vectors close to the principal axis of the training data (fig. 5). More precisely, the set of available models $g_m \mathbf{f}_j$ should match the contours of $\text{pdf} = \text{const}$. For a gaussian signal, the contours of constant pdf are described by the equation

$$\mathbf{x}^t \mathbf{R}^{-1} \mathbf{x} = \text{const} \quad (18)$$

where $\mathbf{R}$ is the covariance matrix of the input signal.

The generalized Lloyd algorithm (fig. 4) may be used for codebooks calculation. The gain centroid is given by the equation (9). The shape centroid may be obtained by minimizing the partial distortion D_j (8) with the constraint $\mathbf{f}_j^t \mathbf{R}^{-1} \mathbf{f}_j = 1$. This leads to a rather complex set of equations which may be solved iteratively. Malone and Fischer[5] used unconstrained minimization(formula 10) for shape centroid. They start the generalized Lloyd algorithm using the gain codebook obtained by clustering of the training sequence $g'(i) = \sqrt{\mathbf{x}^t(i) \mathbf{R}^{-1} \mathbf{x}(i)}$. The values $g'(i)$ are the best gains for codebook vectors $\mathbf{f}_j$ colinear with $\mathbf{x}(i)$ and satisfying $\mathbf{f}_j^t \mathbf{R}^{-1} \mathbf{f}_j = 1$. This initial gain codebook forces the proper lengths of calculated shape vectors (however, the shape vectors only approximatively fulfil the normalizing condition $\mathbf{f}_j^t \mathbf{R}^{-1} \mathbf{f}_j = 1$).

The described above vector quantizers are suitable for stationary signals. In the quantization of speech signals, low-level segments may suffer from quantization noise. To solve this problem, the adaptive gain quantization may be applied. The backward adaptation algorithms perform well for scalar quantizers and do not give very good results for vector quantizers[3].

Therefore we propose the forward adaptation scheme. Each vector is modelled as below:

$$\hat{\mathbf{x}}(i) = g_{m(i)}^{j(i)} \sigma(i) \mathbf{f}_{j(i)} \quad (19)$$

where $\sigma(i)$ is the adaptation parameter.

The parameter $\sigma(i)$ represents the signal level (e.g. mean value of vector norms $\|\mathbf{x}(i)\|$) and is calculated in segments of about 10ms. It is quantized and transmitted to the receiver.

In the encoding algorithm, the shape $\mathbf{f}_{j(i)}$ is chosen for minimum angle $\varphi[\mathbf{x}(i), \mathbf{f}_j]$ then $g_{m(i)}^{j(i)}$ is selected from the gain codebook $\{g_m^{j(i)}\}$, as the nearest neighbour of $g(i)/\sigma(i)$. The adaptation parameter must be taken into consideration in the gain codebooks design: the normalized training data $g(i)/\sigma(i)$ should now be used.

2.4 RESULTS OF SIMULATIONS

The above mentioned product code vector quantizers are tested using 3 phrases of speech signal: the training sequence (used for codebook design), the non-training sequence and the training sequence attenuated by 20 dB. The sampling frequency is 8 kHz and the vector dimension $N=8$ (1 ms). The bit rate of 13 kbit/s is used, except of the SGVQ with forward adaptation for which the bit rate is 13.6 kbit/s due to the quantization of the adaptation parameter.

The segmental signal to noise ratio (SNRseg) is used as a measure of speech quality. SNRseg is a mean value of SNR[dB] calculated in segments of 128 samples.

$$\text{SNRseg} = \frac{1}{S} \sum_{m=1}^S \text{SNR}_m [\text{dB}]$$

where $\text{SNR}_m [\text{dB}] = 10 \log_{10} \frac{\sum_{n=1}^{128} (x_n)^2}{\sum_{n=1}^{128} (x_n - \hat{x}_n)^2}$ is calculated in the m -th segment, S is the

number of segments and $x_n, \hat{x}_n$ are the corresponding input and output signal samples. The non-speech segments of a very small energy are excluded from calculation. The energy threshold is selected 43 dB below the maximum energy of sinusoidal signal that saturates the A/D converter.

The results for the standard SGVQ (sect. 2.2) and the training sequence are given in the tab. 1.

Tab.1

SNRseg[dB] for the standard SGVQ at 13 kbit/s

L_s	L_g	SNRseg
128	64	13.58
256	32	14.16
512	16	13.28

Each of the 3 shape – gain vector quantizers uses 13 bits/ vector. The best shape-gain tradeoff is obtained with 8 bits for a shape ($L_s=256$) and 5 bits for a gain ($L_g=32$). These values are used to obtain results shown in the tab. 2. For a comparison, the results obtained without gain quantization are reported (the non-quantized gain $g(i)$ is used instead of $g_m(i)$).

Two methods of the shape codebook design are tested: the Sabin and Gray algorithm (fig. 4) and the algorithm based on the principal axis calculation (sect. 2.3). Both algorithms give practically the same codebooks (this is because of a rather big number of gain quantization levels $L_g=32$ used in the Sabin and Gray

Tab. 2

SNRseg[dB] for different SGVQs and 3 phrases

Quantizer	bit rate [kbit/ s]	Phrase		
		training sequence	non-training sequence	training sequence attenuated by 20 dB
Standard SGVQ without gain quantization	—	15.48	12.56	15.44
Standard SGVQ	13	14.16	10.74	6.24
SGVQ with many uniform gain quantizers	13	14.97	10.91	8.49
SGVQ with many nonuniform gain quantizers	13	15.25	11.16	13.05
CGVQ	13	14.79	12.02	11.85
SGVQ with adaptive gain quantization	13.6	14.20	10.43	14.17

algorithm). It is to be concluded that the gain codebook design may be decoupled from the shape codebook design without loss of speech quality. The shape codebook may be designed first, using the principal axis extraction algorithm. Then, the gain codebook may be obtained using the Lloyd algorithm with centroids given by the formula (9).

The standard SGVQ performs well for the training sequence, but there is a substantial loss of speech quality for the attenuated sequence. This is because the gain codebook contains too big values to be used for the attenuated sequence. The SGVQ with different gain codebook $\{g_m^i\}$ for each shape f_j is simulated in 2 variants: with the uniform and non—uniform gain quantizers. Both are better than the standard SGVQ, but the SGVQ with many non—uniform gain quantizers is the best, especially for the attenuated sequence. This is because of the special gain codebook design algorithm described in the section 2.3. However this version demands a lot of memory to store 256 gain codebooks.

It may be noticed that not bad results may be obtained using one gain codebook and the contour—gain vector quantizer (CGVQ)[5]. The CGVQ is less phrase—dependent than the SGVQ (note the best SNRseg for the non—training sequence). The best results for the attenuated sequence are obtained with the SGVQ and the forward—adaptive gain quantization. The same adaptation parameter is used for 256 uniform gain quantizers. However, this forward adaptation causes the extra delay of 10ms and requires the extra 600 bit/s for quantization of the adaptation parameter (it is coded once for 10 vectors, i.e. 10ms using 6 bits).

The final conclusion may be formulated as follows: if there is enough memory, the SGVQ with many non—uniform gain quantizers should be used, otherwise the CGVQ should be implemented.

3. MULTISTAGE VECTOR QUANTIZERS

3.1 THE MULTISTAGE SGVQ

The multistage VQ is the PCVQ, in which the signal vector $\mathbf{x}(i)$ is modelled as a sum of K vectors issued from different codebooks. The commonly used, though suboptimal codebook search algorithm consists of K iterations (stages). In each stage, one vector is added to the previously obtained model. At the K -th stage of this successive approximation algorithm the final model $\hat{\mathbf{x}}(i)$ is obtained. In order to cope with various signal levels, the SGVQ may be used at each stage. In the multistage CELP coders known also as the vector excitation (VXC) coders [12], the shape vectors are filtered using an adaptive linear predictive filter.

Practically the number of stages doesn't exceed 4 for the medium bit rate coders and 2 for the low bit rate coders. The most popular is the CELP coder with the long-term prediction (LTP), which may be regarded as a 2 stage SGVQ (with filtered shape vectors). Without filtering, the SGVQ with LTP is obtained (fig.6).

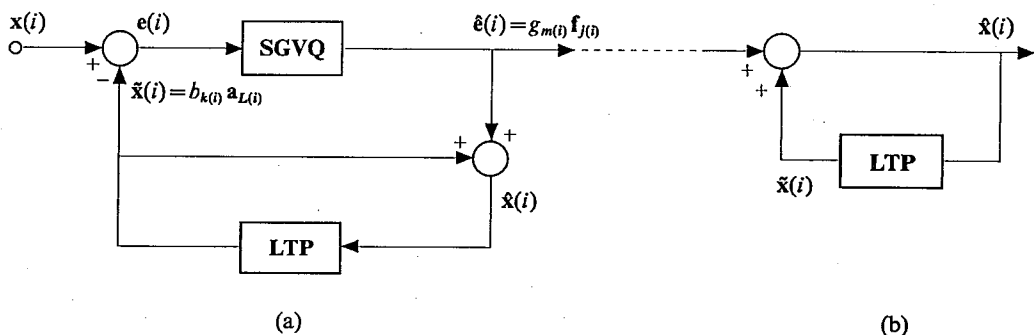


Fig. 6. SGVQ with LTP: (a) coder, (b) decoder

At the first stage of modelling, the first approximation of the input vector $\mathbf{x}(i)$ is obtained as a product $\tilde{\mathbf{x}}(i) = b_{k(i)} \mathbf{a}_{L(i)}$. The shape vector $\mathbf{a}_{L(i)}$ consists of the delayed output signal samples. If the delay equals to the pitch period, very good approximation of the voiced speech vector is obtained. The scalar $b_{k(i)}$ is the LTP gain. The LTP error vector $\mathbf{e}(i)$ is coded using the SGVQ as described in the previous section. Finally the speech signal vector $\mathbf{x}(i)$ is modelled as a linear combination of 2 vectors $\mathbf{a}_{L(i)}$ and: $\mathbf{f}_{j(i)}$

$$\hat{\mathbf{x}}(i) = b_{k(i)} \mathbf{a}_{L(i)} + g_{m(i)} \mathbf{f}_{j(i)}. \quad (20)$$

The LTP shape vector $\mathbf{a}_{L(i)}$ is issued from the predictive codebook $\{\mathbf{a}_L\}$ containing the delayed output signal samples. There is one sample delay between 2 consecutive codebook vectors. The minimum delay equals to the vector length N multiplied by the sampling period, and the maximum delay doesn't exceed 20 ms, which corresponds to the maximum pitch period. This codebook has an adaptive

character: the vectors delayed too much are excluded from the codebook. The LTP gain is issued from the scalar codebook $\{b_k\}$. The final model (20) is described by 4 elements issued from 4 codebooks. The optimal search algorithm, i.e. testing all the combinations b_k, a_L, g_m, f_j is far too complex. Therefore we need suboptimal and less complex search algorithms. Such algorithm, based on the orthogonalization of the codebook $\{f_j\}$ with respect to the LTP vector $a_{L(i)}$, is proposed in the following section. Similar algorithm has been proposed for CELP coders [10]. This orthogonalization influences the process of codebook design. The algorithm for designing of the shape codebook $\{f_j\}$ is described in the sect.3.3 (similar algorithm has been used for VXC coders [13]). Simulation results are given in the sect.3.4.

3.2 CODEBOOK SEARCH ALGORITHMS.

The one-stage modelling algorithm for the SGVQ has been described in sect.2.2. This algorithm may be extended to the multistage SGVQ. For the SGVQ with LTP, the following standard algorithm is thus obtained:

1. The best LTP shape $a_{L(i)}$ is chosen from the adaptive codebook $\{a_L\}$, for the maximum of $\cos^2 \varphi [x(i), a_L]^2$, where

$$\cos^2 \varphi [x(i), a_L] = \frac{[x^t(i) a_L]^2}{x^t(i) x(i) \cdot a_L^t a_L}. \quad (21)$$

2. The unquantized LTP gain $b(i)$ is calculated in such a way that $b(i) a_{L(i)}$ be orthogonal projection of $x(i)$ onto $a_{L(i)}$:

$$b(i) = \frac{x^t(i) a_{L(i)}}{\|a_{L(i)}\|^2}. \quad (22)$$

Then, the scalar $b_{k(i)}$ is selected from the LTP gain codebook $\{b_k\}$, for minimum of $|b(i) - b_k|$. The LTP gain codebook has to contain negative and positive scalars.

3. The LTP error vector $e(i) = x(i) - b_{k(i)} a_{L(i)}$ is calculated. Then modelling continues as described in the sect.2.2: the best shape $f_{j(i)}$ is chosen for $\min \varphi [e(i), f_j]$, the unquantized gain $g(i) = e^t(i) f_{j(i)} / \|f_{j(i)}\|^2$ is calculated and the best gain $g_{m(i)}$ is selected minimizing $|g(i) - g_m|$. It is assumed that the gain codebook $\{g_m\}$ contains only positive scalars.

The described above standard algorithm is rather far from optimality. First of all it doesn't assure that the obtained model (20) is the orthogonal projection of $x(i)$ on the plane span by the chosen vectors $a_{L(i)}$ and $f_{j(i)}$. Therefore the following modelling algorithm with gain correction is proposed:

1. The standard algorithm is performed without gain quantization, i.e. unquantized gains $b(i)$ and $g(i)$ are used.

2. The gains are recalculated in such a way that $x(i) - \hat{x}(i)$ be orthogonal to $a_{L(i)}$ and $f_{j(i)}$,

i.e. the system of 2 linear equations is solved:

$$\begin{cases} \mathbf{a}_{L(i)}^t [\mathbf{x}(i) - \hat{\mathbf{x}}(i)] = 0 \\ \mathbf{f}_{j(i)}^t [\mathbf{x}(i) - \hat{\mathbf{x}}(i)] = 0 \end{cases} \quad (23)$$

where $\hat{\mathbf{x}}(i) = b(i) \mathbf{a}_{L(i)} + g(i) \mathbf{f}_{j(i)}$.

3.3 GAINS ARE QUANTIZED USING GAIN CODEBOOKS $\{b_k\}, \{g_m\}$

The gain correction procedure (23) is performed only once after the shapes are chosen. The choice of shapes $\mathbf{a}_{L(i)}, \mathbf{f}_{j(i)}$ is not optimal. The choice of $\mathbf{f}_{j(i)}$ may be improved if the proper gains (23) are calculated for any candidate $\mathbf{f}_j$ i.e. inside the search algorithm. The calculation of gains may be simplified if the system (23) is decoupled in 2 independent equations, i.e. if $\mathbf{a}_{L(i)}^t \mathbf{f}_j = 0$. This requires the orthogonalization of the shape codebook $\{\mathbf{f}_j\}$ with respect to the chosen LTP vector $\mathbf{a}_{L(i)}$. The resulting modelling algorithm with shape codebook orthogonalization may be described as follows:

1. The best LTP shape $\mathbf{a}_{L(i)}$ and gain $b_{k(i)}$ are selected as in the standard modelling algorithm. (formulas 21,22)

2. The shape codebook $\{\mathbf{f}_j\}$ is orthogonalized with respect to $\mathbf{a}_{L(i)}$, i.e. from each codebook vector $\mathbf{f}_j$ it is subtracted its orthogonal projection on $\mathbf{a}_{L(i)}$:

$$\mathbf{f}_j^{ort} = \mathbf{f}_j - \frac{\mathbf{a}_{L(i)}^t \mathbf{f}_j}{\|\mathbf{a}_{L(i)}\|^2} \mathbf{a}_{L(i)} = \mathbf{f}_j - \frac{\mathbf{a}_{L(i)} \mathbf{a}_{L(i)}^t}{\|\mathbf{a}_{L(i)}\|^2} \mathbf{f}_j = \mathbf{K}_{L(i)} \mathbf{f}_j \quad (24)$$

where $\mathbf{K}_L = \mathbf{I} - \frac{\mathbf{a}_L \mathbf{a}_L^t}{\|\mathbf{a}_L\|^2}$ ($\mathbf{I}$ denotes unit matrix) is the operator of orthogonalization with respect to $\mathbf{a}_L$.

3. The best pair $g_{m(i)}^{ort} \mathbf{f}_{j(i)}^{ort}$ is searched in the second stage of modelling. Since $\{\mathbf{f}_j^{ort}\}$ is orthogonal to $\mathbf{a}_{L(i)}$, modelling is independent on the chosen LTP vector. Thus, the standard approach may be used, i.e. $\min \varphi[\mathbf{x}(i), \mathbf{f}_j^{ort}]$ for the choice of $\mathbf{f}_{j(i)}^{ort}$ and the scalar quantization of $g^{ort}(i) = \mathbf{x}^t(i) \mathbf{f}_{j(i)}^{ort} / \|\mathbf{f}_{j(i)}^{ort}\|^2$ to obtain $g_{m(i)}^{ort}$. Notice that this formula for $g^{ort}(i)$ stems from the second equation of the system (23). In the gain quantization algorithm, any modification described in the sect.2.3 (e.g. many gain codebooks) may be applied.

3.4 CODEBOOK DESIGN ALGORITHMS

The LTP shape codebook $\{\mathbf{a}_L\}$ is not calculated, it is created out of the delayed samples of the decoded (output) signal $\hat{\mathbf{x}}$. The remaining codebooks $\{b_k\}, \{g_m\}, \{\mathbf{f}_j\}$ should be calculated. This is a complex procedure due to the closed loop structure of the coder (fig.6a.). Training data $\{\mathbf{e}(i)\}$ for design of the codebooks $\{g_m\}, \{\mathbf{f}_j\}$ depend on the LTP vector $\hat{\mathbf{x}}(i)$, which depends on the output signal $\hat{\mathbf{x}}$, depending on the

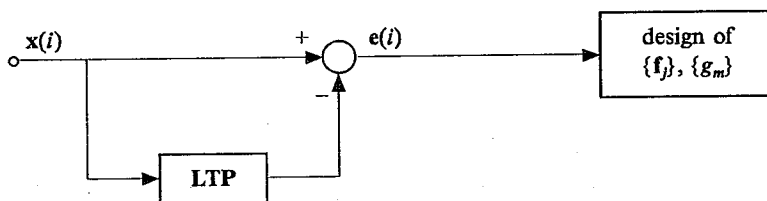


Fig. 7. Codebook design in the open loop scheme

codebooks $\{g_m\}$, $\{f_j\}$ to be designed. Therefore the open loop structure (fig.7) is used for codebook design.

The adaptive codebook $\{a_L\}$ is populated with the delayed input signal x . The unquantized LTP gain values $\{b(i)\}$ (22) are used as training data for LTP gain codebook design $\{b_k\}$. The LTP error vectors $\{e(i)\}$ serve as training data to design of $\{g_m\}$, $\{f_j\}$ as described in the section 2.

If orthogonalization is to be used in codebook search, it has to be taken into consideration in codebook design. The local distortion $D(i)$ may be described as follows:

$$D(i) = \|e(i) - g_{m(i)}^{ort} f_{j(i)}^{ort}\|^2 = \|e(i) - g_{m(i)}^{ort} K_{L(i)} f_{j(i)}\|^2 \quad (25)$$

Using (25) in (7) and (8) results in new formulas for shape and gain centroids. Thus, instead of (9), the following formula should be used for gain centroid:

$$g_m^{new} = \frac{\sum_{\substack{i \\ e(i) \in P(f_j)}} e^T(i) K_{L(i)} f_{j(i)}}{\sum_{\substack{i \\ e(i) \in P(f_j)}} \|K_{L(i)} f_{j(i)}\|}. \quad (26)$$

Instead of (10), the following set of linear equations is obtained for the shape centroid:

$$\left[\sum_{\substack{i \\ e(i) \in P(f_j)}} g_{m(i)}^2 K_{L(i)} \right] f_j^{new} = \sum_{\substack{i \\ e(i) \in P(f_j)}} g_{m(i)} K_{L(i)} e(i). \quad (27)$$

In derivation of (27) symmetry ($K_L^T = K_L$) and idempotence ($K_L K_L = K_L$) of the orthogonalization operator have been used. Codebooks obtained in the open loop scheme may be iteratively recalculated in the closed loop structure (fig.6a.). New codebooks yield new training data $\{e(i)\}$ that are used to obtain new codebooks etc. The closed loop design procedure is computationally complex and rather slowly convergent. There is no theoretical proof for its convergence.

3.5 RESULTS OF SIMULATIONS

In oder to compare several variants of signal modelling and codebook design algorithm, the bit rate 9.6 kbit/s has been chosen. The 9.6 kbit/s speech coders may be used in most telephone lines, in digital radiotelephones and in satellite communications. At this bit rate the vector dimension N must be limited to about 20 in order to keep the reasonable complexity of the modelling algorithm. The available 24 bits per vector are uniformly distributed between the LTP (12 bits) and the SGVQ (12 bits). These 12 bits have to be allocated to shape and gain quantizers. Some adequate experiences have been reported in the Tab. 3.

Tab. 3

SNRseg [dB] for the standard SGVQ with LTP at 9.6 kbit/s

LTP		SGVQ		SNRseg
L_s	L_g	L_s	L_g	
128	32	128	32	11.29
128	32	64	64	10.94
128	32	256	16	11.67
256	64	256	16	11.38

The best shape—gain tradeoff has been obtained for the predictive codebook containing 128 vectors, LTP gain quantization in 5 bits, SGVQ shape codebook containing 256 vectors and SGVQ gain quantization in 4 bits. These values are maintained in the experiences reported in the Tab. 4. In order to compare several codebook search and codebook design algorithms, the training sequence (used for codebook design) and the non—training sequence have been applied. The LTP gain codebook is obtained using Lloyd algorithm and the unquantized values $b(i)$ as training data. Iterations start with the uniform quantizer. The LBG algorithm with splitting [6] has been also tested, but it yielded no better results. For SGVQ gain quantization, two approaches have been studied: a single gain codebook and many gain codebooks (sect. 2.3).

Tab. 4

SNRseg [dB] for different SGVQs with LTP at 9.6 kbit/s

Modelling algorithm	SGVQ codebook design algorithm	single gain		many gains	
		Phrase		Phrase	
		training	non-training	training	non-training
standard	standard open loop	11.67	10.66	12.14	11.26
gain correction	standard open loop	11.95	11.04	12.43	11.67
orthogonalization	standard open loop	12.12	11.30	12.58	11.75
orthogonalization	modified open loop (formulas. 26,27)	12.60	11.72	12.80	11.82
orthogonalization	modified closed loop	12.90	11.72	12.98	12.03

The results justify the proposed modifications. The codebook search algorithm based on the orthogonalization of the SGVQ shape codebook performs better than the standard modelling algorithm and the algorithm with gain correction. The codebook design algorithm taking into account the orthogonalization (new formulas for shape and gain centroids) is better than the standard algorithm. Further improvement is obtained using the closed-loop codebook design algorithm and SGVQ gain coding using many codebooks.

4. ORTHOGONAL TRANSFORMS AND PRODUCT CODE VECTOR QUANTIZERS

4.1 CONCLUSIONS FROM THE HIGH RESOLUTION QUANTIZATION THEORY

Orthogonal transforms such as DFT, DCT are widely used for coding of audio and video signals. Could orthogonal transforms improve performance of vector quantizer? The high resolution quantization theory [14] may give answer to this question.

Let $\{x\}$ be a set of N -dimensional vectors of a zero mean gaussian pdf, and $R(x) = E[xx^t]$ the covariance matrix. These vectors are to be quantized using a vector quantizer of a high resolution b [bits/sample]. High resolution implies a big number of codebook vectors $L = 2^{bN}$. It may be shown [14] that for the optimal codebook, the distortion (energy of quantization error per sample) is given by the following formula:

$$D(b, N) = C(N) [\det R(x)]^{\frac{1}{N}} 2^{-2b} \quad (28)$$

where $C(1) = \frac{\sqrt{3}}{2} \pi$, $C(N) \simeq \frac{C(1)-1}{N} + 1$ and $\det R$ denotes determinant of the matrix R .

The SNR may be calculated as $\sigma_x^2 / D(b, N)$, where σ_x^2 is the signal energy per sample i.e. the diagonal element of $R(x)$:

$$\text{SNR}(b, N) = C^{-1}(N) \frac{\sigma_x^2}{[\det R(x)]^{\frac{1}{N}}} 2^{2b} \quad (29)$$

SNR always grows with N , but its growth is faster if signal is highly correlated. If there are no correlations ($R(x)$ is a diagonal matrix), then $[\det R(x)]^{\frac{1}{N}} = \sigma_x^2$ and SNR attains its minimum. Performance of the VQ may be improved by increasing N but soon the number of codebook vectors $L = 2^{bN}$ becomes so high that the real-time processing is not possible.

Let F^t be an orthogonal transform ($FF^t = I$) and $\bar{x} = F^t x$ the transformed vector. The orthogonal transform preserves signal energy ($\sigma_x^2 = \frac{1}{N} E[x^t x] = \frac{1}{N} E[\bar{x}^t \bar{x}] = \sigma_{\bar{x}}^2$) and determinant of covariance matrix ($\det R(\bar{x}) = \det R(x)$). Therefore, the optimal VQ in transform domain ($\bar{x}$) yields the same SNR as VQ of input signal (x) — see formula (29). The orthogonal transform may be regarded as the rotation of coordinates: the

transformed vector $\bar{x}$ is just x expressed in a new system of coordinates and the optimal VQ in transform domain ($\bar{x}$) uses the transformed codebook obtained in the input signal domain (x). There is no improvement in SNR, but the orthogonal transform may offer some reduction of complexity. It is possible if $\bar{x}$ is less correlated than x . In this case $\bar{x}$ may be split into small vectors of low dimension and quantized using low-complexity VQ's — without substantial loss of SNR [15]. In the extreme case the scalar quantizers may be used for each component of $\bar{x}$. Such decorrelating properties exhibits the Karhunen—Loeve transform (KLT). For the KLT, the matrix F consists of eigenvectors of $R(x)$, reordered using the decreasing tendency of corresponding eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_N$. The covariance matrix $R(\bar{x})$ is a diagonal matrix consisting of eigenvalues, i.e. the components of $\bar{x}$ are not correlated.

Let's assume that $\bar{x}$ is split in two parts: $\bar{x}^1$ of dimension N_1 and $\bar{x}^2$ of dimension N_2 ($N_1 + N_2 = N$). Both parts are coded using separate VQ's. The bit and dimension allocation problem may be formulated as follows [16]: find dimensions N_1 , $N_2 = N - N_1$ and resolutions b_1, b_2 ($b_1 N_1 + b_2 N_2 = bN$) for minimum complexity and possibly low distortion. Minimum complexity will be obtained if two codebooks contain the same number of vectors: $2^{b_1 N_1} = 2^{b_2 N_2}$ i.e. $b_1 N_1 = b_2 N_2 = bN/2$. Distortion per vector component may be obtained using (28):

$$\bar{D}(b_1, b_2, N_1, N_2) = \frac{1}{N} \left\{ N_1 C(N_1) [\det R(\bar{x}^1)]^{\frac{1}{N_1}} 2^{-2b_1} + \right. \\ \left. + N_2 C(N_2) [\det R(\bar{x}^2)]^{\frac{1}{N_2}} 2^{-2b_2} \right\}. \quad (30)$$

The determinants $\det R(\bar{x}^1)$ and $\det R(\bar{x}^2)$ are easily computed as products of corresponding eigenvalues: $\det R(\bar{x}^1) = \prod_{i=1}^{N_1} \lambda_i$, $\det R(\bar{x}^2) = \prod_{i=N_1+1}^N \lambda_i$. Minimization of (30) with a condition $b_1 N_1 = b_2 N_2 = bN/2$ yields the optimum bit and dimension allocation. This PCVQ with concatenation of small vectors in transform domain yields lower SNR than the VQ of full dimension N , but complexity is substantially reduced. The same complexity reduction in the input signal domain (i.e. VQ of vectors of dimension $N/2$) may be obtained at the cost of a bigger drop of SNR:

$$D(b, N) \leq \bar{D}(b_1, b_2, N_1, N_2) \leq D\left(b, \frac{N}{2}\right). \quad (31)$$

4.2 PRODUCT CODE VECTOR QUANTIZERS WITH KLT

In this section a speech coder using KLT of the long-term prediction error and PCVQ in transform domain will be presented. Algorithms for KLT calculation and codebook design will be described in the next section. The scheme of coder is given in the fig. 8.

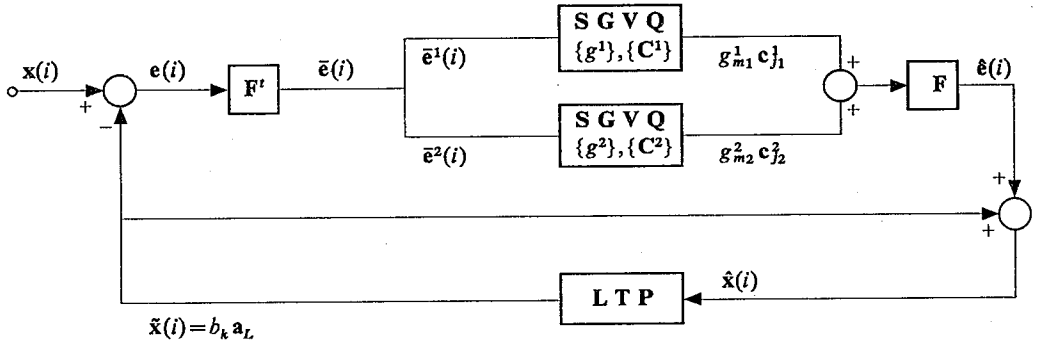


Fig. 8. Speech coder using the LTP, KLT and PCVQ in transform domain

The LTP error vector $\mathbf{e}(i) = \mathbf{x}(i) - \tilde{\mathbf{x}}(i)$ is transformed using the KLT: $\bar{\mathbf{e}}(i) = \mathbf{F}' \mathbf{e}(i)$ then it is split in two parts: $\bar{\mathbf{e}}^1(i)$ of dimension N_1 and $\bar{\mathbf{e}}^2(i)$ of dimension $N_2 = N - N_1$. Each part is quantized using a separate SGVQ with shape codebooks $\{\mathbf{C}^1\}$, $\{\mathbf{C}^2\}$ and gain codebooks $\{g^1\}$, $\{g^2\}$ respectively. Multiple gain codebooks may be used as described in the sect. 2.3.

The signal modelling (codebook search) algorithms described in sect. 3.2 may be extended to three stages and applied to the coder given in the fig. 8.

The standard algorithm may be described as follows:

1. Find the LTP $\tilde{\mathbf{x}}(i) = b_{k(i)} \mathbf{a}_{L(i)}$ using the formulas (21), (22).
2. Calculate the LTP error $\mathbf{e}(i)$, use the KLT to obtain $\bar{\mathbf{e}}(i) = \mathbf{F}' \mathbf{e}(i)$, and split $\bar{\mathbf{e}}(i)$ into $\bar{\mathbf{e}}^1(i)$ and $\bar{\mathbf{e}}^2(i)$.
3. For each part $\bar{\mathbf{e}}^n(i)$; $n=1, 2$, apply the SGVQ algorithm: choose the shape $\mathbf{C}_{j_n(i)}^n$ for $\min \phi[\bar{\mathbf{e}}^n(i), \mathbf{C}_{j_n(i)}^n]$, calculate the unquantized gain $g^n(i) = [\bar{\mathbf{e}}^n(i)]^T \mathbf{C}_{j_n(i)}^n / \|\mathbf{C}_{j_n(i)}^n\|^2$, then choose the gain $g_{m_n(i)}^n$ for $\min |g^n(i) - g_{m_n(i)}^n|$.
4. Concatenate the obtained models $g_{m_1(i)}^1 \mathbf{C}_{j_1(i)}^1$ and $g_{m_2(i)}^2 \mathbf{C}_{j_2(i)}^2$ to obtain the N -dimensional vector $\hat{\mathbf{e}}(i)$, then apply the inverse transform $\hat{\mathbf{e}}(i) = \mathbf{F} \hat{\mathbf{e}}(i)$ and add the LTP to obtain the signal model $\hat{\mathbf{x}}(i) = \hat{\mathbf{e}}(i) + \tilde{\mathbf{x}}(i)$.

The whole modelling procedure may be described in N -dimensional space: N_1 -dimensional vectors $\bar{\mathbf{e}}^1, \mathbf{C}_{j_1}^1$ may be regarded as N -dimensional vectors with N_2 last components equal to zero and N_2 -dimensional vectors $\bar{\mathbf{e}}^2, \mathbf{C}_{j_2}^2$ as N -dimensional vectors with N_1 first components equal to zero. Concatenation may be regarded as vector addition. This is useful for description of the modelling algorithm with gain correction. After selection of the LTP vector $\mathbf{a}_{L(i)}$ and the shapes $\mathbf{C}_{j_1(i)}^1, \mathbf{C}_{j_2(i)}^2$ (using the standard algorithm), gains $b(i)$, $g^1(i)$ and $g^2(i)$ are recalculated. In the transform domain the transformed input vector $\bar{\mathbf{x}}(i) = \mathbf{F}' \mathbf{x}(i)$ is to be modelled using the transformed LTP vector $\bar{\mathbf{a}}_{L(i)} = \mathbf{F}' \mathbf{a}_{L(i)}$ and the shape vectors $\mathbf{C}_{j_1(i)}^1, \mathbf{C}_{j_2(i)}^2$:

$$\hat{\bar{\mathbf{x}}}(i) = b(i) \bar{\mathbf{a}}_{L(i)} + g^1(i) \mathbf{C}_{j_1(i)}^1 + g^2(i) \mathbf{C}_{j_2(i)}^2. \quad (32)$$

The model $\hat{\bar{x}}(i)$ is obtained by orthogonal projection of $\bar{x}(i)$ on the subspace span by $\bar{a}_{L(i)}$, $C_{j_1(i)}^1$ and $C_{j_2(i)}^2$, i.e. gains are calculated as a solution of the following system of linear equations:

$$\begin{cases} [\bar{a}_{L(i)}]^T [\bar{x}(i) - \hat{\bar{x}}(i)] = 0 \\ [C_{j_1(i)}^1]^T [\bar{x}(i) - \hat{\bar{x}}(i)] = 0 \\ [C_{j_2(i)}^2]^T [\bar{x}(i) - \hat{\bar{x}}(i)] = 0. \end{cases} \quad (33)$$

At the end gains are quantized using the proper gain codebooks $\{b_k\}$, $\{g_{m1}^1\}$, $\{g_{m2}^2\}$.

Further improvement may be obtained if the orthogonal projections are used within the modelling procedure. The algorithm may be described as follows:

1. Choose the best LTP vector $a_{L(i)}$ as in the standard algorithm i.e. minimizing distance between $x(i)$ and its orthogonal projection on $a_{L(i)}$.

2. Choose the best shape $C_{j_1(i)}^1$: for each codebook vector $C_{j_1}^1$ calculate the orthogonal projection of $\bar{x}(i) = F^T x(i)$ on the plane span by the vectors $\bar{a}_{L(i)} = F^T a_{L(i)}$ and $C_{j_1}^1$ (this requires solving the system of two linear equations). Quantize obtained gains $b(i)$ and $g^1(i)$ and calculate distance between $\bar{x}(i)$ and the obtained model. Select $C_{j_1}^1$ yielding the minimum distance.

3. Choose the best shape $C_{j_2(i)}^2$: for each candidate $C_{j_2}^2$ calculate the orthogonal projection of $\bar{x}(i)$ on the subspace span by $\bar{a}_{L(i)}$, $C_{j_1(i)}^1$ and $C_{j_2(i)}^2$ (solve the system (33)). Quantize gains $b(i)$, $g^1(i)$ and $g^2(i)$ and calculate distance between $\bar{x}(i)$ and its model. Select $C_{j_2(i)}^2$ yielding the minimum distance.

This algorithm stems from the same idea as the shape codebook orthogonalization algorithm described in the sect. 3.2: at each stage of modelling the orthogonal projection on the subspace is calculated.

4.3 KLT CALCULATION AND CODEBOOK DESIGN

Open-loop scheme (fig. 9) is used for KLT calculation and codebook design for the reasons described in the section 3.3. The estimate of the covariance matrix $R(e)$ is calculated using the LTP error vectors $\{e(i)\}$ of the whole training sequence $\{x(i)\}$:

$$R(e) = \sum_i e(i) e^T(i). \quad (34)$$

Then, eigenvectors are calculated and ordered using the decreasing tendency of the corresponding eigenvalues. Thus columns of the KLT matrix F are obtained. This matrix represents the properties of the whole training phrase. Better results could be obtained by adapting the KLT to the varying properties of speech signal. Such adaptation algorithm would be very complex.

The transformed LTP error vectors $\bar{e}(i)$ are split in 2 parts $\bar{e}^1(i)$ and $\bar{e}^2(i)$ and serve as training data to obtain the shape and gain codebooks using standard methods described in the sect.2. For gain quantization, multiple gain codebooks may be used (sect. 2.3).

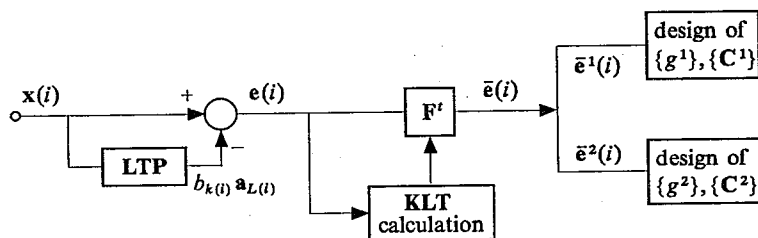


Fig. 9. Open-loop scheme for KLT calculation and codebook design

If orthogonal projections are to be used in the modelling algorithm, the codebook design algorithm should be modified. The optimal codebooks $\{g^1\}$, $\{C^1\}$, $\{g^2\}$, $\{C^2\}$ should be obtained in the same time using the generalized Lloyd algorithm (scheme like in the fig. 4, but with 4 partition and centroid calculation blocks). In order to simplify this procedure, shape codebooks are calculated first, using unquantized gain values. The local distortion may be expressed using (32) as follows:

$$D(i) = \|\bar{x}(i) - b(i) \bar{a}_{L(i)} - g^1(i) C_{j_1(i)}^1 - g^2(i) C_{j_2(i)}^2\|^2. \quad (35)$$

Let's take into consideration some vector from the initial codebook $\{C^1\}$. Partial distortion corresponding to this vector equals:

$$\begin{aligned} D_{j_1} &= \sum_i D(i) = \sum_{\substack{i \\ x(i) \in P(C_{j_1}^1)}} \|\bar{x}(i) - b(i) \bar{a}_{L(i)} - g^1(i) C_{j_1}^1 - g^2(i) C_{j_2(i)}^2\|^2 \\ &= \sum_{\substack{i \\ x(i) \in P(C_{j_1}^1)}} \|\bar{x}(i) - g^1(i) C_{j_1}^1\|^2 \end{aligned} \quad (36)$$

where $\bar{x}(i) \stackrel{\text{def}}{=} \bar{x}(i) - b(i) \bar{a}_{L(i)} - g^2(i) C_{j_2(i)}^2$.

Minimization of D_{j_1} with constraints: $\|C_{j_1}^1\| = 1$ and the last N_2 components of $C_{j_1}^1$ equal zero leads to the following formula for centroid calculation:

$$[C_{j_1}^1]^{\text{new}} = \frac{\sum_{\substack{i \\ x(i) \in P(C_{j_1}^1)}} g^1(i) [\bar{x}(i)]_0^{N_1}}{\left\| \sum_{\substack{i \\ x(i) \in P(C_{j_1}^1)}} g^1(i) [\bar{x}(i)]_0^{N_1} \right\|} \quad (37)$$

where $[\bar{x}(i)]_0^{N_1}$ — vector consisting of N_1 first components of $\bar{x}(i) = \bar{x}(i) - b(i) \bar{a}_{L(i)}$ and N_2 zeros. In the same way the centroid formulas for calculation of the shape codebook $\{C^2\}$ are obtained:

$$[C_{j_2}^2]^{\text{new}} = \frac{\sum_{\substack{i \\ x(i) \in P(C_{j_2}^2)}} g^2(i) [\bar{x}(i)]_{N_2}^0}{\left\| \sum_{\substack{i \\ x(i) \in P(C_{j_2}^2)}} g^2(i) [\bar{x}(i)]_{N_2}^0 \right\|} \quad (38)$$

where $[\tilde{\mathbf{x}}(i)]_{N_2}^0$ — vector consisting of N_1 zeros and N_2 last components of $\tilde{\mathbf{x}}(i)$.

After the calculation of shape codebooks, gain codebooks are calculated using the unquantized gains $g^1(i)$ and $g^2(i)$ as training data.

Codebooks obtained in the open-loop scheme (fig. 9) may be iteratively recalculated using the closed-loop scheme (fig. 8), as proposed in the sect. 3.3.

4.4 RESULTS OF SIMULATIONS

In order to compare the KLT-PCVQ described in the sect.4 with the PCVQ described in the sect.3, the same bit rate 9.6 kbit/s and the same speech phrases have been used. The vector length N has been increased to 30. The available 36 bits per vector are uniformly distributed between the LTP (12 bits: 7 for delay and 5 for gain) and two SGVQs (2×12 bits). Then the bit and dimension allocation problem has to be solved: find dimensions N_1 , $N_2 = N - N_1 = 30 - N_1$ and resolutions b_1 , b_2 that minimize the distortion $\bar{D}(b_1, b_2, N_1, N_2)$ (described with the formula 30, using $\bar{\mathbf{e}}^1$, $\bar{\mathbf{e}}^2$ instead of $\bar{\mathbf{x}}^1$, $\bar{\mathbf{x}}^2$) and fulfil the constraint $b_1 N_1 = b_2 N_2 = bN/2 = 12$. The estimate of the covariance matrix $\mathbf{R}(\mathbf{e})$ has been calculated using (34) and the determinants $\det \mathbf{R}(\bar{\mathbf{e}}^1)$, $\det \mathbf{R}(\bar{\mathbf{e}}^2)$ have been calculated for each N_1 , N_2 using the eigenvalues of $\mathbf{R}(\mathbf{e})$. The minimum distortion has been obtained for $N_1 = 5$ and $N_2 = 25$.

The problem of bit allocation to the shape and gain quantizers has been solved experimentally: one has selected for the 5-dimensional SGVQ 5 bits for gain and 7 bits for shape quantization and for the 25-dimensional SGVQ 4 bits for gain and 8 bits for shape quantization. The results of simulations for 3 modelling algorithms (sect. 4.2), 3 codebook design algorithms (sect. 4.3) and single and multiple gain codebooks (sect. 2.3) are given in the tab.5.

Tab. 5

SNRseg [dB] for different KLT-PCVQs at 9.6 kbit/s

Modelling algorithm	SGVQ codebook design algorithms	single gain codebook		many gain codebooks	
		Phrase		Phrase	
		training	non-training	training	non-training
standard	standard open loop	11.79	10.78	12.19	11.32
gain correction	standard open loop	12.32	11.12	12.39	11.37
orthogonal projections	standard open loop	12.80	11.30	12.96	11.77
orthogonal projections	modified open loop (formulas. 37,38)	13.23	11.72	13.41	11.81
orthogonal projections	modified closed loop	13.51	11.80	13.72	11.88

Modifications proposed in the sect. 4.2 and 4.3 lead to the improvement of speech quality. Codebook search algorithm based on the orthogonal projections within the modelling procedure performs better than the standard algorithm. Further improvement is obtained by re—designing the codebooks using the modified centroid formulas (37,38) and by repeating the codebook design procedure in the closed—loop scheme. The KLT—PCVQs outperform the PCVQs without transformation (compare tab. 4 and 5).

5. CONCLUSIONS

In this work, several PCVQs have been analyzed: the SGVQs, multistage VQs using the LTP and KLT-PCVQs based on the LTP, Karhunen—Loeve transform and concatenation of the SGVQ output vectors in transform domain. Several modifications of the existing algorithms have been proposed, yielding improvement of speech quality without substantial growth of coder's complexity. For the SGVQ, new gain coding algorithm has been proposed, using multiple gain codebooks. For the multistage VQ codebook orthogonalization techniques have proved useful, together with the appropriate codebook design algorithm. The idea of the KLT-PCVQ stems from the high—resolution quantization theory. Due to the improved signal modelling technique, based on the orthogonal projections in transform domain, and the modified codebook design algorithm the KLT-PCVQ outperforms the previously mentioned PCVQs. Based on these algorithms, several 9.6 kbit/s speech coders are proposed. These coders yield acceptable speech quality (SNRseg of 11—13 dB is comparable with the CELP coders at the same bit rate [10]). The informal listening tests lead to the same conclusion. The quantization noise is audible, but not annoying. The proposed coders are of acceptable complexity (KLT-PCVQ with gain correction requires about 3.4 Mflops and KLT-PCVQ with orthogonal projections about 18 Mflops), and may be implemented using the available signal processors.

REFERENCES

- [1] A. Gersho, R.M. Gray: *Vector quantization and signal compression*. Kluwer 1991
- [2] M. Sabin, R.M. Gray: *Product code vector quantizers for waveform and voice coding*. IEEE Trans. on ASSP—32, June 1984
- [3] M. Sabin: *Fixed shape adaptive gain vector quantization for speech waveform coding*. Speech Commun. 1989, 8
- [4] N. Moreau, P. Dymarski: *Successive orthogonalizations in the multistage CELP coder*. Int. Conf. on Acoustics, Speech and Signal Proc. vol. ICASSP'92, San Francisco
- [5] K.T. Malone, T.R. Fischer: *Contour-gain vector quantization* IEEE Trans. on Int. Conf. on Acoustics, Speech and Signal Proc. vol. ASSP—36, June 1988
- [6] Y.L. Linde, A. Buzo, R.M. Gray: *An algorithm for vector quantizer design*. IEEE Trans. on Communications vol. COM—28, Jan. 1980

- [7] T. B e r g e r: *Rate Distortion Theory*. Englewood Cliffs, Prentice-Hall, 1971
- [8] C.E. S h a n n o n: *A mathematical theory of communication*. Bell Syst. Tech. J. 1948 Vol.27, pp.379—423, 623—656
- [9] P. D y m a r s k i, A. C h m i e l e w s k i, S. R a t r e e: *Kodowanie wektorowe sygnału mowy z niezależnym kwantowaniem kształtu i poziomu sygnału*. Krajowe Sympozjum Telekomunikacji, Bydgoszcz 1992, ss. 154—163, pp.154—163
- [10] N. M o r e a u, P. D y m a r s k i: *Selection of excitation vectors for the CELP coder*. IEEE Tran. Speech and Audio Processing, vol.2, no.1, January 1994
- [11] P. D y m a r s k i, S. R a t r e e: *Kwantowanie wektorowe sygnału mowy z wykorzystaniem predykcji długookresowej i transformaty ortogonalnej*. Krajowe Sympozjum Telekomunikacji, Bydgoszcz 1993
- [12] G. D a v i d s o n, M. Y o n g, A. G e r s h o: *Real-time vector excitation coding of speech at 4800bps*. Proc. Int. Conf. on Acoustics, Speech and Signal Proc. ICASSP. 1987
- [13] J.H. Y a o, J.J. S h y n k, A. G e r s h o: *Low-delay VXC at 8kbit/s with interframe coding*. Proc. Int. Conf. on Acoustics, Speech and Signal Proc. ICASSP 1992
- [14] T.D. L o o k a b a u g h, R.M. G r a y: *High-resolution quantization theory and the vector quantizer advantage*. IEEE Trans. on Information Theory, September 1989
- [15] W. L i, Y.Q. Z h a n g: *A Study on the optimal attributes of transform domain vector quantization for image and video compression* IEEE Conf. on Communication, 1993
- [16] P. D y m a r s k i, N. M o r e a u, A. P o p e s c u: *Quantification vectorielle et transformees-etude de quelques gains de codage* Report ENST Dep. SIGNAL, 1993
- [17] C.F. G e r a l d: *Applied Numerical Analysis* 2nd ed., Addison-Wesley, Reading, Massachusetts, 1978

P. DYMARSKI, S. RATREE

KODOWANIE WEKTOROWE SYGNAŁU MOWY
Z WYKORZYSTANIEM DEKOMPOZYCJI, PREDYKCJI DŁUGOOKRESOWEJ
I TRANSFORMATY ORTOGONALNEJ

S t r e s z c z e n i e

Przedstawiono szereg kwantyzerów wektorowych opartych na zasadzie dekompozycji słownika (książki kodowej) na kilka mniejszych zbiorów. Opisano kwantyzerzy z niezależnym kodowaniem kształtu i poziomu sygnału, kwantyzerzy wieloetapowe z wykorzystaniem predykcji długookresowej, a także kwantyzerzy wykorzystujące transformatę Karhunen-Loevego. Zaproponowano szereg nowych algorytmów przeszukiwania słowników kwantyzerów wektorowych i ich projektowania. Połączenie różnych technik dekompozycji słownika pozwoliło na zaprojektowanie kodera sygnału mowy o przepływności 9.6 kbit/s.

WŁAŚCIWOŚCI STRUMIENIA BŁĘDÓW DLA TRANSMISJI BINARNEJ W KANALE Z WOLNYMI ZANIKAMI OPISANYMI ROZKŁADEM CZTEROPARAMETROWYM

KRYSTYNA NOGA

*Wyższa Szkoła Morska
Katedra Automatyki Okrętowej, Gdynia*

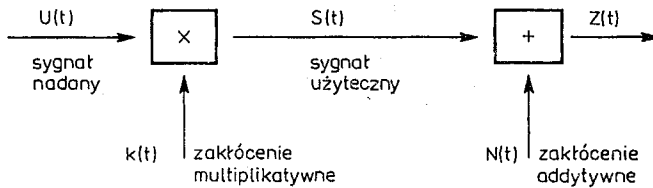
*Otrzymano 1994.07.10
Autoryzowano do druku 1994.11.10*

W artykule przedstawiono dynamiczne prawdopodobieństwo błędu elementowego dla transmisji binarnej przez kanał radiokomunikacyjny z wolnymi zanikami opisanymi rozkładem czteroparametrowym oraz z szumem gaussowskim. Opisano również właściwości strumienia błędów, tj. przedstawiono dynamiczne wagowe widmo błędów, dynamiczne prawdopodobieństwo błędnego bloku oraz dynamiczne prawdopodobieństwo odstępu λ_0 pomiędzy kolejnymi błędami. Analizy dokonano przy założeniu, że do przesyłania informacji wykorzystane są sygnały z modulacją NCFSK oraz DPSK. Uzyskane zależności przedstawiono graficznie.

WSTĘP

W okresie ostatniego trzydziestolecia powstało wiele prac dotyczących oceny jakości transmisji binarnej w kanale radiokomunikacyjnym [1, 2, 4–7, 9–17]. Oddziaływanie ośrodka propagacyjnego, tj. okołoziemskiej przestrzeni wraz z atmosferą, na rozprzestrzenianie się fal radiowych ma w dużym stopniu charakter losowy. Analiza systemu przesyłania informacji binarnych przez kanał radiokomunikacyjny, który obejmuje część systemu od analogowego wejścia nadajnika do analogowego wyjścia odbiornika, za pomocą sygnałów binarnych wymaga więc probabilistycznego opisu. Niezbędne jest uwzględnienie stochastycznych właściwości ośrodka propagacyjnego i jego oddziaływanie na przesyłane sygnały. Warunki propagacji fal radiowych są bardzo złożone. Na jakość transmisji mają wpływ między innymi odbicia, interferencje. Podczas transmisji sygnał nadany ulega losowym zakłóceniom multiplikatywnym $k(t)$ oraz addytywnym $N(t)$. Działanie zakłóceń multiplikatywnych

nych traktujemy jako przypadkowe tłumienie sygnału, zjawisko to nazywamy zanikami sygnału. Zaniki są cechą charakterystyczną przede wszystkim radiokomunikacji wykorzystującej propagację jonosferyczną i troposferyczną. Propagacja jonosferyczna dla fal krótkich (pasmo 3, ..., 30 MHz) wykorzystywana jest między innymi przez radiokomunikację dalekiego zasięgu. Jest ona związana z istnieniem warstw jonosfery D , E , F_1 , F_2 [3]. Zjonizowane w różnym stopniu warstwy jonosfery powodują odbicie w kierunku Ziemi padającej ukośnie fali elektromagnetycznej. Ponadto powierzchnia Ziemi również powoduje odbicie tej fali. Propagacja jonosferyczna polega więc na wielokrotnym odbijaniu się fali elektromagnetycznej od warstw jonosfery i powierzchni Ziemi. W wyniku tego sygnał odbierany jest sumą wielu interferujących ze sobą sygnałów docierających różnymi drogami.



Rys.1. Analogowy model liniowego kanału radiokomunikacyjnego

Dla liniowego modelu kanału radiokomunikacyjnego (rys. 1) sygnał $Z(t)$ odbierany przez odbiornik jest sumą wąskopasmowego sygnału użytecznego $S(t)$ oraz zakłócenia addytywnego $N(t)$, czyli

$$Z(t) = S(t) + N(t) = k(t) \times U(t) + N(t) \quad (1)$$

gdzie $U(t)$ jest sygnałem nadanym.

Jako model zakłóceń addytywnych najczęściej przyjmuje się szum gaussowski o zerowej wartości średniej. Natomiast wartość chwilową sygnału użytecznego można zapisać jako

$$S(t) = R(t) \cos [\omega_0 t + \varphi(t)] \quad (2)$$

gdzie ω_0 jest pulsacją nośnej, $R(t)$ jest obwiednią, a $\varphi(t)$ fazą sygnału użytecznego. Obwiednia sygnału użytecznego zależy wyłącznie od zaników tylko w przypadku sygnałów z modulacją kątową, czyli dla PSK (ang. *phase shift keying*) i FSK (ang. *frequency shift keying*), w dalszej części pracy tylko te modulacje będą rozpatrywane.

2. ROZKŁAD OBWIEDNI SYGNAŁU UŻYTECZNEGO

Model matematyczny zakłóceń multiplikatywnych zależy od zakresu częstotliwości, zasięgu łączności i rodzaju służby. Najczęściej do opisu zaników stosuje się rozkłady Rayleigha, Rice'a, Nakagamiego. Jednoparametrowy rozkład Rayleigha dobrze opisuje obwiednię sygnału użytecznego dochodzącego do odbiornika po

odbiciach od jonosfery w przypadku braku składnika zdeterminowanego oraz po rozproszeniu w troposferze [4, 11, 13, 15]. W przypadku występowania sygnału zdeterminowanego do opisu obwiedni sygnału użytecznego stosuje się rozkład Rice'a. Rozkład ten może być stosowany w przypadku interferencji zdeterminowanego sygnału bezpośredniego z sygnałami o wypadkowym normalnym rozkładzie wartości chwilowej [2, 7, 9, 12, 15, 16]. Znacznie szerszą klasę mechanizmów propagacyjnych odzwierciedla rozkład Nakagami [1, 5, 6, 13, 14, 15, 17]. Rozkład ten w szczególnym przypadku dobrze aproksymuje rozkład Rice'a oraz Rayleigha. Do opisu zaników można również zastosować rozkład czteroparametrowy [5, 6], który odzwierciedla jeszcze szerszą klasę mechanizmów propagacyjnych niż rozkład Nakagami.

W kanale z zanikami opisanymi rozkładem czteroparametrowym obwiednia $R(t)$ sygnału użytecznego, przy założeniu, że składowe ortogonalne tego sygnału mają rozkład normalny o różnej wariancji $\sigma_x^2 \neq \sigma_y^2 \neq 0$ i o różnej wartości średniej $m_x \neq m_y \neq 0$, posiada następującą gęstość prawdopodobieństwa [5, 6]

$$p(r) = \frac{r}{\sigma_x \sigma_y} \exp \left(-\frac{r^2}{2\sigma_x^2} - \frac{m_x^2 \sigma_y^2 + m_y^2 \sigma_x^2}{2\sigma_x^2 \sigma_y^2} \right) \times \sum_{i=0}^{\infty} \sum_{s=0}^{\infty} \frac{(2i+2s-1)!!}{2^i (2s)! i!} \times \frac{(\sigma_y^2 - \sigma_x^2)^i}{\sigma_y^{2i+4s} m_x^{i+s}} \times r^{i+s} I_{i+s} \left(\frac{rm_x}{\sigma_x^2} \right), \quad (3)$$

gdzie $I_i(\alpha)$ jest zmodyfikowaną funkcją Bessela pierwszego rodzaju, rzędu i

$$I_i(\alpha) = \sum_{k=0}^{\infty} \frac{(a/2)^{i+2k}}{k! \Gamma(i+k+1)}$$

$\Gamma(a) = \int_0^{\infty} t^{a-1} \exp(-t) dt$ jest funkcją gamma.

W literaturze, np. [5, 6], można znaleźć jeszcze inne postacie rozkładu czteroparametrowego.

Często zamiast parametrów m_y , m_x , σ_y^2 , σ_x^2 stosuje się inne cztery parametry posiadające określony sens fizyczny [5]:

$$1) \quad q^2 = \frac{m_x^2 + m_y^2}{\sigma_x^2 + \sigma_y^2} = \frac{r_p^2}{\sigma_x^2 + \sigma_y^2}, \quad (4.1)$$

q^2 określa stosunek średniej mocy składowej regularnej sygnału do składowej fluktuującej;

$$2) \quad \beta^2 = \frac{\sigma_x^2}{\sigma_y^2}, \quad (4.2)$$

β^2 określa stosunek wariancji składowych ortogonalnych sygnału $S(t)$; parametr β^2 charakteryzuje więc stopień niesymetrii kanału ze względu na składowe ortogonalne;

$$3) \quad \Psi = \arctg(m_y/m_x), \quad (4.3)$$

Ψ określa kąt fazowy wektora r_p w układzie współrzędnych X, Y , przy czym $m_x = r_p \cos \Psi$ oraz $m_y = r_p \sin \Psi$;

$$4) \quad r_0^2 = \sigma_x^2 + \sigma_y^2 + m_x^2 + m_y^2 \quad (4.4)$$

$r_0^2/2$ określa średnią moc sygnału $S(t)$.

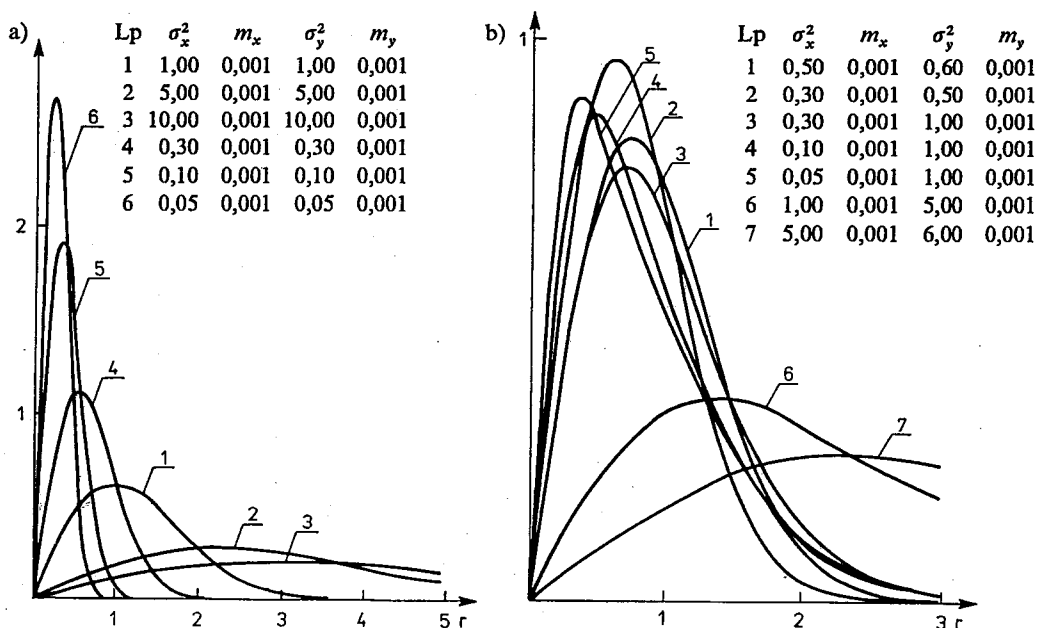
Pomiędzy parametrami r_0^2 , q^2 , β^2 , Ψ oraz σ_x^2 , σ_y^2 , m_x^2 , m_y^2 istnieją następujące zależności:

$$m_x^2 = \frac{q^2 r_0^2 \cos^2 \Psi}{1 + q^2} \quad (5.1); \quad m_y^2 = \frac{q^2 r_0^2 \sin^2 \Psi}{1 + q^2} \quad (5.2)$$

$$\sigma_x^2 = \frac{\beta^2 r_0^2}{(1 + \beta^2)(1 + q^2)} \quad (5.3); \quad \sigma_y^2 = \frac{r_0^2}{(1 + \beta^2)(1 + q^2)} \quad (5.4)$$

Rozkład czteroparametrowy dobrze aproksymuje następujące rozkłady:

- trójparametrowy, gdy $m_x \neq 0$, $m_y = 0$;
- Rice'a, gdy $\sigma_x^2 = \sigma_y^2 = \sigma^2$, $m_x^2 + m_y^2 \neq 0$;
- jednostronny normalny, gdy $m_x = m_y = 0$, $\sigma_x^2 = 0$;
- Rayleigha, gdy $m_x = m_y = 0$, $\sigma_x^2 = \sigma_y^2 = \sigma^2$;
- Hoyota, gdy $m_x = m_y = 0$ oraz $\sigma_x^2 \neq \sigma_y^2 \neq 0$.
- Nakagami, gdy jego parametry m oraz Ω spełniają zależności



Rys.2. Rozkład obwiedni sygnału użytecznego w kanale z zanikami opisanymi rozkładem czteroparametrowym

- dla $m_x = m_y = 0.001$ oraz $\sigma_x^2 = \sigma_y^2 = \sigma^2$ (przypadek ten aproksymuje rozkład Rayleigha)
- dla $m_x = m_y = 0.001$ oraz $\sigma_x^2 \neq \sigma_y^2$ (przypadek ten aproksymuje rozkład Hoyta)

$$m = \frac{(\sigma_x^2 + \sigma_y^2 + m_x^2 + m_y^2)^2}{2(\sigma_x^4 + \sigma_y^4 + 2\sigma_x^2 m_x^2 + 2\sigma_y^2 m_y^2)} = \frac{[(1 + \beta^2)(1 + q^2)]^2}{2[1 + \beta^4 2q^2(1 + \beta^2)(\beta^2 \cos^2 \Psi + \sin^2 \Psi)]}$$

$$\Omega = \sigma_x^2 + \sigma_y^2 + m_x^2 + m_y^2$$

Natomiast dla $\sigma_x^2 = \sigma_y^2 = 0$ otrzymujemy kanał bez zaników.

Na rys. 2 przedstawiono rozkład obwiedni sygnału użytecznego dla różnych wartości parametrów m_y , m_x , σ_y^2 , σ_x^2 .

3. DYNAMICZNE PRAWDOPODOBIENSTWO BŁĘDU ELEMENTOWEGO

Jedną z miar jakości transmisji binarnej jest statyczne prawdopodobieństwo błędu elementowego, które dla kanału bez zaników dla odbioru niekoherentnego można obliczyć na podstawie zależności

$$P_s\left(\frac{r^2}{2N_0}\right) = \frac{1}{2} \exp\left(-\alpha \frac{r^2}{2N_0}\right), \quad (6)$$

gdzie: N_0 jest średnią mocą addytywnego szumu normalnego $N(t)$;

α jest współczynnikiem zależnym od rodzaju modulacji:

$\alpha = 1/2$ dla NCFSK (ang. *noncoherent frequency shift keying*),

$\alpha = 1$ dla DPSK (ang. *differentially coherent phase shift keying*).

W celu uwzględnienia zaników sygnału użytecznego należy obliczyć wartość średnią wyrażenia (6), otrzymujemy wówczas dynamiczne prawdopodobieństwo błędu elementowego określone zależnością

$$P_D(*) = \int_0^{\infty} P_s\left(\frac{r^2}{2N_0}\right) p(r) dr. \quad (7)$$

Dla dłuższych przedziałów czasowych wielkości charakteryzujące cechy sygnału należy traktować jako losowe, w związku z czym obliczone na podstawie (6) statyczne prawdopodobieństwo błędu elementowego jest realizacją zmiennej losowej. Wówczas też dynamiczne prawdopodobieństwo błędu elementowego (7) można przedstawić w równoważnej postaci

$$P_D(*) = E\left[P_s\left(\frac{r^2}{2N_0}\right)\right], \quad (8)$$

gdzie $E(*)$ jest operatorem uśredniania.

Moment w-tęgo rzędu zmiennej losowej reprezentującej statyczne prawdopodobieństwo błędu elementowego dla kanału z wolnymi zanikami opisanymi rozkładem czteroparametrowym, po odpowiednich przekształceniach i po uwzględnieniu zależności przedstawionych w [5], określa wzór:

$$E \left[P_s^w \left(\frac{r^2}{2N_0} \right) \right] = \left(\frac{1}{2} \right)^w \frac{N_0}{[(N_0 + \alpha w \sigma_x^2)(N_0 + \alpha w \sigma_y^2)]^{1/2}} \times \\ \times \exp \left[-\frac{\alpha w m_x^2}{2(N_0 + \alpha w \sigma_x^2)} - \frac{\alpha w m_y^2}{2(N_0 + \alpha w \sigma_y^2)} \right]. \quad (9)$$

Dynamiczne prawdopodobieństwo błędu elementowego w kanale z wolnymi zanikami opisanymi rozkładem czteroparametrowym, dla odbioru niekoherentnego, po uwzględnieniu (8) i (9), można obliczyć na podstawie [10]

$$P_D(m_x, m_y, \sigma_x^2, \sigma_y^2, N_0) = \frac{N_0}{2[(N_0 + \alpha \sigma_x^2)(N_0 + \alpha \sigma_y^2)]^{1/2}} \times \\ \times \exp \left[-\frac{\alpha m_x^2}{2(N_0 + \alpha \sigma_x^2)} - \frac{\alpha m_y^2}{2(N_0 + \alpha \sigma_y^2)} \right]. \quad (10)$$

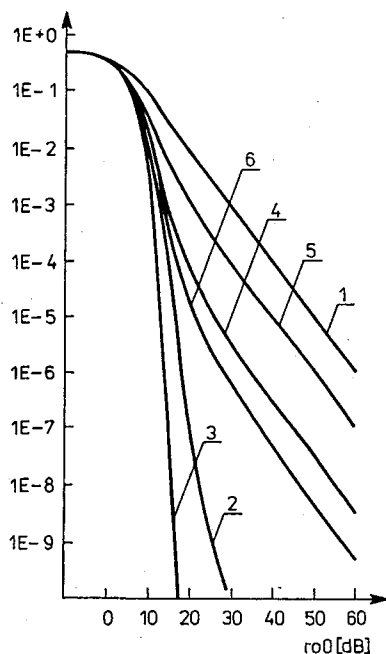
Dla kanału symetrycznego, gdy $\sigma_x^2 = \sigma_y^2 = \sigma^2$ otrzymujemy

$$P_D(m_x, m_y, \sigma^2, N_0) = \frac{N_0}{2(N_0 + \alpha \sigma^2)} \times \exp \left[-\frac{\alpha(m_x^2 + m_y^2)}{2(N_0 + \alpha \sigma^2)} \right]. \quad (11)$$

Ponadto podstawiając we wzorze (11) $m_x = m_y = 0$ otrzymujemy zależność na dynamiczne prawdopodobieństwo błędu elementowego dla kanału Rayleigha.

Dynamiczne prawdopodobieństwo błędu elementowego dla różnych wartości parametrów m_y , m_x oraz σ_y^2 , σ_x^2 zostało przedstawione na rys. 3, przy czym

Lp	σ_x^2	m_x	σ_y^2	m_y	α
1	0,05	0,001	0,05	0,001	0,5
2	0,05	1,000	0,05	1,000	0,5
3	0,05	4,000	0,05	4,000	0,5
4	0,50	2,000	0,50	2,000	0,5
5	1,00	2,000	1,00	2,000	0,5
6	0,10	1,000	0,10	1,000	0,5



Rys.3. Dynamiczne prawdopodobieństwo błędu elementowego dla odbioru NCFSK, gdy $m_x = m_y$ oraz $\sigma_x^2 = \sigma_y^2$

$r_0^2 = \rho_0 = r_0^2/N_0$ jest stosunkiem średniej mocy sygnału użytecznego do średniej mocy szumu addytywnego.

4. DYNAMICZNE WAGOWE WIDMO BŁĘDÓW

Dynamiczne prawdopodobieństwo błędu elementowego (10) nie charakteryzuje ani liczby błędów w bloku N -elementowym, ani ich rozmieszczenia. Podstawową cechą transmisji binarnej przez kanał radiowy jest powstawanie niejednorodnego strumienia błędów, czyli grupowanie się błędów elementowych w serie. Jedną z miar charakteryzujących strumień błędów jest prawdopodobieństwo wystąpienia n błędów w bloku o długości N , które dla NCFSK można określić na podstawie zależności

$$P_s(L_N = n) = \binom{N}{n} \left[P_s\left(\frac{r^2}{2N_0}\right) \right]^n \left[1 - P_s\left(\frac{r^2}{2N_0}\right) \right]^{N-n}, \quad (12)$$

gdzie L_N jest zmienną losową reprezentującą liczbę błędów elementowych w bloku o długości N sygnałów elementarnych.

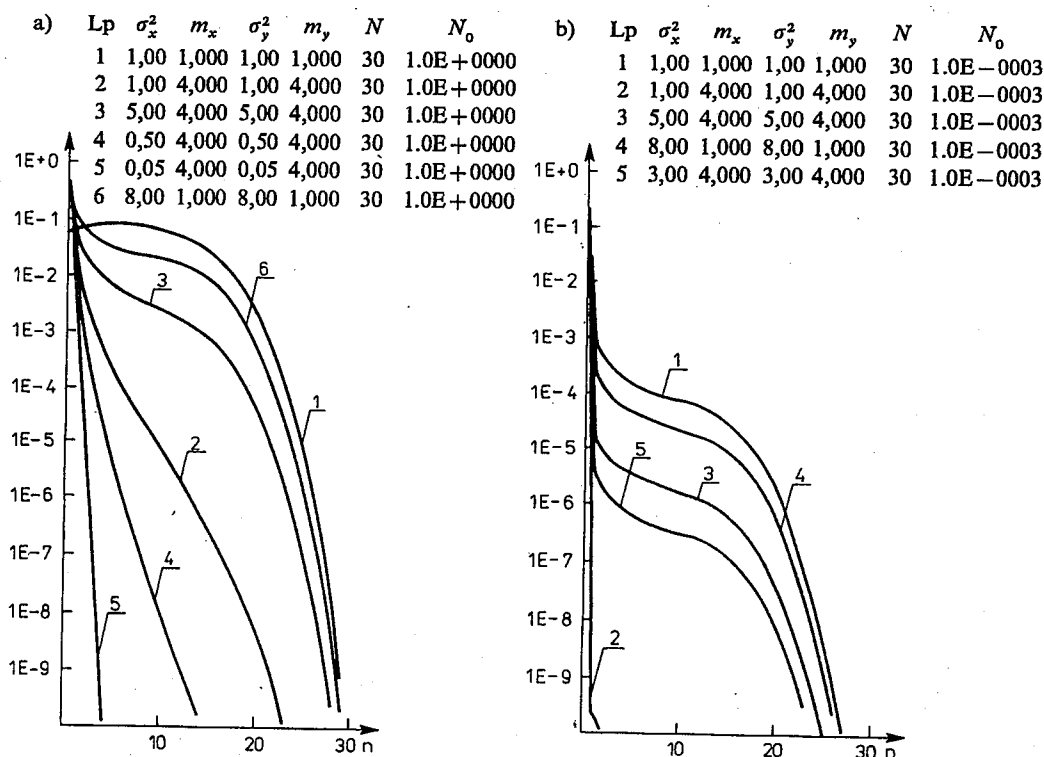
Należy zaznaczyć, że zależność (12) nie jest słuszna dla DPSK, gdyż różnicowa detekcja sąsiednich sygnałów elementarnych powoduje powstawanie par błędów. Prawdopodobieństwo określone wzorem (12) nazywamy wagowym widmem błędów elementowych lub rozkładem krotności błędów. W kanale z zanikami zależność (12) należy uśrednić ze względu na rozkład zmiennej reprezentującej obwiednię sygnału użytecznego. W wyniku uśrednienia uzyskuje się dynamiczne wagowe widmo błędów elementowych

$$\begin{aligned} P_D(L_N = n) &= \int_0^\infty P_s(L_N = n) p(r) dr = \\ &= \binom{N}{n} \sum_{i=0}^{N-n} \binom{N-n}{i} (-1)^i E \left[P_s^{i+n} \left(\frac{r^2}{2N_0} \right) \right]. \end{aligned} \quad (13)$$

Dla kanału z wolnymi zanikami opisanymi rozkładem czteroparametrowym zależność (13), po uwzględnieniu (9), przyjmuje postać [10]

$$\begin{aligned} P_D(L_N = n) &= 2N_0 \binom{N}{n} \sum_{i=0}^{N-n} \binom{N-n}{i} (-1)^i (1/2)^{i+n} \\ &\times \frac{\exp \left\{ -\frac{(i+n)m_x^2}{2[2N_0 + (i+n)\sigma_x^2]} - \frac{(i+n)m_y^2}{2[2N_0 + (i+n)\sigma_y^2]} \right\}}{\{[2N_0 + (i+n)\sigma_x^2][2N_0 + (i+n)\sigma_y^2]\}^{1/2}}. \end{aligned} \quad (14)$$

Dynamiczne wagowe widmo błędów dla różnych wartości parametrów m_y , m_x , σ_y^2 , σ_x^2 zostało przedstawione na rys. 4.



Rys.4. Dynamiczne wagowe widmo błędów dla odbioru NCFSK

a) dla $N = 30$, $N_0 = 0$ dB, $m_x = m_y$ oraz $\sigma_x^2 = \sigma_y^2$,
 b) dla $N = 30$, $N_0 = -30$ dB, $m_x = m_y$ oraz $\sigma_x^2 = \sigma_y^2$,

5. DYNAMICZNE PRAWDOPODOBIENSTWO BŁĘDNego BLOKU

Strumień błędów można również opisać przy pomocy prawdopodobieństwa wystąpienia w bloku N elementowym więcej niż n błędów, które dla kanału bez zaników dla odbioru NCFSK określa zależność

$$\begin{aligned}
 P_s(L_N > n) &= \sum_{k=n+1}^N P_s(L_N = k) = \\
 &= \sum_{k=n+1}^N \binom{N}{k} \left[P_s\left(\frac{r^2}{2N_0}\right) \right]^k \left[1 - P_s\left(\frac{r^2}{2N_0}\right) \right]^{N-k}.
 \end{aligned} \quad (15)$$

Zależność (15) często nazywamy prawdopodobieństwem błędnego bloku. Dla kanału z zanikami, poprzez uśrednienie (15), otrzymujemy dynamiczne prawdopodobieństwo błędnego bloku

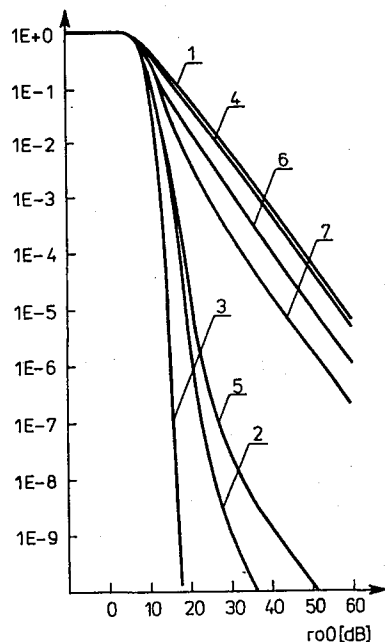
$$\begin{aligned}
 P_D(L_N > n) &= \int_0^{\infty} P_s(L_N > n) p(r) dr = \\
 &= \sum_{k=n+1}^N \binom{N}{k} \sum_{i=0}^{N-k} \binom{N-k}{i} (-1)^i E \left[P_s^{i+k} \left(\frac{r^2}{2N_0} \right) \right].
 \end{aligned} \quad (16)$$

Prawdopodobieństwo (16) dla kanału z wolnymi zanikami opisywanymi rozkładem czterometrowym, po uwzględnieniu (9), przyjmuje postać

$$\begin{aligned}
 P_D(L_N > n) &= 2N_0 \sum_{k=n+1}^N \binom{N}{k} \sum_{i=0}^{N-k} \binom{N-k}{i} (-1)^i (1/2)^{i+k} \times \\
 &\times \frac{\exp \left\{ -\frac{(i+k)m_x^2}{2[2N_0 + (i+k)\sigma_x^2]} - \frac{(i+k)m_y^2}{2[2N_0 + (i+k)\sigma_y^2]} \right\}}{\{[2N_0 + (i+k)\sigma_x^2][2N_0 + (i+k)\sigma_y^2]\}^{1/2}}.
 \end{aligned} \quad (17)$$

Dynamiczne prawdopodobieństwo błędnego bloku zostało przedstawione na rys. 5.

Lp	σ_x^2	m_x	σ_y^2	m_y	N	n
1	0,05	0,001	0,05	0,001	30	0
2	0,05	1,000	0,05	1,000	30	0
3	0,05	4,000	0,05	4,000	30	0
4	1,00	1,000	1,00	1,000	30	0
5	1,00	4,000	1,00	4,000	30	0
6	5,00	4,000	5,00	4,000	30	0
7	3,00	4,000	3,00	4,000	30	0



Rys.5. Dynamiczne prawdopodobieństwo wystąpienia co najmniej jednego błędu ($n = 0$) w bloku o długości $N = 30$ dla $m_x = m_y$ oraz $\sigma_x^2 = \sigma_y^2$.

6. DYNAMICZNE PRAWDOPODOBIEŃSTWO ODSTĘPU λ_0 POMIĘDZY BŁĘDAMI

Grupowanie się błędów w serie można także opisać przy pomocy zmiennej losowej A_0 reprezentującej odstęp pomiędzy kolejnymi błędami w strumieniu błędów. Prawdopodobieństwo to dla kanału bez zaników dla odbioru NCFSK wyraża się wzorem

$$P_s(A_0 = \lambda_0) = P_s\left(\frac{r^2}{2N_0}\right) \left[1 - P_s\left(\frac{r^2}{2N_0}\right) \right]^{\lambda_0}. \quad (18)$$

Z zależności (18) wynika, że odstęp pomiędzy niezależnymi błędami, ma rozkład geometryczny. Zmienna losowa A_0 określa więc długość bloku bezbłędnego, zdefiniowanego jako ciąg kolejnych elementów bezbłędnych występujących między dwoma błędami. Dla kanału z zanikami, po uśrednieniu zależności (18), dynamiczne prawdopodobieństwo odstępów λ_0 pomiędzy kolejnymi błędami przyjmuje postać

$$P_D(A_0 = \lambda_0) = \sum_{i=0}^{\lambda_0} \binom{\lambda_0}{i} (-1)^i E \left[P_s^{i+1} \left(\frac{r^2}{2N_0} \right) \right]. \quad (19)$$

Dla kanału z wolnymi zanikami opisanymi rozkładem czteroparametrowym, po uwzględnieniu (9), zależność (19) przyjmuje postać

$$P_D(A_0 = \lambda_0) = 2N_0 \sum_{i=0}^{\lambda_0} \binom{\lambda_0}{i} (-1)^i (1/2)^{i+1} \times \\ \times \frac{\exp \left\{ -\frac{(i+1)m_x^2}{2[2N_0 + (i+1)\sigma_x^2]} - \frac{(i+1)m_y^2}{2[2N_0 + (i+1)\sigma_y^2]} \right\}}{\{[2N_0 + (i+1)\sigma_x^2][2N_0 + (i+1)\sigma_y^2]\}^{1/2}}. \quad (20)$$

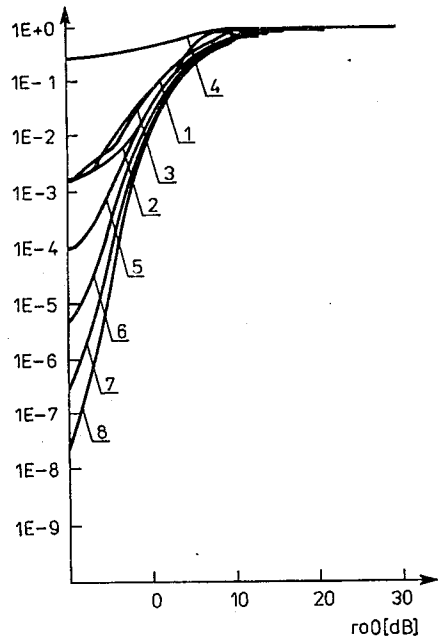
Oczywiście prawdopodobieństwo, że blok bezbłędny będzie nie krótszy niż λ_0 , wyraża się wzorem

$$P_s(A_0 \geq \lambda_0) = \left[1 - P_s\left(\frac{r^2}{2N_0}\right) \right]^{\lambda_0}. \quad (21)$$

Dla kanału z wolnymi zanikami opisanymi rozkładem czteroparametrowym prawdopodobieństwo, że blok bezbłędny będzie nie krótszy niż λ_0 określa zależność

$$P_D(A_0 \geq \lambda_0) = \int_0^\infty P_s(A_0 \geq \lambda_0) p(r) dr = \sum_{i=0}^{\lambda_0} \binom{\lambda_0}{i} (-1)^i E \left[P_s^i \left(\frac{r^2}{2N_0} \right) \right] = \\ = \frac{\exp \left[-\frac{i \times m_x^2}{2(2N_0 + i \times \sigma_x^2)} - \frac{i \times m_y^2}{2(2N_0 + i \times \sigma_y^2)} \right]}{[(2N_0 + i \times \sigma_x^2)(2N_0 + i \times \sigma_y^2)]^{1/2}} = 2N_0 \sum_{i=0}^{\lambda_0} \binom{\lambda_0}{i} (-1)^i (1/2)^i. \quad (22)$$

Lp	σ_x^2	m_x	σ_y^2	m_y	lambda
1	0,05	0,001	0,05	0,001	10
2	0,05	4,000	0,05	4,000	10
3	10,00	4,000	10,00	4,000	10
4	0,05	0,001	0,05	0,001	2
5	0,05	0,001	0,05	0,001	15
6	0,05	0,001	0,05	0,001	20
7	0,05	0,001	0,05	0,001	25
8	0,05	0,001	0,05	0,001	30



Rys.6. Dynamiczne prawdopodobieństwo odstępu nie mniejszego niż λ_0 pomiędzy kolejnymi błędami dla odbioru NCFSK dla $m_x = m_y$ oraz $\sigma_x^2 = \sigma_y^2$

Dynamiczne prawdopodobieństwo odstępu nie mniejszego niż λ_0 pomiędzy kolejnymi błędami zostało przedstawione na rys. 6.

Uzyskane zależności (10), (14), (17), (20), (22) umożliwiają ocenę jakości transmisji binarnej w kanale z wolnymi zanikami. Zastosowanie analizy probabilistycznej oraz rozkładu czteroparametrowego umożliwiło uzyskanie zależności, które uwzględniają szeroką klasę mechanizmów propagacyjnych i w miarę dokładnie opisują warunki panujące w ośrodku propagacyjnym. Zależności opisujące strumień błędów mogą być szczególnie przydatne przy poszukiwaniu nowych kodów detekcyjnych lub korekcyjnych, które odgrywają istotną rolę w systemach transmisji danych. Analiza jakości transmisji binarnej została ograniczona do przypadku kanału z wolnymi zanikami i addytywnym szumem gaussowskim. W przyszłości podobne rozważania należałoby przeprowadzić dla zaników szybkich oraz szumu różnego od gaussowskiego.

BIBLIOGRAFIA

1. E.K. Al-Hussaini, A.A.M. Al Bassiouni: *Performance of MRC diversity systems for the detection of signals with Nakagami fading*. IEEE Trans. on Com., vol. Com-33, No 12, December 1985
2. D. Chasé: *Digital signal design concepts for a time-varying Rician channel*. IEEE Trans. on Com., vol. Com-24, No 2, February 1976
3. M.P. Dołuchanov: *Propagacja fal radiowych*. WKŁ, Warszawa 1975

4. R.E. Eaves, A.H. Levesque: *Probability of block error for very slow Rayleigh fading in Gaussian noise*. IEEE Trans. on Com., vol. Com-25, No 3, March 1977
5. A.A. K ł o w s k i j: *Pieredacza dyskretnych soobszchenij pa radiokanalam*. Swiaz, Moskwa 1969
6. A.A. K ł o w s k i j, W.Sz. Kantorowicz, S.M. Sz i r o k o w: *Modeli nieprierywnych kanalow swiazi na osnovie stochastycznych differencjalnych urawnienij*. Radio i Swiaz, Moskwa 1984
7. W.C. L i n s d e y: *Error probabilities for Rician fading multichannel reception of binary and N-ary signals*. IEEE on IT, vol. IT-10 October 1964
8. M. N a k a g a m i: *The m-distribution a general formula of intensity distribution of fading, in statistical methods in radio wave propagation*. W. C. Hoffman, Ed. London, England: Pergamon, 1960
9. K. N o g a: *Prawdopodobieństwo błędu elementowego oraz właściwości strumienia błędów dla transmisji binarnej w kanale radiokomunikacyjnym z wolnymi zanikami Rice'a i addytywnym szumem gaussowskim*. Rozprawy Elektrotechniczne, 1985, Z. 4
10. K. N o g a: *Jakość transmisji binarnej dla odbioru niekoherentnego w kanale radiokomunikacyjnym z zanikami opisanymi rozkładem czteroparametrowym*. VII Krajowe Sympozjum Nauk Radiowych, Gdańsk, 18-19 luty 1993 r
11. C.E. S u n d b e r g: *Block error probability for noncoherent FSK with diversity for very slow Rayleigh fading in Gaussian noise*. IEEE Trans. on Com., vol. Com-29, No 1, January 1981
12. Z. S y r o k a: *Interferencja sygnałów radiowych w kanale Rice'a*. Przegląd Telekomuniacyjny, 1992, Nr 1
13. Z. S y r o k a, A. W o j n a r: *Systematic analysis of interference of fading radio signals*. International Wrocław Symposium on Electromagnetic Compatibility, Wrocław 1986
14. A. W o j n a r: *Probabilistyczna analiza zasięgów w radiokomunikacji*. Biuletyn Wojskowej Akademii Technicznej, Nr 10, październik 1983
15. A. W o j n a r: *Rayleigh, Rice and Nakagami—in search of efficient models of fading radio channels*. International Wrocław Symposium on Electromagnetic Compatibility, Wrocław 1988
16. A. W o j n a r, Z. S y r o k a: *Graniczna jakość transmisji w kanale Rice'a*. Biuletyn Wojskowej Akademii Technicznej, Nr 4, kwiecień 1985
17. A. W o j n a r: *Unknown bounds on performance in Nakagami channels*. IEEE Trans. on Com. vol. Com-34, No 1, January 1986

K. NOGA

ERROR'S STREAM PROPERTIES FOR BINARY TRANSMISSION IN SLOW FOUR PARAMETERS FADING CHANNELS

S u m m a r y

During the past three decades there have been several studies of communication system performance in Rayleigh, Rice, or Nakagami channels. In this paper the performance of NCFSK or DPSK in a slow four parameters fading channel with additive white Gaussian noise is presented. The bit error ratio is analyzed theoretically and graphically shown. In the open literature the problem of bit error rate performance for digital transmission under fading conditions has been extensively examined but block transmission performance has been largely ignored. In this paper the following expressions for NCFSK receiver are delivered: expression for the probability of n errors occurring in N bits (weighted spectrum of errors), expression for the probability of block error (the probability of more than n errors in a block of N bits), expression for the probability of λ_0 interval between elementary errors is no smaller than λ_0 .

621.317.791
621.3.011.22

Pomiar rezystancji cieplnej bipolarnych tranzystorów mocy Darlingtona

ZBIGNIEW LISIK, TADEUSZ MICHALSKI, ZBIGNIEW SZCZEPANIAK

*Instytut Elektroniki, Politechnika Łódzka**Otrzymano 1994.6.13**Autoryzowano do druku 1994.11.30*

Rezystancja cieplna jest jednym z podstawowych parametrów bipolarnych tranzystorów mocy, a jej wyznaczenie wymaga pomiaru maksymalnej temperatury występującej w pastylce półprzewodnikowej. W tym celu wykorzystuje się zwykle elektryczne metody pośrednie oparte o pomiar zmian wartości elektrycznego parametru termoczułego (EPT). Pomiar taki może być przeprowadzony na wiele sposobów, a wybór metody pomiarowej może mieć istotny wpływ na rezultat pomiaru.

Praca jest poświęcona przedstawieniu problemów występujących przy pomiarach rezystancji cieplnej bipolarnych tranzystorów mocy Darlingtona. Tranzystorów tych nie można traktować jako pojedynczego elementu dyskretnego, jak to ma miejsce w przypadku zwykłych tranzystorów bipolarnych, gdyż są one w rzeczywistości strukturami scalonymi mocy zawierającymi w sobie kilka przyrządów półprzewodnikowych. Przeprowadzone badania eksperymentalne wykazały, że zastosowanie do nich metod korzystających z napięcia przewodzącego złącza baza-emiter jako EPT, powszechnie stosowanych w przypadku zwykłych tranzystorów mocy, nie jest najlepszym rozwiązaniem. W tym przypadku znacznie lepszym EPT jest napięcie na przewodzącym złączu baza-kolektor mierzone w obszarze odpowiadającym diodzie zwrotnej układu zastępczego tranzystora. Uzyskane wyniki badań wykorzystano do opracowania opisanego w pracy miernika rezystancji cieplnej tranzystorów Darlingtona. Miernik ten został wykonany na zamówienie FP TEWA.

1. WSTĘP

Rezystancja cieplna należy od wielu lat do podstawowych parametrów znamionowych deklarowanych przez producentów przyrządów półprzewodnikowych mocy. Określa ona skuteczność odprowadzania ciepła z pastylki półprzewodnikowej przyrządu i jest wykorzystywana w praktyce inżynierskiej do określenia temperatury

struktury półprzewodnikowej, a w dalszej kolejności do określania dopuszczalnych warunków obciążenia przyrządu. Generalnie jest ona zdefiniowana wzorem:

$$R_{thj-r} = \frac{T_j - T_r}{P} \quad (1)$$

w którym T_j — tzw. temperatura złącza reprezentująca temperaturę struktury półprzewodnikowej, T_r — temperatura punktu odniesienia, a P — stała moc rozpraszana w przyrządzie. Wzór (1) definiuje pewną ogólną rezystancję termiczną złącze—punkt odniesienia. W praktyce jako temperaturę odniesienia przyjmuje się temperaturę obudowy T_c lub temperaturę otoczenia T_a , którym odpowiadają, odpowiednio, rezystancja złącze—obudowa R_{thj-c} oraz rezystancja złącze—otoczenie R_{thj-a} .

Temperatura złącza T_j reprezentuje we wzorze (1) temperaturę struktury półprzewodnikowej. Jej nazwa, temperatura złącza, jest związana z genezą pojęcia rezystancji cieplnej. Została ona pierwotnie wprowadzona dla przyrządów jednozłączowych małej mocy, w których zmiany temperatury w pastylce półprzewodnikowej są niewielkie. W przyrządach tych maksymalna temperatura występuje na złączu, jest ona łatwo mierzalna i stanowi podstawę do wyznaczenia ich rezystancji cieplnej. Inaczej wygląda sytuacja w wielozłączowych przyrządach mocy. Temperatury poszczególnych złącz mogą być w nich różne, a rozkład temperatury w płaszczyznach złącz może ulegać dużym zmianom [1–5]. Ponieważ głównym przeznaczeniem rezystancji cieplnej jest badanie czy wzrost temperatury nie przekracza wartości dopuszczalnych, temperaturę T_j traktuje się w tych przyrządach z reguły jako temperaturę maksymalną.

Rezystancja cieplna jest jednym z podstawowych parametrów bipolarnych tranzystorów mocy, a metodom wyznaczania w nich maksymalnej temperatury T_j , niezbędnej do obliczenia R_{th} w oparciu o wzór (1), poświęcono szereg prac zarówno teoretycznych [5–10] jak i eksperymentalnych [1, 2, 10–13]. Producenci bipolarnych tranzystorów mocy prowadzą bieżącą kontrolę wartości rezystancji cieplnej, która obejmuje często 100% wytwarzanych przyrządów i odbywa się już na linii produkcyjnej. Wymaga to dysponowania odpowiednimi automatycznymi miernikami rezystancji cieplnej. Miernik taki, przeznaczony zarówno do prac laboratoryjnych jak i do pracy na linii produkcyjnej został opracowany i wykonany dla potrzeb FP TEWA w Instytucie Elektroniki Politechniki Łódzkiej [14].

Przedstawienie tego miernika jest jednym z celów pracy. Będzie ono poprzedzone omówieniem problemów związanych z samym pomiarem temperatury wewnątrz struktury półprzewodnikowej tranzystora bipolarnego. Pomiar ten może być przeprowadzony na wiele sposobów, a wybór metody pomiarowej nie jest bez znaczenia dla ostatecznego rezultatu pomiaru. Wiele miejsca będzie również poświęcone stronie technicznej pomiaru, ze szczególnym uwzględnieniem jego specyfiki wynikającej z właściwości tranzystorów Darlingтона.

2. POMIARY PARAMETRÓW TERMICZNYCH TRANZYSTORÓW BIPOLARNYCH

2.1. PRZEGLĄD METOD POMIAROWYCH

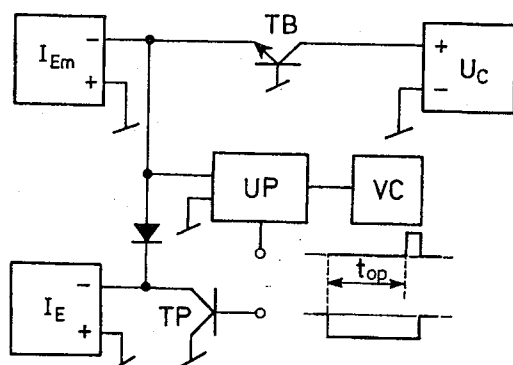
Wyznaczenie rezystancji cieplnej bipolarnego tranzystora mocy w oparciu o wzór (1) wymaga znajomości maksymalnej temperatury występującej w pastylce półprzewodnikowej. Ponieważ rozkład temperatury w pastylce może być silnie nierównomierny [5, 7, 8], najlepiej do tego celu nadawałyby się metody bezpośrednie, oparte o pomiary w podczerwieni [10, 15–18], korzystające z ciekłych kryształów [10, 20, 21] lub past fosforowych [22]. Pomiary takie mogą być jednak przeprowadzone jedynie na otwartych, odpowiednio przygotowanych strukturach i z tego względu nie nadają się do pomiarów na gotowych przyrządach.

W tym przypadku stosuje się elektryczne metody pośrednie oparte o pomiar zmian wartości elektrycznego parametru termoczułego (EPT). Ponieważ prawie wszystkie parametry przyrządu półprzewodnikowego zależą od temperatury, istnieje tu dużo różnych możliwości [23, 24]. Jako kryterium wyboru EPT przyjmuje się zwykle prostotę zależności tego parametru od temperatury, łatwość przeprowadzenia pomiaru oraz warunki pracy przyrządu, dla których pomiar ma być wykonywany. W tranzystorach bipolarnych wykorzystuje się w tym celu:

- spadek napięcia w kierunku przewodzenia na złączu emiter—baza U_{EB} [23–25],
- spadek napięcia w kierunku przewodzenia na złączu kolektor—baza U_{CB} [24, 26],
- prąd upływu złącza kolektor—baza I_{CBO} [23, 24, 27]
- współczynnik wzmocnienia prądowego α_F [23, 24, 28].

W latach osiemdziesiątych dominującą pozycję wśród metod pomiaru temperatury w tranzystorach bipolarnych oraz ich rezystancji cieplnej uzyskały metody wykorzystujące jako EPT napięcie na przewodzącym złączu emiter—baza U_{BE} , a jedna z nich — tzw. metoda przełączanego prądu emitera — została uznana jako standard w normach rekomendowanych przez JEDEC [29] i NBS [30]. Metodzie tej odpowiada układ pomiarowy przedstawiony na rys. 1. Zawiera on źródło napięciowe U_C , zapewniające polaryzację złącza baza—kolektor, źródło prądu grzejnego emitera I_E , źródło prądu pomiarowego emitera I_{Em} , tranzystor pomocniczy TP oraz układ próbkująco-pamiętający UP połączony z woltomierzem VC. Wielkość mocy rozpraszanej w tranzystorze jest określona przez iloczyn napięcia baza—kolektor U_{BC} oraz prądu kolektora I_C , który jest sumą grzejnego i pomiarowego prądu emitera.

Przebieg pomiaru w układzie na rys.1 jest taktowany przez impulsy zegarowe sterujące tranzystorem TP i układem próbkująco-pamiętającym. W trakcie grzania tranzystor TP nie przewodzi i badany tranzystor TB jest grzany mocą $U_{BC}I_C$. Cykl pomiarowy początkuje impuls wprowadzający tranzystor TP w stan nasycenia i kończący okres t_p rozpraszania mocy w tranzystorze. Powoduje to przejście przez



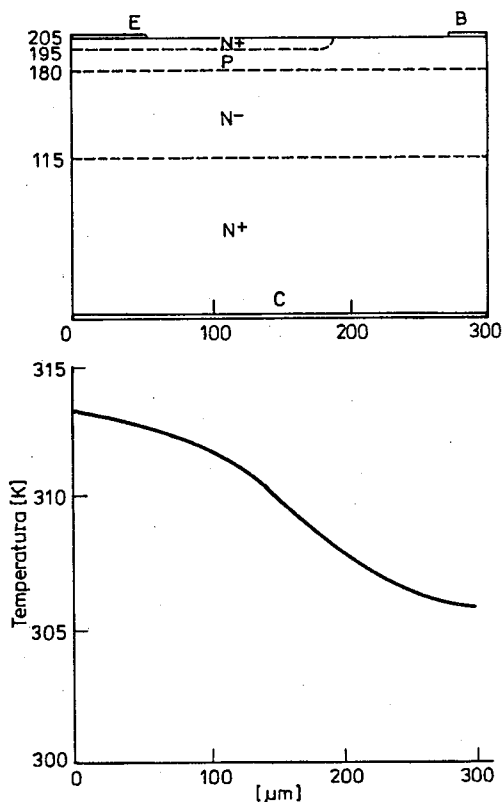
Rys.1 Schemat układu pomiarowego w metodzie przełączanego prądu emitera

ten tranzystor całego prądu grzejnego emitera I_E tak, że przez tranzystor TB płynie tylko prąd pomiarowy I_{Em} . Pomiar napięcia termometrycznego U_{EB} jest wykonywany po upływie czasu opóźnienia t_{op} , rzędu 50–500 μs , niezbędnego do ustalenia się w tranzystorze warunków odpowiadających zmienionej wartości prądu emitera. Natychmiast po zakończeniu pomiaru kończy się impuls wysterowujący tranzystor TP i następuje ponowne przełączenie prądu emitera I_E czyli przejście do cyklu grzania.

Metoda przełączanego prądu emitera może być wykorzystana zarówno do wyznaczania impedancji cieplnej tranzystora jak i do określenia jego rezystancji cieplnej odpowiadającej termicznemu stanowi ustalonemu. W obu przypadkach, warunkiem poprawności pomiaru jest zachowanie odpowiednich relacji pomiędzy czasem grzania t_p i czasem opóźnienia t_{op} w cyklu pomiarowym, tak aby $t_{op} < t_p$.

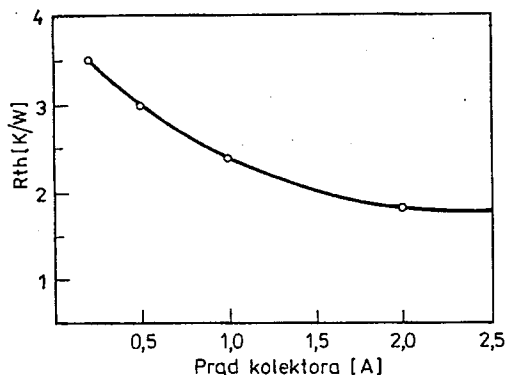
2.2. ZJAWISKA WPŁYWAJĄCE NA WYNIKI POMIARU REZYSTANCJI CIEPLNEJ TRANZYSTORÓW BIPOLARNYCH MOCY

Jak wspomniano we Wstępie, rozkład temperatury wewnątrz przyrządów półprzewodnikowych mocy może być silnie zróżnicowany. Dotyczy to w szczególności tranzystorów bipolarnych, w których obserwuje się nierównomierny rozkład temperatury na powierzchni złącza baza–emiter (rys.2) [1, 5, 7, 11, 19]. W efekcie pomierzona temperatura T_j nie reprezentuje maksymalnej temperatury występującej na złączu, ale jest pewną wartością uśrednioną, zawartą pomiędzy maksymalną i minimalną temperaturą złącza. Aby to uwzględnić należałoby przy pomiarach dokładnych pomnożyć ją przed wprowadzeniem do wzoru (1) przez współczynnik korygujący o wartości z przedziału 1 ÷ 1.4 [1]. Współczynnik ten nie jest jednak dla danego tranzystora stały, ale zmienia się wraz z warunkami obciążenia. Dla zadanych warunków pracy tranzystora może on być wyznaczony na drodze eksperymentalnej, w oparciu o procedurę postępowania zaproponowaną przez Blackburna [12].



Rys.2 Dwuwymiarowa struktura tranzystora bipolarnego (a) oraz rozkład temperatury wzdłuż krawędzi złącza emiter – baza w tej strukturze (b) wyznaczony dla $U_{CE}=400$ V i $U_{BE}=0.6$ V [19]

Należy podkreślić, że wpływ warunków obciążenia na ostateczny wynik pomiaru rezystancji cieplnej tranzystora pracującego w obszarze aktywnym występuje nawet przy stałej wartości mocy rozpraszanej w tranzystorze. Ilustruje to rys. 3, na którym przedstawiono zależność R_{th} od prądu kolektora I_C wyznaczoną dla stałej wartości



Rys.3 Zależność R_{th} od prądu kolektora I_C przy stałej wartości mocy $U_{CE}I_C=20$ W [1]

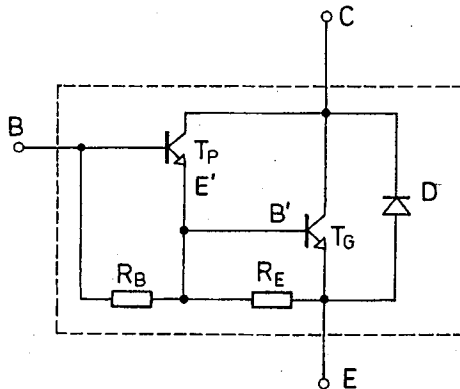
mocy $U_{CE}I_C$ rozpraszanej w przyrządzie. Jest to efektem zmian w lokalizacji rozpraszania mocy wraz ze zmianami napięcia polaryzującego.

Dodatkowym czynnikiem, który może wpływać na przebieg pomiaru w impulsowych metodach pomiarowych jest istnienie pasożytniczych indukcyjności związanych z doprowadzeniami i obudową tranzystora. W przypadku metody przełączanego prądu emitera powoduje to, że po raptownej zmianie prądu emitera z wartości I_E na wartość I_{Em} napięcie obserwowane pomiędzy końcówkami emitera i bazy obrazuje nie tylko zmiany napięcia termometrycznego U_{EB} ale jest sumą tego napięcia i napięć indukowanych w doprowadzeniach. Efekt ten był badany przez Berninga i Blackburna [31]. Stwierdzili oni, że zakłócenia te zamykają się zwykle w przedziale czasowym 5–100 μs liczonym od momentu przełączenia prądu emitera. Wyeliminowanie wpływu tych zakłóceń na pomiar temperatury wymaga opóźnienia momentu pomiaru około 100 μs , co może spowodować zaniżenie wyniku oszacowania temperatury. Niedogodność tą można usunąć stosując dla czasu opóźnienia ekstrapolację wyniku w oparciu o pierwiastkową zależność temperatury od czasu.

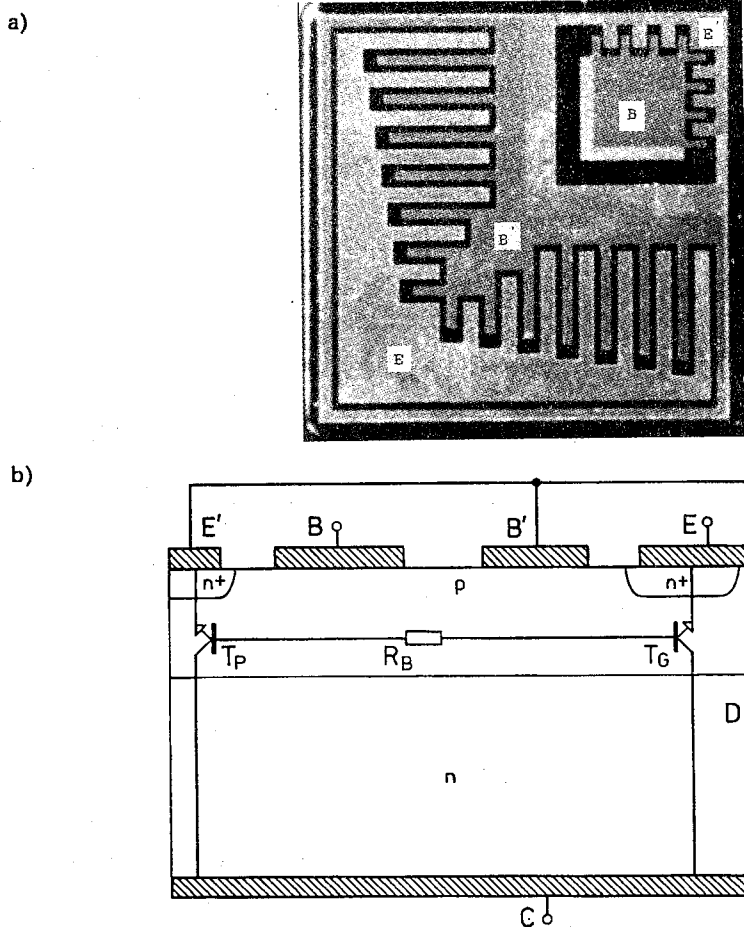
3. TRANZYSTOR DARLINGTONA

3.1. OPIS STRUKTURY

Tranzystor Darlingtona jest strukturą scaloną, której odpowiada elektryczny schemat ideowy przedstawiony na rys.4. Zawiera ona dwa tranzystory bipolarne, tranzystor główny TG i tranzystor pomocniczy TP, diodę zwrotną D oraz dwa rezystory, R_B i R_E . Sposób realizacji takiej struktury jest pokazany na rys. 5a i 5b. Przedstawiają one kolejno widok od góry pastylki półprzewodnikowej tranzystora BU 323 oraz odpowiadający mu przekrój poprzeczny. Na rysunkach zaznaczono lokalizację wszystkich kontaktów i elementów uwzględnionych w schemacie ideowym z rys.4.



Rys.4 Elektryczny schemat ideowy odpowiadający scalonej strukturze tranzystora Darlingtona n-p-n



Rys.5 Budowa scalonej struktury tranzystora Darlingtona n—p—n BU323: widok pastylki półprzewodnikowej od strony kontaktów emitera i bazy (a) oraz schematyczny przekrój poprzeczny przez tę strukturę ukazujący wzajemne położenie elementów uwzględnionych na rys.5 (b)

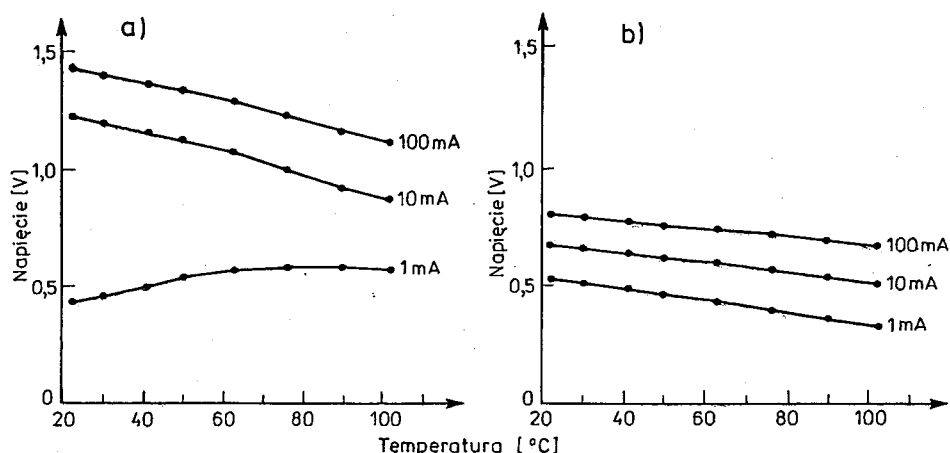
Tranzystory te są wykonywane z reguły jako elementy mocy o wymiarach pastylki półprzewodnikowej rzędu dziesiątek mm². Powoduje to, że nawet w stanie ustalonym występują w nich poprzeczne gradienty temperatury rzędu kilku stopni i wybór obszaru, dla którego będzie mierzony EPT może mieć istotne znaczenie dla oszacowania rezystancji cieplnej przyrządu. Wyboru tego dokonujemy wybierając konkretny EPT. Np. jeżeli jako EPT chcemy wykorzystać napięcie na złączu spolaryzowanym w kierunku przewodzenia, to może to być napięcie U_{EB} , będące sumą spadków napięć na złączach emiterowych obu tranzystorów, napięcie U_{CB} , odpowiadające złączu kolektorowemu tranzystora TG, lub napięcie U_{EC} , odpowiadające napięciu na diodzie zwrotnej D. Zgodnie z rys. 5, w każdym z tych

przypadków lokalizacja obszaru przez które płynie prąd pomiarowy, a więc i obszaru dla którego mierzymy średnią wartość temperatury jest inna, mimo że w dwu ostatnich przypadkach jest mierzone napięcie na tym samym złączu tranzystora.

3.2. WYBÓR METODY POMIARU TEMPERATURY

Jak powiedziano w rozdz. 2, dominującym EPT przy pomiarach rezystancji cieplnej tranzystorów mocy jest napięcie U_{EB} . W tranzystorach Darlingtona jest ono równe sumie spadku napięcia na połączonych szeregowo złączach emiterowych tranzystorów składowych. Powoduje to, że wywołane zmianami temperatury zmiany napięcia obu złącz dodają się czyniąc ten parametr bardziej wrażliwym na zmiany temperatury. Z tego względu wybór właśnie tego parametru wydaje się być szczególnie atrakcyjny i Rubin przedstawił w [32] zmodyfikowaną metodę przełączanego prądu emitera dla tych tranzystorów.

Ocena przydatności napięcia U_{EB} przeprowadzona w oparciu o analizę pracy struktury przedstawionej na rys. 5 nie wypada już tak korzystnie. Po pierwsze droga przepływu prądu pomiarowego jest w tym przypadku wyjątkowo skomplikowana. Jest to prąd poprzeczny płynący poprzez obszary p-bazy o różnych temperaturach, a oba złącza emiterowe są na tyle od siebie odległe, że mogą znajdować się w różnych temperaturach. Już to może znacznie utrudnić właściwą interpretację uzyskanego wyniku, a należy jeszcze wziąć pod uwagę, że oba złącza termometryczne są zbocznikowane przez rezystancje R_B i R_E , reprezentujące rezystancję rozproszoną obszarów p-bazy. Tak więc w tranzystorze istnieje kilka możliwych dróg przepływu prądu od kontaktu B do E. Obejmują one oba złącza emiterowe, aluminiowe połączenia *layoutu* na powierzchni płytki półprzewodnikowej oraz obszary p-bazy tworzące rezystancje R_B i R_E . Są to elementy o odmiennych

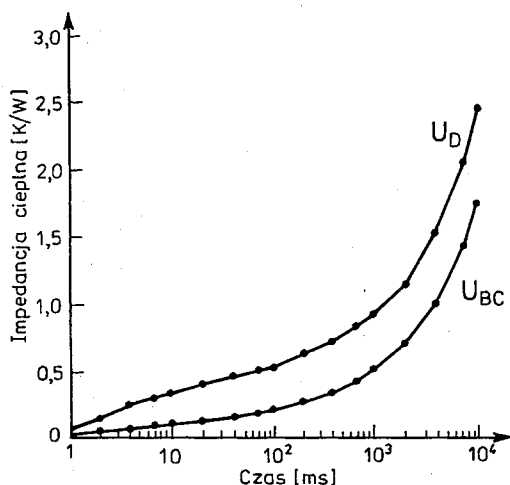


Rys. 6 Krzywe skalowania tranzystorów BU 323 uzyskane dla kilku wartości prądu pomiarowego przy przyjęciu jako napięcia termometrycznego: napięcia emiter-baza U_{EB} (a) i napięcia kolektor-baza U_{CB} (b)

charakterystykach termicznych, co powoduje, że krzywe skalowania, uzyskane dla warunków izotermicznych w tranzystorze, mogą nie być tak liniowe jak to ma miejsce w przypadku innych złącz tranzystora, a ich rozrzut może większy.

Aby ocenić na ile te specyficzne cechy konstrukcji tranzystora Darlingtona wpływają na przydatność poszczególnych złącz jako EPT przeprowadzono badania porównawcze obejmujące szereg typów tranzystorów Darlingtona produkowanych lub przewidywanych do produkcji w FP TEWA. Przykładowe wyniki tych badań, uzyskane dla tranzystorów BU 323, zawiera rys. 6. Porównano na nim krzywe skalowania uzyskane dla tych tranzystorów przy przyjęciu jako EPT, odpowiednio, napięcia U_{EB} i U_{CB} . Jak łatwo zauważyć, krzywe skalowania uzyskane dla napięcia U_{EB} są generalnie mniej liniowe, a krzywa uzyskana przy prądzie pomiarowym 1 mA jest do tego stopnia nieliniowa, że w ogóle nie nadaje się do pomiarów termometrycznych.

Z powyższych względów zdecydowano się przy opracowywaniu miernika na wybór jako EPT spadku napięcia na złączu baza–kolektor. Zgodnie z rys. 5 napięcie to można mierzyć w obszarze znajdującym się pod kontaktem bazy, wybierając jako EPT napięcie U_{BC} , lub w obszarze bliskim krawędzi złącza emiterowego głównego tranzystora, wybierając jako EPT napięcie U_{EC} . W normalnych warunkach pracy tranzystora największa gęstość rozpraszania mocy występuje pod obszarem emitera tranzystora TG i tam należy oczekiwać największego przyrostu temperatury. Aby to sprawdzić przeprowadzono pomiary impedancji i rezystancji cieplnej tranzystorów Darlingtona wykorzystując jako EPT, kolejno, napięcie U_{BC} i napięcie U_{EC} . Przykładowe wyniki, uzyskane dla tranzystora BU 323, są zamieszczone na rys. 7. Przedstawiono na nim fragmenty impedancji cieplnych



Rys.7 Fragmenty impedancji cieplnych tranzystora BU 323 wyznaczone dla impulsu grzejnego 30 W ($U_{CE}=10$ V, $I_C=3$ A) przy wykorzystaniu jako EPT: napięcia kolektor–emiter (U_D) i napięcia kolektor–baza (U_{BC})

tego tranzystora wyznaczone w oparciu o zmiany temperatury w obszarach odpowiadających, kolejno, złączu kolektorowemu (napięcie U_{CB}) i diodzie zwrotnej (napięcie U_{EC}).

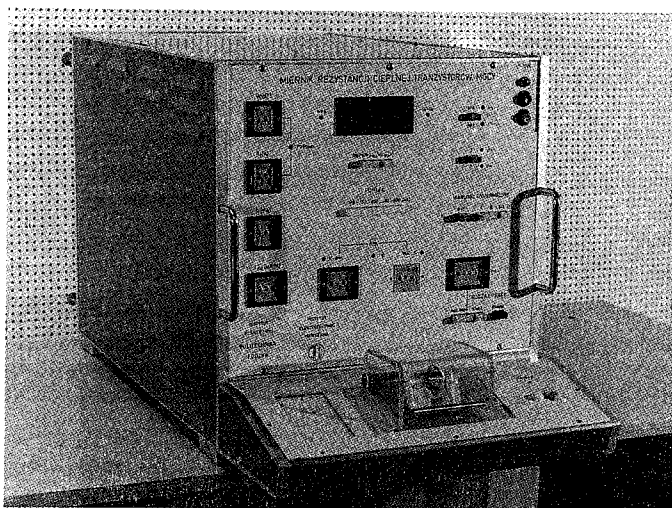
Jak widać, przebieg obu krzywych jest istotnie inny, a przyrosty temperatury wyznaczone z wykorzystaniem napięcia U_{EC} jako EPT dają wyższe wartości impedancji. Różnice pomiędzy oboma oszacowaniami impedancji cieplnej są bardzo wyraźne i w rozpatrywanym przedziale czasowym wynoszą one od 300 do 40 % wartości impedancji wyznaczonej z pomiaru napięcia U_{BC} .

Mając to na uwadze przyjęto ostatecznie jako podstawowy EPT napięcie U_{EC} , czyli spadek napięcia na diodzie zwrotnej D. Postanowiono również zachować w mierniku możliwość pomiaru temperatury w obszarze emitera tranzystora TP w oparciu o zmiany napięcia U_{BC} . Ponieważ rozpraszanie mocy jest w nim stosunkowo niewielkie, a jednocześnie jest on znacznie oddalony od głównego obszaru wydzielania ciepła, pozwoli to na wykorzystanie miernika do kontroli poprzecznych gradientów temperatury w tranzystorze. Możliwość ta może być bardzo pomocna przy ocenie jakości konstrukcji gotowych tranzystorów, jak i na etapie termicznej optymalizacji nowo opracowywanych konstrukcji.

4. OPIS MIERNIKA

4.1. PODSTAWOWE DANE

Widok ogólny miernika przedstawia rys. 8. Jest to miernik przeznaczony zasadniczo do pomiarów rezystancji cieplnej tranzystorów Darlingtona typu p-n-p i n-p-n. Rezystancja jest w nim wyznaczana zgodnie z wzorem (1) w oparciu o pomiar przyrostu temperatury wywołanej stałą mocą rozpraszaną w przyrządzie.



Rys.8 Widok ogólny miernika rezystancji cieplnej tranzystorów Darlingtona

W oparciu o wyniki badań przedstawione w rozdz. 3 przyjęto w nim jako EPT alternatywnie napięcie U_{CE} i U_{BC} . Pozwoli to z jednej strony na wykonywanie rutynowych pomiarów zarówno dla tranzystorów Darlingtona jak i zwykłych tranzystorów mocy bez diody zwrotnej, a z drugiej strony stworzy możliwości przeprowadzania bardziej złożonych pomiarów dla tranzystorów Darlingtona.

Tabela I

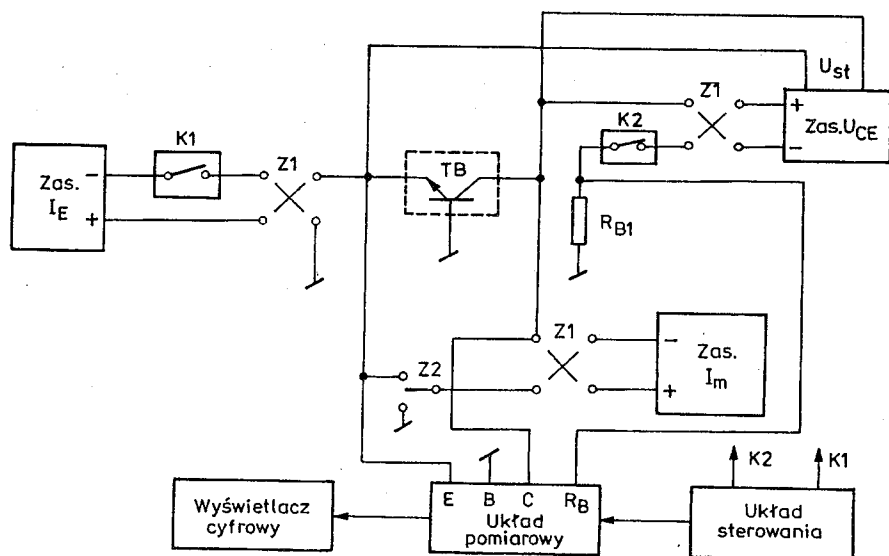
Wielkość	Zakres zmian
Prąd kolektora I_C	0.1 – 10 A
Napięcie kolektor–emiter U_C	1 – 50 V
Prąd złącza termometrycznego I_{CE}	1 – 100 mA
Czas trwania impulsu mocy t_m	0.001 – 10 s
Czas opóźnienia t_{imp}	100 – 150

Podstawowe parametry miernika są zestawione w Tabeli I. Ustalono je mając na uwadze jego przeznaczenie, a w szczególności dane o tranzystorach Darlingtona produkowanych lub przewidywanych do produkcji w FP TEWA oraz wyniki badań testowych egzemplarzy tych tranzystorów. W mierniku przewidziano trzy rodzaje pracy różniące się sposobem przeprowadzenia pomiaru. Może to być:

- praca automatyczna — podczas której pomiar jest wykonywany po osiągnięciu cieplnego stanu ustalonego, odpowiadającego zadanej mocy rozpraszanej w przyrządzie,
- praca impulsowa — podczas której pomiar jest wykonywany po pojedynczym impulsie mocy o zadanym czasie trwania,
- praca kluczowana — podczas której moc jest rozpraszana w tranzystorze przez cały czas, a pomiar odbywa się podczas inicjowanych zewnętrznie przerw w przepływie prądu grzejjego.

Rezystancja cieplna jest mierzona bezpośrednio jedynie podczas pracy automatycznej miernika. W obu pozostałych przypadkach wynikiem pomiaru są wartości chwilowe impedancji cieplnej. Możliwość regulacji w szerokim zakresie czasu trwania impulsu mocy, w przypadku pracy impulsowej, pozwala na wykorzystanie miernika również do wyznaczania pełnej krzywej impedancji cieplnej.

Na rys. 9 przedstawiono uproszczony schemat blokowy miernika. W celu zapewnienia niezależnej regulacji parametrów impulsu grzejjego, przewidziano w nim odrębne zasilacze utrzymujące na zadanym poziomie prąd kolektora, zasilacz prądowy „Zas. I_E ”, i napięcie kolektor–emiter, zasilacz napięciowy „Zas. U_{CE} ”. Zasilacz prądowy „Zas. I_m ” jest źródłem prądu pomiarowego. Konfigurację układu, właściwą dla typu badanego tranzystora (n–p–n lub p–n–p), uzyskuje się poprzez odpowiednie ustawienie kluczy Z_1 . Klucz Z_2 służy do wyboru złącza termometrycznego, a wymagane podczas pomiaru przerwanie przepływu prądu grzejjego i zdjęcie polaryzacji wstecznej ze złącza termometrycznego jest realizowane za pomocą kluczy elektronicznych K_1 i K_1 .



Rys.9 Schemat blokowy miernika

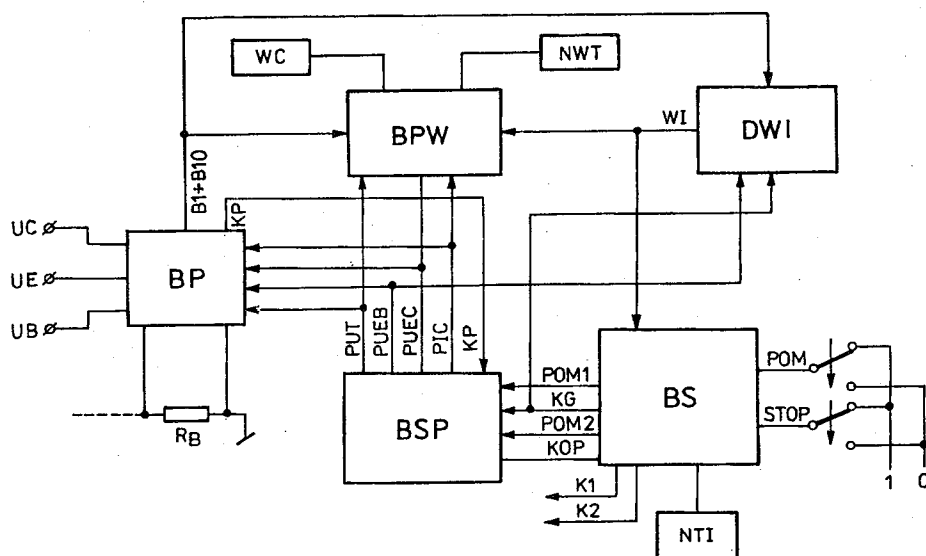
4.2. DZIAŁANIE MIERNIKA

W mierniku przewidziano szereg rozwiązań ułatwiających poprawne przeprowadzenie pomiaru. Wynik pomiaru jest podawany w postaci cyfrowej jako wartość, odpowiednio, rezystancji lub impedancji cieplnej. Wartości napięcia na złączu termometrycznym mierzone przed rozpoczęciem rozpraszania mocy i w momencie wykonywania pomiaru są zapamiętywane i dostępne w postaci cyfrowej. Miernik posiada ponadto układ wykrywający cieplny stan ustalony w tranzystorze przed rozpoczęciem właściwego pomiaru, układy wykrywające i sygnalizujące niepoprawny przebieg pomiaru oraz układy zabezpieczające przed uszkodzeniem.

Na rys. 10 przedstawiono uproszczony schemat blokowy układów sterowania i pomiaru wraz z podstawową siecią połączeń pomiędzy poszczególnymi blokami. Obejmuje on:

A. Blok sterowania pracą miernika BS

Jest on zintegrowany z przełącznikami rodzaju pracy, przełącznikami inicjującymi i kończącymi pomiar (POM, STOP) oraz z nastawnikiem cyfrowym czasu trwania impulsu mocy przy pracy impulsowej (NTI). Blok ten generuje podstawowe sygnały taktujące pracą miernika. Podczas pełnego cyklu pomiarowego jako pierwszy pojawia się sygnał POM1, inicjujący pierwszy etap pomiaru. W etapie tym sprawdza się czy badany tranzystor nie jest uszkodzony i czy został poprawnie podłączony do miernika oraz testuje się czy są spełnione w tranzystorze warunki izotermiczne „na zimno”, tzn. czy temperatura wnętrza tranzystora jest równa temperaturze otoczenia. Jeżeli testy te wypadną pomyślnie jest dokonywany pomiar



Rys.10 Schemat blokowy układów sterowania i pomiaru

napiecia termometrycznego „na zimno” i pojawia się sygnał POM2 inicjujący przejście do drugiego etapu pomiaru. Jest to etap rozpraszania mocy w tranzystorze i rozpoczyna się on zamknięciem kluczy K1 i K2. Okres rozpraszania mocy jest uzależniony od wybranego rodzaju pracy miernika. Gdy odpowiednie warunki zostaną spełnione, następuje otwarcie kluczy K1 i K2 i pojawia się sygnał KG inicjujący pomiar napięcia termometrycznego „na ciepło”. Jest to koniec cyklu pomiarowego dla pracy automatycznej i impulsowej. W przypadku pracy kluczowanej, po zakończeniu pomiaru napięcia termometrycznego „na ciepło” (sygn. KOP) następuje automatyczny powrót do etapu rozpraszania mocy, a zakończenie pomiaru jest inicjowane zewnętrznie sygnałem STOP.

B. Blok sterujący przebiegiem pomiaru BSP

Są w nim generowane podstawowe sygnały synchronizujące pracę bloku pomiarowego BP i bloku przetwarzania wyników BPW. Podczas taktu POM1 generuje on z częstotliwością 1Hz sygnał PUT inicjujący pomiar napięcia termometrycznego „na zimno”, które jest wykorzystywane jako detektor termicznego stanu ustalonego „na zimno”. Po przejściu do taktu POM2 generuje on kolejno sygnały PUEC i PIC inicjujące pomiar napięcia U_{CE} i prądu I_C , a następnie, w przypadku pracy automatycznej, generuje z częstotliwością 1Hz sygnał PUEB inicjujący pomiar napięcia U_{BE} , które jest wykorzystywane do detekcji osiągnięcia przez tranzystor cieplnego stanu ustalonego. Pojawienie się sygnału KG powoduje wygenerowanie impulsu PUT inicjującego pomiar napięcia termometrycznego „na ciepło”. Impuls ten jest opóźniony w stosunku do sygnału KG o nastawiany zewnętrznie czas t_{op} przeznaczony na zakończenie elektrycznego stanu nieustalonego wywołanego zmianą warunków pracy tranzystora.

C. Układ pomiarowy BP

Są w nim dokonywane pomiary poszczególnych wielkości inicjowane odpowiednimi sygnałami z bloku BSP. Sygnał analogowy, odpowiadający mierzonej wielkości, jest standaryzowany i wprowadzany do układu S—H, a następnie przetwarzany w przetworniku analogowo-cyfrowym. Ostateczny wynik pomiaru jest przesyłany w postaci cyfrowej do bloku BPW 10-bitową szyną danych B1—B10. Towarzyszy temu wygenerowanie sygnału KP, który oznacza gotowość układu do przeprowadzenia następnego pomiaru.

D. Blok przetwarzania wyników BPW

Zawiera on specjalizowany arytmometr z rejestrami, w których są przechowywane wyniki pomiarów kolejnych wielkości, zintegrowany ze wskaźnikiem cyfrowym WC, cyfrowym nastawnikiem temperaturowego współczynnika napięcia złącza termometrycznego NWT oraz układem selekcji US. Pracą tego bloku sterują sygnały generowane w BSP. Począwszy od pomiaru napięcia termometrycznego „na zimno” do rejestrów BPW są kolejno wprowadzane wyniki pomiarów wykonanych w BP. Po wprowadzeniu wyniku pomiaru napięcia termometrycznego „na ciepło” rozpoczyna pracę arytmometr wyznaczający rezystancję cieplną. Wynik obliczeń jest wprowadzany do odpowiedniego rejestru i może być wyświetlony na wyświetlaczu cyfrowym przemiennie z napięciem termometrycznym „na zimno” i „na ciepło”. W celu ułatwienia selekcji tranzystorów na linii produkcyjnej, jest on również porównywany w US z nastawianymi na odpowiednich nastawnikach cyfrowych wartościami $R_{th\ max}$ i $R_{th\ min}$, a wynik tego porównania jest sygnalizowany na płycie czołowej miernika.

E. Blok detekcji warunków izotermicznych DWI

Są w nim testowane warunki izotermiczne w badanym tranzystorze w takcie POM1, przed pomiarem napięcia termometrycznego „na zimno” oraz w takcie POM2 dla określenia momentu zakończenia rozpraszania mocy przy pracy automatycznej. Jako kryterium izotermiczności przyjęto w nim wystąpienie w przedziale czasowym t_{wi} zmiany napięcia termometrycznego mniejszej od 6mV. Wielkość t_{wi} jest ustawiana zewnętrznie i w zależności od rodzaju pomiaru oraz typu badanego tranzystora może być równa 0, 4, 8 lub 512s.

4.3. KONTROLA PRZEBIEGU POMIARU

W celu zapewnienia niezawodności i pewności działania miernika, co ma duże znaczenie przy pracy na linii produkcyjnej, wprowadzono do niego szereg układów kontroli i zabezpieczeń. Sygnalizują one gdy badany tranzystor jest uszkodzony, gdy typ tranzystora jest niezgodny z wybraną opcją pracy (pnp/npn) lub gdy nastawiona moc impulsu grzejjego ($I_C U_{CE}$) przekracza maksymalną moc miernika P_{max} . Poniżej zostaną przedstawione nieco szerzej dwa ciekawsze rozwiązania zastosowane w tym mierniku, a mianowicie układ wykrywania i kompensacji wzbudzeń oraz układ kontroli styków Kelvina.

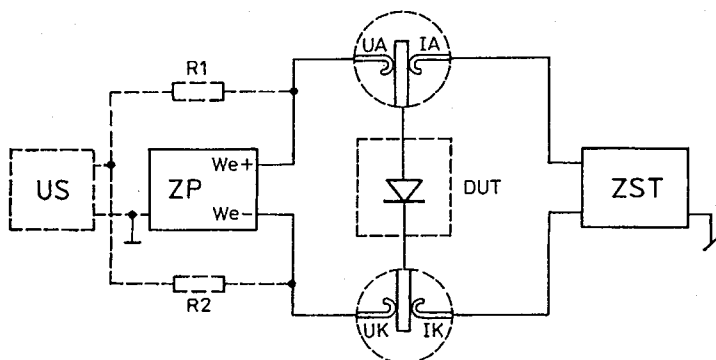
4.3.1. WYKRYWANIE I KOMPENSACJA WZBUDZEŃ

Tranzystory Darlingtona charakteryzują się dużym współczynnikiem wzmocnienia, co było przyczyną dodatkowych trudności pomiarowych, nie występujących przy pomiarach cieplnych zwykłych tranzystorów mocy. Jak można zauważyć na rys. 9, układ pomiarowy zawiera dwa stabilizowane zasilacze, prądu I_C i napięcia U_{CE} , tak wkomponowane w układ pomiarowy, że stabilizują one wielkości mierzone w punktach znacznie oddalonych od samych zasilaczy. Powodowało to, że podczas procesu nagrzewania tranzystorów Darlingtona o dużych współczynnikach wzmocnienia pojawiały się na napięciu U_{CE} oscylacje zakłócające zarówno pomiar tego napięcia jak i sam proces nagrzewania.

W celu wyeliminowania tego niekorzystnego zjawiska wykonano szereg badań z wykorzystaniem prototypowego układu miernika. W efekcie stwierdzono, że wzbudzenia te najskuteczniej można wytłumić wprowadzając niewielką pojemność pomiędzy kolektorowy i emiterowy zacisk napięciowy gniazda pomiarowego, przy czym nie istnieje jedna uniwersalna wartość tej pojemności, odpowiednia dla wszystkich typów tranzystorów Darlingtona. Problem ten rozwiązano w mierniku w ten sposób, że do bloku pomiarowego BP wprowadzono układ wykrywania wzbudzeń, a w gnieździe pomiarowym włączono pojemność dekadową bezpośrednio pomiędzy kolektorowy i emiterowy zacisk napięciowy. Podczas pracy miernika układ wykrywania wzbudzeń sygnalizuje pojawienie się wzbudzeń zapaleniem sygnału kontrolnego na płycie czołowej, co oznacza, że aktualna wartość pojemności jest niewłaściwa i należy zmienić ją na inną. Aby uniknąć wpływu tej pojemności na przebieg pomiaru napięcia termometrycznego „na ciepło”, przed pomiarem jest ona rozładowywana w przedziale czasowym t_{op} w specjalnym układzie rozładowującym.

4.3.2. UKŁAD KONTROLI STYKÓW KELVINA

Dla zapewnienia pewności działania miernika okazało się konieczne wprowadzenie kontroli poprawności działania styków Kelvina w gnieździe pomiarowym. W tym celu, aby operację tę można było przeprowadzić bez rozbudowy miernika, opracowano specjalną metodę pomiarową [32]. Istotę tej metody ilustruje rys. 11, na którym



Rys.11 Układ do testowania styków Kelvina przyłączonych do diody p-n

przedstawiono schematycznie układ do testowania styków Kelvina przyłączonych do diody p-n (DUT). Polega ona na jednoczesnej kontroli styków prądowych IA i IK oraz styków napięciowych UA i UK.

Podczas testu do styków prądowych jest doprowadzany sygnał testujący ze źródła ZST powodujący ustalenie się na diodzie pewnego spadku napięcia U_p . Napięcie to jest mierzone przez układ pomiarowy ZP połączony z diodą poprzez napięciowe styki gniazda pomiarowego, a odseparowany galwanicznie od źródła sygnału testującego ZST. Wynikiem pomiaru jest napięcie różnicowe U_r . W układzie tym brak styczności na którymkolwiek z badanych styków (IA, IK, UA, UK) jest równoważny z odłączeniem od układu pomiarowego ZP źródła napięciowego o napięciu U_p . Pomierzone napięcie U_r będzie w tej sytuacji zależne od przyjętego rozwiązania konstrukcyjnego układu pomiarowego ZP, ale jego wartość będzie zawsze bliska zeru i dużo mniejsza od napięcia U_p . Pozwala to przyjąć za kryterium braku kontaktu na badanych stykach warunek $U_r < U_m < U_p$, gdzie U_m jest napięciem odniesienia przyjętym dla danego układu.

W mierniku przyjęto, że kontrola styków Kelvina będzie się odbywać w dwóch etapach, odpowiednio, dla spolaryzowanego w kierunku przewodzenia złącza kolektor-baza i emiter-baza. Dla pomiarów tych przyjęto $U_m = 0.3 \text{ V}$, a sam układ pomiarowy uzupełniono o oporniki R_1 i R_2 dołączone bezpośrednio do wejść We+ i We- układu pomiarowego ZP oraz źródła napięciowego US o napięciu $U_s = 0.5 \text{ V}$. Miało to na celu zabezpieczenie układu przed ewentualnymi zakłóceniami związanymi z „efektem antenowania” przewodów łączących wejścia We+ i We- ze stykami UA i UK przy braku kontaktu na którymś z testowanych styków.

5. PODSUMOWANIE

Praca jest poświęcona przedstawieniu problemów występujących przy pomiarach rezystancji cieplnej tranzystorów bipolarnych mocy ze szczególnym uwzględnieniem specyfiki pomiarów rezystancji cieplnej tranzystorów Darlingtona. Są to układy scalone mocy o dość złożonej strukturze, co musi być wzięte pod uwagę przy wyborze odpowiedniej metody pomiarowej.

Metodom pomiaru rezystancji cieplnej tranzystorów bipolarnych mocy jest poświęcony rozdz.2. Zestawiono w nim najczęściej stosowane elektryczne parametry termoczułe (EPT) oraz omówiono zjawiska występujące w tranzystorze, które mogą mieć wpływ na przebieg i wynik pomiaru. Nieco dokładniej jest przedstawiona jedynie metoda „przelączanego prądu emitera”, która jest obecnie uważana za metodę standardową, rekomendowaną m.in. przez tak renomowane instytucje normalizacyjne jak JEDEC i NBS.

Rozdz. 3 jest poświęcony omówieniu specyfiki pomiarów przeprowadzanych dla tranzystorów mocy Darlingtona. Z uwagi na to, że nie są one pojedynczym elementem, ale strukturą scaloną mocy zawierającą w sobie kilka przyrządów półprzewodnikowych, zastosowanie do nich metod właściwych dla zwykłych tran-

zystorów mocy nie jest najlepszym rozwiązaniem. Dotyczy to w szczególności metody „przelącanego prądu emitera” z uwagi na niejednoznaczność krzywych skalowania dla napięcia baza — emiter jako EPT. Przeprowadzone badania wykazały, że dla tranzystorów Darlingtona znacznie lepszym EPT będzie napięcie na przewodzącym złączu baza — kolektor mierzone w obszarze odpowiadającym diodzie zwrotnej układu zastępczego tranzystora. W rozdz. 4 opisano automatyczny miernik rezystancji cieplnej tranzystorów Darlingtona opracowany dla FP TEWA [14] w oparciu o wnioski przedstawione w rozdz. 3 oraz wyniki dodatkowych badań przeprowadzonych z wykorzystaniem układu modelowego. Badania te wykazały wielką wrażliwość układu na występowanie oscylacji podczas procesu grzania tranzystorów Darlingtona o dużym współczynniku wzmocnienia. Problem ten rozwiązano wprowadzając do miernika układ wykrywania i kompensacji wzbudzeń. Innym rozwiązaniem specyficznym dla tego miernika jest chroniony patentem układ kontroli styków Kelvina.

Miernik jest przeznaczony zarówno do pracy na linii produkcyjnej jak i prac laboratoryjnych. Przewidziano w nim możliwość pomiaru temperatury z użyciem dwóch EPT — napięcia U_{EC} lub napięcia U_{BC} . Biorąc pod uwagę specyfikę konstrukcji tranzystorów Darlingtona, pozwala to na wykorzystanie miernika do kontroli poprzecznych gradientów temperatury w tranzystorze. Możliwość to może być bardzo pomocna przy ocenie jakości konstrukcji gotowych tranzystorów, jak i na etapie termicznej optymalizacji nowo opracowywanych konstrukcji.

BIBLIOGRAFIA

1. F.F. Oettinger, D.L. Blackburn, S. Rubin: *Thermal characterization of power transistors*, IEEE Trans. Electron Devices, 1976, vol. ED—23, nr. 8, pp. 831—837
2. D.L. Blackburn, F.F. Oettinger: *Transient thermal response measurements of power transistors*, IEEE Trans. Ind. Electron. & Control Instrum., 1975, vol. IECT—22, nr.2, pp.134—141
3. P.R. Gray, D.J. Hamilton: *Analysis of electrothermal integrated circuits*, IEEE J. Solid—St. Circuits, 1971 vol.SC—6, nr.1, 1971, pp.8—14
4. Z. Lisik: *Numeryczne wyznaczenie rozkładu temperatury w tyrystorze podczas procesu załączania*, IV KKTOiUE, 28—33.X.1981, Drzonków, ss.341—346
5. V.C. Alwin, D.H. Navon: *Emitter—junction temperature measurement under nonuniform current and temperature distribution*, IEEE Trans. Electron Devices, 1976, vol.ED—23, nr.1, pp.61—63
6. W.R. Wilcox: *Heat transfer in power transistors*, IEEE Trans. Electron Devices, 1963, vol.ED—10, nr 9, 1963, pp.308—313
7. S.P. Gaur: *Finding temperature distribution in high voltage transistor.*, 1977, IBM Tech. Disc. Bull., vol.19, nr 8, pp.3147—49
8. A. Nowakowski, J. Gajkiewicz: *Computer simulation of power transistor operation as a step to automatic SOA determination*, V Int. Conf. on Control Systems & Computer Science, Bucharest 1983, pp.343—350
9. A. Napieralski, J.M. Dorkel, Ph. Leturcq: *Electro-thermal coupling simulation in high power bipolar transistors*, NASECODE'87 17—19 June Dublin, pp.301—306
10. K.J. Nagus, R.W. Franklin, M.M. Yovanovich: *Thermal modeling and experimental techniques for microwave bipolar devices*, IEEE Trans. Comp. Hybrids Manuf. Technol., 1989, vol.12, nr 4, pp.680—689

11. P.W. Webb: *Measurement of thermal resistance using electrical methods*. IEE Proc., 1987, vol.134, nr 2, pp.51–56
12. D.L. Blackburn: *An electrical technique for the measurement of the peak junction temperature of power transistors*. 13th Ann. Proc. Reliability Phys., 1975, pp.142–150
13. Zommer N.D., Feucht D.L., Heckel R.W.: *Reliability and thermal impedance studies in soft soldered power transistors*, IEEE Trans. Electron Devices, vol.ED–23, nr 8, 1976, pp.843–850
14. Z. Lisik, T. Michalski, J. Powierza, Z. Szczepaniak, H. Świątek, M. Kowalew: *Miernik do wyznaczania rezystancji cieplnej tranzystorów mocy Darlingtona*; praca wykonana dla NCPC CEMI (FP TEWA), 1987
15. A.A. Jaecklin: *Instantaneous temperature profiles inside semiconductor power devices — Part I*. IEEE Trans. Electron Devices, 1974, vol.ED–21, nr 1, pp.50–53
16. L.G. Walshak, W.E. Poole: *Thermal resistance measurement by IR scanning*. Microwave J., 1977, vol.20, nr 2, pp.62–65
17. Ph. Angot: *Modeling and visualization of thermal fields inside electronic systems under operating conditions*. IBM Tech. Report, Bordeaux, 1989
18. B. Więcek, M. Grecki, J. Pacholik: *Thermal measurements of power semiconductor devices using thermographic system*. QIRT'92, 7–9 July 1992 Chatenay–Malabry, pp.291–295
19. S.P. Gaur: *Second breakdown in high-voltage switching transistors*. Electron. Lett., 1976, vol.12, nr 20, 1976, pp.526–527
20. J.L. Ferguson: *Liquid crystals in nondestructive testing*. Appl. Optics, 1968, vol. 7, nr 9, pp.1729–1737
21. M.M. Minot: *Thermal characterisation of microwave power FETs using nematic liquid crystals*. IEEE Int. Microwave Symp. Digest, 1986, pp.495–498
22. D.J. Brenner: *A technique for measuring the surface temperature of transistors by means of fluorescent phosphore*. NBS Tech. Note S91, 1971
23. R. Paul: *Technika pomiarów tranzystorów*. WKiŁ, Warszawa, 1973
24. J. Kołodziejski, L. Spiralski, E. Stolarski: *Pomiary przyrządów półprzewodnikowych*. WKiŁ, Warszawa, 1990
25. P. Coleby: *Thermal resistance of semiconductor devices under steady-state conditions*. Mullard Tech. Comm., June 1963, vol.7, pp.127–140
26. R.F. Gates, R.A. Johnson: *The measurement of thermal resistance — A recommendation for standardization*. Semicon. Prod., July 1959, vol.2, pp.21–26
27. J. Tellerman: *Measuring transistor temperature rise*, Electronics, 1954, vol.27, pp.185–187
28. J.F. Nelson, J.E. Iwersen: *Measurement of internal temperature rise of transistors*. 1958, Proc IRE, vol.46, pp.1207–1208
29. *Thermal resistance measurement of conduction cooled power transistors*. EIA Recommended Standard RS–313–B, 1975
30. S. Rubin, F.F. Oettinger: *Semiconductor Measurement Technology: Thermal Resistance Measurements on Power Transistors*. NBS Special Publication 400–14, 1979
31. D.W. Berning, D.L. Blackburn: *The effect of magnetic package leads on the measurement of thermal resistance of semiconductor devices*. IEEE Trans. Electron Devices, 1981, vol.ED–28, nr 5, pp.609–611
32. S. Rubin: *Thermal resistance measurement on monolithic and hybrid Darlington power transistors*. IEEE PESC, Los Angeles 1975, pp.252–261
33. Z. Lisik: *Sposób i układ do testowania styków Kelvina*. Pat. PL nr 153597

Z. LISIK, T. MICHALSKI, Z. SZCZEPANIAK

THERMAL RESISTANCE MEASUREMENT IN POWER DARLINGTON TRANSISTORS

S u m m a r y

Thermal resistance is one of the basic parameters declared for bipolar power transistors. Its determination requires a measurement of maximal temperature occurring inside the transistor. Usually, such a measurement is performed by means of an thermosensitive electrical parameter (TEP). There is a lot of ways in which the measurement may be performed and the choice of measurement method may essentially effect the results.

In this paper, problems occurring at the measurement of thermal resistance of bipolar power Darlington transistors are presented. The Darlington transistors can not be considered as discrete devices like the ordinary bipolar transistors. In fact, they are some integrate circuits covering a few semiconductor devices. It has been shown that the methods commonly applied in the case of ordinary transistors, in which the voltage of forward biased emitter—base junction is used as TEP, are not convenient for the Darlington transistors. It has been found that in this case, the voltage of forward biased collector—base junction measured in the area of baypass diode is much better as TEP. The results of investigation has been used to work out a thermal resistance meter of Darlington transistors for FP TEWA which is also described.

621.315.992:53.072.1

Badanie struktury energetycznej półprzewodników metodą zewnętrznego zjawiska fotoelektrycznego

JÓZEF WOJAS

*Profesor emerytowany, Politechnika Warszawska**Otrzymano 1994.02.14**Autoryzowano do druku 1994.04.22*

Podano opracowane przez autora sposoby określania wartości fotoelektrycznych parametrów powierzchni i warstw przypowierzchniowych półprzewodników. Uzyskane w ten sposób wyniki umożliwiają konstruowanie modeli energetycznych półprzewodników. W drugiej części pracy wykazano, jak z krzywych rozkładu energii, wydajności fotoelektrycznej, odbicia i gęstości stanów badanych półprzewodników można obliczyć ich strukturę energetyczną.

1. WYZNACZANIE ELEKTRONOWYCH PARAMETRÓW POWIERZCHNI PÓLPRZEWODNIKÓW

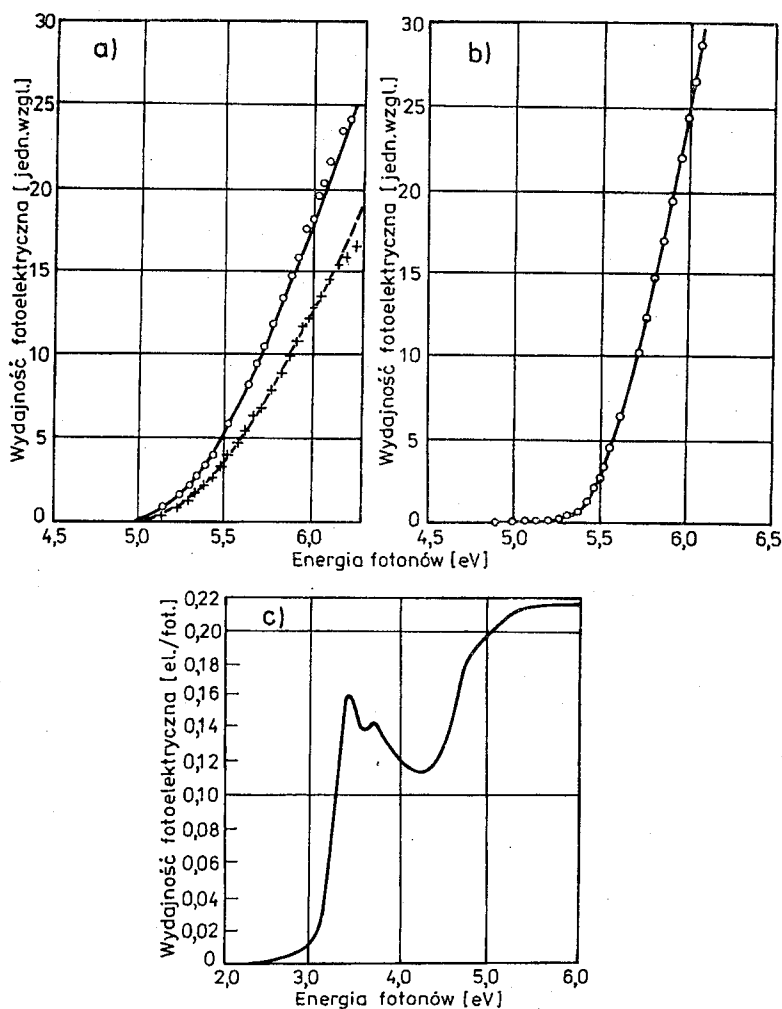
Wystąpienie zjawiska fotoemisji współcześnie traktujemy jako spełnienie się trzech procesów:

- 1) wzbudzenie elektronu po pochłonięciu fotonu,
- 2) wędrówkę wzbudzonego elektronu do powierzchni przez pasmo przewodnictwa, w czasie której może nastąpić rozproszenie elektronu,
- 3) przejście elektronu przez powierzchnię kryształu (pokonanie bariery potencjału powierzchniowego).

Fotoemisja zewnętrzna może być narzędziem badania struktury energetycznej półprzewodników, gdyż przebiegi dwu rodzajów charakterystyk fotoemisji są powiązane z właściwościami objętości i warstw przypowierzchniowych ciał stałych.

Te dwie grupy charakterystyk są następujące:

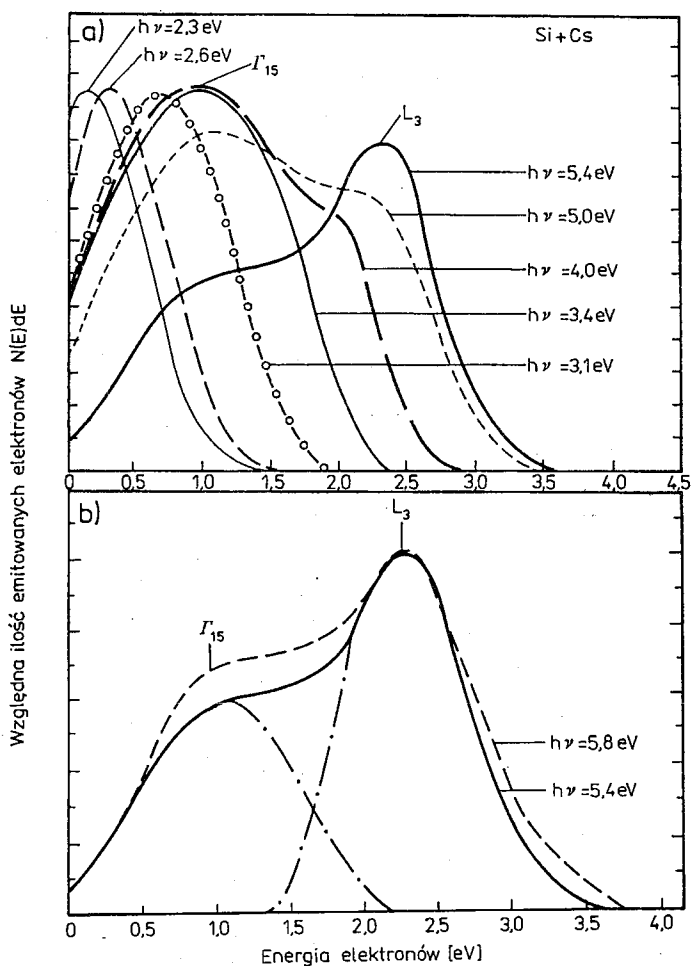
A) Widmowe rozkłady bezwzględnych wydajności kwantowych fotoemitery (w funkcji $h\nu$ lub λ pobudzającego promieniowania), z których otrzymujemy, między innymi, informacje co do położenia długofalowych progów fotoemisji oraz orientujemy się o bezwzględnej czułości fotoemitery (dla przykładu rys. 1 a, b, c) [1].



Rys.1. Fotoemisja i wydajność kwantowa z krzemu

- a) Fotoemisja z kryształów Si typu p domieszkowanego borem o koncentracji $n = 6 \cdot 10^{19}$ atomów/cm³
 (--- obliczono dla $\Delta V = 0,5$ eV, $L = 15 \text{ \AA}$; — obliczono dla $\Delta V = 0,5$ eV, $L = 20 \text{ \AA}$)
- b) Fotoemisja z kryształu Si typu p domieszkowanego galilem o koncentracji $n = 1,6 \cdot 10^{16}$ atomów/cm³
 (O — punkty pomiarowe, — obliczono z równania $I \approx (h\nu - 4,85)^{5/2} + 28,3 (h\nu - 5,36)^{3/2}$)
- c) Rozkład widmowy wydajności kwantowej (poprawiony na odbicie przez Philippa i Tafta [1]. Należy zauważyć maksima przy 3,4, 3,7 i 5,3 eV i minimum przy 4,4 eV. Odpowiadają one dobrze ustalonym przejściom optycznym. Pik przy 3,7 eV jest rozdzielony po raz pierwszy

B) Rozkłady względnych ilości fotoelektronów w funkcji ich energii kinetycznych, otrzymane z charakterystyk prądowo-napięciowych w kondensatorze kulistym (dla przykładu rys. 2) [2]. W tych rozkładach istotne są ich rozpiętości, które określają szerokość pasma stanów zajętych elektronami, lub część pasma, z której emitowane są fotoelektrony, oraz położenia maksimów tych rozkładów.



Rys.2. Krzywe rozkładu energetycznego dla krzemu pokrytego cezem

a) Krzywe rozkładu energetycznego; energia fotonów wzbudzających jest zaznaczona na każdej krzywej.

Krzywe znormalizowano, aby dać równe wartości $N(E)dE$ pikowe

b) Krzywe rozkładu energetycznego. Dla $h\nu = 5,4$ eV krzywą podzielono na dwie krzywe charakterystyczne. Jedna jest scentrowana przy 1,0 eV druga przy 2,3 eV. Symbole Γ_{15} i L_3 wskazują stan końcowy przypisany (przez Ehrenreicha, Philippa i Phillipsa [2]) przejściu optycznemu związanemu z tymi przejściami

Dla każdego półprzewodnika nawet o dużej przerwie energetycznej można uzyskać powyższe informacje pobudzając go fotonami w obszarze nadfioletu, a nawet i w obszarze widzialnym, gdyż można obniżyć pracę wyjścia elektronów, np. przez „powleczenie” powierzchni warstwą aktywatora (cez, bar).

Wartości parametrów fotoemisyjnych, wyznaczone metodą pola hamującego w kondensatorze sferycznym, próbowali po raz pierwszy powiązać z pasmowym modelem półprzewodników Apker, Taft i Dickey [3]. Stwierdzono, że uzyskane

krzywe prądowo-napięciowe odzwierciedlają energetyczne właściwości emitującej warstwy półprzewodnika. Struktura pasmowa może być badana tak głęboko energetycznie, jak głęboko sięga sonda świetlna.

Podstawowym założeniem jest fakt, że główna fotoemisja odbywa się z pasma walencyjnego. A więc tylko ta część pasma walencyjnego może być zbadana, skąd elektrony są wyrwane. Głębokość geometryczna wnikania światła określona jest współczynnikiem absorpcji α . Ze wzoru $I = I_0 e^{-\alpha x}$ wynika, że dla grubości $x = 1/\alpha$ natężenie światła spada do I_0/e , a więc około trzy krotnie. Im α większe, tym głębokość wnikania mniejsza. Dla nadfioletu w wielu półprzewodnikach α jest rzędu 10^6 cm^{-1} , więc światło wnika tam na głębokość około 100 Å. Grubość geometryczna warstw przypowierzchniowych, tzw. obszar zaburzenia zawiera się (zależnie od koncentracji nośników objętościowych w półprzewodnikach) od 10^{-4} cm nawet do 0,1 cm. Np. w przypadku GaAs o koncentracji około 10^{17} cm^{-3} grubość zagięcia przypowierzchniowego pasm (energetycznych) jest rzędu 500 Å. Stąd wniosek, że światło „fotografuje” w GaAs strukturę tylko części warstwy przypowierzchniowej z silnie zagiętymi pasmami. Nie przeszkadza to jednak wyciągnięciu prawidłowych wniosków co do ułożenia pasm i położenia poziomu Fermiego względem poziomu próżni, gdyż odległości międzypasmowe są takie same przy powierzchni, jak i w głębi objętości.

Znając pracę wyjścia kolektora (głównego elementu układu pomiarowego) można określić wielkość pracy wyjścia badanego emitera, czyli energetycznej odległości poziomu Fermiego od poziomu próżni.

Ażeby skonstruować model energetyczny warstwy przypowierzchniowej badanego półprzewodnika należy:

1) wyznaczyć jego ϕ ;

Z charakterystyk prądowo-napięciowych odczytuje się eU_k i znając uprzednio wyznaczoną pracę wyjścia z kolektora stosuje się wzór:

$$\phi_c - \phi = eU_k \quad (2)$$

(czasem pracę wyjścia wyznacza się metodą Kelvina).

To daje nam ϕ i położenie poziomu Fermiego E_F względem poziomu próżni E_p .

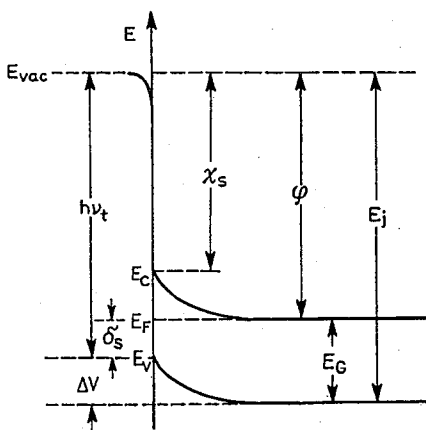
2) następnie trzeba określić położenie poziomu Fermiego względem dna pasma przewodnictwa w objętości, które przyjmujemy w badaniach objętościowych kryształów za poziom odniesienia.

Położenie poziomu Fermiego w przerwie energetycznej w objętości kryształu wyznacza się zależności $E_F = f(n_e)$ energii E_F od koncentracji i masy efektywnej nośników. Koncentrację można wyznaczyć z pomiarów stałej Halla (czyli z pomiarów objętościowych). Mając poziom próżni i poziom Fermiego (względem E_p) kreśli się położenie krawędzi pasma przewodnictwa ($\mathcal{H}_b$). Znając szerokość przerwy energetycznej

$$E_G = |h\nu_t - \phi| + |E_F| \quad (3)$$

odznaczamy ją od dna pasma przewodnictwa w dół i otrzymujemy szczyt pasma walencyjnego i zarazem energię jonizacji elektronów, jako odległość energetyczną szczytu pasma walencyjnego w objętości od poziomu próżni E_p .

3) Na prostopadłej osi oznaczamy – licząc od poziomu próżni w dół próg fotoemisji $h\nu_t$. Jeśli $h\nu_t < E_{j_{obj}}$, to znaczy, że przy powierzchni mamy zagięcie pasm do góry, a gdyby $h\nu_t$ było $>$ od E_j , to mielibyśmy zagięcie pasm energetycznych do dołu (rys. 3). Należy tu dodać, że zagięcie pasm przy powierzchni powstaje na skutek zmiany koncentracji nośników przy powierzchni spowodowanej oddziaływaniem ze stanami powierzchniowymi.



Rys.3. Struktura energetyczna przy powierzchni półprzewodnika

Wartości $h\nu_t$ wyznaczamy ze spektralnych rozkładów Y . Autor omówił to w pracach [4] [5]. Znając pracę wyjścia półprzewodnika ϕ i jeśli również znane jest położenie poziomu Fermiego względem pasma przewodnictwa E_F , to można (z wzoru 3) oszacować:

a) szerokość przerwy energetycznej E_g :

$$E_g = |E_F| + \delta, \quad (4)$$

gdzie δ – energetyczna odległość E_F od najwyższych poziomów emitujących (identyfikowanych przez Apkera i Tafta ze szczytem pasma walencyjnego).

Z pomiarów fotoemisji można określić dalsze parametry struktury fizycznej, jak:

b) wartość fotoelektrycznego progu, czyli odległość szczytu pasma walencyjnego przy powierzchni od poziomu próżni,

c) bezwzględne zagięcie pasm energetycznych przy powierzchni:

dla próbek typu n : $\Delta V = -\mathcal{K}_s$

dla próbek typu p : $\Delta V = h\nu_t - \phi$

d) powinowactwo elektronowe przy powierzchni i w objętości.

($\mathcal{K}$ jest energią potrzebną do wyrwania elektronu z dna pasma przewodnictwa poza kryształ),

e) energię jonizacji objętościowej, (E_j – jest to energia potrzebna do wyrwania elektronu z objętości kryształu, tzn. z objętościowego szczytu pasma walencyjnego poza kryształ – do poziomu próżni),

f) z pomiarów fotoemisji i natężenia światła wzbudzającego ją można liczyć gęstości stanów od ε w emitującym pasmie walencyjnym metodą Kindinga i Spicera [6].

Autor zaproponowanym przez siebie sposobem ustalił i potwierdził dla obszaru przyprogowego emisji z GaAs potęgę 3 w funkcji wykładniczej Kane opisującej fotoemisję z półprzewodników [7] [8]. Sposób ten jest dokładniejszy od informacji z pracy Gobeli i Allena, którzy ustalili tę potęgę metodą prób i błędów; autor natomiast sporządził wykres $Y/Y' = f(h\nu)$, który okazał się prostą; jej przedłużenie do przecięcia osi $h\nu$ dało $h\nu_i$, a

$$\operatorname{ctg} \alpha = \frac{h\nu - h\nu_i}{Y/Y'} = m, \quad \text{gdzie } m = 2,8 - 3,1. \quad (5)$$

Metoda ta potwierdza zasadność przyjęcia potęgi 3 dla związków III–V, a w tym i dla GaAs różnych orientacji. Potęga 3 jest zbliżona do wartości 5/2 może sugerować dominację przejść optycznych skośnych we wspomnianym obszarze przyprogowym.

Opracowana przez autora oryginalna metoda wyznaczania progów $h\nu_i$ i wykładnika m w funkcji potęgowej Kane: $Y = k(h\nu - h\nu_i)^m$ została przedstawiona w pracach [9] [5].

2. OBLICZENIA STRUKTURY ENERGETYCZNEJ Z WYDAJNOŚCI FOTOELEKTRYCZNEJ

2.1. STRUKTURA ELEKTRONOWA ZWIĄZANA Z PRZEJŚCIAMI OPTYCZNYMI W GaAs

Wykresy pasm energetycznych $E(k)$ wyznaczają energię wypadkową pasm wzdłuż głównej osi symetrii. Z półempirycznych pseudopotencjałów opartych na międzypasmowych krawędziach optycznych można (uzyskać obliczyć) modele pasm energetycznych w całym obszarze Brillouin'a. Wykazano, że niektóre struktury krawędzi pasm półprzewodników, jak i mechanizmy ucieczki elektronów można badać nie tylko na drodze optycznej [10], ale również przez fotoemisję. Z kształtu bowiem krzywej wydajności można uzyskiwać informacje o ilości pasm przewodnictwa, do których następują przejścia.

A zatem strukturę pasmową możemy identyfikować nie tylko z optycznego odbicia, ale również z krzywych wydajności fotoemisji. Metoda wydajności $Y(h\nu)$ może być nawet dokładniejsza, gdyż nie tylko uzyskuje się zgodność z odbiciową strukturą, ale przy pomocy krzywych wydajności rozdzielono kilka nowych przejść, które nie uwidaczniały się w poprzednich procesach z przejściami optycznymi. Wiele cech wydajności jest związanych z krawędziami przejść prostych; dostarczają one dokładnego potwierdzenia przyporządkowań optycznych.

Tablica 1 zawiera niektóre energie przejść dla GaAs wraz z ich ostatnimi wynikami doświadczalnymi. Z analizy tej tabeli wynika, że mimo ogólnej zgodności, poziomy Γ_{1c} i Γ_{1sc} są odwrotne; poziomy zaś Γ_{1c} i $\Gamma_{L_{1c}}$ są również wysokie (są one szczególnie wygodne do wyboru potencjału krystalicznego, który określił Bassani [9] analizując elementy macierzowe hamiltonianu H). Potencjał ten można również oszacować doświadczalnie obserwując przesunięcie brzegów pasm energetycznych w przypadku wprowadzenia domieszek [12].

Obserwowany w GaAs pik odbiciowy dla 22 eV [13] przyporządkowano przejściom między nisko leżącymi stanami d a pasmem przewodnictwa. Pik odbiciowy przy 9,6 eV w GaAs [11] przyporządkowano przejściu $L_{3v} \rightarrow L_{1c}^u$ o 10,4 eV na podstawie obliczeń metodą ortogonalizowanych fal płaskich (OPW) [12] i obliczeń Hermana [15]. Dla GaAs pik odbiciowy z maksimumami przy 6,63 i 6,89 eV jest przyporządkowany przejściu $L_{3v} \rightarrow L_{3c}$ z rozszczepieniem spinowo-orbitalnym przy L_{3v} [13]. OWP [14] i obliczenia pseudopotencjalne [16] umiejscawiają te przejścia w 5,9 i 6,0 eV.

Pik w widmie odbicia GaAs przy 4,45 eV przypisuje się przejściu $\Delta_{5v} \rightarrow \Delta_{3c}$ zamykającemu się przejściem do $\Gamma_{15v} \rightarrow \Gamma_{15c}$.

Podobne przyporządkowanie zrobiono [17] dla struktury odbić elektrycznych przy 4,46 eV i 4,64 eV wykazujących rozszczepienie 0,18 eV. W GaAs piki [13] uzyskane przez odbicie, widoczne są przy 2,90 i 3,13 eV. Obliczenia $\varepsilon_2(\omega)$ z GaAs z czynników [16] kształtu pseudopotencjału potwierdza, że piki te powstają z przejść $L_{3v} \rightarrow L_{1c}$ i $\Delta_{3v} \rightarrow \Delta_{1c}$.

Badania fotoemisyjne półprzewodników mogą dać informacje o stanach początkowych i końcowych biorących udział w przejściach optycznych [19] [20]. Wykazano, że w rozkładzie spektralnym wydajności kwantowej fotoemisji (rys. 1c) pojawiają się piki odpowiadające optycznym pikom absorpcyjnym, jeżeli stan końcowy leży powyżej poziomu próżni [17]. Odwrotnie, jeżeli stany końcowe leżą poniżej poziomu próżni, to mogą pojawić się minima w krzywej wydajności fotoemisji. Ponadto, można uzyskać dane odnośnie energii absolutnej stanu końcowego przejścia poprzez pomiar energii emitowanych elektronów. Stąd i z danych absorpcji optycznej można określić energię absolutną stanu końcowego (Tablica 2).

Okazało się, że rozkłady energetyczne dN/dE emitowanych fotoelektronów zawierają dużo więcej informacji niż wydajność $Y(\omega)$; rozkłady energetyczne są bowiem bardziej czułe na szczegółową zależność energetyczną prawdopodobieństw emisji elektronów [7]. Zagadnienie to jest jednak dość zawile.

Badania fotoemisji elektronów z czystych atomowo powierzchni wielu półprzewodników dostarczają informacji o ich elektronowej strukturze w obszarze około 5 eV powyżej krawędzi pasma walencyjnego [21]. Zmniejszając zaś, przez pokrycie powierzchni półprzewodnika monowarstwą cezu, pracę wyjścia ϕ o kilka

Tablica 1

Porównanie obliczonych energii przejścia i piki odbiciowe dla GaAs

Przejście	Obliczona energia przejścia [eV]	Pik odbicia [eV]
$L_{3v} \rightarrow L_{1c}^u$	10,4	9,6
$L_{3v} \rightarrow L_{3c}$	5,9	6,63; 6,89
$\Gamma_{15v} \rightarrow \Gamma_{15c}$	4,2	4,45
$L_{3v} \rightarrow L_{1c}$	5,4	2,90; 3,13
$\Gamma_{15v} \rightarrow \Gamma_{1c}$	5,7	1,54

Dane, których obliczone wartości nie zawierają oddziaływania spin – orbita, lub innych (względnych) efektów

Tablica 2

Porównanie wartości poziomów energetycznych określonych doświadczalnie z wartościami teoretycznymi (patrz odnośniki [34] i [35])

Przejście	Przejście optyczne	Piki w rozkładzie energetycznym	Energia ⁽¹⁾ stanu końcowego ⁽²⁾		Energia ⁽¹⁾ stanu początkowego ⁽²⁾	
			Doświadczenie	Teoria ^[32]	Doświadczenie	Teoria ^[32]
[eV]	[eV]	[eV]	[eV]	[eV]	[eV]	[eV]
$L_{3'} \rightarrow L_3$	5,4	$2,3 \pm 0,2$	$3,8 \pm 0,2$	4,0	$-1,5 \pm 0,2$	-1,4
$\Gamma_{25'} \rightarrow \Gamma_{15}$	3,4	$1,0 \pm 0,2$	$2,5 \pm 0,2$	3,4	$-0,9 \pm 0,2$	0,0
$X_4 \rightarrow X_1$	4,4	...	$1,1 \leq E \leq 1,5$	1,3	$-3,3 \leq E \leq -2,9$	-3,1
$A_3 \rightarrow A_1$				2,55 ⁽³⁾		1.15 ⁽³⁾
	3,7	$0,5 \pm 0,4$	$2,0 \pm 0,4$		$1,7 \pm 0,4$	
$L_3 \rightarrow L_1$				(2,3)		(-1,4)

(1) W odniesieniu do najwyższego maksimum pasma walencyjnego

(2) Dla efektywnego powinowactwa elektronowego 0,4 eV

(3) Zakładając, że różnica energii między A_3 i A_1 taka sama dla Si i dla Ge [33]

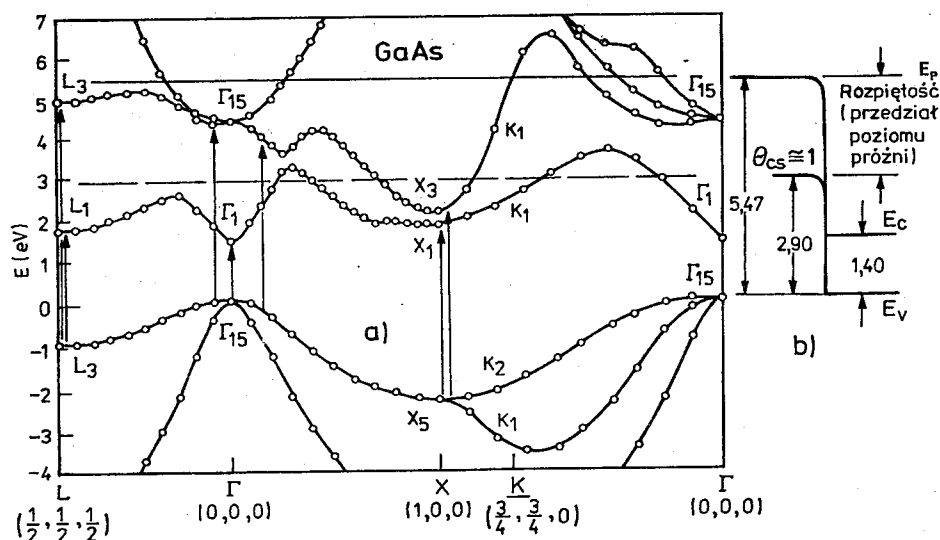
(4) Patrz odnośnik [34]

elektronowoltów w porównaniu z atomowo czystą powierzchnią, można badać elektrony z energetycznie szerszego pasma stanów przewodnictwa.

Rezultaty badań wydajności fotoemisyjnej można porównać z rezultatami metod odbiciowych. W wielu przypadkach badanie wydajności fotoemisji pozwala wyznaczyć bezwzględne energię w pasmie przewodnictwa, czyli dane o kształcie pasm energetycznych. Kształt progów może być wyznaczony przez mechanizmy ucieczki [7].

Przy wysokich energiach fotonów, fotoemisję można stosować do wyznaczania (oprócz gęstości stanów pasm walencyjnych) dokładnych położenia minimów pasmowych półprzewodników, oraz przerwy energetycznej w pasmach walencyjnych. Np. Grobman i Eastman [22] określili krawędź pasma walencyjnego E_v dla GaAs przez ekstrapolację wysokoenergetycznego zbocza krzywej rozkładu energetycznego fotoelektronów $N(E) = f(E)$ do zerowego natężenia emisji.

Odcinek liniowy eksperymentalnej krzywej wydajności spektralnej $Y(h\nu)$ dla GaAs [23] [24] jest interpretowany jako proste przejście optyczne w objętości materiału. Wykresy energii w funkcji wektora falowego $E(k)$ posiadają ekstrema wzdłuż kierunków wysokiej symetrii (np. [111], [110]). Dla półprzewodników grupy III–V wzbudzenie optyczne które powoduje proces emisji zachodzi dla stanów energetycznych mających wektor falowy k zbliżony do kierunku [110]. Stąd też rozsądna jest jakościowa dyskusja wpływu struktury pasmowej na krzywe rozkładu wydajności z krzywych energii elektronów od $k = 0$ do granicy strefy Brillouin'a. Badanie powierzchni energetycznych w przestrzeni k umożliwiło ustalenie struktury położenia i szerokości pików i minimów w wydajności powierzchni GaAs częściowo pokrytej cezem; struktura ta jest jednak dość złożona.



Rys.4. Wykres pasm energetycznych arsenku galu

a) Wykres energetyczny struktury pasmowej GaAs obliczony przez Cohena i Bersiressera metodą pseudopotencjału [16]

b) Poziome energetyczne powierzchni czystych (o orientacji (110)) pokrywanych czernią [24]

W niektórych publikacjach poziom próżni dla atomowo czystych kryształów GaAs (110) umieszczono 4,66 eV powyżej krawędzi dna pasma przewodnictwa [24]. Poziom próżni dla GaAs można obniżyć do zaledwie 1,49 eV powyżej poziomu Γ_{1c} . Rozpatrując pasma energetyczne [25] GaAs wykazano (rys. 4), że dla powierzchni atomowo czystej poziom próżni znajduje się około 0,5 eV powyżej poziomu L_{3c} i 1,8 eV powyżej Γ_{15c} . W wyniku pokrywania powierzchni czernią poziom próżni obniża się i przekracza Γ_{15c} . Przy pewnej minimalnej energii fotonu $\hbar\omega_0$ oczekuje się wzrostu wydajności – związanego z wystąpieniem przejść prostych w pobliżu lub dokładnie dla $\Gamma_{15v} \rightarrow \Gamma_{15c}$.

Rozpatrując rys. 5 [24] widzimy, że fotoelektryczna wydajność zaczyna wzrastać ostro w pobliżu krawędzi $A_1 = \hbar\omega_0 = 4,5$ eV dla poziomu próżni $3,15 \pm 0,3$ eV powyżej krawędzi pasma przewodnictwa Γ_{1c} . Ponieważ $\Gamma_{1c} - \Gamma_{15v}^{3/2} = 1,35$ eV, to jak widać na tym rysunku poziom próżni przechodzi przez poziom Γ_{15c} , który jest około 4,50 eV powyżej $\Gamma_{15v}^{3/2}$. W ten sposób możemy potwierdzić identyfikację krawędzi A_1 umiejscowionej w odbiciu optycznym przy 4,52 eV, jako wkład obszaru energetycznego $\Gamma_{15v} \rightarrow \Gamma_{15c}$ do gęstości stanów złącza pasmo 4 – pasmo 6, oraz można określić absolutne położenie energetyczne, wspomnianych już stanów, początkowego i końcowego (Tablica 2).

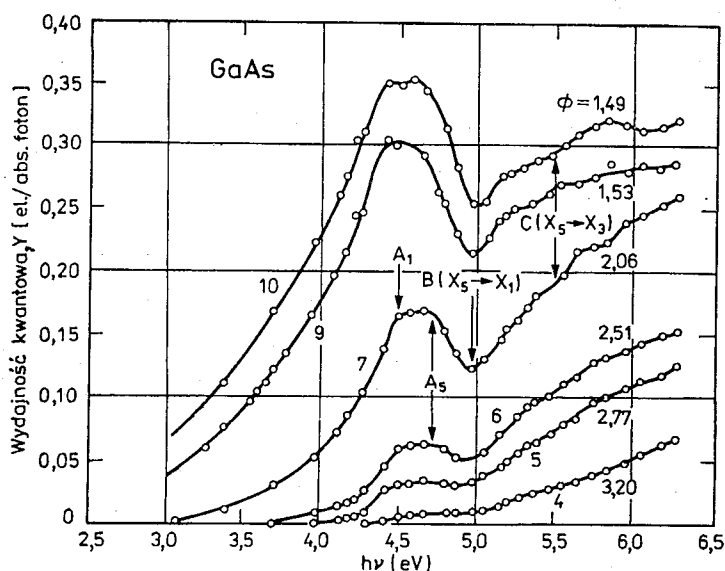
Następną ważną cechą omawianej struktury jest zagłębienie krzywej wydajności w pobliżu 5,0 eV oznaczone literą B na rys. 5. Dla krzywych o małej wydajności, kiedy poziom próżni jest 2,0 eV lub więcej powyżej Γ_{1c} , nie jest jasne czy mamy do czynienia z zagłębieniem, czy też powinniśmy traktować pik A jako nałożony na

łagodnie zmieniające się tło i zanikający w pobliżu 4,7 aż do 5,0 eV. Jakkolwiek trzy ostatnie krzywe (7, 9, 10) na rys. 5 o dużej wydajności, o poziomie próżni od 1,5 do 2,1 eV powyżej Γ_{1c} , wyraźnie wykazują występowanie w okolicy 4,95 eV zagłębienia niższego od tła. (Tego rodzaju zagłębienie było poprzednio znalezione w krzemie [26]). W optycznym odbiciu GaAs między krawędziami międzypasmowymi $X_5 \rightarrow X_1$ występuje pik przy 5,12 eV. Jednakże przed identyfikacją zagłębienia B z krawędzią $X_5 \rightarrow X_1$ należy wprowadzić poprawki w danych odbiciowych i fotoemisyjnych, aby uzyskać prawdziwą energię $X_5 \rightarrow X_1$ związaną z między pasmową krawędzią w ε_2 .

Transformacje Kramersa—Kroniga [27] danych odbiciowych [25] uzyskanych w temperaturze pokojowej umiejscawiają przejście $X_5 \rightarrow X_1$ przy 4,65 eV dla ε_2 niż niż położenie zagłębienia B w fotoelektrycznej wydajności. Uważa się, że X_{1c} znajduje się w GaAs 0,35 eV powyżej Γ_{1c} . Tak więc dla krzywych wydajności, na których zagłębienie jest najwyraźniejsze — poziom próżni położony jest 1,05 do 1,35 eV powyżej X_{1c} .

Charakterystyczną cechą półprzewodników blendy cynkowej jest rozszczepienie na $X_{1c} + X_{3c}$ podwójnie zdegenerowanego poziomu pasma przewodnictwa X_{1c} w kryształach rombowych. Między-pasmowe krawędzie związane z $X_{5v} \rightarrow X_{1c}$, X_{3c} zostały rozdzielone w niskich temperaturach w związkach $A^{III}B^V$, czyli również i w GaAs [28]. Krawędź $X_5 - X_{3c}$ jest znacznie słabsza niż $X_5 \rightarrow X_{1c}$, jak możnaby oczekiwać z krzywych $E(k)$ pokazanych na rys. 4, gdyż zarówno poprzeczna jak i podłużna masa w pobliżu X_{1c} jest znacznie większa niż w pobliżu X_{3c} .

W GaAs wartość optyczna rozszczepienia $X_{3c} \rightarrow X_{1c}$ wynosi 0,43 eV [28]. Dla trzech krzywych o największej wydajności (poziom próżni leży 0,65 i 0,89 eV powyżej X_{1c}) słabe zagłębienie określone literą C obserwuje się pomiędzy 5,5 i 5,6 eV. Środek tego zagłębienia zgadza się) w granicach eksperymentalnego

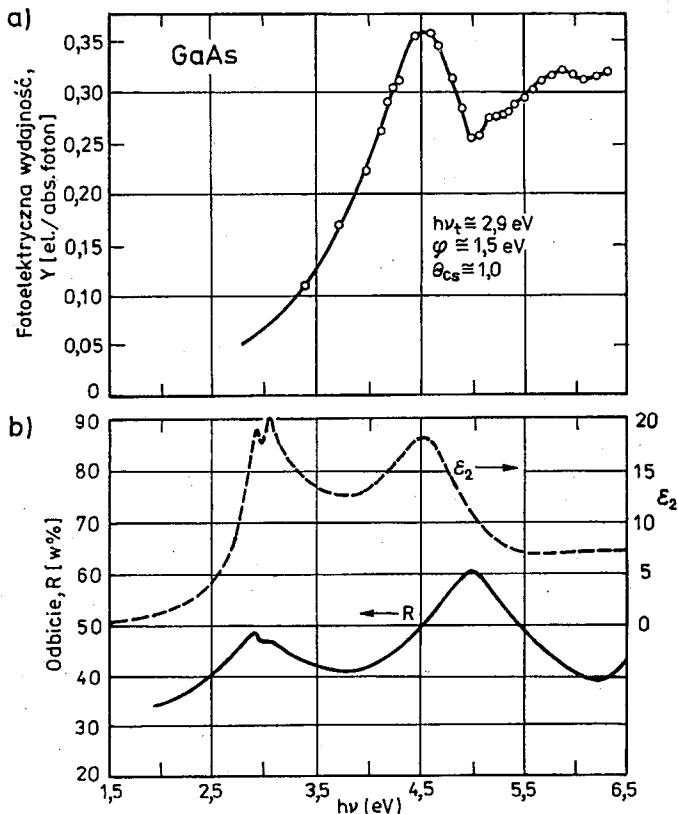


Rys.5. Wydajność fotoelektryczna GaAs w funkcji pokrycia Cs [24]

przedziału) z położeniem przewidywanym z danych optycznych ($4,95 \pm 0,43$ eV. Zagłębienie utrzymuje się w dwu następnych krzywych wydajności) poziom próżni podnosi się do 1,7 eV powyżej X_{3c}).

Omawiane wydajności (rys. 5) są krzywymi teoretycznymi. Warto nadmienić, że aby otrzymać teoretyczną krzywą wydajności trzeba przeprowadzić analizę pseudopotencjalnych pasm energii w $5 \cdot 10^4$ punktach całej sfery Brillouin'a, a aby otrzymać teoretyczny rozkład energetyczny fotoelektronów (przy zdolności rozdzielczej 0,1 eV) trzeba przeprowadzić analizę pseudopotencjalnych pasm energii w $5 \cdot 10^5$ punktach.

Przedyskutujemy zatem jakościowo kształt struktury oznaczonej A_1 i A_2 na krzywych wydajności (rys. 5). Kiedy długość absorpcji jest mała w porównaniu do średniej drogi swobodnej na której elektron emituje lub absorbuje fotony, większość elektronów zostanie wyemitowana bez zderzenia niesprężystego. W tym przypadku zakładając jednolite pokrycie powierzchni cezem i zanedbując efekty zgięcia pasma w pobliżu powierzchni, oczekujemy, że termiczne poszerzenie widma wydajności będzie bardzo podobne do poszerzenia z odbicia (rys. 6).



Rys.6. Badanie parametrów powierzchniowych arsenku galu
a) Wydajność fotoelektryczna GaAs (wg Gobeli'ego i Allena [24])
b) Dane odbicia GaAs wg Philippa i Ehrenreicha [27])

Gdy nie ma rozproszén nieelastycznych, badanie kształtów modeli energetycznych 4–6 na rys. 5 [25] prowadzi do wniosku, że w pobliżu Γ może istnieć struktura pochodząca ze zgrupowania się punktów krytycznych w tym obszarze. Oczywiście struktura tych grup jest zależna od cech szczególnych modelu pasma energetycznego. Niemniej zakres energii punktów krytycznych wynosi około 0,3 eV, co zgadza się jakościowo z różnicą 0,25 eV pomiędzy krawędzią A_1 i drugą krawędzią oznaczoną przez A_5 na rys. 5. W ramach modelu pasma energetycznego pokazanego na rys. 4, bardziej prawdopodobne jest przyporządkowanie $\Gamma_{15c} \rightarrow \Gamma_{15v}$ do A_5 niż do A_1 . Grupa punktów krytycznych znajduje się tak blisko Γ , że trudno jest ustalić położenie tego punktu nawet przy pomocy rozkładu energetycznego.

Porównując dane fotoemisyjne z optycznymi stwierdzono konieczność uwzględnienia wielu małych poprawek, które zmieniały swą wielkość od kryształu do kryształu. Wykazano, że ogólna zgodność pomiędzy optycznymi i fotoemisyjnymi danymi jest lepsza niż 0,2 eV. Prawdopodobnie różnica ta może być wyjaśniona wyłącznie za pomocą przesunięć dyspersyjnych pomiędzy stałą dielektryczną ϵ_2 i odbiciem R . Tak więc ogólna współzależność dla GaAs jest całkiem zadowalająca, jak pokazano w uproszczonej formie na rys. 6.

Przeprowadzona analiza dowodzi poprawności danych Gobeli'ego i Allena w szerszym sensie. Szczegółowe badanie i porównanie kształtów krzywych rozkładów wydajności fotoelektrycznej i odbicia R dla GaAs i dla innych półprzewodników związków III–V, oraz Si i Ge również pokrytych monowarstwą cezu, wykazuje, że rozrzuty nieregularności danych są na ogół mniejsze od 0,1 eV.

Stwierdzono, że przejścia optyczne w kryształach $A^{III}-A^V$ (a więc i w GaAs) są prawdopodobnie w większości proste. Tego można było oczekiwać biorąc pod uwagę duże statyczne stałe dielektryczne, które zmniejszają oddziaływania elektron–elektron i elektron–fonon.

Stwierdzono również występowanie w strukturze pasmowej punktów subtelnych. Punkty te obejmują rozdzielenie w 300 K rozszczepienia $X_{3c} \rightarrow X_{1c}$ w pasmie przewodnictwa, jak również umieszczenie X_{3c} nie tylko względem X_5 , ale także względem Γ_{15v} . Ponadto przy pomocy fotoemisji w wielu przypadkach rozdzielono progi $\Gamma_{15v}^{\alpha} \rightarrow \Gamma_{15c}$ ($\alpha = \frac{3}{2}; \frac{1}{2}$) co nie okazało się dotychczas możliwe w badaniach widm optycznych.

2. ENERGETYCZNE ROZKŁADY FOTOELEKTRONÓW I PRZEJŚCIA OPTYCZNE

Z pośród różnorodnych informacji dostarczanych przez rozkłady energetyczne fotoelektronów wymienimy niektóre. Rozkłady energetyczne fotoelektronów oraz dane wydajności kwantowej wykorzystywane są do obliczania rozkładów gęstości stanów emitujących pasm walencyjnych. Z rozkładów energetycznych wyciąga się wnioski dotyczące rozłożenia stanów powierzchniowych (szerokości powierzchniowego pasma) w przerwie energetycznej. Badając GaAs autor stwierdził, że brak efektu przesuwania się w prawo maksimów rozkładów fotoelektronów z próbek