

POLSKA AKADEMIA NAUK  
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

KWARTALNIK  
ELEKTRONIKI I TELEKOMUNIKACJI

ELECTRONICS AND  
TELECOMMUNICATIONS  
QUARTERLY

TOM 44 — ZESZYT 2



WYDAWNICTWO NAUKOWE PWN  
WARSZAWA 1998

## RADA REDAKCYJNA

### *Przewodniczący*

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

### *Członkowie*

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

## REDAKCJA

### *Redaktor Naczelny*

prof. dr hab. inż. WIESŁAW WOLIŃSKI

### *Zastępca Redaktora Naczelnego*

doc. dr inż. KRYSTYN PLEWKO

### *Sekretarz Odpowiedzialny*

mgr KRYSTYNA LELAKOWSKA

## ADRES REDAKCJI

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

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

Telefony domowe: Redaktora Naczelnego: 12 17 65  
Zastępcy Red. Naczelnego: 826 83 41  
Sekretarza Odpowiedzialnego: 825 29 18

## WYDAWNICTWO NAUKOWE PWN

Warszawa, ul. Miodowa 10

---

Ark. wyd. 11,0 Ark. druk. 9,5

Podpisano do druku w październiku 1998 r.

---

Papier offsetowy kl. III 70 g. B-1

Druk ukończono w listopadzie 1998 r.

---

Skład: *Amel*, Warszawa

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

# A new approach to OCR

WITOLD CZAJEWSKI

*Instytut Podstaw Elektroniki, Politechnika Warszawska*

*Received 1998.02.19*

*Authorized 1998.03.25*

The paper presents an attempt to apply the Rough Sets Theory to Optical Character Recognition with the purpose of accelerating the recognition process and decreasing the database of the characters which is usually very big. Simultaneously, it leads to lower performance and price requirements for automatic identification systems. In this approach specific characters features are referred to as an information system. The Rough Sets Theory allows extracting the most important information from the system, neglecting the other — irrelevant. This process is fully automatic and does not require any human decision in the area of usefulness of certain characters' features. A discernibility matrix, which is built in this way, constitutes a reduced database for classification algorithms. A brief description of Classical Optical Character Recognition Theory and Rough Sets Theory as well as some selected research and experimental results are also presented. As it turns out to be, even 95% of information on the recognized characters may be neglected if certain criteria are met.

*Key words:* optical character recognition, rough sets.

## 1. INTRODUCTION

Automatic identification systems play a great role in our modern civilization. Literally every element of the surrounding world can undergo the process of automatic recognition. Nowadays these elements most often are printed and hand-written characters (addresses, zip codes, signatures), bank notes and coins (e.g. in cash dispensers). A recognition algorithm must be based on a certain, earlier acquired knowledge on the objects it is about to identify. The more objects there are and the more complicated they are, the bigger the knowledge base of the algorithm (information system) is, and thus the required storage space and recognition time of a single object. Therefore it seems important to find an efficient way of reducing the

amount of data in an information system, providing the same (or lower but acceptable) quality of identification. The solution to this problem may be the Rough Sets described in [1].

The analysis of the above problem, discussion on the advantages and disadvantages of applying the Rough Sets Theory may be performed on practically any set of identifiable objects. However due to easiness of construction and acquaintance of the classical recognition methods, in the conducted research printed digits (0-9) were used. An information system based on a classical recognition theory was created in the first part of the research. In the second part, the data was reduced by applying the Rough Sets Theory.

## 2. CLASSICAL OPTICAL CHARACTER RECOGNITION THEORY

### 2.1. Introduction

Optical Character Recognition (OCR) is a field of science dealing with pattern analysis, especially with identification of characters. This problem has not been fully solved yet, which means that there is no universal device allowing perfect (or at least human-like) identification of all the characters.

The research on this problem conducted so far gives a series of methods and algorithms that prepare the text to be recognized and then recognize it. The process of text recognition is highly complex and can be divided into the following stages:

- line separation
- word separation
- single character separation
- single character recognition.

This research is focused on the last stage in which separated characters are treated as objects to be identified. This stage consists of:

- pre-processing
- feature extraction
- classification.

### 2.2. Pre-processing

The pre-processing of a character must be performed prior to its recognition. This process prepares a character in such a way that a representative feature vector can be built allowing identification later on. The pre-processing consists of\*:

- image correction — noise removal, blur removal, etc.
- normalization — applying identical translation, scale and rotation transformations to all the characters
- skeleton generation — replacing a character with thin lines representing its shape

---

\* Pre-processing could also mean the process of loading the character into the computers memory and all the operations necessary to separate a single character.



The only indispensable pre-processing stage is normalization. It eliminates such recognition mistakes in which two identical characters could be classified into two different classes only because of difference in scale or position on a bitmap. All the pre-processing techniques must be used very cautiously, otherwise they may make the recognition process even more difficult.

### 2.3. Feature extraction

There are many identification methods used depending on the characters being recognized. However, they are all based on one common notion of feature space. The feature space is a  $n$ -dimensional space in which all the recognized characters are projected. Each of the characters has its own, unique set of features called a feature vector. Generally there are three types of feature vectors: binary, integer and real. It's worth mentioning that binary features contain the least amount of information, but are computed very fast and require small storage space. On the contrary, real features can store a lot of data, but the trade off is large memory requirements and a long computation time. It is very hard to say what the best type of feature vector is. It depends on the kind of features used and the characters being recognized. From another point of view feature vectors can be divided into constant and variable length vectors. The former are definitely better due to the fact that comparing two vectors belonging to the same space is much easier than comparing two vectors from two different spaces. Besides that, the Rough Sets Theory accepts vectors of one (constant) length. Also from the purely technical point of view (storing vectors in memory) constant length vectors are preferred.

Below a few examples of commonly used feature vectors are given.

1. Profile features — based on certain character dimensions (e.g. width, height, thickness, inclination etc.). The created vector features constant length and binary or integer elements.
2. Skeleton features — calculated from the formerly generated skeleton. There exist two kinds of them. The first one contains number of end- and branching-points in the given region of the skeleton, the second one stores information of shape and orientation of every skeleton's part. The second type of skeleton coding (unlike the first one) creates a variable length vector, which is often unacceptable.
3. Pixel density features — created by dividing a character's bitmap into a certain, constant number of regions and summing up all the lit pixels in each of them. The created vector features constant length and integer elements. Another type of this feature vector is one with elements reaching successive values if number of lit pixels exceeds certain, formerly fixed, threshold values.

### 2.4. Classification

Classification — the last stage of recognition — is the process of qualifying a given character to a certain class. There exist many different classifiers (e.g.

classical, statistical, fuzzy logic or neural network), but the common feature of most of them is that the feature vector of a recognized character is an input, and the output is the number of the membership class. Calculation of the membership class is performed by the classifier employing certain knowledge. This knowledge comprises formerly generated feature vectors of characters from the learning set, which membership classes are known.

Some neural network classifiers do not need feature vectors at all. The bitmap of an unknown character is the input and this is the only information used in calculation of membership class; whereas fuzzy logic classifiers give as the output not a scalar (number of on the membership class), but a vector (fuzzy function of the character's membership to every class).

In order to recognize a character correctly, the closest vector to the character's vector from among learning set vectors must be found. In practical applications a minimum of the norm of vectors' difference is searched for. Although this method is relatively simple and fast, it may cause mistakes. In order to avoid them, another slightly more complicated classification algorithm is used. The K-Nearest Neighbor Classifier finds not one\*, but K nearest vectors to the given one. The membership class of the recognized character is the one, which has the most representatives among the K nearest vectors. In the case that the number of representatives of two or more different classes is equal, a class with minimal mean distance from the given vector is chosen. The optimal value of K parameter can be determined only experimentally, for it depends on the characters themselves (more exactly on their feature vectors).

### 3. ELEMENTS OF THE ROUGH SETS THEORY

#### 3.1. Introduction

Due to the increasing size of databases and huge amount of information stored in them, there occurred a necessity of developing efficient methods of acquiring the most essential and useful knowledge out of databases. This problem was among the most important ones in the field of modern information systems.

In 1982 Zdzisław Pawlak proposed in [1] a new theory of reasoning about data called the Rough Sets Theory.

In this approach knowledge is represented by a set of data organized in a table called an information system. Rows of the table refer to objects (e.g. characters), and columns to their attributes (features). Reduction of knowledge means deleting unnecessary and superfluous attributes and leaving only the most important ones. The efficient methods of knowledge reduction are being developed by the Rough Sets Theory.

---

\* K-Nearest Neighbour classifier with  $K = 1$  is often called a simple classical classifier or minimal distance classifier.

### 3.2. Indiscernibility relation, reduct and core of knowledge

An information systems is a 4-tuple  $S = \langle U, Q, V, f \rangle$  where:

$U$  is a finite set of objects,

$Q$  is a finite set of attributes,

$V = \cup_{q \in Q} V_q$  is a domain of the attribute  $q$ ,

$f : U \times Q \rightarrow V$  is a function such that  $f(x, q) \in V_q$  for every  $q \in Q, x \in U$ , called an information function.

We say that objects  $x, y \in U$  are indiscernible by the set of attributes  $A \subseteq Q$  in  $S$  (denoted by  $x \text{IND}(A)y$ ) if  $f(x, q) = f(y, q)$  for every  $q \in A$ .  $\text{IND}(A)$  is called an indiscernibility relation.

A subset  $A \subseteq Q$  is called a reduct if  $\text{IND}(A) = \text{IND}(Q)$  or if attributes of the subset  $A$  classify objects with the same quality as all the attributes ( $Q$ ). A reduct of knowledge is its essential part, which suffices to define all basic concepts occurring in the considered knowledge. The common part of all reducts is called a core of knowledge and is denoted by  $\text{COR}(Q)$ .

### 3.3. Reduction of consistent decision table\*

Let  $DT$  be a consistent decision table (one without objects with identical condition attributes and different decision attributes\*\*),  $C$  be a set of condition attributes,  $D$  be a set of decision attributes and  $a \in C$ .

- An attribute  $a$  is dispensable in the decision table  $DT$  if the decision table  $DT$  is also consistent without the attribute  $a$ ; otherwise the attribute  $a$  is indispensable in the decision table  $DT$
- A decision table  $DT$  is called independent if all the attributes  $a \in C$  are indispensable in the decision table  $DT$
- The subset of attributes  $R \subseteq C$  is called a reduct of  $C$  in the decision table  $DT$  if the new decision table  $\langle U, R \cup D, V, f \rangle$  is consistent (according to another definition also independent)
- The set of all indispensable attributes in the decision table  $DT$  is called the core of the decision table  $DT$

### 3.4. Discernibility matrix and function

The process of verification if a given subset is a reduct is a process involving a huge number of comparisons. For an information system consisting of  $m$  objects there are  $m^2$  comparisons for a single reduct (according to the definition). When certain conditions are satisfied this number will be reduced by slightly more than a half for each next tested reduct. One should also remember that this number must

\* Due to the fact that information systems have the form of a table, they are often referred to as decision tables

\*\* In other words, one without objects with identical attributes belonging to different class

be still multiplied by the length of the reduct. Moreover, the comparison time of two integer or real numbers must be taken into consideration. Due to the above, the algorithms based on the reduct definition are very slow and inefficient.

Notions of discernibility matrix and function introduced by Rauszer and Skowron in [2] make the reduct search time much shorter. The discernibility matrix stores only information whether the corresponding attributes of any pair of objects are different or not. This information can be stored on single bits grouped in bytes. Hence, there is much less data to be compared than previously and the processing time is shorter.

In an information system consisting of  $m$  objects  $x_i$  described with  $n$  attributes  $a_i \in A$  elements of the discernibility matrix  $c_{ij}$  are defined as follows:

$$c_{ij} = \{a \in A : a(x_i) \neq a(x_j)\} \quad \text{for} \quad ij = 1, \dots, m.$$

Due to symmetry of the discernibility matrix ( $c_{ij} = c_{ji}$ ) and  $c_{ij} = 0$  for  $i = j$ , it can be presented only by elements in the lower triangle of the matrix.

The discernibility function  $f_s$  of an information system  $S$  is a Boolean function of  $n$  Boolean variables  $\bar{a}_1, \dots, \bar{a}_n$  corresponding to attributes  $a_1, \dots, a_n$  respectively, and defined as follows:

$$f_s(\bar{a}_1, \dots, \bar{a}_n) = \bigwedge \{ \bigvee (c_{ij}) : 1 \leq j < i \leq m, c_{ij} \neq 0 \}.$$

In practical applications in order to verify if a given subset is a reduct, it is necessary to check if all the logical products of elements of the subset and elements of discernibility matrix are not empty sets. Otherwise the tested subset is not a reduct.

## 4. RESEARCH AND CONCLUSIONS

### 4.1. Introduction

All the tests conducted within the research were performed on one set of objects. These were printed numbers (0-9) based on MS Windows true-type fonts. Each number was stored on a monochrome  $48 \times 48$ -pixel bitmap. On the whole 2330 characters were tested, 233 of each kind. 155 characters belonged to the teaching set, the other 78 to the testing set.

During the first stage of the research a few types of feature vectors were tested and the best one was established. Profile features and pixel density features (both with and without thresholds) were tested. The highest quality of recognition was achieved with the latter type of feature vector and these features were used in further research. It must be mentioned that it was not the best existing vector that was searched for (it was not the objective of the research), but only several commonly used vectors were tested and the best one from among them was chosen.

In the second stage of the research many attempts of data reduction were performed on an information system built of the above mentioned feature vectors.

The shortest reducts fully describing the system were found, but the recognition results were unsatisfactory. Therefore longer reducts were found and they provided higher recognition quality.

#### 4.2. Establishment of The Best Feature Vector

The table below represents selected recognition results of the testing set for different types of feature vectors

Table 1

Recognition results of the testing set for different types of feature vectors

| Vector No<br>Quality<br>of Recognit. | 1           | 2           | 3           | 4           | 5           | 6           | 7           | 8           | 9           | 10          |
|--------------------------------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
| Minimal                              | 74,4        | 71,8        | 75,6        | 75,6        | 79,3        | 75,6        | 78,2        | 76,9        | 79,5        | 78,2        |
| Maximal                              | 98,8        | 100,0       | 97,4        | 100,0       | 100,0       | 100,0       | 100,0       | 100,0       | 100,0       | 100,0       |
| Average                              | <u>87,1</u> | <u>87,8</u> | <u>90,1</u> | <u>91,7</u> | <u>92,3</u> | <u>92,7</u> | <u>93,0</u> | <u>93,3</u> | <u>93,3</u> | <u>94,1</u> |

- 1 – profile features vector — length 64.
- 2 – pixel density feature vector with thresholds — division  $12 \times 12$ , one threshold at 50%.
- 3 – pixel density feature vector with thresholds — division  $16 \times 16$ , one threshold at 50%.
- 4 – pixel density feature vector with thresholds — division  $8 \times 8$ , five evenly spaced thresholds.
- 5 – pixel density feature vector with thresholds — division  $8 \times 8$ , seven evenly spaced thresholds.
- 6 – pixel density feature vector without thresholds — division  $16 \times 16$ .
- 7 – pixel density feature vector without thresholds — division  $12 \times 8$ .
- 8 – pixel density feature vector without thresholds — division  $6 \times 8$ .
- 9 – pixel density feature vector without thresholds — division  $8 \times 8$ .
- 10 – pixel density feature vector without thresholds — division  $8 \times 12$ .

#### 4.3. The Shortest Reduct

The information system was finally constructed of pixel density features vector without thresholds ( $8 \times 8$ ). On this basis the discernibility matrix was built. It was used for finding reducts. Initially one shortest reduct was found (Fig. 1). The length

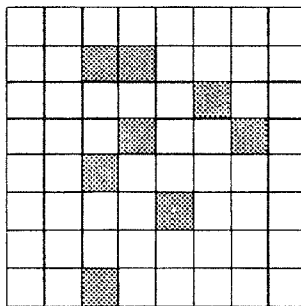


Fig. 1. Geometrical interpretation of the shortest reduct (dark squares represent the features taken into consideration in the recognition process)

of this reduct was only 12,5% of the full feature vector's length. According to the reducts definition it allowed perfect recognition of all the objects (characters) from the learning set. The testing set, however, was recognized with much lower quality, because the shortest reduct was optimally chosen for the known objects from the learning set and no others.

After further examinations the lowest (59,7%), the highest (72,4%) and the average (66,6%) recognition quality for 8-element reducts were found.

The results above clearly indicate that the shortest reduct is unsuitable for the recognition of characters from outside the learning set. The classification quality in this case is much lower (ca 30 percent points on the average) than when all the features are used. The only method of increasing the recognition quality was utilization of longer reducts.

#### 4.4. Longer Reducts

Due to unimaginably big number of all possible subsets of the whole attribute set of the information system ( $2^{64} \approx 1,8 \cdot 10^{19}$ ) verifying every „candidate for reduct” is practically impossible. Therefore a random search algorithm was used. Having tested thousands of „candidates for reducts” some conclusions of statistical nature were reached.

It is easy to notice (Fig. 2) that the number of reducts of a given length (among all the subsets of this length) is decreasing as the length of the tested subsets drops.

The Table 2 presents the worst, the best and the average recognition quality of the characters from the testing set for reducts of a given length.

As one can notice, some of the reducts give better recognition results than the full feature vector! They should be regarded as the sets of features with the least

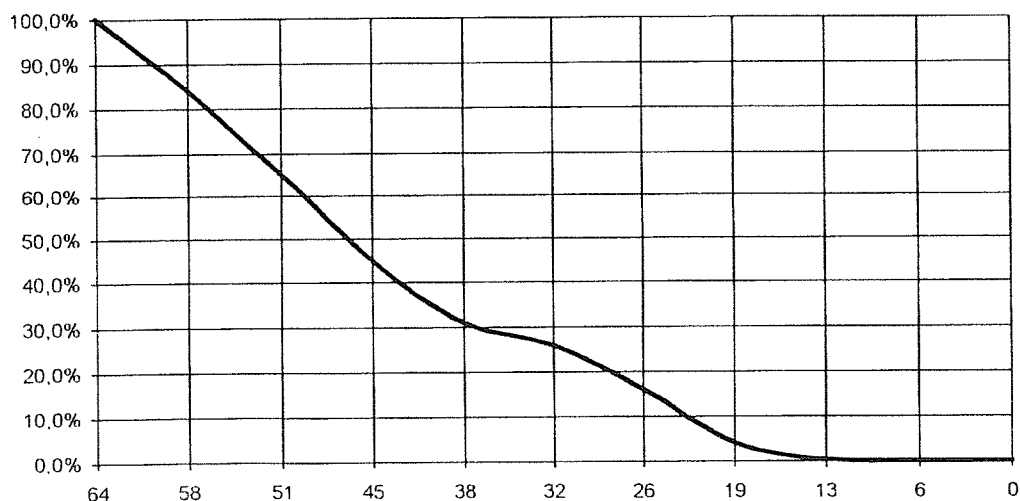


Fig. 2. The ratio of the number of reducts to the number of all the subsets of a given length

Table 2

Recognition results of the testing set for reducts of a given length

| Length of reduct<br>Quality of recognit. | 13          | 19          | 26          | 32          | 38          | 45          | 51          | 58          | 64          |
|------------------------------------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
| Minimal                                  | 70,5        | 71,5        | 79,5        | 83,6        | 86,8        | 87,4        | 89,7        | 91,3        | 93,3        |
| Maximal                                  | 74,0        | 84,0        | 87,7        | 91,3        | 93,5        | 93,2        | 93,6        | 94,1        | 93,3        |
| Average                                  | <u>72,5</u> | <u>80,1</u> | <u>85,3</u> | <u>87,8</u> | <u>89,9</u> | <u>91,1</u> | <u>92,2</u> | <u>92,8</u> | <u>93,3</u> |

important features eliminated. Unfortunately there is no analytical method of finding them. In practical applications one must rely only on statistical results (much less optimistic). As a matter of fact over 27% of 58-element reducts give better recognition results than the full feature vector, but this number decreases quickly as the vector’s length drops (Fig. 3).

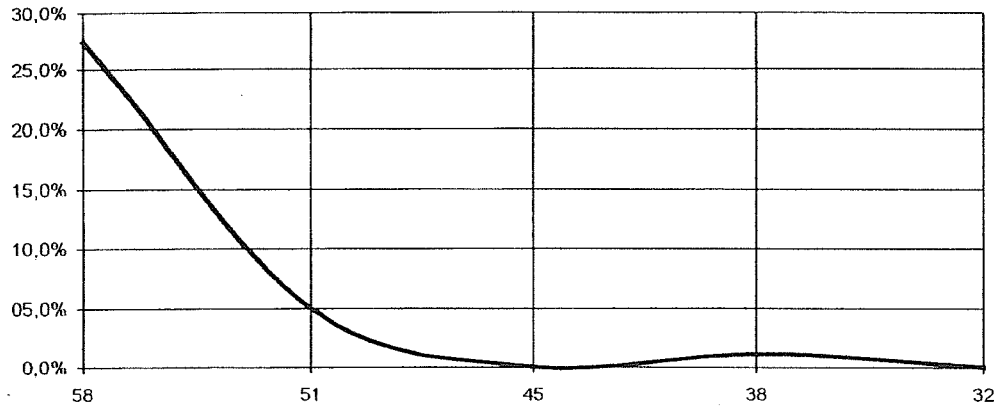


Fig. 3. Statistical distribution of „better” reducts among all the reducts of a given length

The shortest found reduct giving recognition results not worse than the full feature vector had 38 elements (ca. 60% of the full vector’s length). Graphical interpretation of this reduct is shown in Fig. 4a. For comparison, in Fig. 4b the worst found reduct of the same length is shown. The two reducts have 14 different elements.

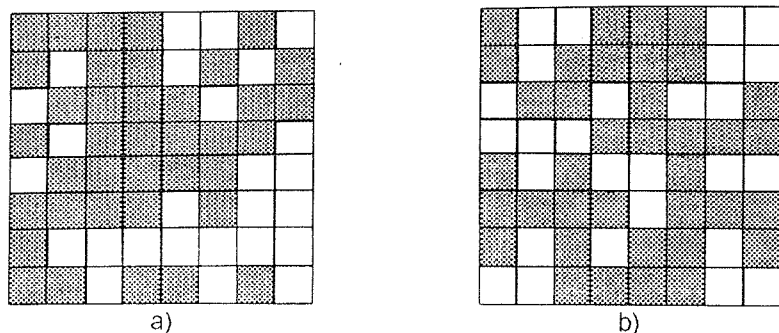


Fig. 4. The distribution of features for a) the best and b) the worst 38-element reduct (white squares represent the features that were neglected in the process of recognition)

## CONCLUSIONS

On the ground of the conducted research it can be stated that:

1. It is enough to use 5-15% of the knowledge from an information system in order to recognize all the characters from this system perfectly. In other words, up to 95% of data in an information system is dispensable.
2. It is possible to find a reduct that will give the same recognition results for characters from outside the information system as the full-length vector. In this case ca. 40-50% of data in an information system is dispensable.

Generally one can notice that the more features a single character has the shorter (in percent, not absolutely) the minimal reduct is. Also less (in percent, not absolutely) data is required to provide the same recognition quality for characters outside the information system.

The obtained results are quite interesting. They mean substantial (even 10 times in the first case) reduction of recognition time and hardware requirements, which leads directly to economical savings.

## BIBLIOGRAPHY

1. Z. Pawlak: *Rough Sets — Theoretical Aspects of Reasoning about Data*. Kluwer Academic Publishers, 1991
2. Meng-Chieh Lin: *Software System For Intelligent Data Processing And Discovering Based On The Fuzzy-Rough Sets Theory*. Thesis, San Diego State University, 1995
3. J. Lichon: *Classical Hand-written Latin Character Recognition Algorithms*. Master thesis, Institute of Control and Industrial Electronics, Warsaw University of Technology, Warszawa (Poland), 15.10.1995 (in Polish)
4. W. Czajewski, J. Żuraw: *Application of the Rough Sets Theory to Optical Character Recognition*. Master thesis, Institute of Control and Industrial Electronics, Warsaw University of Technology, Warszawa (Poland), 19.06.1997 (in Polish)



W. CZAJEWSKI

## ZBIORY PRZYBLIŻONE W ROZPOZNAWANIU ZNAKÓW

## Streszczenie

W artykule zaprezentowano próbę zastosowania teorii zbiorów przybliżonych do rozpoznawania znaków. Ma to na celu przyspieszenie procesu rozpoznawania oraz zmniejszenie (zwykle bardzo dużej) bazy danych o poszczególnych znakach. Tym samym prowadzi to obniżenia wymagań wydajnościowych, a zarazem ceny, stawianych systemom automatycznej identyfikacji znaków. W tym podejściu charakterystyczne cechy znaków (wyznaczone tradycyjnymi metodami) traktowane są jako system informacyjny. Teoria zbiorów przybliżonych pozwala na wyodrębnienie z niego najważniejszych informacji (z punktu widzenia rozpoznawania), a odrzucenie pozostałych — nieistotnych. Proces ten jest w pełni automatyczny i nie wymaga od człowieka podejmowania żadnych decyzji, co do przydatności określonych cech znaków. Utworzona w ten sposób macierz rozróżnialności stanowi zredukowaną bazę danych dla algorytmów klasyfikujących. W artykule zamieszczono także krótki opis teorii rozpoznawania znaków i teorii zbiorów przybliżonych oraz wybrane wyniki przeprowadzonych badań i eksperymentów. Jak się okazuje, przy spełnieniu pewnych warunków, nawet 95% informacji o rozpoznawanych znakach może być pominięte bez uszczerbku dla jakości klasyfikacji.

*Słowa kluczowe:* rozpoznawanie znaków, zbiory przybliżone.



## Własności korelacyjne i widmowe niestacjonarnych sygnałów zmodulowanych\*

IGOR JAWORSKI, ZDZISŁAW DRZYCIMSKI, ZBIGNIEW ZAKRZEWSKI, JACEK MAJEWSKI

*Instytut Telekomunikacji, Akademia Techniczno-Rolnicza, Bydgoszcz*

*Instytut Fizyczno-Mechaniczny NAN Ukrainy, Lwów*

*Otrzymano 1998.01.07*

*Autoryzowano 1998.04.02*

Rozpatrzono probabilistyczne modele sygnałów zmodulowanych w postaci okresowo korelowanych procesów losowych. Przeanalizowano, w ramach teorii drugiego rzędu, właściwości sygnałów, wyprowadzono równania wiążące ich charakterystyki korelacyjne i widmowe oraz charakterystyki procesów modulujących. Opisano struktury korelacyjne i widmowe szczególnych typów okresowo korelowanych procesów losowych, które są najprostszymi modelami amplitudowej i amplitudowo-fazowej modulacji. Omówiono podstawowe sposoby podejścia do szacowania charakterystyk sygnałów.

*Słowa kluczowe:* niestacjonarne sygnały zmodulowane, okresowo korelowane procesy losowe, funkcja autokorelacji, chwilowa gęstość widmowa, komponenty korelacyjne i widmowe.

### 1. WPROWADZENIE

Podstawową zasadą modulacji jest uzmiennienie, zgodnie z odpowiednimi regułami, parametrów pewnej funkcji czasowej zwanej sygnałem nośnym lub falą nośną. Sygnałem nośnym może być funkcja zdeterminowana lub funkcja losowa. W przypadku zdeterminowanym jako nośną stosuje się powszechny sygnał harmoniczny, falę prostokątną, ciąg okresowo powtarzalnych impulsów o różnych kształtach. W ogólności nośna jest ciągłą lub impulsową funkcją okresową. Modulacja wówczas polega na uzmiennieniu jej parametrów. Ponieważ w systemach przesyłania informacji sygnał modulujący ma właściwości stochastyczne, to analizę sygnałów zmodulowanych należy przeprowadzić w ramach pewnej klasy procesów

---

\* Projekt badawczy KBN Nr 8 T11D 021 13

losowych. Chociaż dzięki okresowym zmianom w czasie nośnej właściwości sygnałów zmodulowanych są powtarzalne w czasie, analiza tych sygnałów w większości przypadków jest przeprowadzona w oparciu o teorie i metody analizy stacjonarnych procesów losowych. Badanie właściwości modulacji w ramach przybliżenia stacjonarnego sprowadza się do analizy estymatorów gęstości widmowej, a zadanie poszukiwania regularnych oscylacji określa się jako zadanie wykrycia jej skokowych zmian. Ponieważ takie zadanie może być rozwiązane tylko w poszczególnych przypadkach, dlatego poprawnie należy to sformułować w inny sposób — w oparciu o model niestacjonarny. Opisane właściwości oscylacji stochastycznych z wykorzystaniem gęstości widmowej przybliżenia stacjonarnego umożliwiają tylko analizę rozkładu częstotliwościowego ich uśrednionej w czasie mocy. Taka charakterystyka nie zawiera żadnej informacji o strukturze czasowej, która polega na powtarzalności w czasie właściwości oscylacji, w tym zachowaniu kształtu. Taka powtarzalność jest podstawową cechą procesu oscylującego. Naturalnie można ją opisać za pomocą zmiennych w czasie charakterystyk probabilistycznych, określających pewne typy tych zmian. Modelem probabilistycznym o najprostszej powtarzalności (o jednym okresie) jest okresowo korelowany proces losowy [1, 2, 3]. W oparciu o ten model analizowane są w tej pracy właściwości sygnałów zmodulowanych.

## 2. KORELACYJNA I WIDMOWA ANALIZA SYGNAŁÓW

Okresowo korelowanymi procesami losowymi (OKPL)\* nazywamy takie procesy, których wartość oczekiwana  $m(t) = E\xi(t)$ ,  $E$  — operator uśredniania, i funkcja korelacji  $b(t, u) = E\tilde{\xi}(t+u)\tilde{\xi}(t)$ ,  $\tilde{\xi}(t) = \xi(t) - m(t)$  (procesy  $\xi(t)$  uważamy jako rzeczywiste), są okresowymi funkcjami czasu, tj. dla dowolnych  $t$  i pewnego  $T > 0$  spełniają warunki:

$$m(t) = m(t + T), \quad b(t, u) = b(t + T, u). \quad (1)$$

Wielkość  $T$  uważamy jako okres korelowalności. Jeżeli funkcje  $m(t)$  i  $b(t, u)$  są bezwzględnie całkwalne względem  $t$  w przedziale  $[0, T]$

$$\int_0^T |m(t)| dt < \infty, \quad \int_0^T |b(t, u)| dt < \infty, \quad (2)$$

to mogą być one przedstawione w postaci szeregów Fouriera

$$m(t) = \sum_{k \in \mathbb{Z}} m_k \exp(ik\omega_0 t), \quad b(t, u) = \sum_{k \in \mathbb{Z}} B_k(u) \exp(ik\omega_0 t), \quad (3)$$

gdzie

$$m_k = \frac{1}{T} \int_0^T m(t) \exp(-ik\omega_0 t) dt, \quad B_k(u) = \frac{1}{T} \int_0^T b(t, u) \exp(-ik\omega_0 t) dt \quad (4)$$

\* W literaturze stosowane są też terminy cyklostacjonarne procesy losowe [3, 4], okresowo niestacjonarne [5], okresowe niestacjonarne [6], procesy o okresowej strukturze [7].

oraz  $\omega_0 = \frac{2\pi}{T}$ ,  $Z$  — zbiór liczb całkowitych jak też  $|m_k| \rightarrow 0$ ,  $|B_k(u)| \rightarrow 0$  przy  $k \rightarrow \pm \infty$ . Szeregi (3) dla ciągłych względem  $t$  funkcji  $m(t)$  i  $b(t, u)$  są zbieżne jednostajnie. Wielkości  $B_k(u)$  nazywane są komponentami korelacyjnymi [1, 5].

Istotną rolę przy zrozumieniu istoty OKPL jako modeli sygnałów odgrywa ich przedstawienie [1, 3, 6]

$$\xi(t) = \sum_{k \in Z} \xi_k(t) \exp(ik\omega_0 t), \quad (5)$$

gdzie  $\xi_k(t)$  — łącznie stacjonarne procesy losowe. Z powyższego wyrażenia wynika, że w ogólnym przypadku OKPL jest sumą składowych harmonicznich zmodulowanych amplitudowo i fazowo, których pulsacje są wielokrotnością pulsacji podstawowej harmonicznej  $\omega_0$ . Wartości oczekiwane procesów  $\xi_k(t)$  są współczynnikami Fouriera wartości oczekiwanej OKPL  $m(t)$ . Dla funkcji autokorelacji OKPL  $b(t, u)$ , biorąc pod uwagę przedstawienie (5), otrzymujemy

$$b(t, u) = \sum_{l, n \in Z} \exp[i(l-n)\omega_0 t - il\omega_0 u] E \xi_l(t+u) \xi_n^*(t), \quad (6)$$

gdzie  $\xi_l(t) = \xi_l(t) - m_l$ , a  $\xi_n^*(t)$  — proces sprzężony do  $\xi_n(t)$ . Wprowadźmy funkcję

$$D_{l,n}(u) = E \xi_l(t+u) \xi_n^*(t), \quad (7)$$

która określa powiązania korelacyjne między stacjonarnymi procesami  $\xi_l(t)$  i  $\xi_n(t)$ . Po zamianie wskaźnika sumowania  $n = l - k$  wyrażenie (6) przybiera postać

$$b(t, u) = \sum_{k \in Z} \exp(ik\omega_0 t) \sum_{l \in Z} D_{l, l-k}(u) \exp(il\omega_0 u). \quad (8)$$

Porównując równania (3) i (8), uzyskujemy

$$B_k(u) = \sum_{l \in Z} D_{l, l-k}(u) \exp(il\omega_0 u). \quad (9)$$

Zerowy komponent korelacyjny określony jest, jak widać ze wzoru

$$B_0(u) = \sum_{l \in Z} D_{l, l}(u) \exp(il\omega_0 u) \quad (10)$$

funkcjami autokorelacji procesów  $\xi_l(t)$ , a komponenty wyższego rzędu  $B_k(u)$  — ich wzajemnymi funkcjami korelacji. Przy tym  $k$ -ty komponent korelacyjny określa powiązania korelacyjne procesów  $\xi_l(t)$ , których różnica między wskaźnikami jest równa  $k$ . Stąd wynika, że funkcja autokorelacji procesu (5) będzie zmieniać się okresowo w czasie tylko w przypadku powiązań korelacyjnych między modulującymi procesami  $\xi_l(t)$ . Jeżeli te powiązania nie istnieją, to dla procesu (5) funkcja autokorelacji z biegiem czasu będzie zachowywać się jak stała.

Zauważmy, że przedstawienie (5) uogólnia znane szczególne przypadki powtarzalności stochastycznej. Jeżeli

$$\xi_k(t) = a_k + \eta_k(t), \quad (11)$$

gdzie  $\eta_k(t)$  — nieskorelowane procesy losowe, to otrzymujemy model addytywny:

$$\xi(t) = \sum_{k \in \mathbb{Z}} a_k \exp(ik\omega_0 t) + \sum_{k \in \mathbb{Z}} \eta_k(t) \exp(ik\omega_0 t) = \eta(t) + f(t) \quad (12)$$

Jeżeli komponenty mają postać

$$\xi_k(t) = a_k \eta(t), \quad (13)$$

to OKPL przechodzi w model multiplikatywny:

$$\xi(t) = \eta(t) \sum_{k \in \mathbb{Z}} a_k \exp(ik\omega_0 t) = \eta(t)f(t). \quad (14)$$

Model poliharmoniczny (okresowy proces losowy) otrzymujemy w przypadku, gdy  $\xi_k(t) \equiv \eta_k$ , gdzie  $\eta_k$  — zmienne losowe:

$$\xi(t) = \sum_{k \in \mathbb{Z}} \eta_k \exp(ik\omega_0 t). \quad (15)$$

Jeżeli wszystkie współczynniki szeregu (15) są zdeterminowane, to  $\xi(t)$  jest prosto funkcją okresową.

Korzystając z zależności (5) można także otrzymać inne szczególne typy OKPL, które mogą być przedstawione w postaci szeregu trygonometrycznego:

$$\xi(t) = \xi_0(t) + \sum_{k \in \mathbb{N}} [\xi_k^c(t) \cos k\omega_0 t + \xi_k^s(t) \sin k\omega_0 t], \quad (16)$$

gdzie  $\xi_k^c(t) = [\xi_k(t) + \xi_{-k}(t)]$ ,  $\xi_k^s(t) = i[\xi_k(t) - \xi_{-k}(t)]$ ,  $\xi_{-k}(t)$  i  $\xi_k(t)$  — procesy sprzężone  $\xi_{-k}(t) = \xi_k^*(t)$ ,  $\mathbb{N}$  — zbiór liczb naturalnych.

Wówczas

$$m(t) = m_0 + \sum_{k \in \mathbb{N}} (m_k^c \cos k\omega_0 t + m_k^s \sin k\omega_0 t), \quad (17)$$

$$b(t, u) = B_0(u) + \sum_{k \in \mathbb{N}} [B_k^c(u) \cos k\omega_0 t + B_k^s(u) \sin k\omega_0 t], \quad (18)$$

przy tym

$$m_0 = \frac{1}{T} \int_0^T m(t) dt, \quad m_k^c = \frac{2}{T} \int_0^T m(t) \cos k\omega_0 t dt, \quad m_k^s = \frac{2}{T} \int_0^T m(t) \sin k\omega_0 t dt, \quad (19)$$

$$B_0(u) = \frac{1}{T} \int_0^T b(t, u) dt, \quad (20)$$

$$B_k^c(u) = \frac{2}{T} \int_0^T b(t, u) \cos k\omega_0 t dt, \quad B_k^s(u) = \frac{2}{T} \int_0^T b(t, u) \sin k\omega_0 t dt. \quad (21)$$

Oczywiste jest, że

$$m_{\pm k} = \frac{1}{2}(m_k^c \mp i m_k^s), \quad B_{\pm k}(u) = \frac{1}{2}[B_k^c(u) \mp i B_k^s(u)]. \quad (22)$$

Funkcje  $B_k^c(u)$  i  $B_k^s(u)$  nazywane są odpowiednio kosinusowymi i sinusowymi komponentami korelacyjnymi.

Wartość oczekiwana (17) opisuje regularny przebieg sygnału, a wariancja  $d(t) = b(t, 0)$  — jego chwilową moc. Własności okresowych zmian w czasie powiązań korelacyjnych między wartościami procesu w chwilach  $t$  i  $t+u$  określa się funkcją autokorelacji  $b(t, u)$ . Współczynniki Fouriera szeregów (17) i (18) są parametrami kształtów krzywych czasowych zmian odpowiednich charakterystyk.

Funkcja autokorelacji jest symetryczna względem argumentów  $t$  i  $t+u$ , tj.  $b(t, u) = b(t+u, -u)$  i  $b(t, -u) = b(t, -u, u)$ . Z tego powodu

$$\sum_{k \in Z} B_k(-u) \exp(ik\omega_0 t) = \sum_{k \in Z} B_k(u) \exp[ik\omega_0(t-u)] \quad (23)$$

i

$$B_k(-u) = B_k(u) \exp(-ik\omega_0 u). \quad (24)$$

Z ostatniego wzoru otrzymujemy

$$B_0(-u) = B_0(u), \quad B_k(-u) = B_k^c(u) \cos \omega_0 u - B_k^s(u) \sin k\omega_0 u, \quad (25)$$

$$B_k^s(-u) = B_k^s(u) \cos k\omega_0 u + B_k^c(u) \sin k\omega_0 u. \quad (26)$$

Jak widać, zerowy komponent korelacyjny jest parzystą funkcją przesunięcia  $u$ . Jest on równy uśrednionej w czasie funkcji autokorelacji OKPL:

$$B_0(u) = \mu_t \{b(t, u)\} = \lim_{\theta \rightarrow \infty} \frac{1}{2\theta} \int_{-\theta}^{\theta} b(t, u) dt = \frac{1}{T} \int_0^T b(t, u) dt. \quad (27)$$

Zerowy komponent korelacyjny  $B_0(u)$  jest funkcją dodatnio określoną, tj. funkcją spełniającą nierówność

$$\sum_{i,k=1}^L B_0(t_i - t_k) \alpha_k \alpha_i^* \geq 0, \quad (28)$$

gdzie  $\alpha_k$  i  $\alpha_i$  — dowolne liczby zespolone. Nierówność tę można wykazać biorąc pod uwagę przemienność operacji uśrednienia w czasie względem operacji liniowych:

$$\begin{aligned} \sum_{i,k=1}^L B_0(t_i - t_k) \alpha_k \alpha_i^* &= \sum_{i,k=1}^L \mu_t \{b(t + t_k, t_i - t_k)\} \alpha_i \alpha_k^* = \\ &= \mu_t \left\{ \sum_{i,k=1}^L b(t + t_k, t_i - t_k) \alpha_i \alpha_k^* \right\} = \mu_t E \left\{ \left| \sum_{k=1}^L \xi(t + t_k) \alpha_k \right|^2 \right\} \geq 0 \end{aligned} \quad (29)$$

Tak więc, zerowy komponent korelacyjny ma własności funkcji autokorelacji stacjonarnego procesu losowego i zgodnie z twierdzeniem Wienera-Chinczyna jeżeli jest całkowny bezwzględnie

$$\int_{-\infty}^{\infty} |B_0(u)| du < \infty, \quad (30)$$

może być przedstawiony w postaci całki

$$B_0(u) = \int_{-\infty}^{\infty} e^{i\omega u} f_0(\omega) d\omega, \quad (31)$$

gdzie funkcja  $f_0(\omega)$ , tak zwany zerowy komponent widmowy, jest funkcją rzeczywistą, parzystą i nieujemną:  $f_0(\omega) = f_0^*(\omega)$ ,  $f_0(-\omega) = f_0(\omega)$ ,  $f_0(\omega) \geq 0$ . Własności te umożliwiają interpretowanie komponentu zerowego jako gęstości widmowej uśrednionej w czasie mocy sygnału.

Widać więc, że parzystą funkcją przesunięcia  $u$  jest funkcja

$$K_l(u) = B_l(u) \exp\left(-il\omega_0 \frac{u}{2}\right), \quad (32)$$

gdych

$$K_l(-u) = B_l(-u) \exp\left(il\omega_0 \frac{u}{2}\right) = B_l(u) \exp\left(-il\omega_0 \frac{u}{2}\right) = K_l(u), \quad (33)$$

oraz funkcje

$$K_{\pm l}^c(u) = B_l^c(u) \cos l\omega_0 \frac{u}{2} - B_l^s(u) \sin l\omega_0 \frac{u}{2}, \quad (34)$$

$$K_{\pm l}^s(u) = \pm B_l^c(u) \sin l\omega_0 \frac{u}{2} \pm B_l^s(u) \cos l\omega_0 \frac{u}{2}, \quad l \in \mathbb{Z} \quad (35)$$

przy tym

$$K_l(u) = \frac{1}{2} [K_l^c(u) - K_l^s(u)]. \quad (36)$$

Utwórzmy szereg

$$k(t, u) = \sum_{l \in \mathbb{Z}} K_l(u) \exp(il\omega_0 t) = \sum_{l \in \mathbb{Z}} B_l(u) \exp\left[il\omega_0 \left(t - \frac{u}{2}\right)\right]. \quad (37)$$

Porównując równania (3) i (37), mamy

$$k(t, u) = b\left(t - \frac{u}{2}, u\right) = E \xi\left(t - \frac{u}{2}\right) \xi^*\left(t + \frac{u}{2}\right). \quad (38)$$

Stąd wynika, że funkcje  $K_l(u)$  są komponentami Fouriera momentu centralnego  $k(t, u)$ , który określa powiązania korelacyjne między wartościami procesu w chwili



lach  $t - \frac{u}{2}$  i  $t + \frac{u}{2}$ . Oczywiście taka funkcja w odróżnieniu od  $b(t, u)$  jest funkcją parzystą względem zmiennej  $u$ .

Uwzględniając przedstawienie OKPL poprzez procesy stacjonarne, dla funkcji autokorelacji  $k(t, u)$  mamy

$$k(t, u) = \sum_{m, n \in \mathbb{Z}} \exp \left[ i(n-m)\omega_0 t + i(m+n)\omega_0 \frac{u}{2} \right] E \xi_m^* \left( t - \frac{u}{2} \right) \xi_m \left( t + \frac{u}{2} \right) \quad (39)$$

Wprowadzając nowy wskaźnik sumowania  $l = n - m$ , mamy

$$k(t, u) = \sum_{l \in \mathbb{Z}} \exp \left[ il\omega_0 \left( t - \frac{u}{2} \right) \right] \sum_{n \in \mathbb{Z}} D_{n, n-1}(u) \exp(in\omega_0 u). \quad (40)$$

Z powyższego wzoru uzyskujemy zależności komponentów  $K_l(u)$  od auto- i interkorelacyjnych funkcji stacjonarnych składowych OKPL  $\xi_k(t)$ :

$$K_l(u) = \exp \left( -il\omega_0 \frac{u}{2} \right) \sum_{n \in \mathbb{Z}} D_{n, n-1}(u) \exp(in\omega_0 u). \quad (41)$$

Wydzielimy z funkcji autokorelacji  $b(t, u)$  część parzystą i nieparzystą

$$b^{(p)}(t, u) = \frac{1}{2} [b(t, u) + b(t, -u)] = \frac{1}{2} [b(t, u) + b(t - u, u)], \quad (42)$$

$$b^{(n)}(t, u) = \frac{1}{2} [b(t, u) - b(t, -u)] = \frac{1}{2} [b(t, u) - b(t - u, u)]. \quad (43)$$

Jeżeli proces  $\xi(t)$  jest stacjonarny, tj. szereg Fouriera funkcji autokorelacji zawiera tylko zerowy komponent korelacyjny  $B_0(u)$ , to część nieparzysta jest równa zeru  $b^{(n)}(t, u) = 0$ . Tak więc obecność składowej nieparzystej jest dowodem na to, że proces jest niestacjonarny, zaś składowa ta może być wykorzystana przy opisie niestacjonarności. Ze wzoru (43) wynika, że wartości części nieparzystej są określone przez zależność między szybkością zmian powiązań korelacyjnych z czasem  $t$  i przesunięciem  $u$ . Jeżeli funkcja korelacji wraz ze wzrostem przesunięcia szybko maleje, to część nieparzysta jest bardzo mała. W przypadku, gdy funkcja korelacji różni się od opisanej powyżej, to dla każdej chwili czasu będą istnieć tzw. korelacje przejściowe. Własności takich procesów losowych opisuje funkcja  $b^{(n)}(t, u)$ . Analiza struktury probabilistycznej sygnałów przy pomocy autokorelacji parzystej  $k(t, u)$  już tego nie umożliwia.

Parzystą i nieparzystą część funkcji autokorelacji można przedstawić w postaci

$$\begin{aligned} b^{(p)}(t, u) &= B_0(u) + \frac{1}{2} \sum_{k \in \mathbb{N}} [(B_k^c(u) + B_k^c(-u)) \cos k\omega_0 t + [B_k^s(u) + B_k^s(-u)] \sin k\omega_0 t] = \\ &= B_0(u) + \sum_{k \in \mathbb{N}} \cos k\omega_0 \frac{u}{2} \left[ B_k^c(u) \cos k\omega_0 \left( t - \frac{u}{2} \right) + B_k^s(u) \sin k\omega_0 \left( t - \frac{u}{2} \right) \right], \end{aligned}$$

$$\begin{aligned}
 b^{(n)}(t, u) &= \frac{1}{2} \sum_{k \in N} [[B_k^c(u) - B_k^c(-u)] \cos k\omega_0 t + [B_k^s(u) - B_k^s(-u)] \sin k\omega_0 t] = \\
 &= \sum_{k \in N} \sin k\omega_0 \frac{u}{2} \left[ B_k^s(u) \cos k\omega_0 \left( t - \frac{u}{2} \right) - B_k^c(u) \sin k\omega_0 \left( t - \frac{u}{2} \right) \right]. \quad (44)
 \end{aligned}$$

Wychodząc z równości (34) i (35) uzyskujemy następujące zależności

$$\begin{aligned}
 b^{(p)}(t, u) &= B_0^-(u) + \sum_{k \in N} \cos k\omega_0 \frac{u}{2} \left[ \cos k\omega_0 t \left[ B_k^c(u) \cos k\omega_0 \frac{u}{2} - B_k^s(u) \sin k\omega_0 \frac{u}{2} \right] + \right. \\
 &\quad \left. + \sin k\omega_0 t \left[ B_k^c(u) \sin k\omega_0 \frac{u}{2} + B_k^s(u) \cos k\omega_0 \frac{u}{2} \right] \right] = \\
 &= B_0(u) + \sum_{l \in N} \cos l\omega_0 \frac{u}{2} [K_l^c(u) \cos l\omega_0 t + K_l^s(u) \sin l\omega_0 t], \quad (45)
 \end{aligned}$$

$$\begin{aligned}
 b^{(n)}(t, u) &= \sum_{k \in N} \sin k\omega_0 \frac{u}{2} \left[ \cos k\omega_0 t \left[ B_k^s(u) \cos k\omega_0 \frac{u}{2} + B_k^c(u) \sin k\omega_0 \frac{u}{2} \right] + \right. \\
 &\quad \left. + \sin k\omega_0 t \left[ B_k^s(u) \sin k\omega_0 \frac{u}{2} - B_k^c(u) \cos k\omega_0 \frac{u}{2} \right] \right] = \\
 &= \sum_{l \in N} \sin l\omega_0 \frac{u}{2} [K_l^s(u) \cos l\omega_0 t - K_l^c(u) \sin l\omega_0 t], \quad (46)
 \end{aligned}$$

które określają związek parzystej i nieparzystej  $b^{(n)}(t, u)$  części z parzystymi komponentami Fouriera  $K_l^c(u)$  i  $K_l^s(u)$ .

Określmy za pomocą przekształcenia Fouriera funkcji autokorelacji  $b(t, u)$  chwilową gęstość widmową  $f(\omega, t)$ :

$$f(\omega, t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} b(t, u) \exp(-i\omega u) du. \quad (47)$$

Funkcja  $f(\omega, t)$  jest funkcją zespoloną:  $f(\omega, t) = \operatorname{Re} f(\omega, t) - i \operatorname{Im} f(\omega, t)$ . Po podstawieniu do (47) funkcji autokorelacji  $b(t, u)$  jako sumy parzystej  $b^{(p)}(t, u)$  i nieparzystej  $b^{(n)}(t, u)$  części i uproszczeniu otrzymujemy

$$\operatorname{Re} f(\omega, t) = \frac{1}{\pi} \int_0^{\infty} b^{(p)}(t, u) \cos \omega u du, \quad \operatorname{Im} f(\omega, t) = \frac{1}{\pi} \int_0^{\infty} b^{(n)}(t, u) \sin \omega u du. \quad (48)$$

Oznacza to, że część rzeczywista chwilowej gęstości widmowej jest kosinusowym przekształceniem Fouriera parzystej części funkcji autokorelacji, a część urojona — sinusowym przekształceniem części nieparzystej. Widać, że  $\operatorname{Re} f(\omega, t)$  jest parzystą funkcją częstotliwości, a  $\operatorname{Im} f(\omega, t)$  — nieparzystą:

$$\operatorname{Re} f(-\omega, t) = \operatorname{Re} f(\omega, t), \quad \operatorname{Im} f(-\omega, t) = -\operatorname{Im} f(\omega, t). \quad (49)$$

Wyrażenie dla funkcji autokorelacji

$$b(t, u) = \int_{-\infty}^{\infty} f(\omega, t) \exp(i\omega u) d\omega \quad (50)$$

zapiszemy w postaci

$$b(t, u) = 2 \int_0^{\infty} [\operatorname{Re} f(\omega, t) \cos \omega u + \operatorname{Im} f(\omega, t) \sin \omega u] d\omega, \quad (51)$$

stąd

$$b^{(n)}(t, u) = 2 \int_0^{\infty} \operatorname{Re} f(\omega, t) \cos \omega u d\omega, \quad b^{(n)}(t, u) = 2 \int_0^{\infty} \operatorname{Im} f(\omega, t) \sin \omega u d\omega. \quad (52)$$

Dla wariancji sygnału  $d(t) = b(t, 0)$  mamy

$$d(t) = 2 \int_0^{\infty} \operatorname{Re} f(\omega, t) d\omega. \quad (53)$$

Z tego równania wynika, że całkując w całym zakresie częstotliwości rzeczywistą część chwilowej gęstości widmowej otrzymujemy chwilową moc procesu. Jednak funkcja  $\operatorname{Re} f(\omega, t)$  nie może być interpretowana jako gęstość widmowa mocy, ponieważ nie zawsze jest ona nieujemną dla wszystkich  $(\omega, t)$ . Warunek  $\operatorname{Re} f(\omega, t) \geq 0$  jest zawsze spełniony dla lokalnie stacjonarnego procesu losowego\* [8], gdyż powiązania korelacyjne szybko maleją wraz ze wzrostem przesunięcia  $u$  w porównaniu ze zmianami w czasie i wtedy w przybliżeniu  $b(t - u, u) \approx b(t, u)$ . Wówczas część nieparzysta jest równa zero  $b^{(n)}(t, u) \approx 0$ . Warunek dodatniej określoności dla funkcji autokorelacji

$$\sum_{j,k=1}^L b(t_j, t_k - t_j) \alpha_j \alpha_k^* \geq 0 \quad (54)$$

może być wówczas przepisany w postaci

$$\sum_{j,k=1}^L b(t, t_k - t_j) \alpha_j \alpha_k^* \geq 0. \quad (55)$$

W tym wypadku, dla wszystkich  $t$ , funkcja autokorelacji  $b(t, u)$  ma własności funkcji autokorelacji procesu stacjonarnego, wobec tego  $\operatorname{Re} f(\omega, t) \geq 0$  i  $\operatorname{Re} f(\omega, t)$  można interpretować jako chwilową gęstość widmową mocy. Dla procesu lokalnie stacjonarnego  $\operatorname{Im} f(\omega, t) \approx 0$ . W ogólnym przypadku sinusowe przekształcenia części

\* Lokalnie stacjonarne procesy losowe mogą być zdefiniowane z wykorzystaniem pojęcia  $\tau$  — skorelowanych procesów losowych [9].  $\tau$  — skorelowanymi procesami losowymi nazywane są procesy, których funkcja korelacji  $b(t, u)$  dla dowolnych  $t$  jest równa zero, jeżeli  $|u| > \tau$ . Lokalnie stacjonarne procesy są to takie  $\tau$  — skorelowane procesy, których funkcja korelacji w dodatku spełnia warunek:  $b(t + x, u) \approx b(t, u)$  dla  $|x| < \tau$ .

nieparzystej funkcji autokorelacji opisuje własności częstotliwości procesów przejściowych.

Gęstość widmowa  $f(\omega, t)$  jest funkcją okresową zatem

$$f(\omega, t) = \sum_{k \in Z} f_k(\omega) \exp(ik\omega_0 t), \quad (56)$$

gdzie

$$f_k(\omega) = \frac{1}{T} \int_0^T f(\omega, t) \exp(-ik\omega_0 t) dt. \quad (57)$$

Podstawiając wartość (47) do wzoru (57) otrzymujemy wyrażenie

$$\begin{aligned} f_k(\omega) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \exp(-i\omega u) \left[ \frac{1}{T} \int_0^T b(t, u) \exp(-ik\omega_0 t) dt \right] du = \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k(u) \exp(-i\omega u) du, \end{aligned} \quad (58)$$

określające związek między komponentami korelacyjnymi i widmowymi.

Po uwzględnieniu równania (9) mamy

$$f_k(\omega) = \frac{1}{2\pi} \sum_{l \in Z} \int_{-\infty}^{\infty} D_{l,l-k}(u) \exp[-i(\omega - l\omega_0)u] du = \sum_{l \in Z} f_{l,l-k}(\omega - l\omega_0). \quad (59)$$

Stąd wynika, że każdy komponent widmowy jest równy sumie gęstości widmowych odpowiednich procesów modulujących, przesuniętych wzdłuż osi częstotliwości o wielkość  $l\omega_0$ . Zerowy komponent widmowy zależy tylko od gęstości widmowych mocy procesów modulujących  $f_{l,l}(\omega + l\omega_0)$ , a komponenty widmowe rzędu  $k$  — od wzajemnych gęstości widmowych procesów, których różnica między wskaźnikami jest równa  $k$ .

Komponenty widmowe określają również stopień korelowalności składowych harmonicznych formułujących pewny sygnał. Ten wynik otrzymujemy analizując rozwinięcie harmoniczne OKPL. Takie rozwinięcie można wyprowadzić na podstawie przedstawienia (5). Rozwinięcie harmoniczne procesów stacjonarnych ma postać

$$\xi_k(t) = \int_{-\infty}^{\infty} \exp(i\omega t) dz_k(\omega), \quad (60)$$

gdzie  $z_k(\omega)$  — funkcje losowe o przyrostach

$$Edz_k(\omega) = m_k \delta(\omega) d\omega, \quad Edz_k(\omega_1) d\dot{z}_l^*(\omega_2) = f_{k,l}(\omega_1 - \omega_2) d\omega_1 d\omega_2, \quad (61)$$

$$d\dot{z}_k(\omega) = dz_k(\omega) - m_k \delta(\omega) d\omega \quad (62)$$

$\delta(\omega)$  — funkcja delta Diraca.

Po uwzględnieniu wzoru (60) szereg (5) przybiera postać:

$$\xi(t) = \sum_{k \in Z} \int_{-\infty}^{\infty} e^{i(\omega + k\omega_0)t} dz_k(\omega) = \int_{-\infty}^{\infty} e^{i\omega t} \sum_{k \in Z} dz(\omega - k\omega_0) \quad (63)$$

Wprowadźmy nową funkcję losową  $Z(\omega)$ , której przyrosty są równe

$$dz(\omega) = \sum_{k \in Z} dz_k(\omega - k\omega_0). \quad (64)$$

Dla OKPL wtedy mamy

$$\xi(t) = \int_{-\infty}^{\infty} e^{i\omega t} dz(\omega). \quad (65)$$

Zauważmy, że przyrosty  $dz(\omega)$  spełniają warunki:

$$Edz(\omega) = \sum_{k \in Z} Edz_k(\omega - k\omega_0) = \sum_{k \in Z} m_k \delta(\omega - k\omega_0) d\omega \quad (66)$$

$$\begin{aligned} Edz^2(\omega_1) d\hat{z}^*(\omega_2) &= \sum_{n, l \in Z} Edz_l(\omega_1 - l\omega_0) d\hat{z}_n^*(\omega_2 - n\omega_0) = \\ &= \sum_{n, j \in Z} f_{l, n}(\omega_1 - l\omega_0) \delta[\omega_1 - \omega_2 - (l - n)\omega_0] d\omega_1 d\omega_2 = \\ &= \sum_{k \in Z} f_k(\omega_1) \delta(\omega_1 - \omega_2 - k\omega_0) d\omega_1 d\omega_2, \end{aligned} \quad (67)$$

gdzie

$$d\hat{z}^2(\omega) = dz(\omega) - \sum_k m_k \delta(\omega - k\omega_0) d\omega, \quad (68)$$

$$f_k(\omega_1) = \sum_{l \in Z} f_{l, l-k}(\omega - l\omega_0). \quad (69)$$

Wzory (66) i (67) jak widać zbiegają się. Tak więc, wartość oczekiwana przyrostów częstotliwościowej funkcji losowej w odróżnieniu od procesu stacjonarnego ma skoki nie tylko w punkcie  $\omega = 0$  lecz również w punktach  $\omega = k\omega_0$ . Składowe harmoniczne o częstotliwościach  $\omega$  i  $\omega = k\omega_0$  są skorelowane. Uwzględniając wyrażenia (20) i (21) łatwo otrzymujemy

$$m(t) = \int_{-\infty}^{\infty} \exp(i\omega t) Edz(\omega) = \sum_{k \in Z} m_k \exp(ik\omega_0 t), \quad (70)$$

$$b(t, u) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \exp[i(\omega_1 - \omega_2)t - i\omega_2 u] Ed\hat{z}(\omega_1) d\hat{z}^*(\omega_2) = \sum_{k \in Z} B_k(u) \exp(ik\omega_0 t). \quad (71)$$

W oparciu o rozwinięcie (65) może być analizowane przetwarzanie sygnałów przy pomocy układów o różnych charakterystykach częstotliwościowych.

Przyjmijmy, że  $f_k(\omega) = \operatorname{Re} f_k(\omega) - i \operatorname{Im} f_k(\omega)$ . Wtedy wzór (56) przybiera postać

$$f(\omega, t) = \sum_{k \in \mathbb{Z}} [[\operatorname{Re} f_k(\omega) \cos k\omega_0 t + \operatorname{Im} f_k(\omega) \sin k\omega_0 t] - i[\operatorname{Im} f_k(\omega) \cos k\omega_0 t - \operatorname{Re} f_k(\omega) \sin k\omega_0 t]] \quad (72)$$

$$\operatorname{Re} f(\omega, t) = \sum_{k \in \mathbb{Z}} [\operatorname{Re} f_k(\omega) \cos k\omega_0 t + \operatorname{Im} f_k(\omega) \sin k\omega_0 t], \quad (73)$$

$$\operatorname{Im} f(\omega, t) = \sum_{k \in \mathbb{Z}} [\operatorname{Im} f_k(\omega) \cos k\omega_0 t - \operatorname{Re} f_k(\omega) \sin k\omega_0 t]. \quad (74)$$

Rzeczywiste i urojone części komponentów widmowych są wyrażone za pośrednictwem rzeczywistej i urojonej części chwilowej gęstości widmowej przez równania

$$\operatorname{Re} f_k(\omega) = \frac{1}{T} \int_0^T [\operatorname{Re} f(\omega, t) \cos k\omega_0 t - \operatorname{Im} f(\omega, t) \sin k\omega_0 t] dt, \quad (75)$$

$$\operatorname{Im} f_k(\omega) = \frac{1}{T} \int_0^T [\operatorname{Im} f(\omega, t) \cos k\omega_0 t + \operatorname{Re} f(\omega, t) \sin k\omega_0 t] dt. \quad (76)$$

Ponieważ  $B_{-k}(u) = B_k^*(u)$ , to

$$f_{-k}(\omega) = \frac{1}{2\pi} \int_{-\infty}^{\infty} B_{-k}(u) \exp(-i\omega u) du = \left[ \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k(u) \exp(i\omega u) du \right]^* = f_k^*(-\omega). \quad (77)$$

Uwzględniając równość  $B_k(-u) = B_k(u) \exp(-i\omega_0 u)$  mamy

$$\begin{aligned} f_k(-\omega) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k(-u) \exp(-i\omega u) du = \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k(u) \exp[-i(\omega + k\omega_0)u] du = \\ &= f_k(\omega + k\omega_0). \end{aligned} \quad (78)$$

Tak więc

$$f_{-k}(\omega) = f_k^*(-\omega), \quad f_k(-\omega) = f_k(\omega + k\omega_0). \quad (79)$$

Równania te dla komponentów widmowych są wynikiem symetrii funkcji autokorelacji względem argumentów  $t$  i  $t + u$ . Dla części rzeczywistej i urojonej mamy

$$\operatorname{Re} f_{-k}(\omega) = \operatorname{Re} f_k(-\omega), \quad \operatorname{Im} f_{-k}(\omega) = -\operatorname{Im} f_k(-\omega), \quad (80)$$

$$\operatorname{Re} f_{-k}(-\omega) = \operatorname{Re} f_k(\omega + k\omega_0), \quad \operatorname{Im} f_{-k}(-\omega) = \operatorname{Im} f_k(\omega + k\omega_0). \quad (81)$$

Po podstawieniu tych równości do wyrażeń (73) i (74) otrzymujemy

$$\begin{aligned} \operatorname{Re} f(\omega, t) &= f_0(\omega) + \sum_{k \in N} [[\operatorname{Re} f_k(\omega) + \operatorname{Re} f_k(-\omega)] \cos k\omega_0 t + \\ &+ [\operatorname{Im} f_k(\omega) + \operatorname{Im} f_k(-\omega)] \sin k\omega_0 t] = f_0(\omega) + \sum_{k \in N} [[\operatorname{Re} f_k(\omega) + \\ &+ \operatorname{Re} f_k(\omega + k\omega_0)] \cos k\omega_0 t + [\operatorname{Im} f_k(\omega) + \operatorname{Im} f_k(\omega + k\omega_0)] \sin k\omega_0 t]. \end{aligned} \quad (82)$$

$$\begin{aligned} \operatorname{Im} f(\omega, t) &= \sum_{k \in N} [[\operatorname{Im} f_k(\omega) - \operatorname{Im} f_k(-\omega)] \cos k\omega_0 t - \\ &- [\operatorname{Re} f_k(\omega) - \operatorname{Re} f_k(-\omega)] \sin k\omega_0 t] = \sum_{k \in N} [[\operatorname{Im} f_k(\omega) - \\ &- \operatorname{Im} f_k(\omega + k\omega_0)] \cos k\omega_0 t - [\operatorname{Re} f_k(\omega) - \operatorname{Re} f_k(\omega + k\omega_0)] \sin k\omega_0 t]. \end{aligned} \quad (83)$$

Oczywiste jest, że funkcje

$$\operatorname{Re} f_k(\omega) + \operatorname{Re} f_k(\omega + k\omega_0) \quad \text{ i } \quad \operatorname{Im} f_k(\omega) + \operatorname{Im} f_k(\omega + k\omega_0) \quad (84)$$

są funkcjami parzystymi, a funkcje

$$\operatorname{Im} f_k(\omega) - \operatorname{Im} f_k(\omega + k\omega_0) \quad \text{ i } \quad \operatorname{Re} f_k(\omega) - \operatorname{Re} f_k(\omega + k\omega_0) \quad (85)$$

nieparzystymi funkcjami częstotliwości  $\omega$ .

Podobnie jak zostały zdefiniowane komponenty widmowe  $f_k(\omega)$ , można również zdefiniować komponenty widmowe  $g_k(\omega)$ , będące przekształceniami Fouriera parzystych komponentów korelacyjnych  $K_l(u)$ :

$$g_l(\omega) = \frac{1}{2\pi} \int_{-\infty}^{\infty} K_l(u) \exp(-i\omega u) du. \quad (86)$$

Ponieważ  $K_l(-u) = K_l(u)$ , więc funkcje  $g_l(\omega)$  są parzystymi funkcjami częstotliwości, a całka (86) sprowadza się do postaci

$$g_l(\omega) = \frac{1}{\pi} \int_0^{\infty} K_l(u) \cos \omega u du, \quad (87)$$

stąd

$$g_l^e(\omega) = \frac{1}{2\pi} \int_0^{\infty} K_l^e(u) \cos \omega u du, \quad g_l^s(\omega) = \frac{1}{2\pi} \int_0^{\infty} K_l^s(u) \cos \omega u du, \quad (88)$$

przy tym

$$g_l(\omega) = g_l^e(\omega) - i g_l^s(\omega), \quad g_{-l}(\omega) = g_l^*(\omega). \quad (89)$$

Uwzględniając, że  $K_l(u) = B_l(u) \exp\left(-i\omega_0 \frac{u}{2}\right)$ , otrzymujemy następujące zależności, wiążące komponenty widmowe obu typów

$$g_k(\omega) = \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k(u) \exp \left[ -iu \left( \omega + k \frac{\omega_0}{2} \right) \right] du = f_k \left( \omega + k \frac{\omega_0}{2} \right). \quad (90)$$

Z powyższego wypływa wniosek, że  $f_k \left( \omega + k \frac{\omega_0}{2} \right)$  — jest parzystą funkcją częstotliwości. Łatwo o tym przekonać się korzystając bezpośrednio ze wzoru  $f_k(-\omega) = f_k(\omega + k\omega_0)$ :

$$f_k \left( -\omega + \frac{k\omega_0}{2} \right) = f_k \left( \omega - \frac{k\omega_0}{2} + k\omega_0 \right) = f_k \left( \omega + \frac{k\omega_0}{2} \right). \quad (91)$$

Szeregi (42) i (43) dla parzystej i nieparzystej części funkcji autokorelacji przedstawmy w następującej postaci:

$$b^{(p)}(t, u) = \sum_{l \in \mathbb{Z}} K_l(u) \cos l\omega_0 \frac{u}{2} \exp(il\omega_0 t), \quad (92)$$

$$b^{(n)}(t, u) = i \sum_{l \in \mathbb{Z}} K_l(u) \sin l\omega_0 \frac{u}{2} \exp(il\omega_0 t). \quad (93)$$

Na ich podstawie łatwo uzyskujemy

$$\operatorname{Re} f(\omega, t) = \frac{1}{2} \sum_{l \in \mathbb{Z}} \left[ g_l \left( \omega + l \frac{\omega_0}{2} \right) + g_l \left( \omega - l \frac{\omega_0}{2} \right) \right] \exp(il\omega_0 t), \quad (94)$$

$$\operatorname{Im} f(\omega, t) = \frac{i}{2} \sum_{l \in \mathbb{Z}} \left[ g_l \left( \omega - l \frac{\omega_0}{2} \right) - g_l \left( \omega + l \frac{\omega_0}{2} \right) \right] \exp(il\omega_0 t), \quad (95)$$

które mogą być przepisane w postaci

$$\begin{aligned} \operatorname{Re} f(\omega, t) = g_0(\omega) + \sum_{l \in \mathbb{Z}} \left[ \left[ g_l^c \left( \omega + l \frac{\omega_0}{2} \right) + g_l^c \left( \omega - l \frac{\omega_0}{2} \right) \right] \cos l\omega_0 t + \right. \\ \left. + \left[ g_l^s \left( \omega + l \frac{\omega_0}{2} \right) + g_l^s \left( \omega - l \frac{\omega_0}{2} \right) \right] \sin l\omega_0 t \right], \end{aligned} \quad (96)$$

$$\begin{aligned} \operatorname{Im} f(\omega, t) = \sum_{l \in \mathbb{N}} \left[ \left[ g_l^s \left( \omega - l \frac{\omega_0}{2} \right) - g_l^s \left( \omega + l \frac{\omega_0}{2} \right) \right] \cos l\omega_0 t + \right. \\ \left. + \left[ g_l^c \left( \omega + l \frac{\omega_0}{2} \right) - g_l^c \left( \omega - l \frac{\omega_0}{2} \right) \right] \sin l\omega_0 t \right]. \end{aligned} \quad (97)$$

Przedstawienia (82)-(83) i (96)-(97) oczywiście są identyczne. Od jednej postaci funkcji  $\operatorname{Re} f(\omega, t)$  i  $\operatorname{Im} f(\omega, t)$  łatwo przechodzimy do drugiej stosując zależności między składnikami komponentów obu typów:

$$g_l^s \left( \omega - l \frac{\omega_0}{2} \right) = \operatorname{Im} f_l(\omega), \quad g_l^s \left( \omega + l \frac{\omega_0}{2} \right) = \operatorname{Im} f_l(\omega + l\omega_0), \quad (98)$$



$$g_l^p\left(\omega + l\frac{\omega_0}{2}\right) = \operatorname{Re} f_l(\omega + l\omega_0), \quad g_l^c\left(\omega - l\frac{\omega_0}{2}\right) = \operatorname{Re} f_l(\omega). \quad (99)$$

Komponenty widmowe  $g_l^c(\omega)$  i  $g_l^s(\omega)$  są kosinusowymi i sinusowymi współczynnikami Fouriera gęstości widmowej  $g(\omega, t)$ , określonej przez funkcję  $k(t, u)$  równaniem

$$g(\omega, t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} k(t, u) e^{-i\omega u} du = \frac{1}{\pi} \int_0^{\infty} k(t, u) \cos \omega u du. \quad (100)$$

To znaczy, że

$$g(\omega, t) = g_0(\omega) + 2 \sum_{l \in N} [g_l^c(\omega) \cos l\omega_0 t + g_l^s(\omega) \sin l\omega_0 t], \quad (101)$$

i

$$g_0(\omega) = \frac{1}{T} \int_0^T g(\omega, t) dt, \quad g_l^c(\omega) = \frac{1}{T} \int_0^T g(\omega, t) \cos l\omega_0 t dt, \quad (102)$$

$$g_l^s(\omega) = \frac{1}{T} \int_0^T g(\omega, t) \sin l\omega_0 t dt. \quad (103)$$

Gęstość widmowa  $g(\omega, t)$  w odróżnieniu od  $f(\omega, t)$  jest funkcją rzeczywistą i parzystą.

Znajdźmy teraz związki między rzeczywistą  $\operatorname{Re} f_k(\omega)$  i urojoną  $\operatorname{Im} f_k(\omega)$  częścią komponentów widmowych oraz kosinusowymi i sinusowymi komponentami korelacyjnymi. Przekształcenie wzoru (58) daje

$$f_0(\omega) = \frac{1}{2\pi} \int_{-\infty}^{\infty} B_0(u) \cos \omega u du, \quad (104)$$

$$\operatorname{Re} f_{\pm k}(\omega) = \frac{1}{4\pi} \int_{-\infty}^{\infty} [B_k^c(u) \cos \omega u \mp B_k^s(u) \sin \omega u] du, \quad (105)$$

$$\operatorname{Im} f_{\pm k}(\omega) = \frac{1}{4\pi} \int_{-\infty}^{\infty} [B_k^c(u) \sin \omega u \pm B_k^s(u) \cos \omega u] du, \quad (16)$$

a po uwzględnieniu warunków symetryczności komponentów korelacyjnych (25)-(26) mamy

$$f_0(\omega) = \frac{1}{\pi} \int_0^{\infty} B_0(u) \cos \omega u du, \quad (107)$$

$$\begin{aligned} \operatorname{Re} f_k(\omega) &= \frac{1}{4\pi} \int_0^{\infty} [B_k^c(u) [\cos \omega u + \cos(\omega - k\omega_0)u] - \\ &\quad - B_k^s(u) [\sin \omega u - \sin(\omega - k\omega_0)u]] du, \end{aligned} \quad (108)$$

$$\begin{aligned} \operatorname{Im} f_k(\omega) = & \frac{1}{4\pi} \int_0^{\infty} [B_k^s(u) [\cos \omega u + \cos(\omega - k\omega_0)u] + \\ & + B_k^c(u) [\sin \omega u - \sin(\omega - k\omega_0)u]] du, \end{aligned} \quad (109)$$

Odwrotne przekształcenia mają postać

$$B_0(u) = \int_{-\infty}^{\infty} f_0(\omega) \cos \omega u du, \quad (110)$$

$$B_k^c(u) = 2 \int_{-\infty}^{\infty} [\operatorname{Re} f_k(\omega) \cos \omega u + \operatorname{Im} f_k(\omega) \sin \omega u] d\omega, \quad (111)$$

$$B_k^s(u) = 2 \int_{-\infty}^{\infty} [\operatorname{Im} f_k(\omega) \cos \omega u - \operatorname{Re} f_k(\omega) \sin \omega u] d\omega, \quad (112)$$

albo

$$B_0(u) = 2 \int_0^{\infty} f_0(\omega) \cos \omega u d\omega, \quad (113)$$

$$\begin{aligned} B_k^c(u) = & 2 \int_0^{\infty} [[\operatorname{Re} f_k(\omega) + \operatorname{Re} f_k(\omega + k\omega_0)] \cos \omega u + \\ & + [\operatorname{Im} f_k(\omega) - \operatorname{Im} f_k(\omega + k\omega_0)] \sin \omega u] d\omega, \end{aligned} \quad (114)$$

$$\begin{aligned} B_k^s(u) = & 2 \int_0^{\infty} [[\operatorname{Im} f_k(\omega) + \operatorname{Im} f_k(\omega + k\omega_0)] \cos \omega u + \\ & + [\operatorname{Re} f_k(\omega + k\omega_0) - \operatorname{Re} f_k(\omega)] \sin \omega u] d\omega. \end{aligned} \quad (115)$$

Korzystając ze wzorów (108)-(109) łatwo można wykazać słuszność następujących równości

$$\begin{aligned} \operatorname{Re} f_l(\omega) + \operatorname{Re} f_l(-\omega) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} B_l^c(u) \cos \omega u du, \\ \operatorname{Im} f_l(\omega) + \operatorname{Im} f_l(-\omega) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} B_l^s(u) \cos \omega u du, \\ \operatorname{Re} f_l(\omega) - \operatorname{Re} f_l(-\omega) &= -\frac{1}{2\pi} \int_{-\infty}^{\infty} B_l^s(u) \sin \omega u du, \\ \operatorname{Im} f_l(\omega) - \operatorname{Im} f_l(-\omega) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} B_l^c(u) \sin \omega u du. \end{aligned} \quad (116)$$

Wówczas rozwinięcie rzeczywistej i urojonej części chwilowej gęstości widmowej  $f(\omega, t)$  przyjmuje postać

$$\operatorname{Re} f(\omega, t) = \frac{1}{\pi} \int_0^{\infty} B_0(u) \cos \omega u \, du + \\ + \sum_{k \in N} \left[ \cos k \omega_0 t \left[ \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k^c(u) \cos \omega u \, du \right] + \sin k \omega_0 t \left[ \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k^s(u) \cos \omega u \, du \right] \right], \quad (117)$$

$$\operatorname{Im} f(\omega, t) = \sum_{k \in N} \left[ \cos k \omega_0 t \left[ \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k^c(u) \sin \omega u \, du \right] + \right. \\ \left. + \sin k \omega_0 t \left[ \frac{1}{2\pi} \int_{-\infty}^{\infty} B_k^s(u) \sin \omega u \, du \right] \right]. \quad (118)$$

Jak widać współczynniki Fouriera części rzeczywistej chwilowej gęstości widmowej określa się tylko przez przekształcenie kosinusowe komponentów korelacyjnych, a urojonej tylko przez przekształcenie sinusowe. Jeżeli komponenty korelacyjne przedstawimy jako sumy parzystych i nieparzystych części

$$B_k^c(u) = P_k^c(u) + N_k^c(u), \quad B_k^s(u) = P_k^s(u) + N_k^s(u), \quad (119)$$

to otrzymamy

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} B_k^{c,s}(u) \cos \omega u \, du = \frac{1}{\pi} \int_0^{\infty} P_k^{c,s}(u) \cos \omega u \, du, \quad (120)$$

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} B_k^{c,s}(u) \sin \omega u \, du = \frac{1}{\pi} \int_0^{\infty} N_k^{c,s}(u) \sin \omega u \, du, \quad (121)$$

Stąd wynika, że część rzeczywista chwilowej gęstości widmowej zależy tylko od parzystych części komponentów korelacyjnych, a urojona — tylko od nieparzystych części tych komponentów. Przy tym kosinusowe i sinusowe składniki każdej charakterystyki są wyrażone za pośrednictwem, odpowiednio, tylko kosinusowych i sinusowych komponentów korelacyjnych.

Zauważmy, że urojona i rzeczywista część komponentów widmowych też bezpośrednio może być określona przez kosinusowe i sinusowe składowe parzystych komponentów korelacyjnych. uwzględniając równania (34)-(35) ze wzorów (108) i (109) znajdujemy następujące zależności

$$\operatorname{Re} f_l(\omega) = \frac{1}{2\pi} \int_0^{\infty} K_l^c(u) \cos \left( \omega - l \frac{\omega_0}{2} \right) u \, du, \quad (122)$$

$$\operatorname{Im} f_l(\omega) = \frac{1}{2\pi} \int_0^{\infty} K_l^s(u) \cos \left( \omega - l \frac{\omega_0}{2} \right) u \, du. \quad (123)$$

wówczas

$$\operatorname{Re} f_l(\omega) + \operatorname{Re} f_l(-\omega) = \frac{1}{\pi} \int_0^{\infty} K_l^c(u) \cos l\omega_0 \frac{u}{2} \cos \omega u \, du, \quad (124)$$

$$\operatorname{Im} f_l(\omega) + \operatorname{Im} f_l(-\omega) = \frac{1}{\pi} \int_0^{\infty} K_l^s(u) \cos l\omega_0 \frac{u}{2} \cos \omega u \, du, \quad (125)$$

$$\operatorname{Re} f_l(\omega) - \operatorname{Re} f_l(-\omega) = \frac{1}{\pi} \int_0^{\infty} K_l^c(u) \sin l\omega_0 \frac{u}{2} \sin \omega u \, du, \quad (126)$$

$$\operatorname{Im} f_l(\omega) - \operatorname{Im} f_l(-\omega) = \frac{1}{\pi} \int_0^{\infty} K_l^s(u) \sin l\omega_0 \frac{u}{2} \sin \omega u \, du. \quad (127)$$

Po podstawieniu tych wyrażeń do szeregów (82) i (83) otrzymujemy takie związki:

$$\begin{aligned} \operatorname{Re} f(\omega, t) = & \frac{1}{\pi} \int_0^{\infty} B_0(u) \cos \omega u \, du + \frac{1}{\pi} \sum_{l \in N} \left[ \cos l\omega_0 t \left[ \int_{-\infty}^{\infty} K_l^c(u) \cos l\omega_0 \frac{u}{2} \cos \omega u \, du \right] + \right. \\ & \left. + \sin l\omega_0 t \left[ \int_{-\infty}^{\infty} K_l^s(u) \cos l\omega_0 \frac{u}{2} \cos \omega u \, du \right] \right] \end{aligned} \quad (128)$$

$$\begin{aligned} \operatorname{Im} f(\omega, t) = & \frac{1}{\pi} \sum_{l \in N} \left[ \cos l\omega_0 t \left[ \int_{-\infty}^{\infty} K_l^c(u) \sin l\omega_0 \frac{u}{2} \sin \omega u \, du \right] - \right. \\ & \left. - \sin l\omega_0 t \left[ \int_{-\infty}^{\infty} K_l^s(u) \sin l\omega_0 \frac{u}{2} \sin \omega u \, du \right] \right]. \end{aligned} \quad (129)$$

Przekształcenie Fouriera parzystej funkcji autokorelacji ma postać

$$k(t, u) = 2 \int_0^{\infty} g(\omega, t) \cos \omega u \, d\omega. \quad (130)$$

Stąd wynika, że wariancję procesu można określać przez całkowanie, w całym zakresie częstotliwości, gęstości widmowej  $g(\omega, t)$ :

$$k(t, 0) = 2 \int_0^{\infty} g(\omega, t) \, d\omega. \quad (131)$$

Jednak funkcja  $g(\omega, t)$  jak i funkcja  $\operatorname{Re} f(\omega, t)$  nie może być interpretowana jako gęstość widmowa mocy gdyż nie dla wszystkich  $(\omega, t)$  jest spełniony warunek  $g(\omega, t) \geq 0$ .

Oczywiste jest, że  $k(t, 0) = b(t, 0)$  i

$$\int_0^{\infty} g(\omega, t) \, d\omega = \int_0^{\infty} f(\omega, t) \, d\omega. \quad (132)$$

Z równości tych całek jednak nie wynika równość samych funkcji  $g(\omega, t)$  i  $\operatorname{Re} f(\omega, t)$ . Zgodnie z rozwinięciem (96) i (101) w przypadku ogólnym mają one

różne współczynniki Fouriera. Współczynniki te, a co za tym idzie gęstości widmowe, są równe tylko dla lokalnie stacjonarnych procesów losowych, gdyż

$$K_l^s(u) = B_l^s(u), \quad K_l^c(u) = B_l^c(u), \quad \operatorname{Re} f_l(\omega) = g_l^c(\omega), \quad \operatorname{Im} f_l(\omega) = g_l^s(\omega).$$

Warto zauważyć, że chociaż komponenty korelacyjne  $K_l^c(u)$ ,  $K_l^s(u)$  i komponenty widmowe  $g_l^s(\omega)$  i  $g_l^c(\omega)$  są odpowiednio współczynnikami Fouriera funkcji  $k(t, u)$  i  $g(\omega, t)$ , które opisują nieparzyste powiązania korelacyjne i procesy przejściowe, a które mogą być wykorzystane do skonstruowania potrzebnych charakterystyk. Na podstawie parzystych komponentów korelacyjnych może być określona parzysta i nieparzysta część funkcji autokorelacji  $b(t, u)$  (równania (45) i (46)), a także rzeczywista i urojona część gęstości widmowej  $f(\omega, t)$  (równania (128) i (129)). Przy zastosowaniu wzorów (96) i (97) i znanych komponentów widmowych  $g_l^c(\omega)$  i  $g_l^s(\omega)$  mogą być wyliczone funkcje  $\operatorname{Re} f(\omega, t)$  i  $\operatorname{Im} f(\omega, t)$ .

### 3. WYBRANE MODELE SYGNAŁÓW ZMODULOWANYCH

Podstawą wyprowadzenia wyrażeń dla charakterystyk probabilistycznych konkretnych typów sygnałów zmodulowanych są związki między tymi charakterystykami i charakterystykami stacjonarnych procesów modulujących. Komponenty korelacyjne sygnałów wyraża się za pośrednictwem funkcji auto- i interkorelacji tych procesów, komponenty widmowe — za pośrednictwem ich gęstości widmowych. Dobierając auto- i interkorelacje procesów w postaci znanych funkcji, zależnych od pewnych parametrów, można wyprowadzić równania, określające komponenty sygnału jako funkcji tych parametrów, a na ich podstawie — podobne równania dla funkcji autokorelacji, jak też rzeczywistej i urojonej części gęstości widmowej sygnału.

Modelem sygnału zmodulowanego amplitudowo jest proces multiplikatywny  $\xi(t) = \eta(t)f(t)$ , gdzie  $f(t)$  — okresowa funkcja czasowa, opisująca nośną a  $\eta(t)$  — rzeczywisty proces stacjonarny. Jeżeli nośna jest funkcją harmoniczną, to funkcja autokorelacji tego sygnału ma postać

$$b(t, u) = \frac{1}{2} R_\eta(u) \cos \omega_0 u + \frac{1}{2} R_\eta(u) \cos \omega_0 (2t + u), \quad (133)$$

gdzie  $R_\eta(u) = E \dot{\eta}(t) \dot{\eta}(t + u)$ ,  $\dot{\eta}(t) = \eta(t) - m_\eta$ ,  $m_\eta = E \eta(t)$  — wartość oczekiwana procesu  $\eta(t)$ .

Komponenty korelacyjne w danym przypadku określone są wzorami

$$B_0(u) = \frac{1}{2} R_\eta(u) \cos \omega_0 u, \quad B_1^c(u) = B_1^s(u) = 0 \quad (134)$$

$$B_2^c(u) = \frac{1}{2} R_\eta(u) \cos \omega_0 u, \quad B_2^s(u) = -\frac{1}{2} R_\eta(u) \sin \omega_0 u, \quad (135)$$

$$B_k^c(u) = B_k^s(u) = 0 \quad \text{dla } k > 2. \quad (136)$$

W oparciu o równania, określające część parzystą i nieparzystą funkcji autokorelacji

$$b^{(n)}(t, u) = \frac{1}{2} R_{\eta}(u) \cos \omega_0 u (1 + \cos 2\omega_0 t), \quad (137)$$

$$b^{(n)}(t, u) = -\frac{1}{2} R_{\eta}(u) \sin \omega_0 u \sin 2\omega_0 t, \quad (138)$$

znajdujemy część rzeczywistą i urojoną gęstości widmowej

$$\operatorname{Re} f(\omega, t) = \frac{1}{2} (1 + \cos 2\omega_0 t) [J(\omega + \omega_0) + J(\omega - \omega_0)], \quad (139)$$

$$\operatorname{Im} f(\omega, t) = \frac{1}{2} \sin 2\omega_0 t [J(\omega + \omega_0) - J(\omega - \omega_0)], \quad (140)$$

gdzie

$$J(\omega) = \frac{1}{2\pi} \int_0^{\infty} R_{\eta}(u) \cos \omega u du. \quad (141)$$

Jeżeli funkcja autokorelacji  $R_{\eta}(u)$  ma postać  $R_{\eta}(u) = De^{-\alpha|u|}$ , to

$$J(\omega) = \frac{D\alpha}{2\pi(\alpha^2 + \omega^2)}. \quad (142)$$

Wówczas zerowy komponent widmowy jest określony następująco:

$$f_0(\omega) = \frac{1}{2} [J(\omega + \omega_0) + J(\omega - \omega_0)], \quad (143)$$

Dla rzeczywistej i urojonej części drugiego komponentu mamy

$$\operatorname{Re} f_2(\omega) = \frac{1}{4\pi} \int_{-\infty}^{\infty} [B_2^c(u) \cos \omega u - B_2^s(u) \sin \omega u] du \quad (144)$$

$$\operatorname{Im} f_2(\omega) = \frac{1}{4\pi} \int_{-\infty}^{\infty} [B_2^c(u) \sin \omega u + B_2^s(u) \cos \omega u] du. \quad (145)$$

Jak widać ze wzorów (134)-(135), drugi komponent kosinusowy jest parzystą funkcją przesunięcia, a sinusowy — nieparzystą funkcją. Wtedy funkcja pod całką, wyrażenia (144), jest parzysta względem przesunięcia, zaś funkcja pod całką wyrażenia (145) — jest nieparzysta względem przesunięcia. Stąd wynika, że

$$\operatorname{Re} f_2(\omega) = \frac{1}{2\pi} \int_0^{\infty} [B_2^c(u) \cos \omega u - B_2^s(u) \sin \omega u] du, \quad (146)$$

$$\operatorname{Im} f_2(\omega) = 0. \quad (147)$$

Po podstawieniu wzorów (134)-(135) do równania (145) uzyskujemy

$$\operatorname{Re} f_2(\omega) = \frac{1}{2\pi} \int_0^{\infty} R_{\eta}(u) \cos(\omega - \omega_0) u \, du = \frac{1}{2} J(\omega - \omega_0). \quad (148)$$

Rzeczywista część gęstości widmowej, opisana wzorem (139), jest nieujemna. Funkcje  $b^{(n)}(t, u)$  i  $\operatorname{Im} f(\omega, t)$  zawierają tylko drugie harmoniczne składowe sinusowe i przy  $u = 0$  i  $\omega = 0$  odpowiednio są zerowymi. Wartości amplitud tych harmonicznych zależą od stosunku między wielkościami  $\alpha$  i  $\omega_0$ . Pierwsza z nich  $\alpha$  określa szybkość zanikania powiązań korelacyjnych ze wzrostem przesunięcia, a druga  $\omega_0$  szybkość zmian tych powiązań w czasie  $t$ . Amplitudy są duże przy małych  $\frac{\alpha}{\omega_0}$ , a ze wzrostem  $\frac{\alpha}{\omega_0}$  maleją. Dla  $\frac{\alpha}{\omega_0} \gg 1$   $b^{(n)}(t, u) \cong 0$ ,  $\operatorname{Im} f(\omega, t) \cong 0$ . Sygnał wówczas będzie lokalnie stacjonarny. W takim wypadku można założyć, że  $\sin \omega_0 u_k \cong 0$ ,  $\cos \omega_0 u_k \cong 1$ , gdzie  $u_k$  — przedział korelacji procesu  $\xi(t)$ , skąd wynikają następujące zależności

$$B_0(u) = \frac{1}{2} R I_{\eta}(u), \quad B_2^c(u) = \frac{1}{2} R_{\eta}(u), \quad B_2^s(u) = 0, \quad (149)$$

$$f_0(\omega) = J(\omega), \quad \operatorname{Re} f_2(\omega) = \frac{1}{2} J(\omega). \quad (150)$$

Funkcja korelacji takiego sygnału

$$b(t, u) = \frac{1}{2} R_{\eta}(u) (1 + \cos 2\omega_0 t) \quad (151)$$

i gęstość widmowa

$$f(\omega, t) = \frac{1}{2} J(\omega) (1 + \cos 2\omega_0 t) \quad (152)$$

dla dowolnych  $t$  mają wszystkie właściwości procesu stacjonarnego. Funkcja korelacji  $b(t, u)$  jest parzystą funkcją przesunięcia i spełnia nierówność (55).

Gęstość widmowa jest funkcją rzeczywistą, parzystą i nieujemną:

$$f^*(\omega, t) = f(\omega, t), \quad f(-\omega, t) = f(\omega, t), \quad f(\omega, t) \geq 0. \quad (153)$$

To znaczy, że może być ona traktowana jako gęstość widmowa mocy. Tak więc, sygnał zmodulowany amplitudowo w tym przypadku w każdej chwili  $t$  może być rozważony jako proces stacjonarny.

Dla funkcji autokorelacji  $k(t, u)$  mamy

$$k(t, u) = \frac{1}{2} R_{\eta}(u) \cos \omega_0 u + \frac{1}{2} R_{\eta}(u) \cos 2\omega_0 t, \quad (154)$$

wówczas

$$g(\omega, t) = \frac{1}{2}[J(\omega + \omega_0) + J(\omega - \omega_0)] + J(\omega) \cos 2\omega_0 t. \quad (155)$$

Gęstość widmowa  $g(\omega, t)$ , jak widać, może przybierać wartości ujemne. Komponenty korelacyjne i widmowe są zdefiniowane wzorami

$$K_0(u) = \frac{1}{2}R_\eta(u) \cos \omega_0 u, \quad K_2^c(u) = \frac{1}{2}R_\eta(u), \quad K_2^s(u) = 0, \quad (156)$$

$$g_0(\omega) = \frac{1}{2}[J(\omega + \omega_0) + J(\omega - \omega_0)], \quad g_2^c(\omega) = \frac{1}{2}J(\omega), \quad g_2^s(\omega) = 0.$$

Komponenty zerowe  $B_0(u)$ ,  $f_0(\omega)$  i  $K_0(u)$ ,  $g_0(\omega)$  są jednakowe, zaś drugie komponenty korelacyjne  $K_2^c(u)$ ,  $K_2^s(u)$  i widmowe  $g_2^c(\omega)$  i  $g_2^s(\omega)$  odpowiednio równają się  $B_2^c(u)$ ,  $B_2^s(u)$  oraz  $\text{Re}f_2(\omega)$ ,  $\text{Im}f_2(\omega)$  tylko przy spełnieniu warunków lokalnej stacjonarności.

W modelu multiplikatywnym założono, że własności wszystkich składowych harmonicznych, tworzących OKPL, zmieniają się jednakowo. Zmiany te mają charakter modulacji amplitudowej. Model, który rozważymy poniżej, opisuje już modulacje amplitudowe i fazowe i jest najprostszym modelem tego typu. Wzięto w nim pod uwagę tylko jedną harmoniczną nośnej, która została poddana zmianom modulacyjnym. Z wyrażenia (5) ten model otrzymujemy w przypadku, gdy nierównymi zeru są tylko procesy modulujące  $\xi_1(t)$  i  $\xi_{-1}(t)$ . Przedstawimy je w postaci

$$\xi_1(t) = \frac{1}{2}[\xi_c(t) - i\xi_s(t)], \quad \xi_{-1}(t) = \frac{1}{2}[\xi_c(t) + i\xi_s(t)]. \quad (157)$$

Wówczas mamy

$$\xi(t) = \xi_c(t) \cos \omega_0 t + \xi_s(t) \sin \omega_0 t. \quad (158)$$

Wartość oczekiwana tego procesu jest określona wzorem

$$m(t) = m_c \cos \omega_0 t + m_s \sin \omega_0 t.$$

gdzie  $m_c = E\xi_c(t)$ ,  $m_s = E\xi_s(t)$ . Funkcja autokorelacji zawiera zerowy

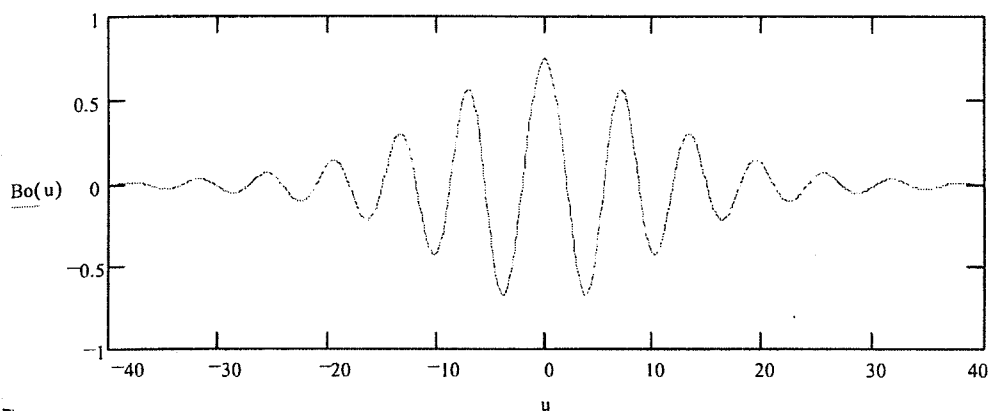
$$B_0(u) = \frac{1}{2}[R_c(u) + R_s(u)] \cos \omega_0 u + \frac{1}{2}[R_{cs}(u) - R_{sc}(u)] \sin \omega_0 u \quad (159)$$

i drugie komponenty korelacyjne

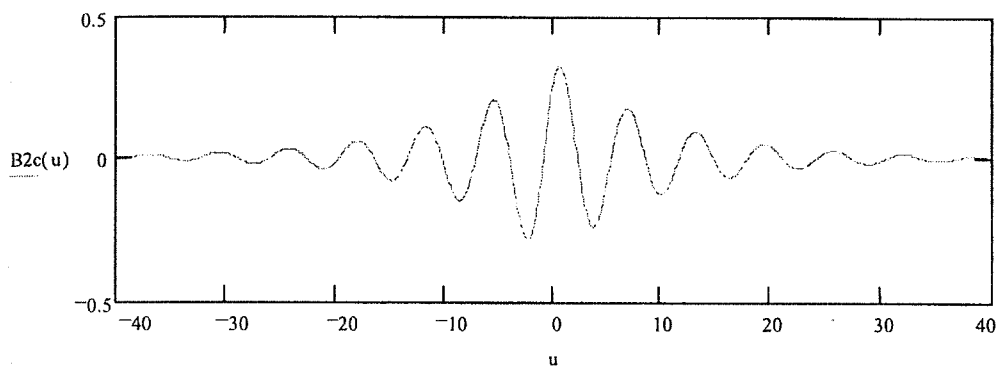
$$B_2^c(u) = \frac{1}{2}[R_c(u) - R_s(u)] \cos \omega_0 u + \frac{1}{2}[R_{cs}(u) + R_{sc}(u)] \sin \omega_0 u, \quad (160)$$

$$B_2^s(u) = \frac{1}{2}[R_{cs}(u) + R_{sc}(u)] \cos \omega_0 u + \frac{1}{2}[R_s(u) - R_c(u)] \sin \omega_0 u. \quad (161)$$

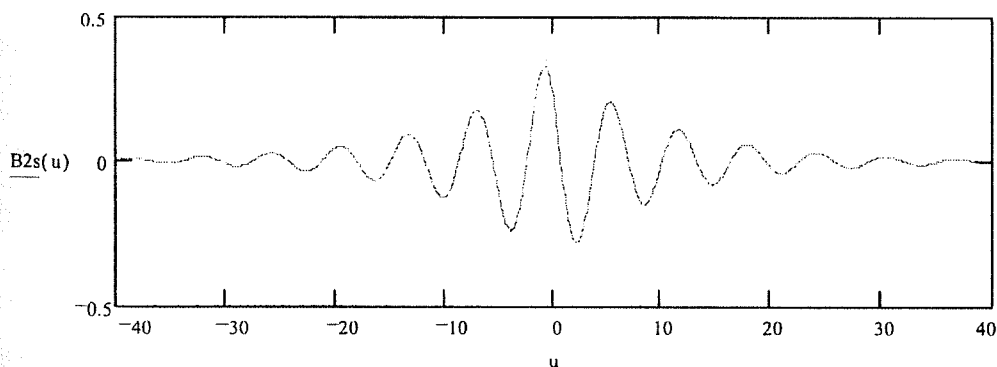




Rys. 1. Zerowy komponent korelacyjny  $B_0(u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$



Rys. 2. Kosinusowy drugi komponent korelacyjny  $B_2^c(u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$



Rys. 3. Sinusowy drugi komponent korelacyjny  $B_2^s(u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$

W wyrażeniach (159)-(160)  $R_c(u) = E \dot{\xi}_c(t) \dot{\xi}_c(t+u)$ ,  $R_s(u) = E \dot{\xi}_s(t) \dot{\xi}_s(t+u)$ ,  $R_{cs}(u) = E \dot{\xi}(t) \dot{\xi}(t+u)$ .

Utwórzmy funkcje

$$R_{cs}^{(p)}(u) = \frac{1}{2}[R_{cs}(u) + R_{cs}(-u)] = \frac{1}{2}[R_{cs}(u) + R_{sc}(u)], \quad (162)$$

$$R_{cs}^{(n)}(u) = \frac{1}{2}[R_{cs}(u) - R_{cs}(-u)] = \frac{1}{2}[R_{cs}(u) - R_{sc}(u)]. \quad (163)$$

Pierwsza z nich jest parzystą, a druga — nieparzystą funkcją przesunięcia. Ze wzorów wynika, że

$$R_{cs}(u) = R_{cs}^{(p)}(u) + R_{cs}^{(n)}(u), \quad R_{sc}(u) = R_{cs}^{(p)}(u) - R_{cs}^{(n)}(u). \quad (164)$$

Uwzględniając te zależności, wnioskujemy, że zerowy komponent korelacyjny, który jest funkcją autokorelacji przybliżenia stacjonarnego procesu (158), określa się sumą funkcji autokorelacji składowych kwadraturowych  $\xi_c(t)$  i  $\xi_s(t)$  oraz części nieparzystej ich korelacji wzajemnej. Drugie komponenty korelacyjne, określające właściwości niestacjonarności procesu, zależą tylko od części parzystej funkcji korelacji wzajemnej składowych kwadraturowych i różnicy funkcji autokorelacji. Tak więc, gdy  $R_c(u) = R_s(u)$  i  $R_{cs}^{(p)}(u) = 0$ , to  $B_2^c(u) = B_2^s(u) = 0$  i proces (158) jest stacjonarny.

Stosując wyrażenia (159)-(161) uzyskujemy następujące zależności dla komponentów widmowych

$$f_0(\omega) = \frac{1}{2\pi} \int_0^\infty [[R_c(u) + R_s(u)] \cos \omega_0 u + 2R_{cs}^{(n)}(u) \sin \omega_0 u] \cos \omega u du, \quad (165)$$

$$\operatorname{Re} f_2(\omega) = \frac{1}{4\pi} \int_0^\infty [R_c(u) - R_s(u)] \cos(\omega - \omega_0) u du, \quad (166)$$

$$\operatorname{Im} f_2(\omega) = \frac{1}{4\pi} \int_0^\infty R_{cs}^{(p)}(u) \cos(\omega - \omega_0) u du. \quad (167)$$

Każda charakterystyka widmowa, jak widać, opisuje inne własności składowych kwadraturowych — sygnałów modulujących. Należy zauważyć, że część rzeczywista drugiego komponentu widmowego określa się tylko różnicą autokorelacji  $R_c(u)$  i  $R_s(u)$ , a urojona — tylko częścią parzystą funkcji korelacji wzajemnej. Oczywiście własności sygnału, które są związane z drugimi komponentami widmowymi, w żaden sposób nie mogą być otrzymane przy analizie tylko zerowego komponentu, tj. w ramach przybliżenia stacjonarnego.

Wprowadźmy całki

$$J_{c,s}^{(1)}(\omega) = \frac{1}{2\pi} \int_0^\infty R_{c,s}(u) \cos \omega u du, \quad J_{cs}^{(1)}(\omega) = \frac{1}{2\pi} \int_0^\infty R_{cs}^{(p)}(u) \cos \omega u du, \quad (168)$$

$$J_{cs}^{(2)}(\omega) = \frac{1}{2\pi} \int_0^{\infty} R_{cs}^{(n)}(u) \sin \omega u \, du. \quad (169)$$

Wówczas mamy

$$f_0(\omega) = \frac{1}{2} [J_c^{(1)}(\omega + \omega_0) + J_c^{(1)}(\omega - \omega_0) - J_s^{(1)}(\omega + \omega_0) + J_s^{(1)}(\omega - \omega_0)] + \\ + J_{cs}^{(2)}(\omega + \omega_0) + J_{cs}^{(2)}(\omega - \omega_0), \quad (170)$$

$$\operatorname{Re} f_2(\omega) = \frac{1}{2} [J_c^{(1)}(\omega - \omega_0) - J_s^{(1)}(\omega - \omega_0)], \quad (171)$$

$$\operatorname{Im} f_2(\omega) = J_{cs}^{(1)}(\omega - \omega_0). \quad (172)$$

Jeżeli funkcje autokorelacyjne, parzystej i nieparzystej części funkcji korelacji wzajemnej mają postacie

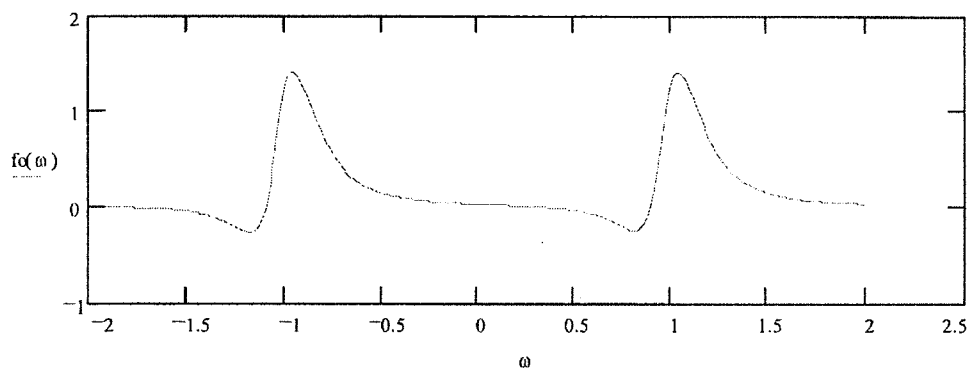
$$R_{c,s}(u) = D_{c,s} e^{-\alpha_{c,s}|u|}, \quad R_{cs}^{(p)}(u) = D_{cs}^{(p)} e^{-\alpha_{cs}^{(p)}|u|}, \quad R_{cs}^{(n)}(u) = D_{cs}^{(n)} u e^{-\alpha_{cs}^{(n)}|u|}, \quad (173)$$

$$J_{c,s}^{(1)}(\omega) = \frac{D_{c,s} \alpha_{c,s}}{2\pi(\alpha_{c,s}^2 + \omega^2)}, \quad J_{cs}^{(1)}(\omega) = \frac{D_{cs}^{(p)} \alpha_{cs}^{(p)}}{2\pi[(\alpha_{cs}^{(p)})^2 + \omega^2]}, \quad (174)$$

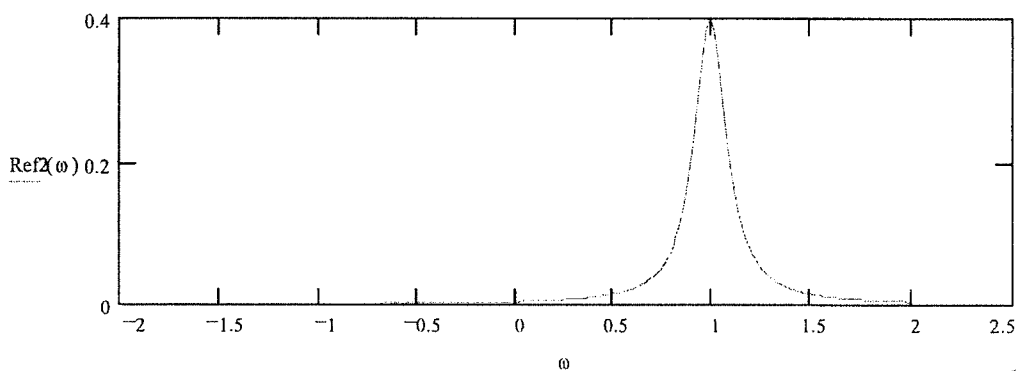
$$J_{cs}^{(2)}(\omega) = \frac{D_{cs}^{(n)} \alpha_{cs}^{(n)} \omega}{\pi[(\alpha_{cs}^{(n)})^2 + \omega^2]^2}.$$

Zerowy komponent widmowy (170) jest parzystą funkcją częstotliwości  $f_0(-\omega) = f_0(\omega)$ , a charakterystyki (171)-(172) spełniają równości  $\operatorname{Re} f_2(-\omega) = \operatorname{Re} f_2(\omega + 2\omega_0)$ ,  $\operatorname{Im} f_2(-\omega) = \operatorname{Im} f_2(\omega + 2\omega_0)$ . Rzeczywista i urojona część drugiego komponentu jest symetryczna względem punktu  $\omega = \omega_0$ , to znaczy, że funkcje  $\operatorname{Re} f_2(\omega + \omega_0)$  i  $\operatorname{Im} f_2(\omega + \omega_0)$  — parzyste funkcje częstotliwości. Jeżeli procesy losowe  $\xi_c(t)$  i  $\xi_s(t)$  zmieniają się w czasie bardzo wolno w porównaniu z oscylacjami o częstotliwości  $\omega_0$  ( $\alpha \ll \omega_0$ ), to zerowy komponent widmowy będzie istotnie wyróżniać się od zera tylko w wąskim pasmie wokół częstotliwości  $\pm \omega_0$ . Punkty jego maksimum będą zbiegać się w przybliżeniu z punktami  $\pm \omega_0$ . Zależność częstotliwościowa funkcji  $f_0(\omega)$  ma symetryczną i niesymetryczną składową względem punktów  $\pm \omega_0$ . Niesymetryczność kształtu jest uwarunkowana tylko parzystą częścią korelacji wzajemnej. Charakterystyki (171)-(172) w tym wypadku będą skoncentrowane tylko w dziedzinie dodatnich częstotliwości oraz będą mieć wierzchołki w punkcie  $\omega_0$ . OKPL wówczas będzie wąskopasmowym procesem. Maksima charakterystyk rozszerzają się wraz ze wzrostem szybkości zanikania powiązań korelacyjnych, tj. wartości wielkości  $\alpha$ , i przemieszczają się do punktu  $\omega = 0$ . Gdy  $\alpha \gg \omega_0$ , to

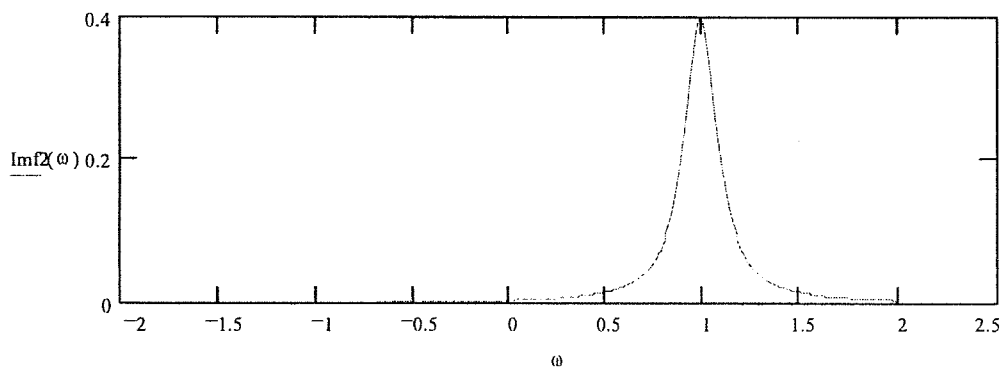
$$J_{c,s}^{(1)}(\omega \pm \omega_0) = J_{c,s}^{(1)}(\omega), \quad J_{cs}^{(1,2)}(\omega \pm \omega_0) = J_{cs}^{(1,2)}(\omega) \quad (175)$$



Rys. 4. Zerowy komponent widmowy  $f_0(\omega)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$



Rys. 5. Część rzeczywista drugiego komponentu widmowego  $\text{Re} f_2(\omega)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$



Rys. 6. Część urojona drugiego komponentu widmowego  $\text{Im} f_2(\omega)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$

i

$$B_0(u) = \frac{1}{2}[R_c(u) + R_s(u)], \quad B_2^e(u) = \frac{1}{2}[R_c(u) - R_s(u)], \quad B_2^s(u) = R_{cs}^{(p)}(u), \quad (176)$$

$$f_0(\omega) = J_c^{(1)}(\omega) + J_s^{(1)}(\omega) + 2J_{cs}^{(2)}(\omega), \quad \operatorname{Re} f_2(\omega) = \frac{1}{2}[J_c^{(1)}(\omega) - J_s^{(1)}(\omega)], \quad (177)$$

$$\operatorname{Im} f_2(\omega) = J_{cs}^{(1)}(\omega).$$

Parzysta i nieparzysta część funkcji autokorelacji procesu (158) mają postać

$$b^{(p)}(t, u) = \frac{1}{2}[R_c(u) + R_s(u)] \cos \omega_0 u + R_{cs}^{(n)}(u) \sin \omega_0 u + \\ + \frac{1}{2}[R_c(u) - R_s(u)] \cos \omega_0 u \cos 2\omega_0 t + R_{cs}^{(p)}(u) \cos \omega_0 u \sin 2\omega_0 t, \quad (178)$$

$$b^{(n)}(t, u) = R_{cs}^{(p)}(u) \sin \omega_0 u \cos 2\omega_0 t + \frac{1}{2}[R_s(u) - R_c(u)] \sin \omega_0 u \sin 2\omega_0 t. \quad (179)$$

Biorąc pod uwagę równania (168)-(169), (178)-(179) otrzymujemy następujące zależności

$$\operatorname{Re} f(\omega, t) = f_0(\omega) + \frac{1}{2}[J_c^{(1)}(\omega + \omega_0) + \\ + J_c^{(1)}(\omega - \omega_0) - J_s^{(1)}(\omega + \omega_0) - J_s^{(1)}(\omega - \omega_0)] \cos 2\omega_0 t + \\ + [J_{cs}^{(1)}(\omega + \omega_0) + J_{cs}^{(1)}(\omega - \omega_0)] \sin 2\omega_0 t, \quad (180)$$

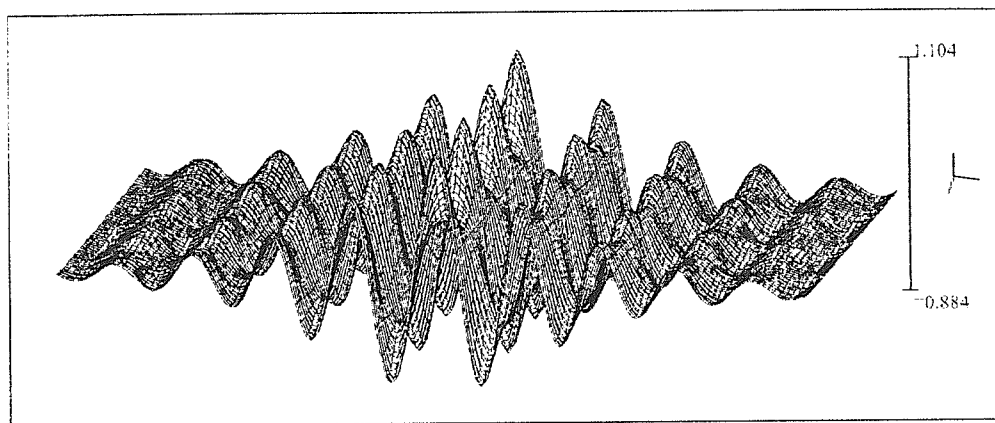
$$\operatorname{Im} f(\omega, t) = [J_{cs}^{(1)}(\omega - \omega_0) - J_{cs}^{(1)}(\omega + \omega_0)] \cos 2\omega_0 t + \\ + \frac{1}{2}[J_s^{(1)}(\omega - \omega_0) - J_s^{(1)}(\omega + \omega_0) - J_c^{(1)}(\omega - \omega_0) + J_c^{(1)}(\omega + \omega_0)] \sin 2\omega_0 t. \quad (181)$$

Część nieparzysta  $b^{(n)}(t, u)$  i jej sinusowe przekształcenie  $\operatorname{Im} f(\omega, t)$  określają cechy niestacjonarności sygnału. Jeżeli powiązania korelacyjne składowych kwadraturowych maleją szybko w porównaniu ze zmianami w czasie nośnej, to  $b^{(n)}(t, u) \approx 0$  i  $\operatorname{Im} f(\omega, t) \approx 0$ . W tym wypadku zmienia się postać parzystej części  $b^{(p)}(t, u)$  i odpowiednio części rzeczywistej chwilowej gęstości widmowej  $\operatorname{Re} f(\omega, t)$ :

$$b^{(p)}(t, u) = \frac{1}{2}[R_c(u) + R_s(u)] + \frac{1}{2}[R_c(u) - R_s(u)] \cos 2\omega_0 t + R_{cs}^{(p)}(u) \sin 2\omega_0 t, \quad (182)$$

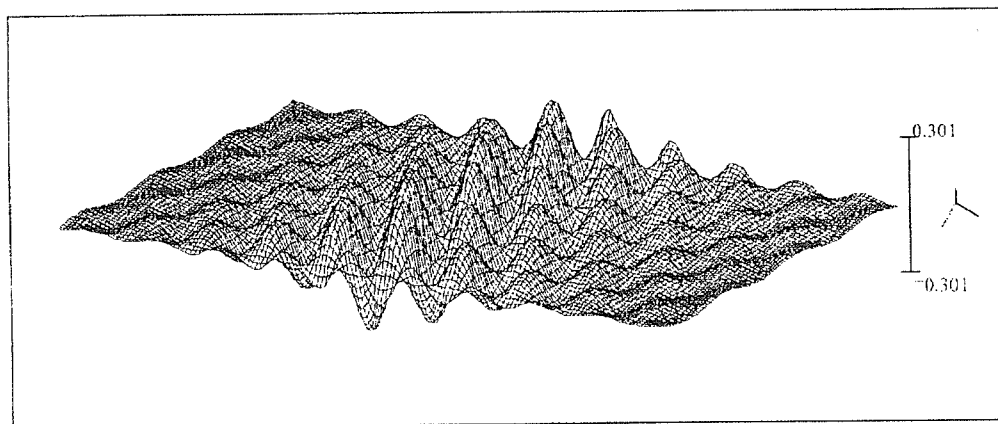
$$\operatorname{Re} f(\omega, t) = J_c^{(1)}(\omega) + J_s^{(1)}(\omega) + [J_c^{(1)}(\omega) - J_s^{(1)}(\omega)] \cos 2\omega_0 t + 2J_{cs}^{(1)}(\omega) \sin 2\omega_0 t \quad (183)$$

Funkcja  $\operatorname{Re} f(\omega, t)$  będzie nieujemną dla wszystkich  $(\omega, t)$  i może być traktowana jako chwilowa gęstość widmowa mocy takiego lokalnie stacjonarnego OKPL.



bp

Rys. 7. Parzysta część funkcji autokorelacji  $b^{(p)}(t, u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $u = -30 \div 30$



bn

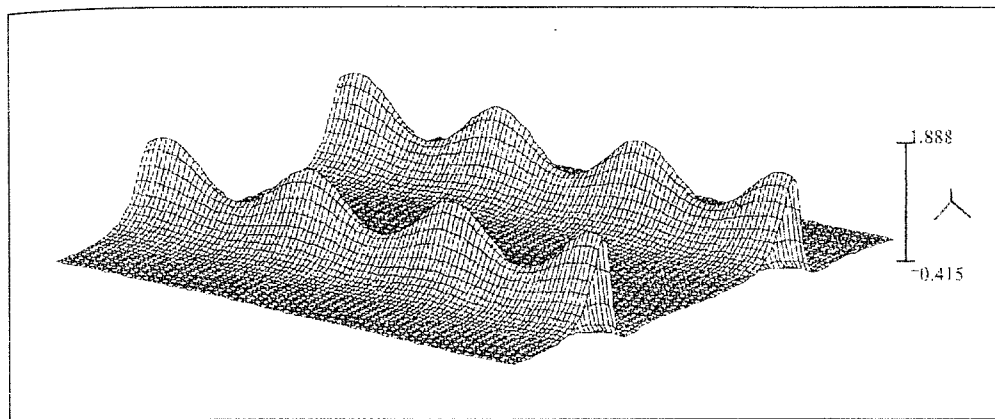
Rys. 8. Nieparzysta część funkcji autokorelacji  $b^{(n)}(t, u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $u = -30 \div 30$

Oczywiście, że proces losowy (158) w tym wypadku nie może być wąkopasmowy.

Parzysta funkcja autokorelacji  $k(t, u)$  i odpowiednia jej chwilowa gęstość widmowa  $g(\omega, t)$  są określone wzorami:

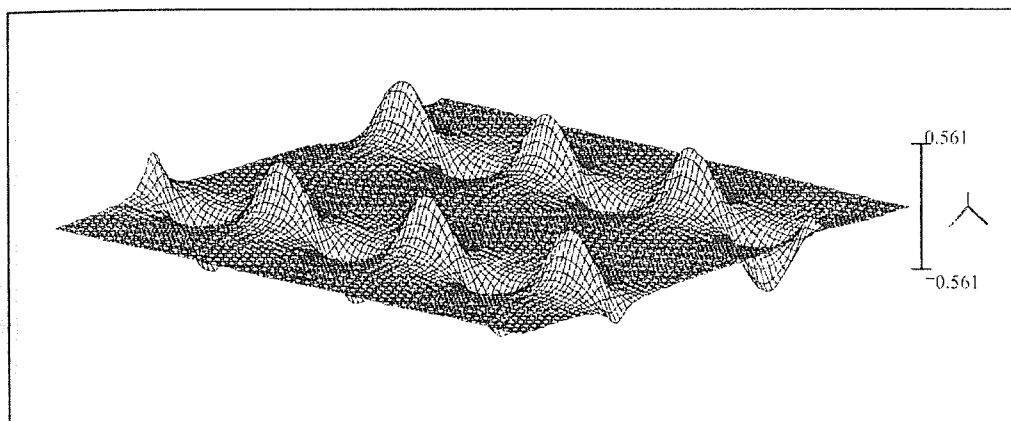
$$k(t, u) = K_0(u) + K_2^c(u) \cos 2\omega_0 t + K_2^s(u) \sin 2\omega_0 t, \quad (184)$$

$$g(\omega, t) = g_0(\omega) + g_2^c(\omega) \cos 2\omega_0 t + g_2^s(\omega) \sin 2\omega_0 t, \quad (185)$$



Ref

Rys. 9. Część rzeczywista chwilowej gęstości widmowej  $\text{Ref}(\omega, t)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(0)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(0)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $\omega = -2 \div 2$



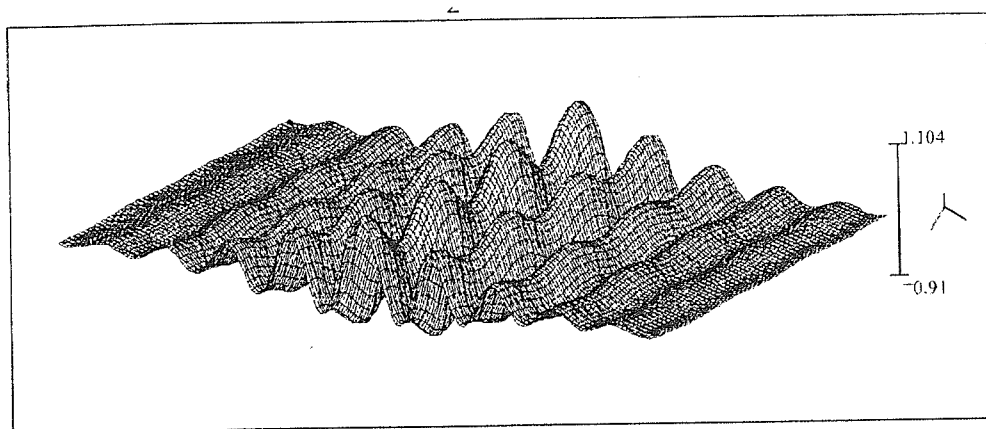
Imf

Rys. 10. Część urojona chwilowej gęstości widmowej  $\text{Imf}(\omega, t)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(0)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(0)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $\omega = -2 \div 2$

gdzie

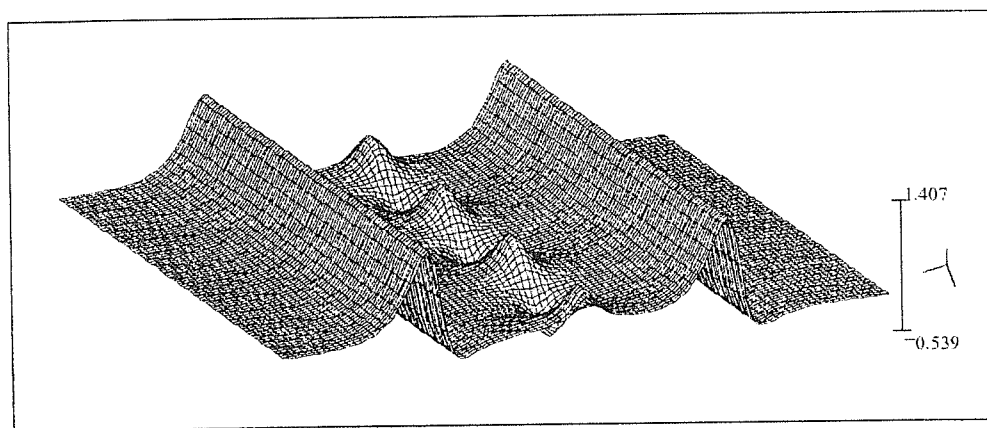
$$K_0(u) = B_0(u), \quad K_2^c(u) = \frac{1}{2}[R_c(u) - R_s(u)], \quad K_2^s(u) = R_{cs}^{(0)}(u), \quad (186)$$

$$g_0(\omega) = f_0(\omega), \quad g_2^c(\omega) = \frac{1}{2}[J_c^{(1)}(\omega) - J_s^{(1)}(\omega)], \quad g_2^s(\omega) = J_{cs}^{(1)}(\omega). \quad (187)$$



k

Rys. 11. Parzysta funkcja autokorelacji  $k(t, u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $u = -30 \div 30$



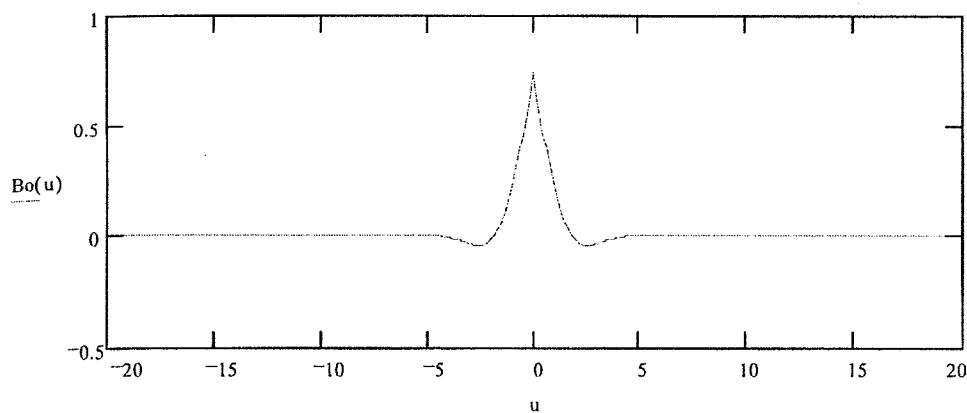
g

Rys. 12. Chwilowa gęstość widmowa  $g(\omega, t)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,25$ ,  $D_{cs}^{(n)} = 0,25$ ,  $\alpha_s = 0,1$ ,  $\alpha_c = 0,1$ ,  $\alpha_{cs}^{(p)} = 0,1$ ,  $\alpha_{cs}^{(n)} = 0,2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $\omega = -2 \div 2$

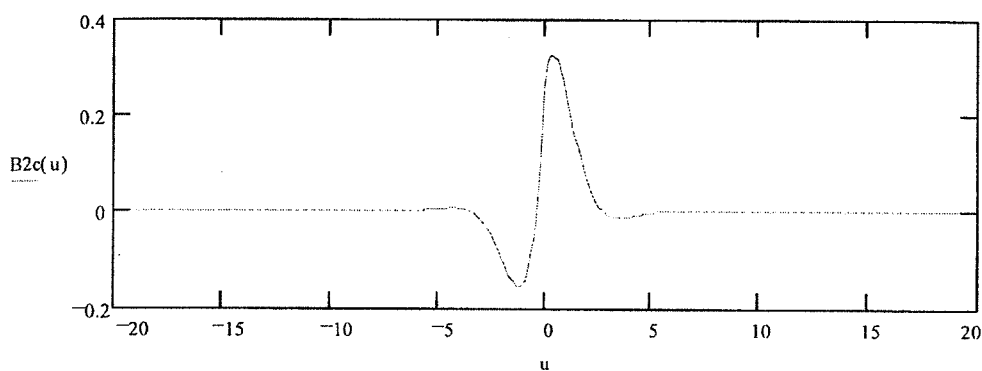
Drugie komponenty korelacyjne i widmowe mają postać charakterystyczną dla procesu lokalnie stacjonarnego. Nie zależą one od częstotliwości nośnej  $\omega_0$ . Właśnie to odróżnia je od charakterystyk sygnałów w przypadku gdy zanik auto- i interkorelacji nie jest taki szybki.

Poniżej przedstawiono przebiegi porównawcze określające właściwości charakterystyk sygnału (158) w przypadku szybkich zaników korelacji.

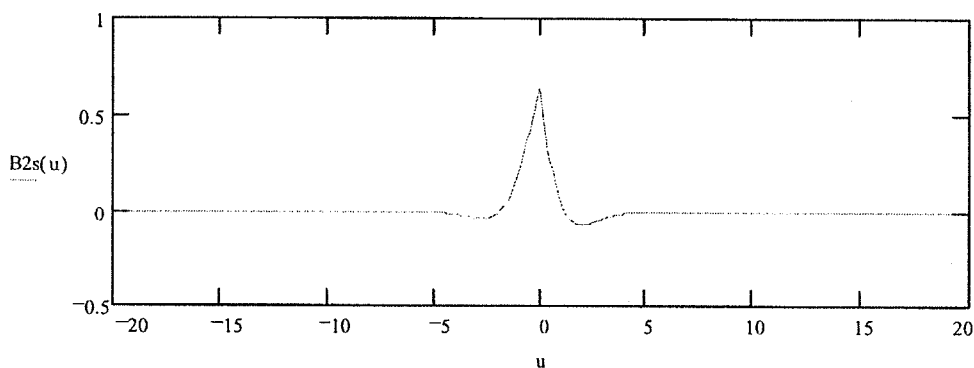




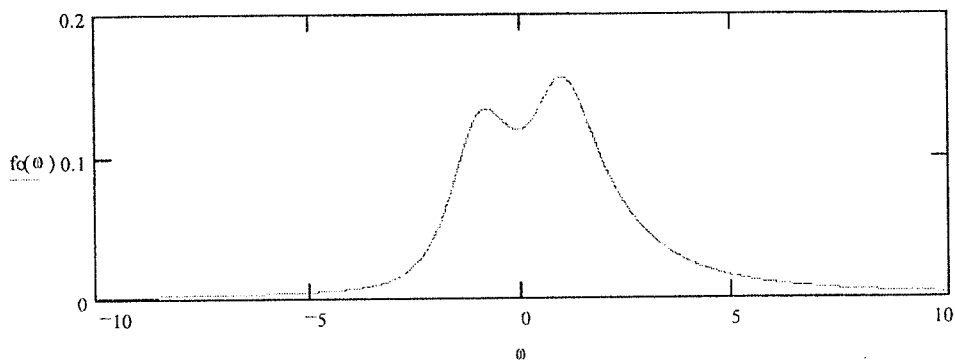
Rys. 13. Zerowy komponent korelacyjny  $B_0(u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(p)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,



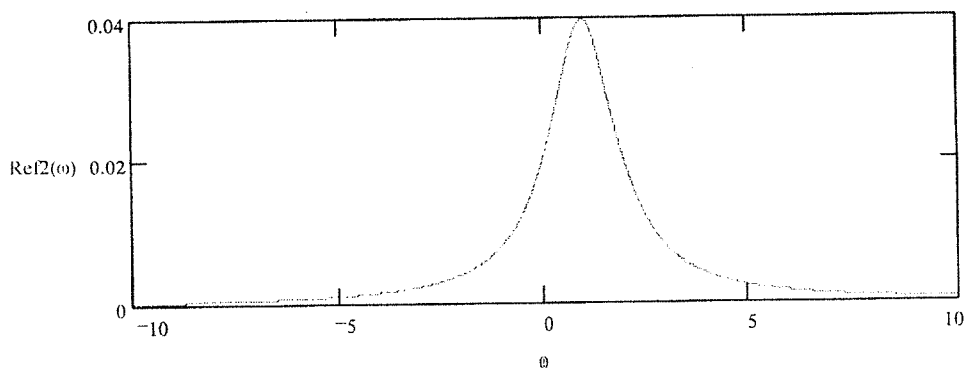
Rys. 14. Kosinusowy drugi komponent korelacyjny  $B_2^c(u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(p)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,



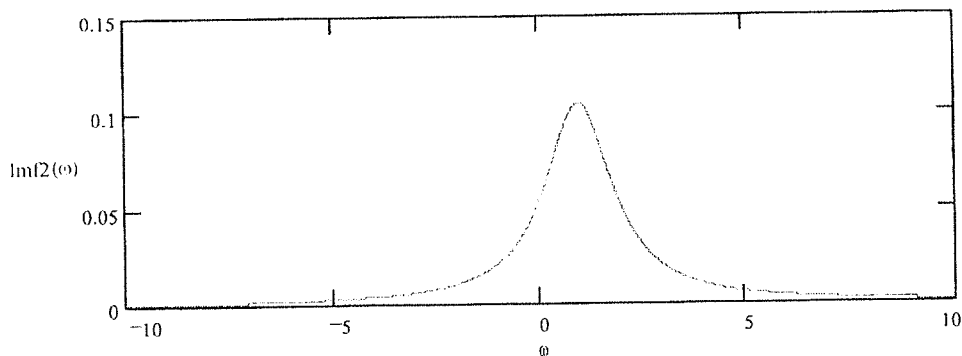
Rys. 15. Sinusowy drugi komponent korelacyjny  $B_2^s(u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,65$ ,  $D_{cs}^{(n)} = 0,5$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(p)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,



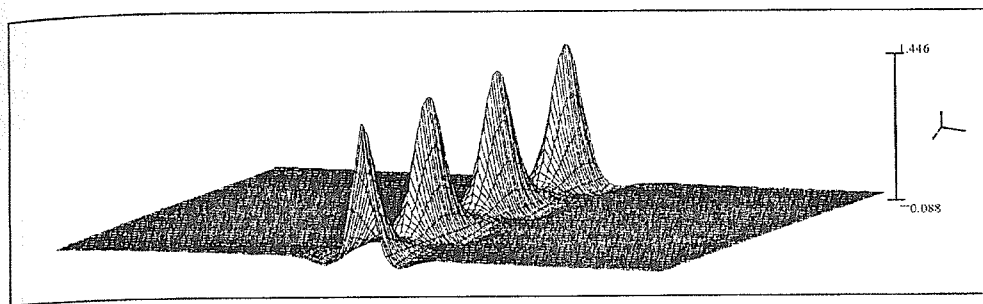
Rys. 16. Zerowy komponent widmowy  $f_0(\omega)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(n)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(n)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,  $\omega_0 = 1$



Rys. 17. Część rzeczywista drugiego komponentu widmowego  $\text{Re } f_2(\omega)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(n)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(n)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,  $\omega_0 = 1$

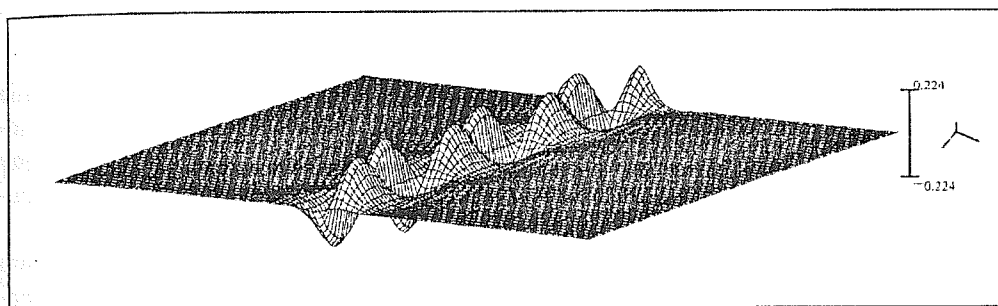


Rys. 18. Część urojona drugiego komponentu widmowego  $\text{Im } f_2(\omega)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(n)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(n)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,  $\omega_0 = 1$



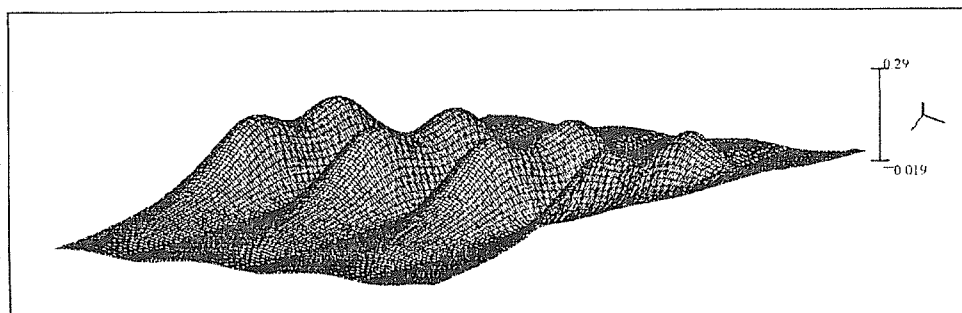
bp

Rys. 19. Parzysta część funkcji autokorelacji  $b^{(n)}(t, u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(p)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $u = -20 \div 20$



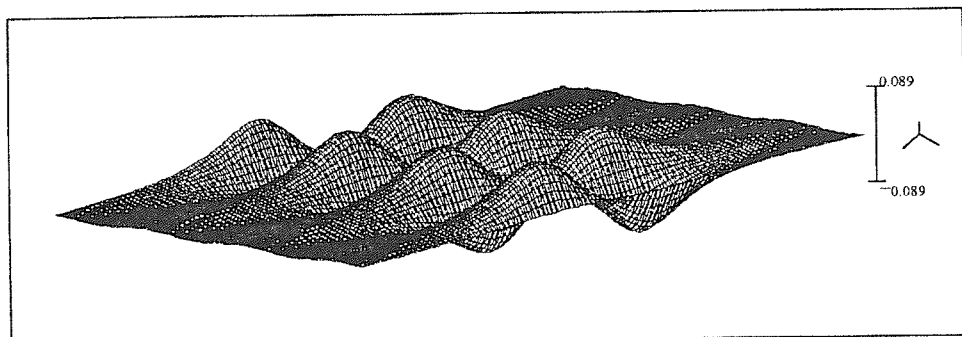
bn

Rys. 20. Nieparzysta część funkcji autokorelacji  $b^{(n)}(t, u)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(p)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $u = -20 \div 20$



Ref

Rys. 21. Część rzeczywista chwilowej gęstości widmowej  $\text{Re } f(\omega, t)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(p)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $\omega = -5 \div 5$



Imf

Rys. 22. Część urojona chwilowej gęstości widmowej  $\text{Im } f(\omega, t)$  dla  $D_c = 1$ ,  $D_s = 0,5$ ,  $D_{cs}^{(p)} = 0,65$ ,  $D_{cs}^{(n)} = 0,65$ ,  $\alpha_s = 1$ ,  $\alpha_c = 1$ ,  $\alpha_{cs}^{(p)} = 1$ ,  $\alpha_{cs}^{(n)} = 2$ ,  $\omega_0 = 1$ ,  $t = 0 \div 10$ ,  $\omega = -5 \div 5$

#### 4. ESTYMACJA CHARAKTERYSTYK

Okresowość charakterystyk probabilistycznych OKPL umożliwia stosowanie naturalnych metod ich szacowania — koherentnej i komponentnej. Metoda koherentna polega na uśrednieniu próbek procesów, odstęp między którymi jest równy okresowi korelowalności  $T$ , natomiast komponentna — na całkowitych przekształceniach realizacji procesów [1, 3].

Algorytmy obróbki statystycznej sygnałów typu OKPL mogą ulec uproszczeniu jeżeli w pierwszym etapie będą szacowane charakterystyki, które opisują ich własności z wykorzystaniem parzystej funkcji autokorelacji  $k(t, u)$ , a następnie na ich podstawie zostaną obliczone inne wartości. Koherentny estymator parzystej funkcji autokorelacji ma postać

$$\hat{k}(t, u) = \frac{1}{2N+1} \sum_{n=-N}^N \xi\left(t - \frac{u}{2} + nT\right) \xi\left(t + \frac{u}{2} + nT\right) - \hat{m}\left(t - \frac{u}{2}\right) \hat{m}\left(t + \frac{u}{2}\right), \quad (188)$$

gdzie

$$\hat{m}(t) = \frac{1}{2N+1} \sum_{n=-N}^N \xi(t + nT). \quad (189)$$

Estymator wartości oczekiwanej (189) jest nieobciążony, a jego wariancja opisana jest wzorem

$$D[\hat{m}(t)] = \frac{1}{2N+1} \sum_{n=-N}^N \left(1 - \frac{|n|}{2N+1}\right) k\left(t - \frac{u}{2}, nT\right). \quad (190)$$

Jeżeli funkcje autokorelacji sygnału maleją wraz ze wzrostem przesunięcia  $u$

$$\lim_{|u| \rightarrow \infty} k(t, u) = 0, \quad (191)$$

to  $D[\hat{m}(t)] \rightarrow 0$  przy  $N \rightarrow \infty$ , skąd wnioskujemy, że estymator (189) jest zgodny. Warunek (191) zapewnia też asymptotyczną nieobciążoność i zgodność estymatora funkcji autokorelacji. Obciążenie i wariancja tego estymatora w pierwszym przybliżeniu przy skończonych  $N$  są odpowiednio równe

$$\varepsilon[\hat{k}(t, u)] = -\frac{1}{2N+1} \sum_{n=-N}^N \left(1 - \frac{|n|}{2N+1}\right) k(t, u + nT), \quad (192)$$

$$D[\hat{k}(t, u)] = \frac{1}{2N+1} \sum_{n=-2N}^{2N} \left(1 - \frac{|n|}{2N+1}\right) [k(t, nT)k(t+u, nT) + k(t, u+T)k(t, u-nT)] \quad (193)$$

Na podstawie statystyk (188) i (190) mogą być utworzone komponenty Fouriera wartości oczekiwanej i funkcji autokorelacji:

$$\hat{m}(t) = \hat{m}_0 + \sum_{l=1}^{N_1} (\hat{m}_l^c \cos l\omega_0 t + \hat{m}_l^s \sin l\omega_0 t), \quad (194)$$

$$\hat{k}(t, u) = \hat{K}_0(u) + \sum_{l=1}^{N_2} [\hat{K}_l^c(u) \cos l\omega_0 t + \hat{K}_l^s(u) \sin l\omega_0 t], \quad (195)$$

przy tym

$$\hat{m}_0 = \frac{1}{\theta} \int_0^\theta \xi(t) dt, \quad \hat{m}_l^c = \frac{1}{\theta} \int_0^\theta \xi(t) \cos l\omega_0 t dt, \quad \hat{m}_l^s = \frac{1}{\theta} \int_0^\theta \xi(t) \sin l\omega_0 t dt, \quad (196)$$

$$\hat{K}_0(u) = \frac{1}{\theta} \int_0^\theta \xi^0\left(t - \frac{u}{2}\right) \xi^0\left(t + \frac{u}{2}\right) dt, \quad (197)$$

$$\hat{K}_l^c(u) = \frac{1}{\theta} \int_0^\theta \xi^0\left(t - \frac{u}{2}\right) \xi^0\left(t + \frac{u}{2}\right) \cos l\omega_0 t dt, \quad (198)$$

$$\hat{K}_l^s(u) = \frac{1}{\theta} \int_0^\theta \xi^0\left(t - \frac{u}{2}\right) \xi^0\left(t + \frac{u}{2}\right) \sin l\omega_0 t dt, \quad (199)$$

gdzie  $\xi^0(t) = \xi(t) - \hat{m}(t)$ , a  $N_1$  i  $N_2$  są numerami największych komponentów. Jeżeli powiązania korelacyjne OKPL maleją wraz ze wzrostem przesunięcia  $u$ , tj. jeżeli jest spełniony warunek (191), to estymatory (194)-(195) są asymptotycznie nieobciążone i zgodne. Wartości obciążeń i wariancji estymatorów komponentnych dla pewnej długości wycinka realizacji  $\theta$  są mniejsze od wartości tych charakterystyk estymatorów koherentnych w wypadku szybkiego zanikania powiązań korelacyjnych i małej ilości szacowanych komponentów.

Stosując metodę koregramową Blackmana-Tuckeya [9] możemy skonstruować następujące statystyki dla szacowania charakterystyk widmowych

$$\hat{g}(\omega, t) = \frac{1}{\pi} \int_0^{U_m} \hat{k}(t, u) \lambda(u) \cos \omega u du, \quad \hat{g}_0(\omega) = \frac{1}{\pi} \int_0^{U_m} \hat{K}_0(u) \lambda(u) \cos \omega u du, \quad (200)$$

$$\hat{g}_i^c(\omega) = \frac{1}{\pi} \int_0^{u_m} \hat{K}_i^c(u) \lambda(u) \cos \omega u \, du, \quad \hat{g}_i^s(\omega) = \frac{1}{\pi} \int_0^{u_m} \hat{K}_i^s(u) \lambda(u) \cos \omega u \, du, \quad (201)$$

Tutaj  $\lambda(u)$  — funkcja wagowa:  $\lambda(-u) = \lambda(u)$ ,  $\lambda(0) = 1$ ,  $\lambda(u) = 0$  dla  $|u| > u_m$ . Za pomocą statystyk (197)–(199) i otrzymanych powyżej równań (45)–(46) oraz (128)–(129) można oszacować funkcje  $b(t, u)$  i  $f(\omega, t)$ , opisujące ważne cechy niestacjonarności sygnałów zmodulowanych. Wiarygodność estymatorów charakterystyk pewnych typów sygnałów może być obliczona z wykorzystaniem wzorów dla obciążenia i wariancji tych estymatorów oraz wyrażeń opisujących ich konkretne struktury probabilistyczne, które zostały wyprowadzone powyżej.

### PODSUMOWANIE

Własności sygnałów zmodulowanych w ogóle mogą być opisane w ramach OKPL. Charakterystyki probabilistyczne tej klasy procesów niestacjonarnych opisują zarówno stochastyczność jak i powtarzalność w czasie własności sygnałów. Model sygnałów w postaci stacjonarnych procesów losowych jest szczególnym przypadkiem OKPL i może być traktowany jako jego przybliżenie stacjonarne. Wartość oczekiwana, funkcja autokorelacji, gęstość widmowa mocy są wartościami uśrednionych w czasie charakterystyk OKPL, a więc nie zawierają żadnej informacji o strukturze czasowej sygnałów i powiązanej z jej własnościami procesów modulujących. Założenie stacjonarności przy analizie sygnałów zmodulowanych w większości przypadków są formalne. Weryfikacja modelu może być dokonana tylko z wykorzystaniem metod analizy statystycznej OKPL. Ignorowanie cech niestacjonarności sygnałów prowadzi do mylnych wniosków przy badaniach statystycznych ponieważ dokładność estymatorów określa się nie tylko uśrednionymi w czasie wartościami charakterystyk. Uwzględnienie tych cech jest konieczne przy analizie przetwarzania sygnałów, szczególnie nieliniowych, gdy operacje uśredniania i przetwarzania są nieprzemienne. Są one ważne przy badaniach na podstawie danych eksperymentalnych własności fizycznych układów, w których takie sygnały są wytwarzane. Model sygnałów zmodulowanych w postaci OKPL umożliwia poprawne sformułowanie i rozwiązanie powyższych zagadnień.

### BIBLIOGRAFIA

1. Я.П. Драган: И.Н. Яворский: Методы вероятностного анализа ритмики океанологических процессов. Гидрометеиздат, Ленинград, 1987
2. W. Gardner: Statistical spectral analysis: A nonprobabilistic theory. Prentice Hall, Englewood Cliffs, NJ, 1987
3. I. Jaworski, Z. Drzycimski, W. Mychajlyszyn: Probabilistyczne modele i analiza statystyczna sygnałów telemetrycznych o okresowej niestacjonarności. Krajowe Sympozjum Telekomunikacji '96, A, Bydgoszcz, 1996, s. 53-60
4. I. Jaworski, W. Pogribny, Z. Drzycimski: Analiza stochastyczna sygnałów telemetrycznych o okresowo wolnozmienniej niestacjonarności. Krajowe Sympozjum Telekomunikacji, A, Bydgoszcz, 1996, s. 61-66

5. L.I. Gu z d e n k o: *On periodically nonstationary processes*. Radiotek. Elektron., vol 4, no 6, 1959, pp. 1062-1064
6. H. O g u r a: *Spectral representation of periodic nonstationary random processes*. IEEE Trans. Inform. Theory, vol. 17, 1971, pp. 143-149
7. W.M. B r e l s f o r d and R.H. J o n e s: *Time series with periodic structure*. Biometrika, vol. 54, 1967, pp. 403-407
8. A. P a p o u l i s: *Signal Analysis*. Mc Graw—Hill, New York, 1997
9. J.. B e n d a t, A.G. P i e r s o l: *Metody analizy i pomiaru sygnałów losowych*. PWN, Warszawa 1976

J. JAWORSKI, Z. DRZYCIMSKI, Z. ZAKRZEWSKI, J. MAJEWSKI

## CORRELATION AND SPECTRAL PROPERTIES OF NONSTATIONARY MODULATED SIGNALS

### S u m m a r y

Probabilistic models of modulated signals in the form of periodically correlated stochastic processes are considered. The second order signal properties are analysed. The formulas for signal correlation and spectral characteristics dependence on characteristics of modulating signals are obtained. The correlation and spectral structure of amplitude and phase modulated signals are described. The basic approaches of signal characteristics estimation are discussed.

*Key words:* nonstationary modulated signals, periodically correlated stochastic processes, autocorrelation functions, variable spectral density, correlation and spectral components.





# Coding capacity of programmable transcoder

DARIUSZ KANIA

*Instytut Elektroniki, Politechnika Śląska, Gliwice*

*Received 1998.01.07*

*Authorized 1998.04.20*

Programmable circuits have limited internal resources. The problem of implementation digital circuits in programmable devices is strictly connected with decomposition. Decomposition means breaking a large logic blocks into relatively smaller ones. This method of reducing the number of module inputs, outputs, product terms often uses different transcoders. This paper presents a coding capacity of PLE, PAL, PLA-based programmable transcoder. The result can be used to implement circuits in PLE, PAL, PLA-based CPLDs. The decomposition and word coding algorithms presented in this paper are implemented in Decomp system which has been developed at Silesian Technical University.

*Key word:* decomposition, partitioning, coding, Programmable Logic Devices.

## 1. INTRODUCTION

The set of combinational functions often cannot be made to fit into any single PLD or internal logic block of CPLD, FPGA devices. The only solution is to decompose the set of functions to be implemented requires a logic block with N-input, M-output and K-product terms in such a way that the requirement can be met by a network of two or more subcircuits with ninputs, moutputs and k product terms ( $n < N$ ;  $m < M$ ;  $k < K$ ). The classical methods of functional decomposition are based on Ashenhurst and Curtis theory [1,5].

Curtis proved the fundamental theorem:

$$v(X_2|X_1) \leq 2^p \Leftrightarrow f(X_2, X_1) = F[g_1(X_1), g_2(X_1), \dots, g_p(X_1), X_2] \quad (1)$$

If the column multiplicity  $v(X_2|X_1)$  is less or equal  $2^p$ , then the function  $f(X_2, X_1)$  can be decomposed into  $F$  form as shown in Fig.1.

Different algorithms of functional decomposition often based on Curtis theory are presented in [2, 3, 6, 8, 9, 10, 11, 12, 12, 13, 15]. The complexity of the propose

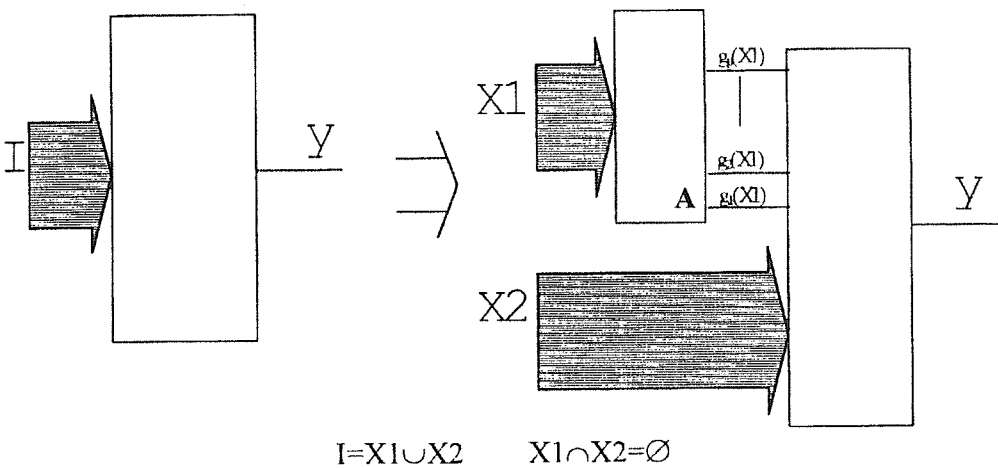


Fig. 1. Curtis decomposition

algorithms often preclude their practical application. The simplest method of reducing the number inputs required is input coding [3, 9]. This form of decomposition is presented in Fig. 2.

Systematic methods of detecting this type of decomposition are known as column partitioning that is truth table columns are allocated to partitions which contain rows with a limited number of combinations.

From Fig. 1 and 2, we observe that, after decomposition, the big block is broken into smaller subblocks. Block A (see Fig. 1, Fig.2 ) is a specific transcoder. The greater part of existing PLDs are using one of the three base structures (PLE, PAL,

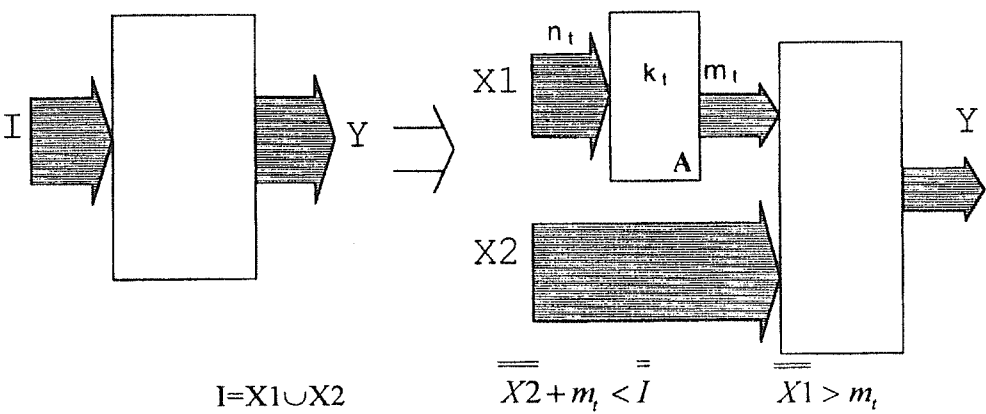


Fig. 2. The principle of partitioning the inputs based on input coding

PLA). Due to the variety of programmable structures partitioning the design and selecting appropriate PLD circuits has become one of the major issues in the development of digital devices. What is a coding capacity of PLE, PAL, PLA-based programmable transcoder? The present paper is a answer to these question.

The decomposition and word coding algorithms presented in this paper are implemented in Decomp system which has been developed at Silesian Technical University. Decomp includes different input, output, terms decomposition algorithms dedicated for PLE, PAL, PLA-based devices.

## 2. THEORETICAL BACKGROUND

Various existing methods of address modification allow the length of the address word to be reduced [3]. This principle can be used to partition a digital circuit for implementation in separate logical blocks. In this paper an input partitioning algorithm using an external transcoder will be demonstrated using a simple example.

**Example 1.** Let an exemplary function  $y$  be defined over the set  $I = \{I_{11}, \dots, I_1, I_0\}$ . Its description is given by the file  $y.pla$ , shown in Column A of Table 1. Let us try to find a partition of the function arguments using an 10-input transcoder.

First we shall introduce the necessary terminology:

**Definition 1** The complexity coefficient  $w_k(I_{s-1}, \dots, I_{p+1}, I_p)$  is the number of different words formed by the input variables  $I_{s-1}, \dots, I_{p+1}, I_p$ .

If a set  $X_1$  (see Fig.2) exists such that  $X_1 \subset I$  and  $w_k(X_1) < 2^{(X_1-1)}$ , then it is possible to limit the number of inputs by using an external transcoder. If we assume that the transcoder has  $n_t$ -inputs, then our search for  $X_1$  will of course be limited to such subsets for which  $X_1 \leq n_t$ . Additionally, we want to maximize the expression  $X_1 - [lg_2(w_k(X_1))]$ , because in this way we will most strongly limit the number of inputs for the  $y$  function circuit.

In our case  $n_t = 10$ , so our search will begin with those subsets for which  $X_1 = 10$ . A subset  $X_1 = \{I_{11}, I_{10}, I_9, I_8, I_7, I_6, I_5, I_4, I_1, I_0\}$  meeting all the conditions does indeed exist; it is therefore possible to propose the partitioning presented in Column B,C of Table 1 and Fig. 3.

Programmable circuits have limited internal resources. The problem of implementing a transcoder is strictly related to these resources. In the general case a programmable transcoder has a limited coding capacity.

A transcoder (PLE, PAL, PLA) is described by  $n_i$  — the number of inputs,  $m_i$  — the number of outputs and  $k_i$  — the number of code words which can be obtained in a given programmable structure (Fig.2). The  $k_i$  coefficient depends exactly on the base structure of the transcoder.

Table 1

Pla files defining the function y before (y.pla) and after argument partitioning

| y.pla file                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             | transcoder                                                                                                                                                                                                                                                                                                                                                                                                                                | y function circuits                                                                                                                                                                                                                                                                                                                                     |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <pre> .i 12 .o 1 .ilb I11,I10,I9,I8,I7,I6,I5,I4,I3,I2,I1,I0 .ob y .p 28 000000000000 1 000111110111 1 000000010100 1 011000010100 1 000000010111 1 011000011000 1 000000011011 1 011001010000 1 0000000100100 1 011001100000 1 000000101000 1 011001101100 1 000001000100 1 011000000111 1 000001001000 1 101010100101 1 000001010010 1 101010101001 1 000001100001 1 101010000111 1 000001110100 1 101100010100 1 000010000100 1 101100011000 1 000010001000 1 111111010010 1 000101110011 1 111111100100 1 .e </pre> | <pre> .i10 .o5 .ilb I11,I10,I9,I8,I7,I6,I5,I4,I1,I0 .ob a4,a3,a2,a1,a0 .p20 0000000000 00001 0110000011 01111 0000000100 00010 1010101001 10000 0000000111 00011 1010100011 10001 0000001000 00100 1011000100 10010 0000010000 00101 1111110110 10011 0000010110 00110 1111111000 10100 0000011001 00111 .e 0000011100 01000 0000100000 01001 0001011111 01010 0001111111 01011 0110000100 01100 0110010100 01101 0110011000 01110 </pre> | <pre> .i 7 .o 1 .ilb I3,I2,a4,a3,a2,a1,a0 .ob y .p 28 0000001 1 0101011 1 0100010 1 0101100 1 0100011 1 1001100 1 1000011 1 0001101 1 0100100 1 0001110 1 1000100 1 1101110 1 0100101 1 0101111 1 1000101 1 0110000 1 0000110 1 1010000 1 0000111 1 0110001 1 0101000 1 0110010 1 0101001 1 1010010 1 1001001 1 0010011 1 0001010 1 0110100 1 .e </pre> |
| A                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | B                                                                                                                                                                                                                                                                                                                                                                                                                                         | C                                                                                                                                                                                                                                                                                                                                                       |

$I = \{I_{11}, \dots, I_1, I_0\}$

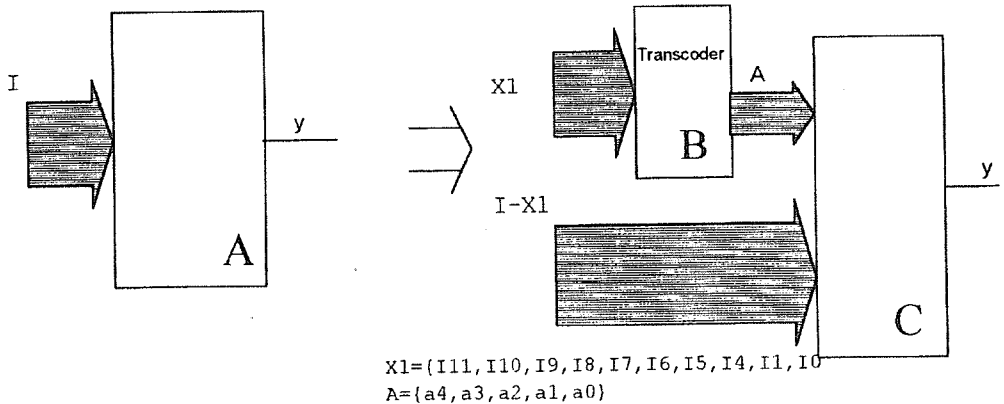


Fig. 3. The example of partitioning the inputs using an external transcoder

### 3. CODING CAPACITY OF PLE-BASED TRANSCODER

The important feature of the PLE architecture, is that the inputs are fully decoded by a fixed AND array which drives a programmable OR array. This means that every combination of inputs is represented by a separate AND gate. Since there are  $2^{n_i}$  combinations possible from  $n_i$ -inputs.

The above allows us to state that

$$k_{i(PLE)} = 2^{n_i}. \quad (2)$$

**Example 2** Let us try to implement the functions  $y_3, y_2, y_1, y_0$  described by the file *y.pla* (Column A of Table 2) using Configurable Logic Blocks (CLB) contained in an FPGA structure by XILINX, series 3000. A CLB is based on PLE structure and can implement one of the followings:

- any single output function of up to five input variables,
- any two output function of up to five input variables with each output depending on at most four input variables.

If the CLB is used as a transcoder, it will have the following parametres:

- $n_i = 5; m_i = 1; k_i = 2$  or
- $n_i = 4; m_i = 2; k_i = 4$ .

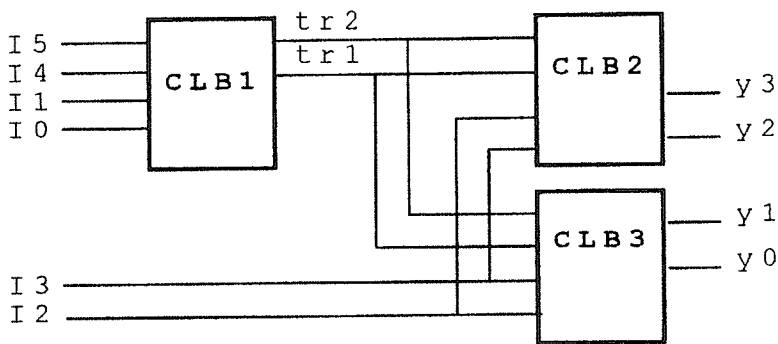
Table 2

The *y.pla* file and functions implemented in the individual CLBs

| <i>y.pla</i>           | CLB1 (transcoder) | CLB2               | CLB3               |
|------------------------|-------------------|--------------------|--------------------|
| i 6                    | i 4               | i 4                | i 4                |
| .o 4                   | .o 2              | .o 2               | .o 2               |
| .ilb 15,I4,I3,I2,I1,I0 | .ilb 15,I4,I1,I0  | .ilb 13,I2,tr2,tr1 | .ilb 13,I2,tr2 tr1 |
| .ob y3 y2 y1 y0        | .ob tr2 tr1       | .ob y3,y2          | .ob y1,y0          |
| .p 7                   | .p 7              | .p 5               | .p 6               |
| 000001 1011            | 0001 11           | 0011 10            | 0011 11            |
| 001001 0100            | 1011 10           | 1011 01            | 1111 11            |
| 001101 0011            | 1111 01           | 0110 01            | 0110 01            |
| 100111 0101            | 0000 00           | 1010 11            | 1010 11            |
| 101011 1111            | 0-1- 00           | 1101 10            | 0101 11            |
| 110111 0011            | 1-0- 00           | .e                 | 1101 10            |
| 111111 1010            | -10- 00           |                    | .e                 |
| .e                     | .e                |                    |                    |
| A                      | B                 | C                  | D                  |

Because a subset  $X1 = \{I5, I4, I1, I0\}$  meeting all the conditions does indeed exist; it is therefore possible to propose the partitioning presented in Fig. 4.

Using the decomposition system Decomp it is possible to implement the function in three CLBs, as shown in Fig. 4. The files describing each CLB are presented in Column B, C, D of Table 2. The PLE-based transcoder is described in Column B of Table 2.



$$X_1 = \{I5, I4, I1, I0\}; \quad X_2 = \{I3, I2\}$$

CLB1 - Configurable Logic Block used as transcoder

Fig. 4. Implementing the  $y$  function in PLE-based CLBs of the XILINX FPGA series 3000

#### 4. CODING CAPACITY OF PAL-BASED TRANSCODER

The PAL-based structure is exactly the opposite of a PLE. In these circuits the AND array is programmable and the OR array is fixed and therefore issue of coding becomes essential if a programmable PAL structure is to be used as an external transcoder.

PAL-based programmable chips have a limited number of terms connected to their OR gates. Let  $m_{PAL}$  be the number of PAL-based transcoder outputs and  $k_{PAL}$  — the number of terms connected to each output. The number  $k_{PAL}$  of terms is considerably inferior to  $k_{PAL} \ll 2^n$ , where  $n$  is the number of inputs. If binary coding is used, then, due to uneven term allocation, only  $2k_{PAL} + 1$  different words can be coded. It is of course possible to expand the number of terms [3, 9] and increase the number of code words by using additional outputs. The above analysis indicates the need of a different coding which will evenly use the products connected to the output OR gates.

How many different words can be coded using  $m_{PAL}$ -outputs, if  $k_{PAL}$ -products are connected to every output?

Optimal coding means using a minimal number of terms while keeping this number evenly distributed between the individual outputs. Let us consider what maximum number  $k_{i(PAL)}$  of code words can be obtained using a circuit with given parameters ( $m_{PAL}$ ,  $k_{PAL}$ ). To use the terms optimally, we should first use all the combinations "0 of  $m_{PAL}$ ", "1 of  $m_{PAL}$ ", "2 of  $m_{PAL}$ " etc. active outputs.

Each of the blocks "i of  $m_{PAL}$ ", where  $0 \leq i \leq m_{PAL}$ , makes use of a constant number of terms connected to each output (Table 3).

If, for a given programmable transcoder, there exists a number  $j$  such that  $\sum_{i=0}^j \binom{m_{PAL}-1}{i} = k_{PAL}$  and  $w_k(X1) < \sum_{i=0}^{j+1} \binom{m_{PAL}}{i}$ , then the coding algorithm is relative-

Table 3  
Number of combination and number of terms per output for different type of output code

| Type of output code        | 0 z $m_{PAL}$ | 1 z $m_{PAL}$ | 2 z $m_{PAL}$        | 3 z $m_{PAL}$          | .... $m_{PAL}$ z $m_{PAL}$          |
|----------------------------|---------------|---------------|----------------------|------------------------|-------------------------------------|
| Number of combination      | 1             | $m_{PAL}$     | $\binom{m_{PAL}}{2}$ | $\binom{m_{PAL}}{2}$   | .... $\binom{m_{PAL}}{m_{PAL}}$     |
| Number of terms per output | 0             | 1             | $m_{PAL}-1$          | $\binom{m_{PAL}-1}{2}$ | .... $\binom{m_{PAL}-1}{m_{PAL}-1}$ |

vely simple. It consists in choosing an arbitrary set of combinations “i of  $m_{PAL}$ ”; this set will have  $w_k(X1) + 1$  elements, where  $0 \leq i \leq j + 1$  (the combination 0...00 corresponds to all the other words, not included in the number  $w_k(X1)$ ).

The algorithm becomes more complicated if no number j meets the condition  $\sum_{i=0}^j \binom{m_{PAL}-1}{i} = k_{PAL}$  and  $w_k(X1) < \sum_{i=0}^{j+1} \binom{m_{PAL}}{i}$ . In this case, we can present differently the concept of optimal coding. In the first step we find a value of j which meets the inequalities  $\sum_{i=0}^{j-1} \binom{m_{PAL}-1}{i} < k_{PAL} < \sum_{i=0}^j \binom{m_{PAL}-1}{i}$ . Now a part of the required number of words can be obtained using all the combinations “i of  $m_{PAL}$ ”, where  $0 \leq i \leq j$ .

This step yields  $\sum_{i=0}^j \binom{m_{PAL}}{i}$  different combinations. When these have been coded, the transcoder still has a certain number  $k_r = k_{PAL} - \sum_{i=0}^{j-1} \binom{m_{PAL}-1}{i}$  of unused terms at each output. These terms can be used to code more words. With  $k_r$  unused terms connected to every output we can code a maximum number of “j + 1 of  $m_{PAL}$ ” combinations. The total number of terms used by every combination is j + 1. With  $m_{PAL}$  outputs and  $k_r$  unused terms at every output we can code a maximum of  $\left\lfloor \frac{m_{PAL} k_r}{j+1} \right\rfloor$  additional words. The total number of possible words is therefore  $\sum_{i=0}^j \binom{m_{PAL}}{i} + \left\lfloor \frac{m_{PAL} k_r}{j+1} \right\rfloor$ .

The above allows us to state that

$$k_{I(PAL)} = \sum_{i=0}^j \binom{m_{PAL}}{i} + \left\lfloor \frac{m_{PAL} \left( k_{PAL} - \sum_{i=0}^{j-1} \binom{m_{PAL}-1}{i} \right)}{j+1} \right\rfloor, \quad (3)$$

where j is a number satisfying the inequality

$$\sum_{i=0}^{j-1} \binom{m_{PAL}-1}{i} < k_{PAL} \leq \sum_{i=0}^j \binom{m_{PAL}-1}{i}. \quad (4)$$

The optimal output coding algorithm allowing an external PAL-based programmable transcoder was presented in [7].

**Example 3** We shall now try to implement the transcoder of Example 1 in a programmable structure, in which 8 terms are connected to each output (e.g. a macrocell of GAL16V8). Let the outputs of the programmable transcoder are active in the high state. Under this assumption, a "1" in a output code means that a term was used. It is not possible to implement the transcoder directly in such macrocells. It is necessary to expand the number of terms [7, 9], which leads to using additional 3 outputs. The individual words can be coded better, using fewer outputs.

Because we want to code  $w_k(X1) = 20$  words using a PAL-based programmable transcoder with  $k_{PAL} = 8$  products per output (Table 1; Fig.2) we must using transcoder for which  $k_t > w_k(X1)$ . (Note: a ">" is required because the combination 0...00 is reserved for the remaining words, not included in the number  $w_k(X1)$ ).

From equations 3 and 4 we can determine the parameter  $j$  and the number of outputs necessary to code a given number of words. The creation of code words

Table 4

Description of a PAL-based programmable transcoder before/after optimal coding

| transcoder<br>(binary coding)        | transcoder<br>(optimal coding)       | transcoder optimal coding $k_{PAL} = 8$<br>(PAL description) |
|--------------------------------------|--------------------------------------|--------------------------------------------------------------|
| .i10                                 | .i10                                 | .i10                                                         |
| .o5                                  | .o5                                  | .o5                                                          |
| .ilb i11,i10,i9,i8,i7,i6,i5,i4,i1,i0 | .ilb i11,i10,i9,i8,i7,i6,i5,i4,i1,i0 | .ilb i11,i10,i9,i8,i7,i6,i5,i4,i1,i0                         |
| .ob a4,a3,a2,a1,a0                   | .ob a4,a3,a2,a1,a0                   | .ob a3,a2,a1,a0                                              |
| .p20                                 | .p20                                 | .p40                                                         |
| 0000000000 00001                     | 0000000000 00001                     | 0000000000 00001 1111110110 00100                            |
| 0000000100 00010                     | 0000000100 00010                     | 1010101001 00001 0001011111 00100                            |
| 0000000111 00011                     | 0000000111 00100                     | 1010100011 00001 0110010100 00100                            |
| 0000001000 00100                     | 0000001000 01000                     | 0000011001 00001 0110011000 00100                            |
| 0000010000 00101                     | 0000010000 10000                     | 0000011001 00001 0000001000 01000                            |
| 0000010110 00110                     | 0000010110 00011                     | 0000011100 00001 1010101001 01000                            |
| 0000011001 00111                     | 0000011001 00101                     | 1111110110 00001 0000011100 01000                            |
| 0000011100 01000                     | 0000011100 01001                     | 0000100000 00001 1011000100 01000                            |
| 0000100000 01001                     | 0000100000 10001                     | 0000000100 00010 1111110110 01000                            |
| 0001011111 01010                     | 0001011111 00110                     | 1010101001 00010 0001111111 01000                            |
| 0001111111 01011                     | 0001111111 01010                     | 1010100011 00010 0110010100 01000                            |
| 0110000100 01100                     | 0110000100 10010                     | 1111111000 00010 0110000011 10000                            |
| 0110010100 01101                     | 0110010100 01100                     | 0000010110 00010 0000010000 10000                            |
| 0110011000 01110                     | 0110011000 10100                     | 0001011111 00010 1010100011 10000                            |
| 0110000011 01111                     | 0110000011 11000                     | 0001111111 00010 1111111000 10000                            |
| 1010101001 10000                     | 1010101001 01011                     | 0110000100 00010 0110000011 10000                            |
| 1010100011 10001                     | 1010100011 10011                     | 0000000111 00100 1011000100 10000                            |
| 1011000100 10010                     | 1011000100 11100                     | 1111111000 00100 0000100000 10000                            |
| 1111110110 10011                     | 1111110110 01101                     | 0000011001 00100 0110000100 10000                            |
| 1111111000 10100                     | 1111111000 10110                     | 1011000100 00100 0110011000 10000                            |
| .e                                   | .e                                   | .e                                                           |



occurs in two stages. In the first stage, all “ $i$  of  $m_{PAL}$ ” combinations are generated, where  $0 \leq i \leq j$ . In the second stage, the remaining words are coded as “ $j+1$  of  $m_{PAL}$ ” combinations. The first stage requires no special explanations; the purpose of the second stage is to find a maximum number of “ $j+1$  of  $m_{PAL}$ ” combinations when  $k_r$  terms per output are available. The optimal PAL-based programmable transcoder structure is described in Table 4.

## 5. CODING CAPACITY OF PLA-BASED TRANSCODER

The PLA allows both the AND array and the OR array to be programmed. This gives the PLA additional coding capability over PALs. Every term can be used for coding one input combination. Coding capacity is limited by the total number of AND-gates allocated to all outputs. Let  $k_{PLA}$  be the total number of terms. Knowing the value of  $k_{PLA}$  we can determine the parameter  $k_{i(PLA)}$  according to the following dependence

$$k_{i(PLA)} = k_{PLA} + 1. \quad (5)$$

(Note: a ‘+ 1’ is required because the combination 0...00 is reserved for the remaining words, not included in the number  $w_k(X1)$ ).

## 6. CONCLUSIONS

The problem of implementing a transcoder is strictly related to resources of programmable structures. A PLE-based programmable transcoder hasn't a limited coding capacity (eqn.2). Coding capacity of PLA-based transcoder is related to number of terms (eqn.5).

Equation 3 indicates that the coding capacity of PAL-based devices depends on two parameters  $k_{i(PAL)} = f(m_{PAL}, k_{PAL})$ . This fact has practical implications. If, after partitioning the arguments  $w_k(X1) \geq k_r$ , then we should either use a programmable transcoder with a greater number of terms per output or use additional outputs from additional blocks or macrocells. Table 5 shows the values of the  $k_{i(PAL)}$  coefficient for various values of  $m_{PAL}$  and  $k_{PAL}$  (note the difference between optimal coding and binary coding, for which  $k_{i(PAL)} = 2k_{PAL} + 1$ ).

The word coding method presented above was implemented in a prototype program Decomp assisting the decomposition of combinational circuits. It is one of the algorithms used to optimally code words in a PLE, PAL, PLA-based external transcoder. In the present paper the algorithm was described for PAL-based structures with an equal number of terms connected to each output and with high-active outputs. It is, however, also directly applicable to the coding of words in circuits with varied number of terms and programmable output type (e.g. 22V10).

With slight modifications, the method can also be used to code the states of sequential Mealy automata being implemented in PLE, PLA, PAL-based structures (e.g. GAL, MAPL, MAX etc.).

Table 5

Coefficient values  $k_{i(PAL)} = f(m_{PAL}, k_{PAL})$ 

| $k_{PAL} \backslash m_{PAL}$ | 3 | 4  | 5  | 6  | 7  | 8  | 9  | 10 |
|------------------------------|---|----|----|----|----|----|----|----|
| 3                            | 7 | 9  | 11 | 13 | 15 | 17 | 19 | 21 |
| 4                            | 8 | 11 | 13 | 16 | 18 | 21 | 23 | 26 |
| 5                            | — | 12 | 16 | 19 | 22 | 25 | 28 | 31 |
| 6                            | — | 13 | 17 | 22 | 25 | 29 | 32 | 36 |
| 7                            | — | 15 | 19 | 24 | 29 | 33 | 37 | 41 |
| 8                            | — | 16 | 21 | 26 | 31 | 37 | 41 | 46 |

## REFERENCES

1. R.L. Ashenurst: *The decomposition of switching functions*. Proceedings of an International Symposium on the Theory of Switching, April 1957
2. B. Babba, M. Crastes, G. Saucier: *Input driven synthesis on PLDs and PGAs*. The European Conference on Design Automation, Brussels (Belgium), March 1992
3. M. Bolton: *Digital Systems Design with Programmable Logic*. Addison-Wesley Publishing Company, 1990, p. 133-140
4. S.D. Brown, R.J. Francis, J. Rose, Z.G. Vranesic: *Field Programmable Gate Arrays*. Kluwer Academic Publishers, Boston/London/Dordrecht 1993, p.45-86
5. H.A. Curtis: *The Design of switching Circuits*. D.van Nostrand Company, Inc., Princeton, New Jersey, Toronto, New York, 1962
6. F. Dresig, Ph. Lanches, O. Rettig, U.G. Baitiger: *Functional Decomposition for Universal Logic Cells using Substitution*. The European Conference on Design Automation, Brussels (Belgium), March 1992
7. E. Hryniewicz, K. Pucher, and D. Kania: *The Input Partitioning and Coding Problem in PAL-based CPLDs*. XX National Conference Circuits Theory and Electronic Networks, Krynica 1997,
8. D. Kania: *Complex decomposition of multiple output-functions*. Proc. International Conference on Programmable Devices and Systems, Ostrava 1996
9. D. Kania: *Table decomposition methods of digital circuits. Implementation of these circuits in selected PLD structures*. PhD Thesis, Silesian Technical University, Gliwice 1995
10. T. Luba, H. Selveraj, A. Krasniewski: *A new approach to FPGA-based logic synthesis*. Workshop on Design Methodologies for Microelectronics and Signal Processing, Gliwice-Cracow, 1993
11. M. Rawski, M. Nowicka, P. Tomaszewicz, T. Luba: *Decomposition-based logic synthesis and its application in FPGA-oriented technology mapping*. Proc. International Conference on Programmable Devices and Systems, Ostrava 1996
12. T. Sasa o: *Logic Synthesis and Optimization*, Kluwer Academic Publishers, Boston/London/Dordrecht, 1993
13. T. Sasa o: *FPGA Design by Generalized Functional Decomposition, in Logic Synthesis and Optimization*. Kluwer Academic Publishers, Boston/London/Dordrecht, 1993

14. H. Selveraj, T. Łuba, M. Nowicka, B. Bignall: *Multiple-valued decomposition and its applications in data compression and technology mapping*. ICCIMA'97
15. W. Wan, M.A. Perkowski: *A new Approach to the Decomposition of Incompletely Specified Multi-Output Functions Based on Graph Coloring and Local Transformations and Its Applications to FPGA Mapping*. Proceedings EDAC'92, Bruesels

D. KANIA

## POJEMNOŚĆ KODOWA PROGRAMOWALNYCH TRANSKODERÓW

### Streszczenie:

Struktury programowalne mają ograniczone zasoby wewnętrzne. Problem realizacji układów cyfrowych w oparciu o struktury programowalne jest ściśle związany z dekompozycją, oznaczającą podział dużych bloków logicznych na mniejsze. Metody redukcji liczby wejść, wyjść, termów często wykorzystują zewnętrzne transkodery. W niniejszym artykule jest analizowany problem pojemności kodowej programowalnych transkoderów zbudowanych na bazowych strukturach PLE, PAL i PLA. Wyniki przedstawionych rozważań mogą być wykorzystywane do realizacji układów cyfrowych w strukturach CPLD bazujących na strukturach PLE, PAL, PLA. Algorytm dekompozycji i kodowania słów został zaimplementowany w systemie Decomp opracowanym w Politechnice Śląskiej.

*Słowa kluczowe:* dekompozycja, podział, kodowanie, Układy Logiki Programowalnej.



## Semi-empirical boundary conditions at p-n junctions for device simulation

ANDRZEJ PFITZNER, BEATA JUCHT-SUCHTA, ALFRED ŚWIT

*Instytut Mikroelektroniki i Optoelektroniki, Politechnika Warszawska*

*Received 1998.02.19*

*Authorized 1998.04.08*

Semi-empirical formulae have been obtained, which approximate the pn product at depletion layer edges in p-n junction versus bias voltage. These formulae have been derived using full numerical solutions of the semiconductor transport equations, assumed as true distributions of the physical quantities of interest. This modification of the classic boundary conditions allows to expand the applicability of analytical expressions, like the classic current-voltage characteristics, to moderate and high forward bias, where the neutrality condition may be violated and the voltage drop on resistances of the p and n regions is important. In the case of reverse bias, the exact numerical solutions show, that the quasi fermi-levels and the difference between them are not constant inside the depletion layer. Moreover, the separation of the fermi-levels does not depend on the applied voltage for a high enough reverse bias, but is a function of the impurity concentration. In consequence, formulae are proposed, which approximate the pn product at the depletion layer edges, stabilized at the value depending on the impurity concentration for the reverse bias voltages above  $10 V_T$ . This approach avoids the time consuming numerical solutions, maintaining a high accuracy in the modeling of devices containing p-n junctions. Owing to this, the boundary conditions obtained in the paper, can be used in statistical process/device simulators.

**Key words:** p-n junctions, transport in semiconductors, c-v characteristics, process/device simulators.

### 1. INTRODUCTION

CAD tools for IC device modeling create rigorous requirements as to accuracy and speed of simulation. The most accurate and versatile modeling of semiconductor devices based on full numerical solution of the carrier transport equations is extremely time consuming. This method is impractical for complex circuits, particularly in the case of statistical process/device simulators, which repeat the analysis of the

device many times for the randomly disturbed material and design parameters. From this point of view, fast analytical models are more appropriate but unfortunately, in many applications they are too simplified and inaccurate. Good results can be obtained using exact numerical solutions to provide such modifications of analytical formulae that allow a simulation in an acceptable long computing time and retaining a high accuracy. Analysis of the boundary conditions at p-n junctions is the basis for many analytical models.

In this paper we propose a semi-empirical formula describing the  $pn$  product at the space charge layer edges in the p-n junction. This modification of the classic boundary condition allows to expand the applicability of analytical expressions, like the classic current-voltage characteristics, to moderate and high forward bias, where the neutrality condition may be violated and the resistances of the p and n regions are important.

## 2. THE APPROACH

A semi-empirical formula for the  $pn$  product at the depletion layer boundaries versus bias voltage has been derived using "accurate" numerical solutions of the semiconductor transport equations [1], [2], assumed as true spatial distributions of the physical quantities of interest, e.g. carrier concentrations and electric field. Analysis of these functions allows to derive "exact" values of the fermi-levels in the transition region and the  $pn$  product at the edges of this region (see section 3).

Our modification of the classic boundary conditions consists of matching an approximate function to the voltage changes of the  $pn$  product at the edges of the space charge layer (section 4).

Numerical solutions of the transport equation set do not introduce the classic sectional approach i.e. the partitioning of the device into space-charge and neutral regions. Hence, at the begining of our considerations, it is necessary to define the coordinates of the depletion layer edges. In order to derive the abrupt space charge edges (ASCE), the compensation method was used [3]. This method assumes that the charge of ionized impurities compensates the true charge existing in the given part of the transition region (Fig. 1).

Thus, the following relations should be satisfied:

$$\int_{-d_p}^0 q N(x) dx = \int_{-x_p}^0 \rho(x) dx, \quad (1a)$$

$$\int_{d_n}^0 q N(x) dx = \int_0^{x_p} \rho(x) dx, \quad (1b)$$

where:  $-d_p, d_n$  — coordinates of the abrupt space charge edges (ASCE),  
 $x_p, x_n$  — depths of the charge penetration into p and n regions,  
 $-x_{kp}, x_{kn}$  — coordinates of the electric contacts,  
 $\rho(x)$  — charge density,  
 $N(x) = N_d - N_a$  — net impurity concentration,  
 $N_d, N_a$  — concentrations of the donors and acceptors.

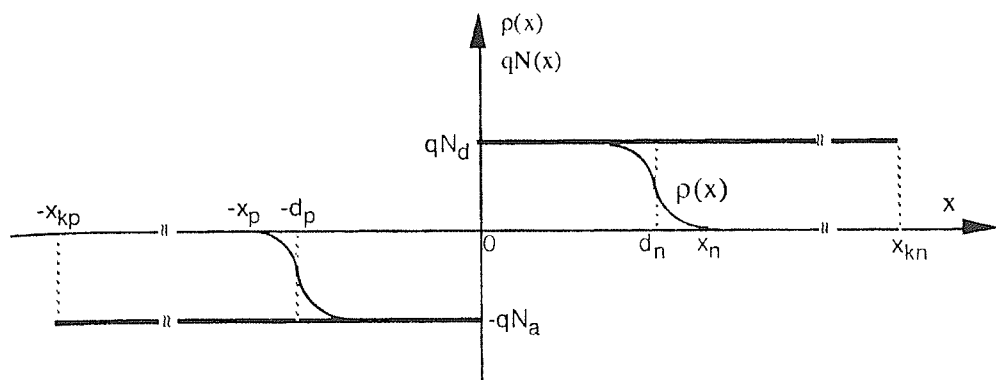


Fig. 1. The charge density distribution and the ASCE approximation in an abrupt p-n junction

Using the numerical solutions of the transport equations, the right sides of equations (1) can be easily evaluated and then a fit can be made for  $d_n$  and  $d_p$ . The above method of derivation of the space charge layer edges is simple and universal (arbitrary impurity distributions and arbitrary biases). Moreover, this method is verified experimentally through measurements of the junction capacitance [3].

### 3. RESULTS OF THE NUMERICAL ANALYSIS

In the following, an abrupt junction is used as an example. This example is numerically difficult but clear in an analytical notation. Computations were realised using a program in which the algorithm presented in papers [1], [2] was applied. A few results in the form of fermi-level plots are shown ( $L_D$  in the following figures means intrinsic Debye length).

In the case of forward bias, for high voltages, a nonzero slope of the fermi-levels is observed in the neutral regions (Fig. 2).

In the space charge layer the fermi-levels are practically constant (Fig. 3). In our example, in the investigated range of bias voltages ( $U = 0 \div 800$  mV), the relative (to  $U$ ) slope does not exceed 1% even though the current of the carriers which recombine in the depletion layer is a dominant component of the total current. (This result does not confirm information published in [4], but is conformable to the paper [5]).

In the case of reverse bias the fermi-levels can be assumed constant in the depletion layer only for major carriers (i.e. for electrons and holes in the metallurgical  $n$  and  $p$  part of the depletion region, respectively) and for low voltages (Fig. 4).

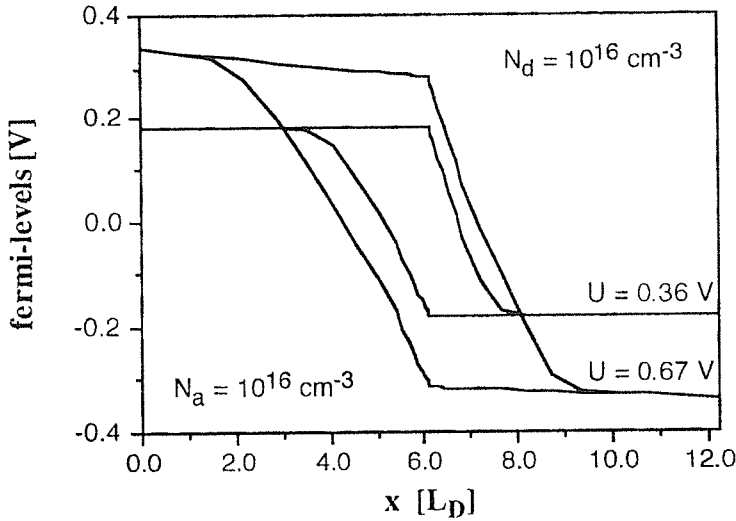


Fig. 2. Spatial variations of the fermi-levels in a step p-n junction under forward bias

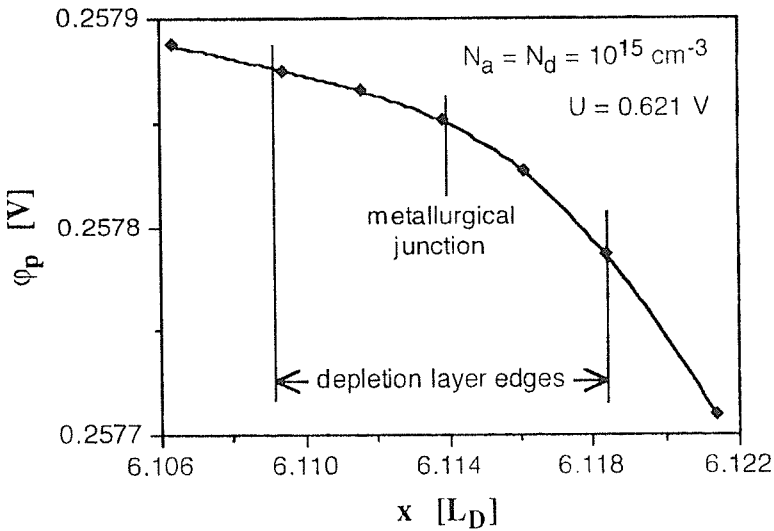


Fig. 3. Spatial variation of the fermi-level for holes in the depletion layer of a step p-n junction (forward bias)

Changes of the fermi-levels for voltages exceeding  $10 kT/q$  are shown in Fig. 5. For voltages  $U_R \geq 60 kT/q$ , the relative (to  $U_R$ ) change of the fermi-levels in the depletion layer is over 90%. Therefore, in the case of higher reverse biases practically the whole slope of the fermi-levels lies in the space charge layer.



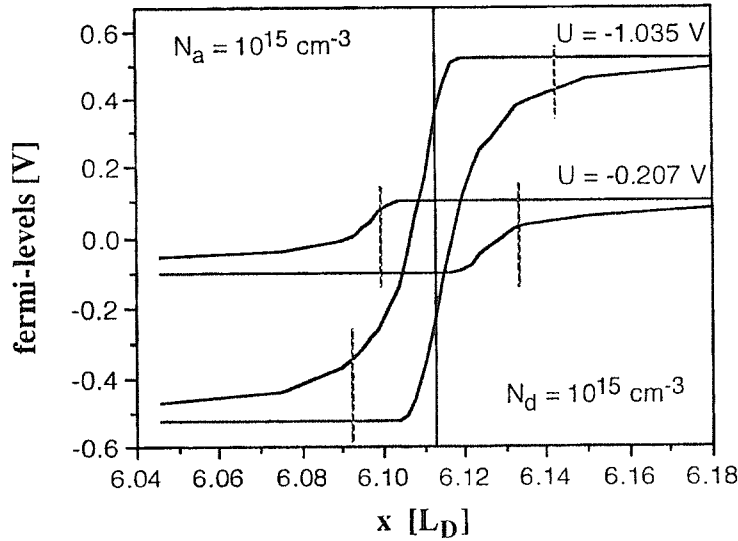


Fig. 4. Spatial variations of the fermi-levels in a step p-n junction under reverse bias (dashed lines indicate the depletion layer edges)

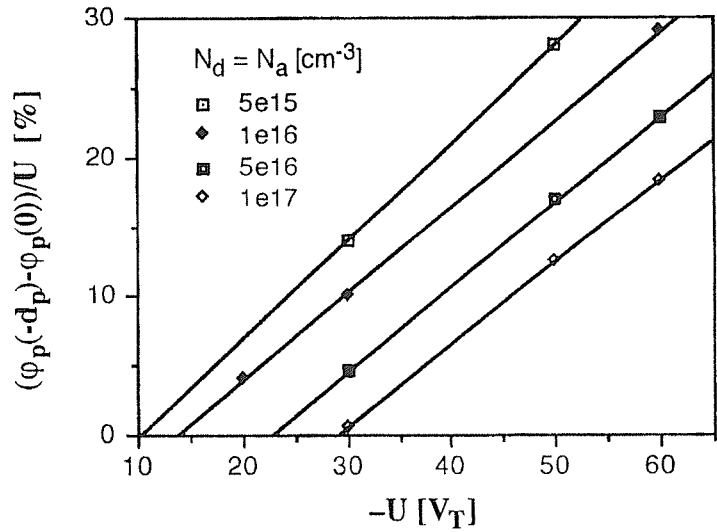


Fig. 5. Relative change of the fermi-level for holes in the metallurgical p part of the depletion layer versus reverse applied bias

4. DERIVATION OF THE BOUNDARY CONDITIONS

A splitting of the fermi-level into two levels: for electrons and holes, characterizes quantitatively a disturbance of the thermal equilibrium due to the bias of the p-n junction. The difference between the fermi-levels for electrons and holes allows to

evaluate the  $pn$  product:

$$np = \exp(\varphi_p - \varphi_n), \quad (2)$$

where:  $n, p$  — concentrations of electrons and holes in units of the intrinsic concentration  $n_i$ ,

$\varphi_n, \varphi_p$  — fermi-levels in units of the thermal potential  $V_T = kT/q$ .

In analytical considerations, the values of the  $pn$  product at the space charge edges versus bias voltages define the boundary conditions for the carrier distributions in neutral regions and, in consequence, they allow to find the current-voltage characteristics of the p-n junction. The classic theory of the p-n junction assumes the fermi-levels to be constant inside the depletion layer and the fermi-level for major carriers — constant in quasi-neutral regions. This means, that:

$$n(-d_p)p(-d_p) = n(d_n)p(d_n) = \exp(U), \quad (3)$$

where  $U$  is the applied voltage in units of  $V_T$ . The current-voltage characteristics obtained according to condition (3) are applicable only for low forward bias, when a finite conductivity of the quasi-neutral regions and a noncomplete compensation of the electric charge of the excess minority carriers by the majority carriers in these regions can be ignored.

Numerical calculations indicate, that the assumption of the constant  $pn$  product in the depletion layer is valid for forward bias (Fig. 6). The separation of the fermi-levels assumes maximum at the metallurgical junction, but changes of its magnitude in the whole depletion layer are negligible (hundredth parts of a percent) in the investigated range of bias voltages ( $U = 0 \div 800$  mV). Therefore, this assump-

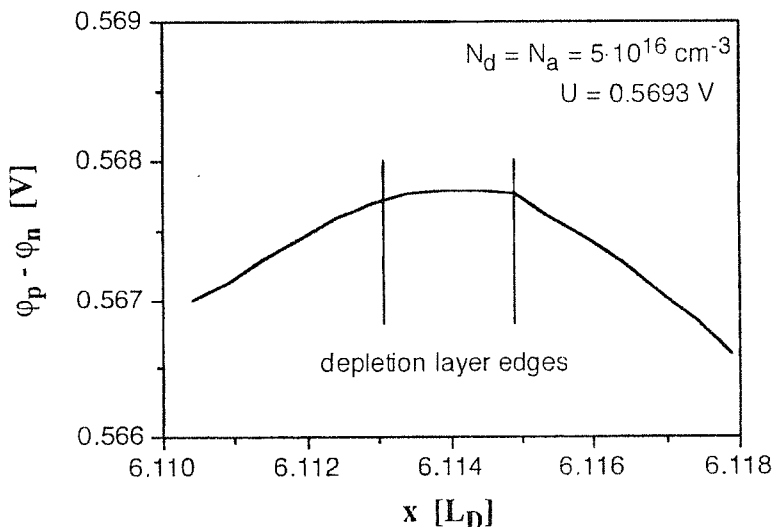


Fig. 6. Separation of the fermi-levels in the transition region of the forward biased p-n junction

tion is satisfied better than the assumption of the constant fermi-levels in the depletion layer (section 3).

However, the right hand equality in the expression (3) is not satisfied for arbitrary bias. For higher bias voltages, a difference between both fermi-levels is smaller than the applied voltage  $U$ : The voltage drop on the quasi-neutral regions are not negligible because of the finite conductivity of these regions.

In the following, a simplified evalutaion of this voltage drop is made for the worst case:

$$\Delta U_r = I(\rho_p x_{kp} + \rho_n x_{kn}), \quad (4)$$

i.e. when the modulation of the resistivity of the quasi-neutral regions ( $\rho_p$  and  $\rho_n$ ) and the penetration of the depletion layer into these regions are ignored. It was found that the relation:

$$\Delta \varphi_F = \varphi_p - \varphi_n = U - \Delta U_r, \quad (5)$$

is satisfied only until a characteristic value of the bias voltage  $U_{min}$ , above which the "true" separation of the fermi-levels (Fig. 7) and, thus, the  $pn$  product are smaller. The voltage  $U_{min}$  is a function of the logarithm of impurity concentrations (Fig. 8).

In the range of higher voltages (above  $U_{min}$ ), a noncomplete compensation of the electric charge of the excess minority carriers in quasi-neutral regions is significant. A quantitative description of these phenomena is complicated and analytical

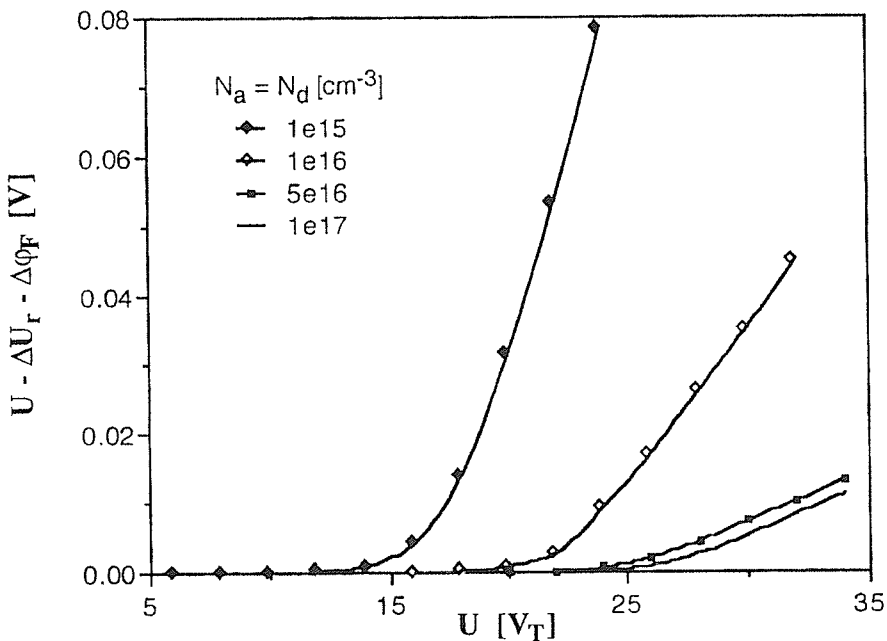
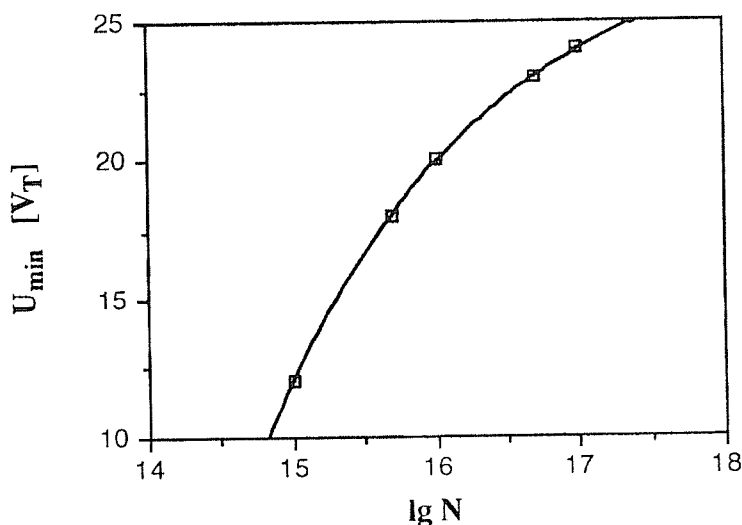
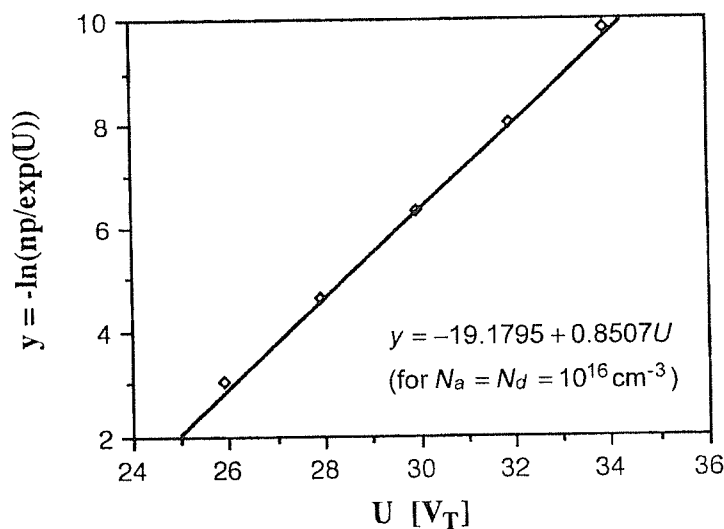


Fig. 7. Difference between the forward applied voltage reduced by  $\Delta U_r$ , and the separation of the fermi-levels

Fig. 8.  $U_{min}$  versus carrier concentration

considerations are simplified [5], [6]. For this reason, we have searched for a semi-empirical formula, which includes phenomena occurring at higher bias voltages and based on information obtained from the fully numerical solution of the transport equations in the p-n junction. Our approximation is derived on the basis of the plot of  $-\ln(np/\exp(U))$  versus the applied voltage (Fig. 9).

Fig. 9. Derivation of approximation of the function  $-\ln(np/\exp(U))$  for high voltages

At voltages high enough, this function can be approximated by a straight line:

$$-\ln(np/\exp(U)) = w + mU,$$

whereas at low voltages it is equal to zero (the classical formula (3) is satisfied). Therefore, for arbitrary forward bias voltages this function can be approximated by the formula:  $\ln(1 + \exp(w + mU))$  (Fig. 10).

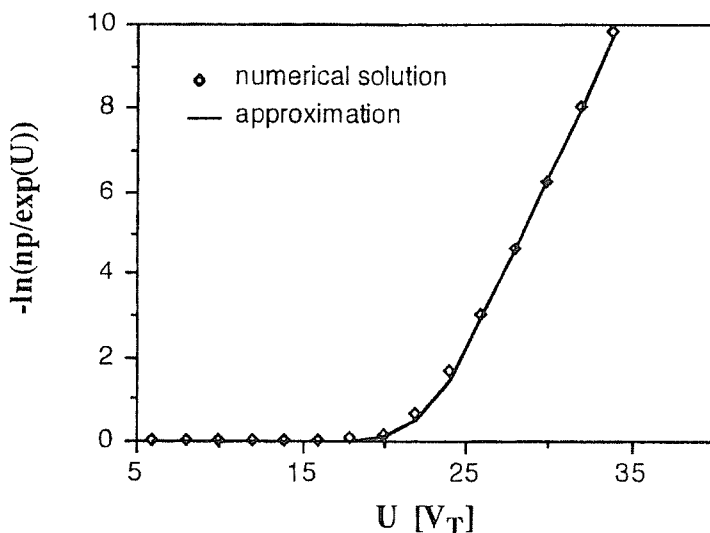


Fig. 10. Approximation of the function  $-\ln(np/\exp(U))$  (example for  $N_a = N_d = 10^{16} \text{ cm}^{-3}$ )

Finally, the following semi-empirical expression, giving the  $pn$  product at the depletion layer edges, was obtained for the case of forward bias of a p-n junction:

$$n(-d_p)p(-d_p) = n(d_n)p(d_n) = \exp(U)/(1 + \exp(w + mU)). \quad (6)$$

An example of the comparison between the numerically evaluated values of the  $pn$  product at the depletion layer edges and those obtained using the semi-empirical formula (6) is shown in Fig. 11 [7].

From the practical point of view it is important that the factors in this formula are nearly second power functions of the logarithm of doping concentrations (Fig. 12).

Exact boundary conditions at a p-n junction under reverse bias are not of great practical importance, but have also been derived.

It is known that equation (3) is satisfied only at low reverse bias voltages. Considering the difference between the separation of the fermi-levels and the applied voltage, it is found that the voltage drop on resistances of the quasi-neutral regions can be ignored because of the small value of the reverse current. Now, in opposite to the forward bias case, it is important that the quasi fermi-levels and the difference

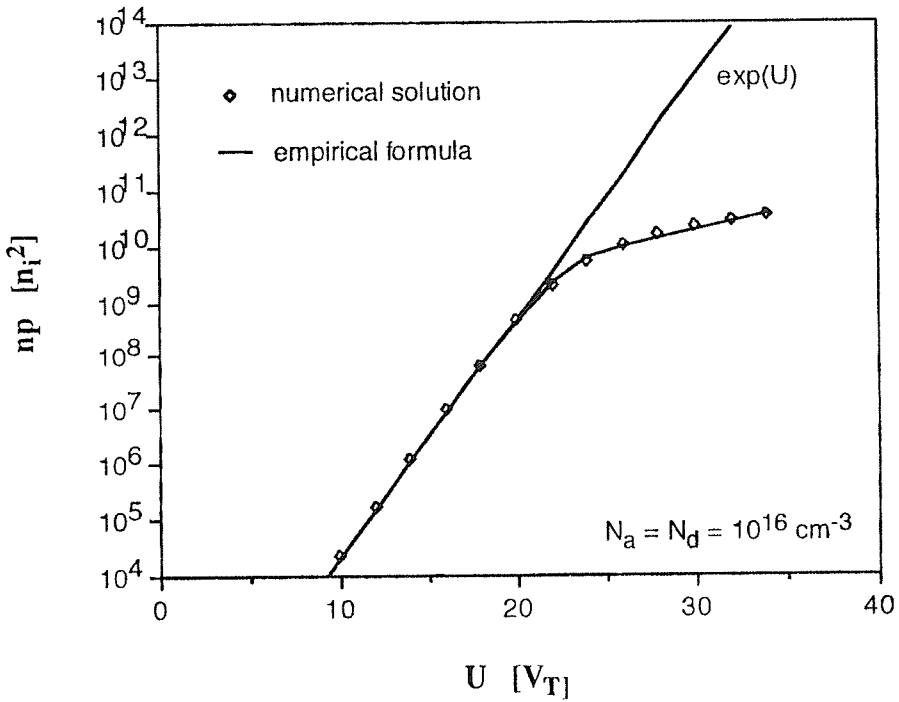


Fig. 11. Comparison between the numerically evaluated values of the  $pn$  product at the depletion layer edges and those obtained using the semi-empirical formula (6)

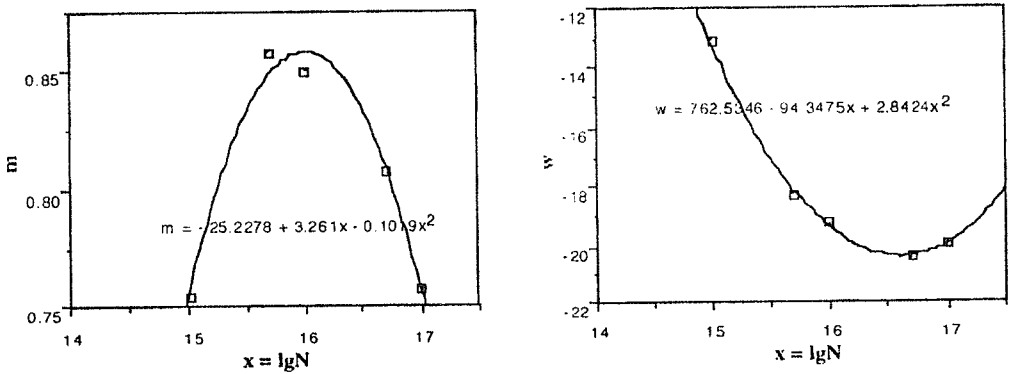


Fig. 12. Factors  $m$  and  $w$  (Eq. 6) vs. doping concentration [8]

between them are not constant inside the depletion layer. Moreover, the exact numerical solutions show, that the separation of the fermi-levels, greatest at the metallurgical junction, does not depend on the applied voltage for a high enough reverse bias, but is a function of the impurity concentration (Fig. 13).

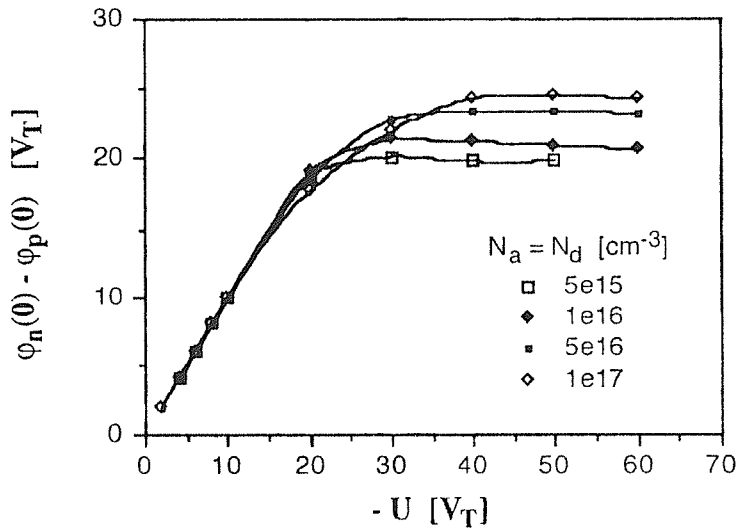


Fig. 13. Separation of the fermi-levels at the metallurgical junction versus reverse applied bias

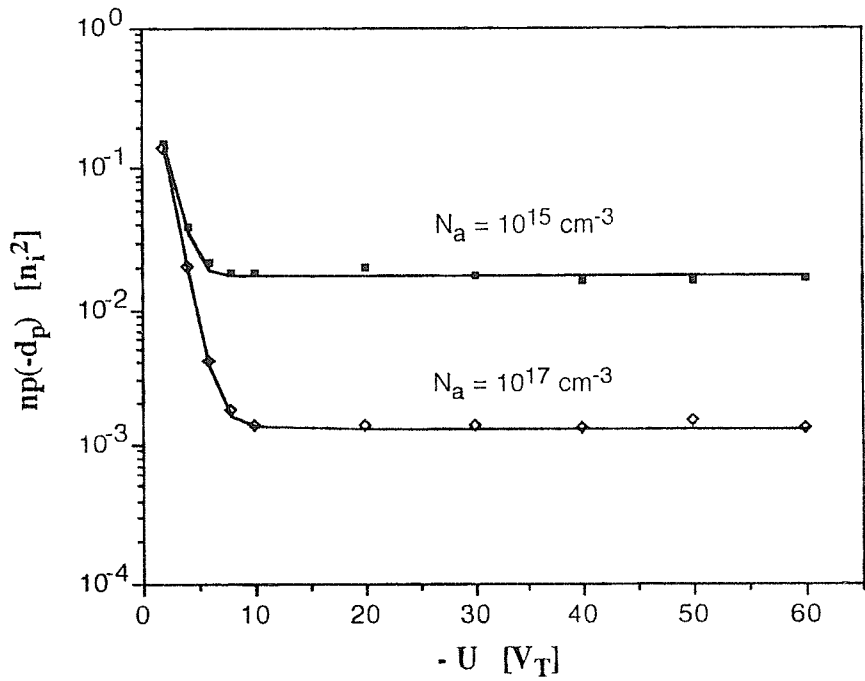


Fig. 14. Comparison between the numerically evaluated values of the  $pn$  product at the depletion layer edge and those obtained using the semi-empirical formula (8)

In consequence, the  $pn$  product at the depletion layer edges becomes stabilized at the value depending on the impurity concentration for the reverse bias voltages above  $10 V_T$ . This function can be approximated by the following expression:

$$np = \exp(U) + aN^{-b}. \quad (7)$$

An example of the comparison between the numerically evaluated values of the  $pn$  product at the depletion layer edge ( $x = -d_p$ ) and those obtained using the semi-empirical formula:

$$n(-d_p)p(-d_p) = \exp(U) + 5.18 \cdot 10^6 \cdot N^{-0.565}, \quad (8)$$

is shown in Fig. 14.

## CONCLUSIONS

We have obtained new semi-empirical formulae, which approximate the  $pn$  product versus bias voltage at the space charge layer edges in the p-n junction. This approach avoids the time consuming numerical solution of the semiconductor transport equations, maintaining a high accuracy in the analysis of diodes, bipolar transistors and other devices containing p-n junctions operating at high bias conditions. Analytical formulae obtained in such a way can find application in IC modeling, particularly in statistical process/device simulators.

## REFERENCES

1. W. Kuźmicz, A. Pfitzner: *A highly accurate numerical solution of semiconductor transport equations*. NASECODE IV Proc., pp. 372-377, 1985
2. A. Pfitzner: *Numerical solution of the one-dimensional phenomenological transport equation set in semiconductors*. Electron Technology 1977, vol. 10, pp. 3-21
3. W. Kuźmicz, A. Świt: *The built-in voltage and space charge layer capacitance of p-n junctions*. Solid State Electronics, 1974, vol. 17, no. 4, pp. 457-463
4. A. Nussbaum: *The modified Fletcher boundary conditions*. Solid State Electronics, 1975, vol. 18, pp. 107-109
5. E.I. Heasell: *Boundary conditions at p-n junctions*. Solid State Electronics, 1979, vol. 22, pp. 853-856
6. C.T. Sah: *The spatial variation of the quasi-Fermi potentials in p-n junctions*. IEEE Trans. on Electron Devices, 1966, vol. ED-13, no. 12, pp. 809-845
7. A. Pfitzner: *Dokładne modelowanie charakterystyki I-U złącza p-n w statystycznej symulacji układów scalonych*, XIII KKTOiUE, 449-452, Bielsko-Biała, październik 1990
8. A. Pfitzner: *A Methodology of Building Hybrid Models for Statistical Process/Device Simulators in the IC Design Systems*. MIXDES, 183-185, Łódź 1996



A. PFITZNER, B. JUCHT-SUCHTA, A. ŚWIT

PÓLEMPIRYCZNE WARUNKI BRZEGOWE W ZŁĄCZU P-N  
DLA SYMULACJI PRZYRZĄDÓW PÓŁPRZEWODNIKOWYCH

## Streszczenie

Sformułowano wzory półempiryczne, które aproksymują iloczyn koncentracji np na krawędziach warstwy zaporowej złącza p-n w funkcji napięcia polaryzacji. Związki te zostały określone na podstawie pełnych numerycznych rozwiązań półprzewodnikowych równań transportu, traktowanych jako rzeczywiste rozkłady interesujących wielkości fizycznych. Zaproponowana modyfikacja klasycznych warunków brzegowych pozwala rozszerzyć zakres stosowności wyrażań analitycznych, takich jak klasyczne charakterystyki prądowo-napięciowe, do średnich i wysokich napięć polaryzacji przewodzenia, kiedy warunek neutralności nie jest spełniony i spadek napięcia na rezystancjach obszarów p i n może być znaczny. Dla polaryzacji zaporowej, dokładne rozwiązania numeryczne wskazują, że quasi poziomy Fermiego i różnica pomiędzy nimi nie są stałe wewnątrz warstwy zaporowej, a ponadto ich rozszczepienie nie zależy od napięcia polaryzacji dla odpowiednio dużego napięcia wstecznego, lecz jest funkcją koncentracji domieszek. W rezultacie zaproponowano wzór aproksymujący iloczyn koncentracji np na krawędziach warstwy zaporowej, zapewniający jego stabilizację na wartości zależnej od koncentracji domieszek, dla napięć polaryzacji zaporowej powyżej  $10 V_T$ . Przedstawione podejście unika czasochłonnych rozwiązań numerycznych zachowując jednak wysoką dokładność modelowania przyrządów zawierających złącza p-n. Dzięki temu otrzymane warunki brzegowe mogą być wykorzystane w modelach stosowanych w statystycznych symulatorach procesów wytwarzania i elementów układów scalonych, gdy obliczenia powtarzane są wielokrotnie dla zaburzonych wartości parametrów materiałowych i konstrukcyjnych.

*Słowa kluczowe:* złącza p-n, transport w półprzewodnikach, charakterystyki prądowo-napięciowe, symulatory przyrządów półprzewodnikowych.



# Varying capacitances with the use of current conveyors controlled by a resistor

LECH TOMAWSKI

*Uniwersytet Śląski, Katowice*

*Received 1998.02.24*

*Authorized 1998.03.20*

Four circuits realising positive capacitances, the values of which can be set by a grounded resistor are described. Computer simulation shows that dependences of the capacitance vs. frequency are very similar for all circuits but dependences of the parallel resistance and the dissipation factor are different. One of these circuits has such a low dissipation factor which cannot be obtained in any physical capacitors.

*Key words:* controlled capacitances, current conveyor CCII

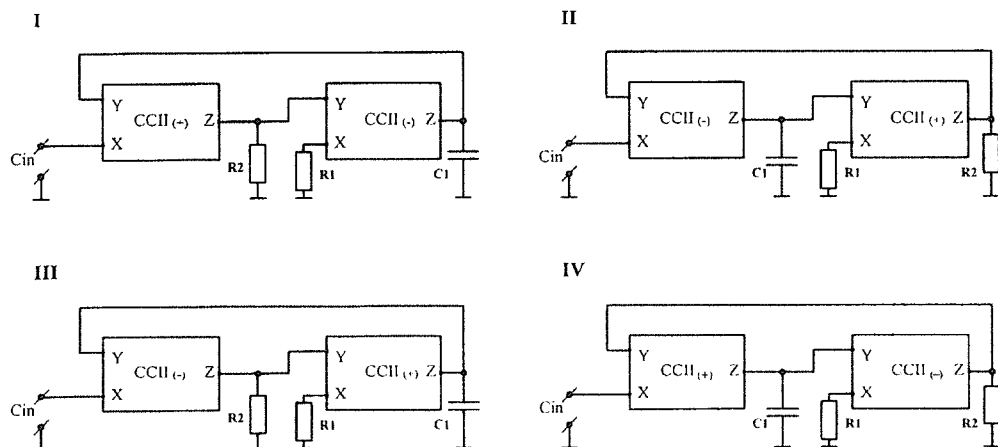
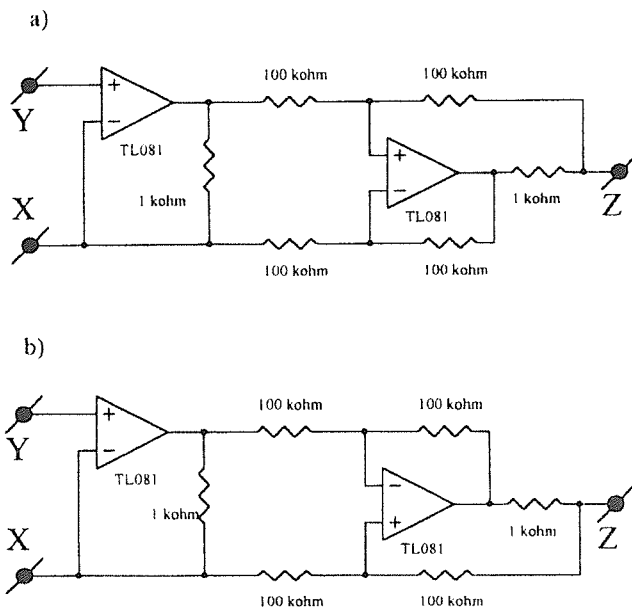
## INTRODUCTION

Four circuits presented in this paper, shown in Fig. 1, realise positive capacitances which can be easily set by a grounded resistor. They have a low dissipation factor, especially in the region of high values of the simulated capacitance. Current conveyors CCII (–) and CCII (+) realised according to Huertas proposition [1] are shown in Fig. 2 a and b. In comparison with circuits, known in literature [2, 3], the described circuits and the five circuits reported in [4] exhibit low sensitivity values.

## 1. BASIC PROPERTIES

By direct analysis it can be shown from Fig. 1 that capacitance  $C_{in}$  at the input terminals of all circuits is:

$$C_{in} = C_1 \frac{R_1}{R_2}. \quad (1)$$

Fig. 1. Circuits simulating positive capacitance  $C_{in}$  at the input terminalsFig. 2. Current conveyors  $CCII(+)$  – a, and  $CCII(-)$  – b

However, the block diagrams of these circuits, apart from  $C_{in}$  also include some other elements, among them resistance  $R_{in}$ ,  $FDNC^*$  ( $Y(s) = s^2 D$ ) and  $FDNCap^{**}$  ( $Y(s) = s^3 F$ ) are essential. Two resonance phenomena appear in these

\*  $FDNC$  = Frequency Dependent Negative Conductance.

\*\*  $FDNCap$  = Frequency Dependent Negative Capacitance.

circuits. The first one between positive resistance  $R_{in}$  and negative resistance  $-1/\omega^2 D$  at frequency:

$$f_R = \frac{1}{2\pi \sqrt{R_{in} D}}. \quad (2)$$

The second one between positive capacitance  $C_{in}$  and negative capacitance  $-1/\omega^2 F$  at frequency:

$$f_C = \frac{1}{2\pi \sqrt{\frac{C_{in}}{F}}}. \quad (3)$$

Since computer simulations of the electronic circuits with the use of standard programs do not operate in the complex frequency domain and do not define the values of  $C_{in}$ ,  $R_{in}$ ,  $D$ ,  $F$ , a simple simulation has been carried out with the admittance  $Y(j\omega)$  at the input terminals described as:

$$Y_{in}(j\omega) = 1/R_{in} + j\omega C_{in}. \quad (4)$$

From now on resistance  $R_{in}$  and capacitance  $C_{in}$  become dependent on the frequency and the formula (1) holds true for a low frequency region only (below  $f_C$  frequency).

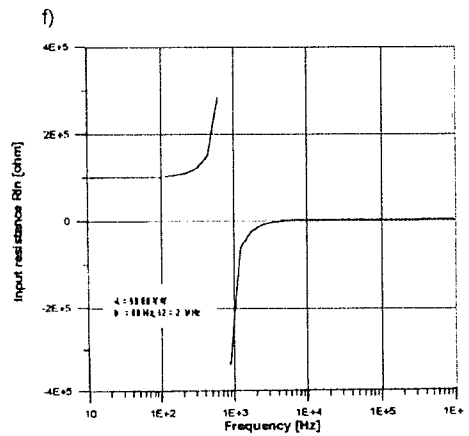
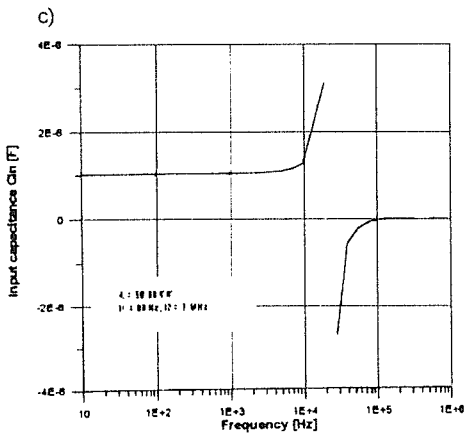
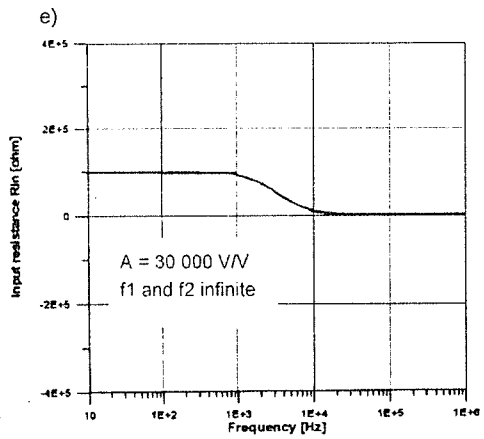
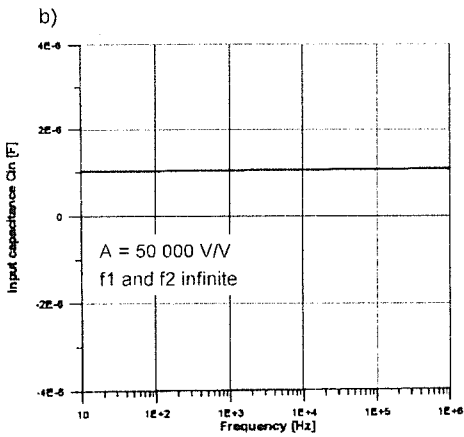
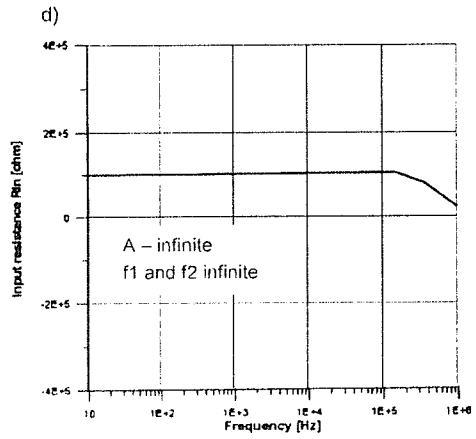
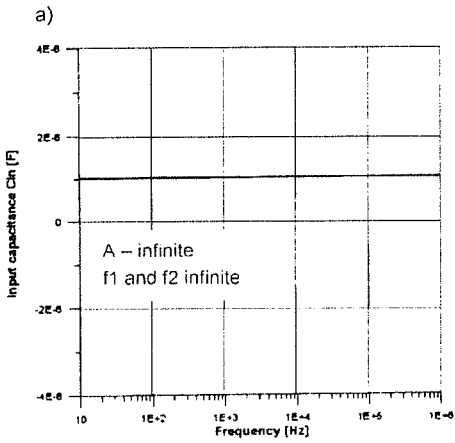
The presence of FDNC causes the appearance of a discontinuity point on the  $R_{in}$  dependence vs. frequency, whereas the FDNCap leads to the appearance of a similar point on the  $C_{in}$  dependence vs. frequency.

The shape of the dependences  $R_{in}$  and  $C_{in}$  vs. frequency depends on the construction of the circuit and on the operational amplifiers used. Curves for  $C_{in}$ ,  $R_{in}$  and dissipation factor  $\tan \delta$ , calculated as  $1/(\omega R_{in} C_{in})$ , for circuit I at  $C_1 = 1 \mu F$ ,  $R_1 = R_2 = 2 \text{ k}\Omega$ , are shown in Fig. 3.

Input capacitance  $C_{in}$  is equal  $1 \mu F$ , (according to (1)) when the operational amplifiers are ideal- see Fig. 3a. In this case input resistance  $R_{in}$  (Fig. 3d) has the value  $100 \text{ k}\Omega$  which decreases at frequency above  $200 \text{ kHz}$ . This drop in  $R_{in}$  is caused by capacitance  $C_1$ . If  $C_1 = 0$  this curve has a constant value  $100 \text{ k}\Omega$  independent of frequency.

Effects of nonideality of applied operational amplifiers in the CCII construction are shown in Fig. 3b-3i. Finite gain value ( $50\,000$  for TL081 op. amp.) caused the decrease of  $R_{in}$  from  $100 \text{ k}\Omega$  to a few  $\Omega$  (see Fig. 3e).

Both finite gain and finite frequencies of the operational amplifier poles caused the occurrence of discontinuity points in all dependences (see Fig. 3c, f, i). Above these points  $C_{in}$  and  $R_{in}$  have negative signs and hence the dissipation factor becomes positive again.



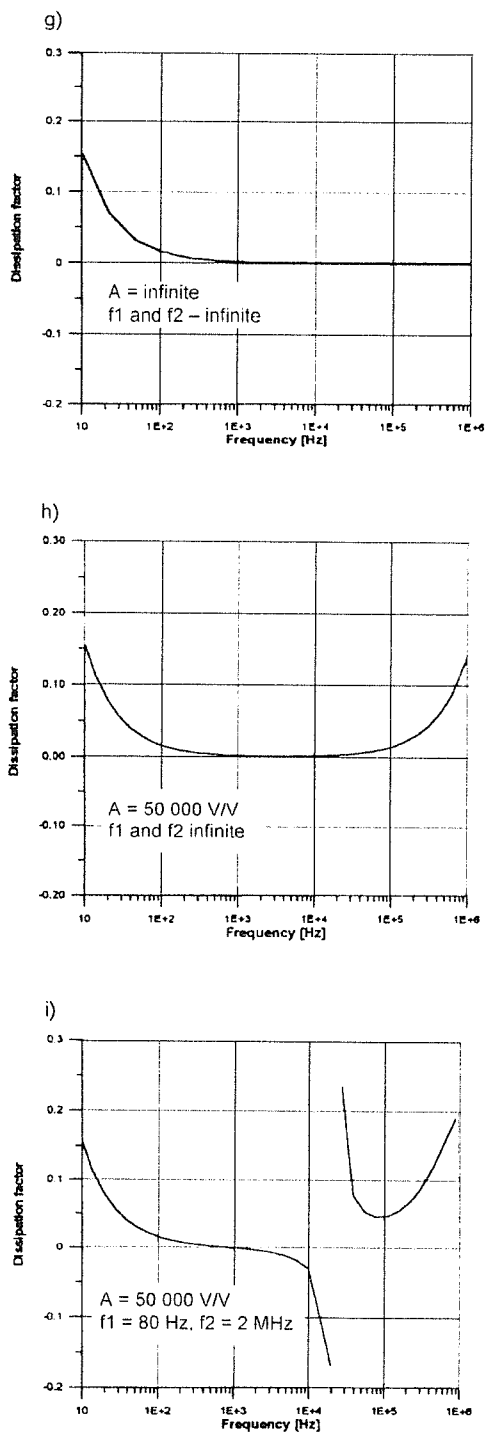


Fig. 3. Basic dependences for  $C_{in}$ ,  $R_{in}$  and the dissipation factor calculated for circuit I with capacitance  $C_1 = 1 \mu\text{F}$  and  $R_1 = R_2 = 2 \text{ k}\Omega$ . Curves a, d, g—operational amplifiers are ideal, curves b, e, h, c, f, i—operational amplifiers are nonideal

## 2. COMPUTER SIMULATION

Computer simulation of described circuits was carried out in the MicroCap program for TL081 operational amplifiers represented by the following data: input resistance  $1\text{E}9$ , gain  $50\,000\text{ V/V}$ , frequency of first pole  $80\text{ Hz}$ , frequency of second pole  $2\text{ MHz}$ , slew rate  $1.2\text{ V}/\mu\text{s}$ . Capacitance  $C_1$  varied from  $1\text{ nF}$  to  $10\text{ }\mu\text{F}$ . Resistors  $R_1$  and  $R_2$  have constant values equal  $2\text{ k}\Omega$ . Input admittance  $Y_{in}(j\omega) = 1/R_{in} + j\omega C_{in}$  was calculated from the following formulas:

$$R_{in} = 1/(G_{10} + \sum_{i=2}^n (1 - A_{1i} \cos \varphi_{1i}) G_{1i}) \quad (5)$$

$$C_{in} = C_{10} - \sum_{i=2}^n (G_{1i} A_{1i} \sin \varphi_{1i})/2\pi f, \quad (6)$$

where  $G_{10}$  and  $C_{10}$  denote conductance and capacitance between the hot input terminal (node 1) and the ground (node 0);  $G_{1i}$  the conductance between node 1 and other node;  $A_{1i}$  and  $\varphi_{1i}$  gain and phase angle between nodes 1 and  $i$  obtained from MicroCap simulation. All nodes which are connected to node 1 by any conductance different from zero must be taken into account.

## 3. RESULTS OF COMPUTER SIMULATIONS

Dependences of input capacitance  $C_{in}$ , input resistance  $R_{in}$  and dissipation factor  $\tan \delta$  vs. frequency are shown in Fig. 4 for all described circuits I-IV. The comparison of all dependence families is only possible with the use of logarithmic scales but then the negative regions of plotted curves cannot be shown.

$C_{in}$  curves for all circuits have the same shapes in Fig. 4 since the differences between them are of the range of few per mille. Marked differences appear between  $R_{in}$  curves. For circuits I and III  $R_{in}$  has the value of  $100\text{ k}\Omega$  in the low frequency region but for circuit II it is  $25\text{ M}\Omega$  and for circuit IV:  $50\text{ k}\Omega$ . Hence the dissipation coefficients for circuit II are lower from the remaining ones. It is noteworthy that the dissipation factor lower than  $10^{-4}$  cannot be obtained in physical capacitors. Above frequency  $f_R$  (formula (2)) the dissipation factor becomes negative but since the parallel resistance  $R_{in}$  has a very high value it becomes meaningless. It is always possible to connect an additional resistor to the input terminals, the resistance value of which can be matched in such a way that the equivalent resistance at the input terminals becomes positive at the right frequency.

In Table 1 the border frequencies where the modulus of the dissipation factor is lower than  $10^{-2}$  are given in order to compare the properties of the described circuits.

From Table 1 it can be seen, that circuit II has the widest region in which the simulated capacitance  $C_{in}$  exhibits low losses.



Table 1

The frequencies limiting the range of the dissipation factor where its modulus is lower than  $10^{-2}$ 

| Capacitance C1 | Circuit I |         | Circuit II |         | Circuit III |         | Circuit IV |         |
|----------------|-----------|---------|------------|---------|-------------|---------|------------|---------|
| 1 nF           | 21 kHz-   | 27 kHz  | 600 Hz-    | 4.8 kHz | 21 kHz-     | 27 kHz  | 30 kHz-    | 36 kHz  |
| 10 nF          | 5.5 kHz-  | 9.5 kHz | 60 Hz-     | 4.8 kHz | 5.5 kHz-    | 9.5 kHz | 8 kHz-     | 14 kHz  |
| 0.1 $\mu$ F    | 1.2 kHz-  | 5.2 kHz | 6 Hz-      | 4.8 kHz | 1.2 kHz-    | 5.2 kHz | 1.8 kHz-   | 6.2 kHz |
| 1 $\mu$ F      | 130 Hz-   | 4.8 kHz | 0.6 Hz-    | 4.8 kHz | 130 Hz-     | 4.8 kHz | 250 Hz-    | 4.8 kHz |
| 10 $\mu$ F     | 15 Hz-    | 4.8 kHz | 0.06 Hz-   | 4.8 kHz | 15 Hz-      | 4.8 kHz | 30 Hz-     | 4.8 kHz |

#### 4. RESULTS OF MEASUREMENTS

Measurements were carried out for 5 values of capacitor C1 vs. resistance R2. According to formula (1) the increase of R2 caused the decrease of Cin. Capacitance Cin was measured by the standard multimeter up to the value of about 20  $\mu$ F and for higher values it was calculated from the voltage drop on the capacitance.

Fig. 5 shows the families of curves for all described circuits. It can be seen, that for circuit II the measured dependences are regular in a wide region. The dissipation factor for circuit II has very high values especially in the region of high input capacitances.

Fig. 6 shows curves obtained from measurements for two parallel resonant circuits consisting of:

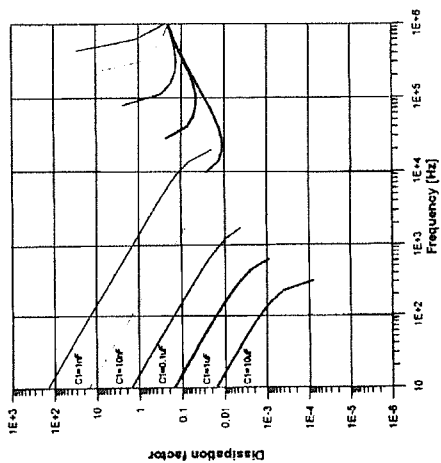
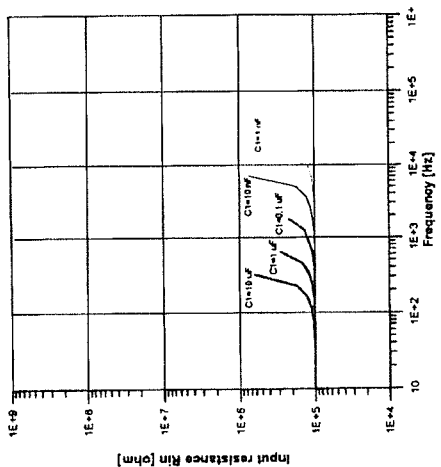
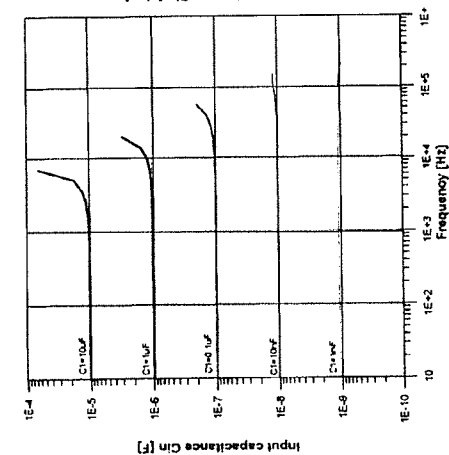
- inductor 66 mH and capacitor 10 nF (curve a),
- inductor 66 mH and capacitance 10 nF realised by circuit I in which the same capacitor 10 nF was used as C1, at resistances  $R1 = R2$  (curve b)

Despite the negative sign of resistance Rin and the dissipation factor of circuit I at frequency 6.2 kHz this circuit can have a practical use. A slight disadvantage can appear when this circuit is connected by a series capacitor. Then negative resistance Rin charges the capacitance Cin to a voltage which can overdrive the operational amplifiers. Connecting a suitable resistor to input terminals neutralises this effect.

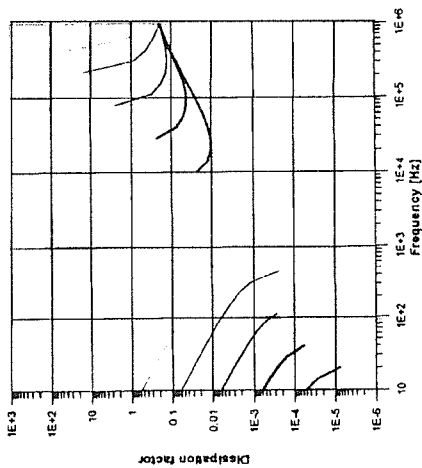
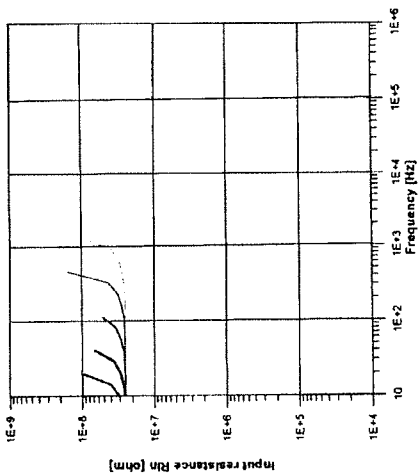
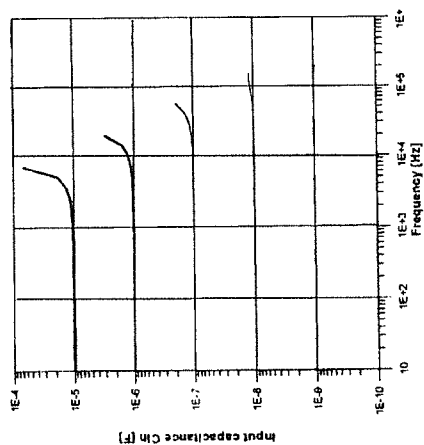
#### SUMMARY

Each circuit I-IV can be used as a positive grounded capacitance which is easily controlled by resistor R1 or R2. The controlling signal can be analogue — when the resistors R1 or R2 are substituted by a transconductance amplifier [3], or a digital one — when these resistors are substituted by a digital potentiometer chip, for example DS 1267 [5]. The described circuits have interesting properties in the region

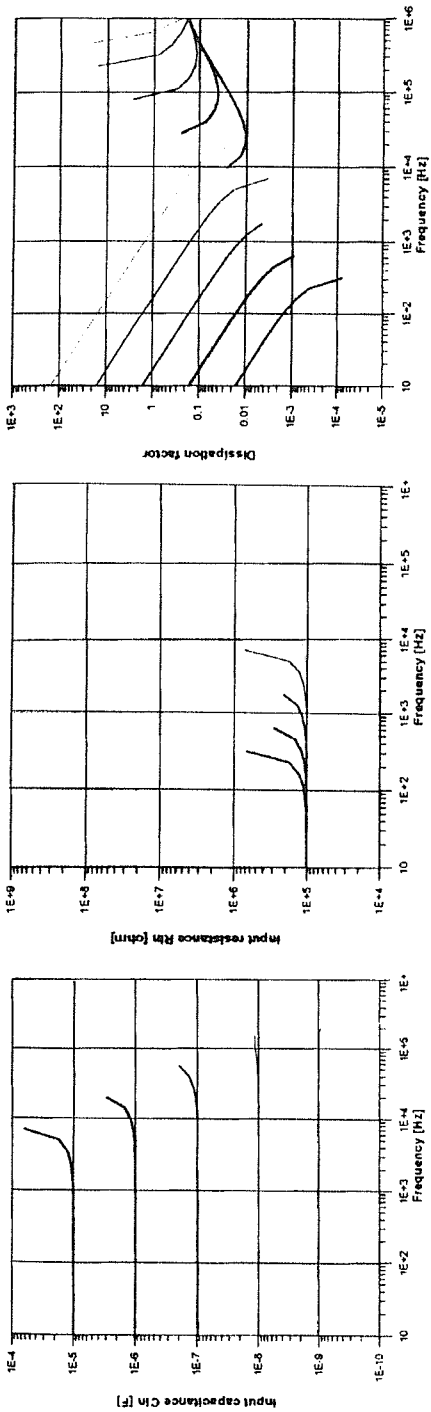
## Circuit I



## Circuit II



### Circuit III



### Circuit IV

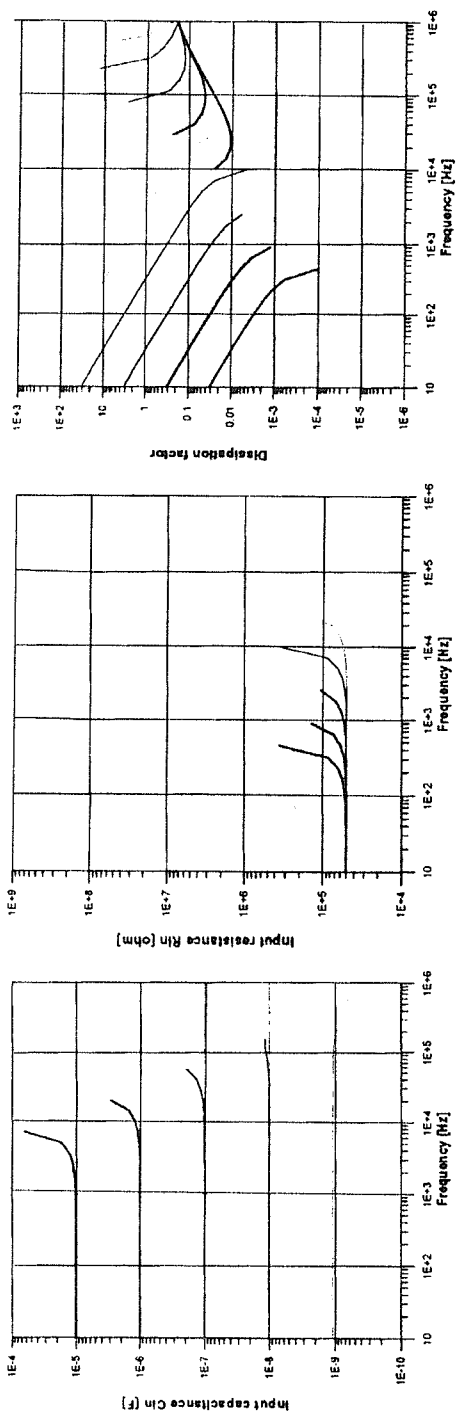
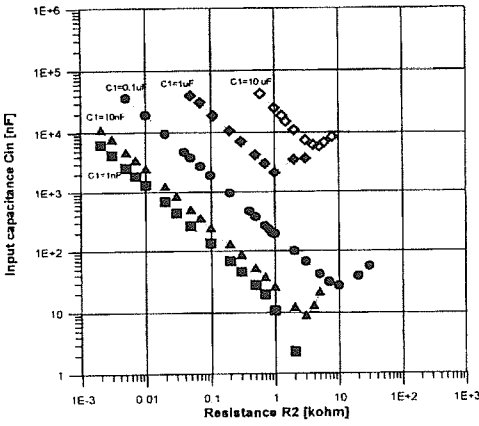
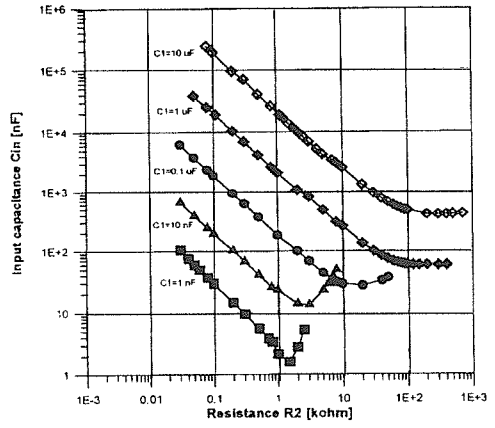


Fig. 4. Results from computer simulation for circuits I-IV vs. frequency in log-log scales

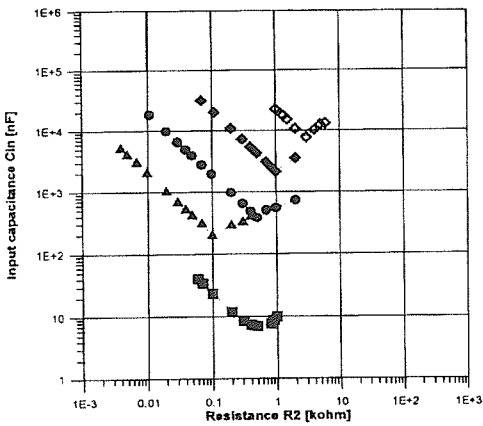
Circuit I



Circuit II



Circuit III



Circuit IV

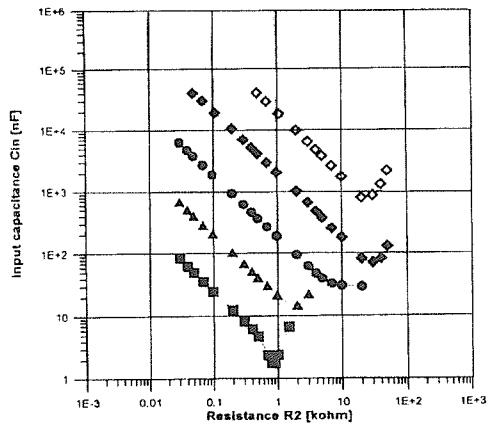


Fig. 5. Input capacitance  $C_{in}$  obtained from measurements for circuits I-IV vs. resistance  $R_2$

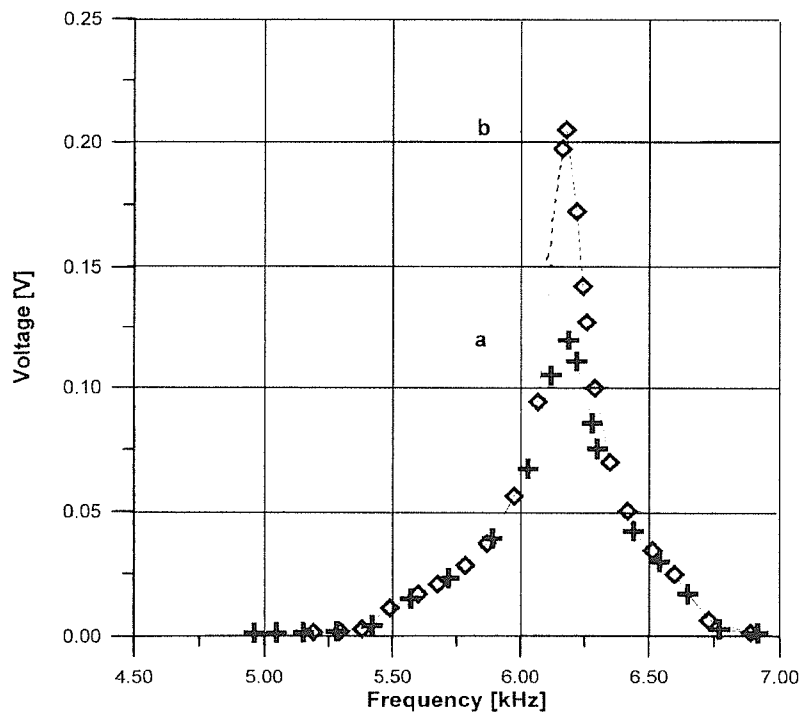


Fig. 6. The comparison of two resonant circuits: the first one consists of inductor 66 mH and capacitor 10 nF — a, the second one consists of the same inductor 66 mH and capacitance  $C_{in}$  realised by circuit I, in which the same capacitor 10 nF was used as C1

of high simulated capacitances because the dissipation factors drop in this range. Circuit II has the lowest dissipation factor and moreover its measurement dependences are regular and its frequency range starts from mHz. Hence circuit II should be considered as the best.

#### REFERENCES

1. J.L. Huertas: *Circuit implementation of current conveyor*. Electron. Lett., 1980, vol. 16, no. 6, pp. 225-226
2. J. Dunn: *Vary capacitance to positive or negative*. Electronic Design., 1991, vol. 5, p. 113
3. A.R. Al-Ali, M.T. Abuelma'atti: *Bipolar programmable capacitor*. Electronics World + Wireless World, July 1995, p. 602
4. L. Tomawski: *Variable grounded capacitances*. XX th National Conference Circuit Theory and Electronic Networks, Kołobrzeg, Poland, October 21-24, 1997, pp. 37-41
5. L. Tomawski, M. Janota: *Computer controllable band-pass filter*. Kwartalnik Elektroniki i Telekomunikacji, 1998, vol. 44, nr 1

L. TOMAWSKI

#### ZMIENNA POJEMNOŚĆ Z UŻYCIEM KONWEJORÓW PRĄDOWYCH STEROWANA PRZEZ REZYSTOR

##### Streszczenie

Opisano cztery układy wytwarzające pojemności, których wartości mogą być łatwo nastawiane za pomocą uziemionego rezystora. Komputerowa symulacja pokazuje, że przebiegi pojemności w funkcji częstotliwości są dla wszystkich układów bardzo podobne, natomiast przebiegi równoległej rezystancji i współczynnika strat są różne. Jeden z układów ma współczynnik strat niższy niż możliwy do uzyskania w fizycznych kondensatorach.

*Słowa kluczowe:* sterowane pojemności, konwejtory prądowe CCH.

## Projektowanie cyfrowego korektora fazy NOI przy użyciu metod optymalizacji nieliniowej

FELICJA WYSOCKA-SCHILLAK

*Zakład Elektroniki i Informatyki, Wyższa Szkoła Pedagogiczna, Bydgoszcz*

*Otrzymano 1998.01.26*

*Autoryzowano 1998.03.20*

W pracy przedstawiono metodę projektowania cyfrowego korektora fazy NOI przy założeniu równomiernie falistego przebiegu funkcji błędu. W metodzie tej zadanie zaprojektowania korektora fazy jest przekształcane w odpowiednie zadanie programowania nieliniowego z ograniczeniami, do rozwiązania którego jest stosowana metoda przesuwanej funkcji kary. Otrzymane zadania programowania nieliniowego bez ograniczeń rozwiązywano metodą gradientu sprzężonego Polaka-Ribiery. Praca zawiera ponadto opis opracowanego programu EQUA oraz przykład obliczeniowy ilustrujący działanie tego programu.

*Słowa kluczowe:* filtry cyfrowe, korektory fazy, optymalizacja.

### 1. WSTĘP

W wielu zastosowaniach jest wymagane, aby sygnał niesinusoidalny przy przejściu przez określony filtr nie ulegał zniekształceniu. Spełnienie tego wymagania jest możliwe w przypadku, gdy charakterystyka fazowa rozpatrywanego filtru jest liniowa w pasmie przepustowym, lub co jest równoważne, gdy charakterystyka opóźnienia grupowego jest w tym pasmie stała. Przebieg charakterystyki fazowej w pasmie zaporowym nie jest istotny ze względu na duże tłumienie filtru w tym zakresie częstotliwości [1].

Dokładnie liniową charakterystykę fazową można uzyskać stosując filtry cyfrowe o skończonej odpowiedzi impulsowej (w skrócie SOI). Istnieją różne metody projektowania tych filtrów, zarówno bez uwzględnienia, jak i z uwzględnieniem dodatkowych warunków ograniczających [1-9]. W przypadku filtrów SOI o liniowej fazie opóźnienie grupowe jest związane z długością  $N$  odpowiedzi impulsowej filtru i wynosi  $(N-1)/2$  próbek [1, 2]. Do uzyskania filtrów o dobrej selektywności

(o stromym zboczu charakterystyki amplitudowej i wąskim pasmie przejściowym) wymagane jest jednak zastosowanie dużych wartości  $N$ , co pociąga za sobą odpowiednio duże opóźnienie grupowe filtru. Ponadto, realizacja filtrów SOI w podstawowych strukturach dla dużych wartości  $N$  (rzędu kilkudziesięciu i więcej) jest nieekonomiczna [3].

Jeżeli zależy nam na uzyskaniu mniejszego opóźnienia grupowego przy jednoczesnym zachowaniu w przybliżeniu liniowej charakterystyki fazowej, alternatywnym rozwiązaniem jest zastosowanie filtru o nieskończonej odpowiedzi impulsowej (w skrócie NOI) połączonego kaskadowo z układem wszechprzepustowym pełniącym rolę korektora fazy. Odpowiednio dobierając liczbę biegunów układu wszechprzepustowego oraz ich położenie wewnątrz okręgu jednostkowego można dowolnie kształtować przebieg charakterystyki fazowej układu, nie zmieniając przy tym jego charakterystyki amplitudowej. Korektory fazy są z reguły realizowane jako kaskadowe połączenie układów pierwszego i drugiego rzędu [3].

W literaturze opisane są różne metody projektowania układów wszechprzepustowych [1-3, 10-16]. Metody te są z reguły stosunkowo skomplikowane, a ich zastosowanie jest niekiedy dość ograniczone [10]. Oparte są one najczęściej na minimalizacji błędu średniokwadratowego w dziedzinie częstotliwości lub równomiernie falistej aproksymacji stałego opóźnienia grupowego.

W tej pracy zaproponowano metodę projektowania korektora fazy przy założeniu równomiernie falistego przebiegu funkcji błędu. Zaprojektowanie korektora fazy o zadanej charakterystyce fazowej jest przekształcanie w odpowiednie zadanie programowania nieliniowego z ograniczeniami, które jest następnie rozwiązywane metodą przesuwanej funkcji kary. Artykuł zawiera również opis opracowanego programu EQUA oraz przykład obliczeniowy.

## 2. SFORMUŁOWANIE PROBLEMU

Rozważmy układ wszechprzepustowy, zwany również filtrem wszechprzepustowym, posiadający stałą, niezależną od częstotliwości charakterystykę amplitudową. W przypadku tego układu dla wszystkich częstotliwości zachodzi następująca zależność [1, 2, 16]:

$$|H_E(e^{j\omega})|^2 = 1, \quad (1)$$

gdzie:  $H_E(e^{j\omega})$  jest charakterystyką częstotliwościową (amplitudowo-fazową) układu, a  $\omega$  — pulsacją unormowaną względem częstotliwości próbkowania  $F_p$ .

Transmitancja układu wszechprzepustowego ma następującą postać [16]:

$$H_E(z) = e^{j\varphi} \prod_{k=1}^M \frac{z_k^* - z^{-1}}{1 - z_k z^{-1}}. \quad (2)$$

Należy zauważyć, że bieguny i zera tej transmitancji występują w punktach odpowiadających sprzężonym odwrotnościom określonych liczb zespolonych.



Warunek stabilności układu o transmitancji opisanej powyższym wzorem jest następujący:

$$|z_k| < 1, \quad k = 1, 2, \dots, M, \quad (3)$$

czyli wszystkie bieguny transmitancji muszą znajdować się wewnątrz okręgu jednostkowego.

Aby transmitancja  $H_E(z)$  była funkcją rzeczywistą, zespolone zera muszą być zwierciadlanym odbiciem (względem okręgu jednostkowego) biegunów, a ponadto musi zachodzić równość  $\varphi = 0$  lub  $\varphi = \pi$ . Transmitancja ta może być wówczas wyrażona w następującej postaci:

$$H_E(z) = \frac{a_M + \dots + a_1 z^{-(M-1)} + a_0 z^{-M}}{a_0 + a_1 z^{-1} + \dots + a_M z^{-M}} = \frac{z^{-M} D(z^{-1})}{D(z)}. \quad (4)$$

Należy zauważyć, że wielomian znajdujący się w liczniku jest utworzony z wielomianu stanowiącego mianownik poprzez odwrócenie porządku współczynników. Transmitancję wyrażoną powyższym wzorem wygodnie jest przedstawić w postaci iloczynu transmitancji układów 2-go rzędu o rzeczywistych współczynnikach  $c_{ik}$  [1, 14]:

$$H_E(z) = \prod_{k=1}^K \frac{1 + c_{1k} z + c_{2k} z^2}{c_{2k} + c_{1k} z + z^2}. \quad (5)$$

Kaskadowa struktura filtru realizująca taką transmitancję jest mało wrażliwa na wartości współczynników [2].

Opóźnienie grupowe  $\tau_E(\omega)$  układu wszechprzepustowego zdefiniowane jest następująco:

$$\tau_E(\omega) = -\frac{d}{d\omega} [\arg H_E(e^{j\omega})] \quad (6)$$

i może być ono wyrażone wzorem [14]:

$$\tau_E(\omega) = 2T \sum_{k=1}^K \frac{1 - c_{2k}^2 + c_{1k}(1 - c_{2k}) \cos \theta}{1 + c_{2k}^2 + c_{1k}^2 + 2c_{2k}(2 \cos^2 \theta - 1) + 2c_{1k}(1 + c_{2k}) \cos \theta} \quad (7)$$

gdzie:  $\theta = \omega T$ , a  $T = 1/F_p$  — jest okresem próbkowania.

Aby układ o transmitancji danej wzorem (5) był stabilny, współczynniki  $c_{ik}$  muszą spełniać następujące warunki [14, 17]:

$$\left. \begin{aligned} c_{2k} &< 1 \\ c_{1k} - c_{2k} &< 1 \\ c_{1k} + c_{2k} &> -1 \end{aligned} \right\} \quad k = 1, 2, \dots, K. \quad (8)$$

W przypadku stabilnego układu wszechprzepustowego  $M$ -tego rzędu o współczynnikach rzeczywistych zachodzi następująca zależność [16, 17]

$$\int_0^{\pi} \tau_E(\omega) d\omega = M\pi \quad (9)$$

czyli przy zmianie częstotliwości od zera do  $\pi$  następuje zmiana fazy od zera do  $-M\pi$  radianów.

Kaskadowe połączenie filtru NOI i układu wszechprzepustowego (pełniącego rolę korektora fazy) o odpowiednio dobranych elementach umożliwia kształtowanie charakterystyki fazowej bez zmiany charakterystyki amplitudowej filtru. Transmitancja filtru NOI może być przedstawiona w postaci:

$$H_F(z) = H \prod_{k=1}^L \frac{a_{0k} + a_{1k}z + a_{2k}z^2}{b_{0k} + b_{1k}z + b_{2k}z^2}, \quad (10)$$

gdzie:  $H > 0$  — stała,

$a_{ik}, b_{ik}, i = 0, 1, 2$  — współczynniki rzeczywiste.

Charakterystyka częstotliwościowa  $H(e^{j\omega})$  układu będącego kaskadowym połączeniem filtru NOI o charakterystyce częstotliwościowej  $H_F(e^{j\omega})$  i korektora fazy jest wyrażona następująco [1]:

$$H(e^{j\omega}) = H_F(e^{j\omega}) H_E(e^{j\omega}). \quad (11)$$

Ponieważ  $|H_E(e^{j\omega})|^2 = 1$ , zachodzą następujące zależności:

$$|H(e^{j\omega})| = |H_F(e^{j\omega})|, \quad (12)$$

$$\arg[H(e^{j\omega})] = \arg[H_F(e^{j\omega})] + \arg[H_E(e^{j\omega})], \quad (13)$$

a ponadto:

$$\tau(\omega) = \tau_F(\omega) + \tau_E(\omega), \quad (14)$$

gdzie  $\tau_F(\omega)$  i  $\tau(\omega)$  jest odpowiednio opóźnieniem grupowym filtru NOI i opóźnieniem grupowym kaskadowego połączenia filtru NOI i korektora fazy.

Zaprojektowanie korektora fazy sprowadza się do takiego doboru współczynników  $c_{ik}$  występujących we wzorze (5), aby  $\arg[H(e^{j\omega})]$  był możliwie najlepszą aproksymacją funkcji liniowej względem zmiennej  $\omega$  lub, co jest równoważne, aby  $\tau(\omega)$  stanowiło możliwie najlepszą aproksymację pewnego stałego opóźnienia grupowego  $\tau_0$ . W dalszej części pracy będziemy zajmować się aproksymacją stałego opóźnienia grupowego.

Oznaczmy przez  $\mathbf{C} = [c_{11}, c_{21}, c_{12}, c_{22}, \dots, c_{1K}, c_{2K}]^T$  wektor współczynników  $c_{ik}$ ,  $k = 1, 2, \dots, K$ ,  $i = 1, 2$ , a przez  $\tau(\omega, \mathbf{C})$  charakterystykę opóźnienia grupowego otrzymaną przy przyjęciu we wzorze (7) współczynników o wartościach określonych przez wektor  $\mathbf{C}$ . Funkcja błędu  $E(\omega, \mathbf{C})$  może być przedstawiona w postaci:

$$E(\omega, \mathbf{Y}) = \tau(\omega, \mathbf{Y}) - \tau_0. \quad (15)$$

Zadanie zaprojektowania korektora fazy sformułujemy w sposób następujący: Dla zadanej liczby  $K$ , zadanej charakterystyki opóźnienia grupowego  $\tau_F(\omega)$  filtru NOI oraz dla zadanych pulsacji krańcowych pasma przepustowego  $\omega_{p1}$ ,

$\omega_{p2}$  znaleźć wektor  $\mathbf{C}$ , dla którego funkcja błędu  $E(\omega, \mathbf{C})$  ma przebieg równomiernie falisty w pasmie przepustowym  $P = [\omega_{p1}, \omega_{p2}]$ , przy jednoczesnym spełnieniu warunków (8).

### 3. PRZEKSZTAŁCENIE PROBLEMU W ZADANIE OPTIMALIZACJI

Rozpatrywane zadanie zaprojektowania korektora fazy można przekształcić w odpowiednie zadanie optymalizacji. W tym celu, przy założonych wartościach początkowych współrzędnych wektora  $\mathbf{C}$ , podzielimy pasmo przepustowe  $P = [\omega_{p1}, \omega_{p2}]$  na podprzedziały  $\theta_j$ ,  $j = 1, 2, \dots, n$ , zdefiniowane następująco:

$$\begin{aligned}\theta_1 &= [\omega_{p1}, \omega_1) \\ \theta_k &= [\omega_{k-1}, \omega_k) \quad k = 2, 3, \dots, n-1 \\ \theta_n &= [\omega_{n-1}, \omega_{p2}],\end{aligned}\tag{16}$$

gdzie  $\omega_1 < \omega_2 < \dots < \omega_{n-1}$  są unormowanymi pulsacjami, dla których spełnione są odpowiednio warunki:

$$(-1)^{m+k} \left. \frac{d\tau(\omega, \mathbf{C})}{d\omega} \right|_{\omega=\omega_k} < 0, \quad k = 1, 2, \dots, n-1,\tag{17}$$

gdzie:

$$m = \begin{cases} 0 & \text{— gdy dla pulsacji } \omega_{p1} \text{ występuje maksimum lokalne,} \\ 1 & \text{— gdy dla pulsacji } \omega_{p1} \text{ występuje minimum lokalne.} \end{cases}$$

Podprzedziały  $\theta_j$ ,  $j = 1, 2, \dots, n$ , definiujemy tak, aby otrzymana w wyniku aproksymacji charakterystyka opóźnienia grupowego  $\tau(\omega, \mathbf{C})$  miała w każdym z nich dokładnie jedno ekstremum lokalne.

Określimy następnie wartości największych odległości  $\Delta\tau_i(\mathbf{C})$ ,  $i = 1, 2, \dots, n+2$  pomiędzy przyjętym stałym opóźnieniem grupowym  $\tau_0$  i charakterystyką  $\tau(\omega, \mathbf{C})$  w sposób następujący:

— w przypadku, gdy dla pulsacji  $\omega_{p1}$  występuje minimum lokalne

$$\Delta\tau_1(\mathbf{C}) = \tau_0 - \tau(\omega_{p1}, \mathbf{C});\tag{18}$$

$$\Delta\tau_i(\mathbf{C}) = \max_{\omega \in \theta_{i-1}} (\tau(\omega, \mathbf{C}) - \tau_0), \quad i = 2, 4, 6, \dots, (i \leq n+1);\tag{19}$$

$$\Delta\tau_i(\mathbf{C}) = \max_{\omega \in \theta_{i-1}} (\tau_0 - \tau(\omega, \mathbf{C})), \quad i = 3, 5, 7, \dots, (i \leq n+1);\tag{20}$$

$$\Delta\tau_{n+2}(\mathbf{C}) = (-1)^n (\tau(\omega_{p2}, \mathbf{C}) - \tau_0);\tag{21}$$

— w przypadku, gdy dla pulsacji  $\omega_{p1}$  występuje maksimum lokalne

$$\Delta\tau_1(\mathbf{C}) = \tau(\omega_{p1}, \mathbf{C}) - \tau_0;\tag{22}$$

$$\Delta\tau_i(\mathbf{C}) = \max_{\omega \in \theta_{i-1}} (\tau_0 - \tau(\omega, \mathbf{C})), \quad i = 2, 4, 6, \dots, (i \leq n+1); \quad (23)$$

$$\Delta\tau_i(\mathbf{C}) = \max_{\omega \in \theta_{i-1}} (\tau(\omega, \mathbf{C}) - \tau_0), \quad i = 3, 5, 7, \dots, (i \leq n+1); \quad (24)$$

$$\Delta\tau_{n+2}(\mathbf{C}) = (-1)^n (\tau_0 - \tau(\omega_{p2}, \mathbf{C})). \quad (25)$$

Należy zauważyć, że wartość  $\tau_0$  nie może być przyjmowana zupełnie dowolnie. Maksymalna wartość opóźnienia grupowego  $\tau_0^{\max}$  w przypadku korektora fazy połączonego z filtrem NOI o pasmie przepustowym  $P = [\omega_{p1}, \omega_{p2}]$  jest określona zależnością [10, 11]:

$$\tau_0^{\max} = \frac{M\pi - [\Theta(\omega_{p2}) - \Theta(\omega_{p1})]}{\omega_{p2} - \omega_{p1}}, \quad (26)$$

gdzie:  $\Theta(\omega_{p1}) = \arg[H_F(e^{j\omega_{p1}})]$  i  $\Theta(\omega_{p2}) = \arg[H_F(e^{j\omega_{p2}})]$ .

Przyjęcie zbyt małej wartości  $\tau_0$  prowadzi do rozwiązania, które odpowiada układowi niestabilnemu [15]. Zgodnie za wskazówkami podanymi w pracy [11], jako początkowe przybliżenie przyjęto  $\tau_0 = 0,8 \tau_0^{\max}$ .

Zadanie zaprojektowania korektora fazy przy założeniu równomiernie falistego przebiegu funkcji błędu, zdefiniowanej wzorem (15), czyli zadanie równomiernie falistej aproksymacji stałego opóźnienia grupowego  $\tau_0$ , sprowadza się do znalezienia takiego wektora  $\mathbf{C} = [c_{11}, c_{21}, c_{12}, c_{22}, \dots, c_{1K}, c_{2K}]^T$ , dla którego spełniony jest warunek

$$\Delta\tau_i(\mathbf{C}) = \Delta\tau_k(\mathbf{C}), \quad k, i = 1, 2, \dots, n+2. \quad (27)$$

W celu rozwiązania sformułowanego powyżej zadania wprowadźmy funkcję celu  $\tilde{X}(\Delta\tau_1, \Delta\tau_2, \dots, \Delta\tau_{n+2})$  o następującej postaci:

$$\tilde{X}(\Delta\tau_1, \Delta\tau_2, \dots, \Delta\tau_{n+2}) = \sum_{i=1}^{n+2} (\Delta\tau_i - \tilde{S})^2, \quad (28)$$

przy czym

$$\tilde{S} = \frac{1}{n+2} \sum_{i=1}^{n+2} \Delta\tau_i \quad (29)$$

jest średnią arytmetyczną odległości  $\Delta\tau_i$ .

Tak określona funkcja  $\tilde{X}$ :

- jest nieujemną funkcją ciągłą zmiennych  $\Delta\tau_1, \Delta\tau_2, \dots, \Delta\tau_{n+2}$ ,
- przyjmuje wartość zero gdy  $\Delta\tau_1 = \Delta\tau_2 = \dots = \Delta\tau_{n+2}$ ,
- nie posiada lokalnych minimów.

Ponieważ odległości  $\Delta\tau_1, \Delta\tau_2, \dots, \Delta\tau_{n+2}$  są funkcjami wektora  $\mathbf{C}$ , więc równoważne zadanie optymalizacji z ograniczeniami ma postać:

Znaleźć wektor  $\mathbf{C}$ , który minimalizuje funkcję celu  $X(\mathbf{C}) = \tilde{X}(\Delta\tau_1(\mathbf{C}), \Delta\tau_2(\mathbf{C}), \dots, \Delta\tau_{n+2}(\mathbf{C}))$ , przy warunkach ograniczających zdefiniowanych zależnościami (8) oraz

$$\Delta\tau_k(\mathbf{C}) > 0, \quad k = 1, 2, \dots, n + 2. \quad (30)$$

Sformułowane powyżej zadanie optymalizacji można rozwiązać jedną z metod poszukiwania minimum z ograniczeniami. Z podanych w literaturze [18, 19] metod, do rozwiązywania rozpatrywanego zadania wybrano metodę przesuwanej funkcji kary. Metoda ta uważana jest bowiem za jedną z najbardziej efektywnych metod poszukiwania minimum z ograniczeniami. Umożliwia ona sprowadzenie zadania z ograniczeniami do ciągu zadań bez ograniczeń poprzez odpowiednią modyfikację funkcji celu. Zadania programowania nieliniowego bez ograniczeń rozwiązywano metodą gradientu sprzężonego Polaka-Ribiery [18, 19].

Należy zauważyć, że dla określonego przyjętego zbioru danych wejściowych i przy uwzględnieniu ograniczeń (8), zadanie równomiernie falistej aproksymacji stałego opóźnienia grupowego może nie mieć rozwiązania. Ponadto należy zwrócić uwagę na fakt, że w przypadku układów wszechprzepustowych rozwiązanie zadania równomiernie falistej aproksymacji stałego opóźnienia grupowego nie jest równoznaczne z uzyskaniem rozwiązania zadania aproksymacji w sensie czebyszewowskim. Szczegółowe rozważania na ten temat można znaleźć w literaturze [12]. W pracy tej podany jest również przykład rozwiązania zadania czebyszewowskiej aproksymacji stałego opóźnienia grupowego, w przypadku którego funkcja aproksymująca nie ma przebiegu równomiernie falistego.

Przy rozwiązywaniu zadań optymalizacji, podobnie jak przy rozwiązywaniu innych zagadnień metodami numerycznymi, istotnym problemem jest wybór odpowiedniego punktu startowego. Punkt ten powinien być w ogólnym przypadku niezbyt odległy od poszukiwanego rozwiązania, gdyż przyjęcie zbyt odległego punktu startowego może niekiedy powodować numeryczną rozbieżność algorytmu. W rozpatrywanej procedurze projektowania korektora fazy, dla przyjętego punktu startowego charakterystyka  $\tau(\omega, \mathbf{C})$  powinna mieć występujące na przemian po sobie minima i maksima lokalne, co jest niezbędne do określenia przedziałów  $\theta_j$ ,  $j = 1, 2, \dots, n$ . Wyznaczenie takiego punktu startowego stanowi więc pierwszy, istotny krok w procedurze projektowania.

Jako punkt startowy przyjęto wektor  $\tilde{\mathbf{C}}$  będący rozwiązaniem zadania minimalizacji błędu średniokwadratowego  $E_2(\omega, \mathbf{C})$ . Błąd ten jest zdefiniowany zależnością:

$$E_2(\omega, \mathbf{C}) = \sum_{j=1}^J [\tau_0 - \tau(\omega_j, \mathbf{C})]^2, \quad (31)$$

gdzie  $\omega_j$  jest jedną z  $J$  równoodległych wartości pulsacji z przedziału  $P = [\omega_{p1}, \omega_{p2}]$ .

Zadanie minimalizacji błędu  $E_2(\omega, \mathbf{C})$  jest zadaniem optymalizacji nieliniowej bez ograniczeń, które również rozwiązywano metodą gradientu sprzężonego Polaka-Ribiery. Przeprowadzone próby wykazały, że procedura rozwiązywania tego zadania jest szybko zbieżna.

Należy zauważyć, że wyznaczony wektor  $\tilde{\mathbf{C}}$  jest niezbyt odległy od poszukiwanego rozwiązania zadania równomiernie falistej aproksymacji stałego opóźnienia

grupowego  $\tau_0$ . Znając współrzędne tego wektora można następnie przystąpić do drugiego kroku procedury projektowania, czyli do rozwiązania zadania minimalizacji funkcji  $X(C)$  określonej wzorem (28).

Do określenia wartości funkcji  $X(C)$  konieczna jest znajomość wartości największych odległości  $\Delta\tau_i(C)$  pomiędzy stałym opóźnieniem grupowym  $\tau_0$  i charakterystyką  $\tau(\omega, C)$ . Wartości tych odległości w przedziałach  $\theta_j, j = 1, 2, \dots, n$ , wyznaczano posługując się metodą złotego podziału [18, 19]. Metoda ta umożliwia znalezienie odciętej  $a$  oraz rzędnej  $f(a)$  minimum rozpatrywanej funkcji  $f(a)$  w określonym przedziale, przy założeniu, że w tym przedziale funkcja  $f(a)$  jest ciągła i ma dokładnie jedno minimum.

#### 4. OPIS PROGRAMU EQUA

Na podstawie zaproponowanej metody projektowania korektorów fazy został opracowany program EQUA. Program ten umożliwia zaprojektowanie korektora fazy, którego kaskadowe połączenie z filtrem NOI o zadanej charakterystyce częstotliwościowej  $H_F(e^{j\omega})$ , zapewnia otrzymanie w przybliżeniu stałej charakterystyki opóźnienia grupowego  $\tau(\omega, C)$  układu, przy założeniu równomiernie falistego przebiegu funkcji błędu. Program EQUA został napisany w języku Fortran 77, przy czym do kompilacji wykorzystano kompilator PowerStation firmy Microsoft (wersja 4.0).

Danymi wejściowymi dla programu są:

- liczba  $L$  oraz wartości współczynników  $a_{ik}$  i  $b_{ik}, k = 1, 2, \dots, L, i = 0, 1, 2$ , występujących we wzorze (10) określającym transmitancję filtru NOI;
- liczba  $K$ ;
- wartości unormowanych częstotliwości  $f_{p1} = \omega_{p1}/2\pi$  oraz  $f_{p2} = \omega_{p2}/2\pi$ ;
- wartość parametru IW informującego, czy chcemy otrzymać wydruki kolejnych współrzędnych wyznaczonej charakterystyki opóźnienia grupowego  $\tau(\omega, C)$  oraz funkcji błędu  $E(\omega, C)$  (IW = 1), czy też nie (IW = 0);
- wartość parametru EPS, zwanego kryterium zakończenia procedury minimalizacji bez ograniczeń [19] (procedura ta kończy swoje działanie, gdy  $\|\nabla X(C)\| < \text{EPS}$ );
- wartość parametru IT informującego, czy podajemy wartość  $\tau_0$  z klawiatury (IT = 1), czy też ma być ona wyznaczana przez program według wzoru (26).

W programie EQUA wykorzystano odpowiednio zmodyfikowaną procedurę GSPR1E opracowaną w Instytucie Automatyki Politechniki Warszawskiej [19]. Procedura ta umożliwia rozwiązywanie zadań programowania nieliniowego bez ograniczeń metodą gradientu sprzężonego Polaka-Ribiery, przy czym gradient funkcji celu jest estymowany numerycznie.

Po przeprowadzeniu obliczeń przy użyciu programu EQUA jako wyniki końcowe otrzymujemy wartości współczynników  $c_{ik}, k = 1, \dots, 2K, i = 1, 2$ , występujących we wzorze (5).

Istnieje ponadto możliwość drukowania:

- kolejnych wartości współrzędnych wyznaczonej charakterystyki opóźnienia grupowego  $\tau(\omega, C)$  oraz funkcji błędu  $E(\omega, C)$ ;

- wartości współrzędnych ekstremów wyznaczonej funkcji  $E(\omega, C)$  oraz wartości największych odległości  $\Delta\tau_i(C)$ ,  $i = 1, 2, \dots, n+2$ .

Przy użyciu programu EQUA przeprowadzono projektowanie szeregu korektorów fazy przy założeniu równomiernie falistego przebiegu funkcji błędu  $E(\omega, C)$ . Obliczenia wykonywano dla  $K < 6$  przy użyciu mikrokomputera IBM PC Pentium z zegarem 120 MHz. Czas wykonywania obliczeń wahał się od kilkunastu do kilkudziesięciu sekund, przy czym czas ten wyraźnie zwiększał się w miarę wzrostu liczby  $K$  oraz w miarę zmniejszania się wartości parametru EPS. Działanie programu zilustrujemy następującym przykładem.

**Przykład.** Dany jest eliptyczny filtr cyfrowy NOI 6-tego rzędu o transmitancji posiadającej współczynniki podane w tablicy 1 [14]. Pasma przepustowe filtru mieści się w przedziale częstotliwości od  $f_{p1} = 0$  do  $f_{p2} = 0,175$ . Zaprojektować korektor fazy, którego kaskadowe połączenie z danym filtrem umożliwi otrzymanie w przybliżeniu stałej charakterystyki opóźnienia grupowego. Korektor ten ma składać się z trzech układów drugiego rzędu.

Tablica 1

Wartości współczynników transmitancji filtru NOI

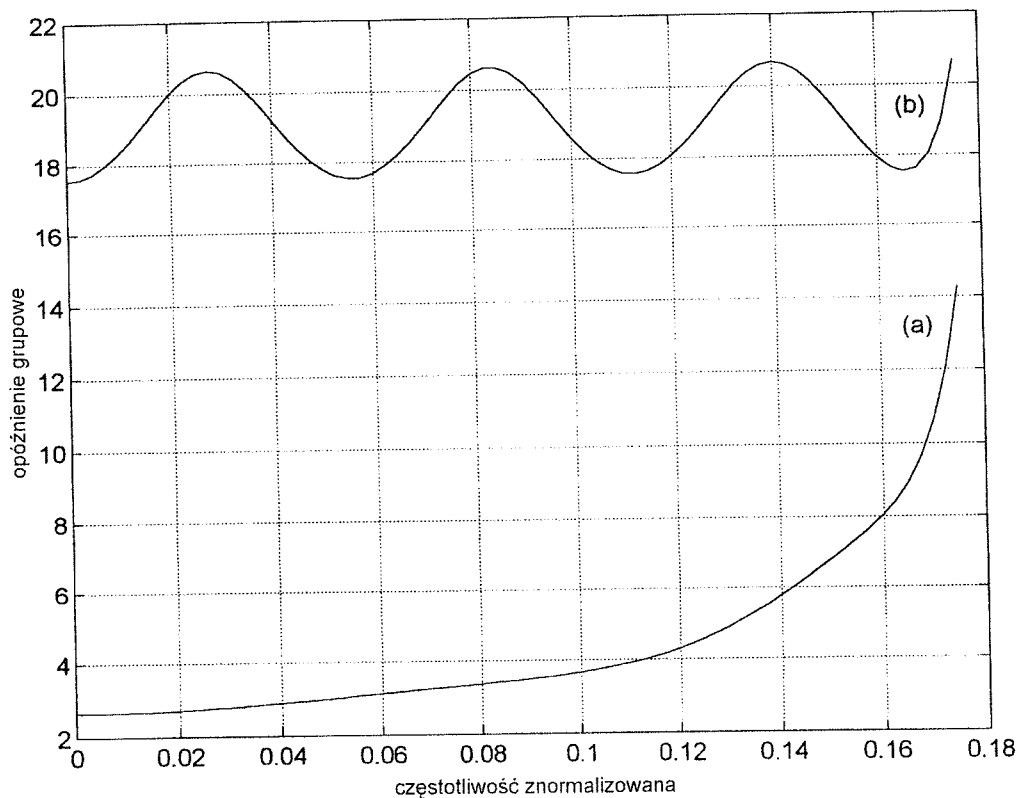
| $k$ | $a_{0k}$ | $a_{1k}$   | $a_{2k}$ | $b_{0k}$  | $b_{1k}$   | $b_{2k}$ |
|-----|----------|------------|----------|-----------|------------|----------|
| 1   | 4,681699 | 7,0020200  | 4,681699 | 0,3378687 | -1,0235090 | 1,0      |
| 2   | 1,328189 | 0,1544323  | 1,328189 | 0,6436865 | -0,8582590 | 1,0      |
| 3   | 1,144349 | -0,3799679 | 1,144349 | 0,9002355 | -0,7684310 | 1,0      |

$$H = 0,3110289$$

Dla powyższych danych przeprowadzono projektowanie korektora fazy przy użyciu programu EQUA. Do obliczeń przyjęto wartość parametru EPS =  $10^{-6}$ . Otrzymane wyniki projektowania są podane w tablicy 2. Przebieg charakterystyki opóźnienia grupowego  $\tau_F(\omega)$  rozpatrywanego filtru NOI oraz przebieg charakterystyki  $\tau(\omega, C)$  kaskadowego połączenia filtru NOI i zaprojektowanego korektora fazy są pokazane na rysunku 1.

W przypadku, gdy wartość odległości  $\Delta\tau_{n+2}(C)$  dla pulsacji  $\omega_{p2}$  nie jest dla nas szczególnie istotna i może być nieco większa od pozostałych odległości, można uzyskać wyraźnie mniejsze odległości  $\Delta\tau_i(C)$ ,  $i = 1, 2, \dots, n+1$  w pozostałej części pasma przepustowego niż w przypadku aproksymacji równomiernie falistej. Aby przeprowadzić projektowanie korektora dla takiego przypadku, we wzorze (28) należy przeprowadzić sumowanie odległości  $\Delta\tau_i(C)$  nie do  $n+2$ , a do  $n+1$ , a wartość  $\tilde{S}$  przyjmie wówczas postać następującą:

$$\tilde{S} = \frac{1}{n+1} \sum_{i=1}^{n+1} \Delta\tau_i(C). \quad (32)$$

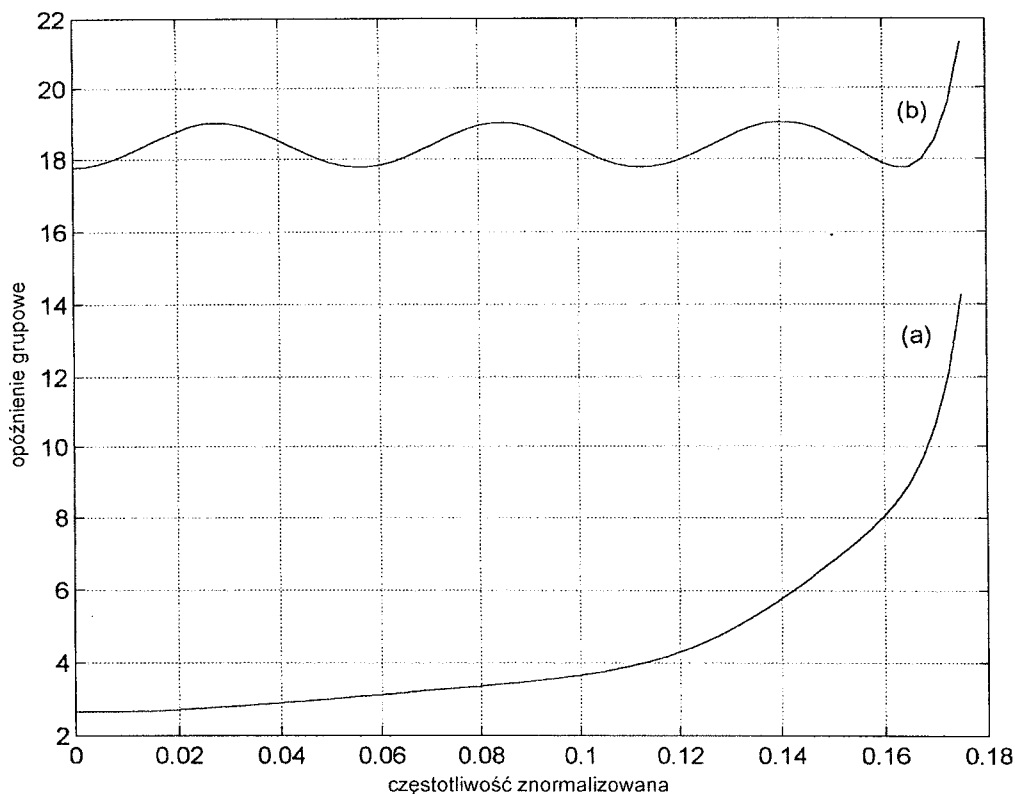


Rys. 1. Przebieg charakterystyki opóźnienia grupowego  $\tau_f(\omega)$  rozpatrywanego filtra NOI (a) oraz charakterystyki  $\tau(\omega, C)$  kaskadowego połączenia filtra NOI i zaprojektowanego korektora fazy (b)

Tablica 2

| Otrzymane wyniki obliczeń                       |                    |                                  |
|-------------------------------------------------|--------------------|----------------------------------|
| Wartości współczynników transmitancji korektora |                    |                                  |
| c11 = -.10694410E+01                            |                    |                                  |
| c21 = .69875940E+00                             |                    |                                  |
| c12 = -.16540130E+01                            |                    |                                  |
| c22 = .70493810E+00                             |                    |                                  |
| c13 = -.14475140E+01                            |                    |                                  |
| c23 = .70180620E+00                             |                    |                                  |
| Częstotliwości ekstremów                        | Wartości ekstremów | Wartości największych odległości |
| .00000000                                       | 17.53934           | 1.551177                         |
| .05555738                                       | 20.64169           | 1.551178                         |
| .11147760                                       | 17.53293           | 1.557588                         |
| .16797880                                       | 20.62525           | 1.534739                         |
| .22384700                                       | 17.55086           | 1.539651                         |
| .28030480                                       | 20.66103           | 1.570511                         |
| .33081880                                       | 17.52216           | 1.568353                         |
| .35000000                                       | 20.66413           | 1.573610                         |





Rys. 2. Przebieg charakterystyki opóźnienia grupowego  $\tau_F(\omega)$  rozpatrywanego filtra NOI (a) oraz charakterystyki  $\tau(\omega, C)$  kaskadowego połączenia filtra NOI i korektora fazy dla przypadku sumowania  $n + 1$  odległości  $\Delta\tau_i(C)$  (b)

Przebieg charakterystyki opóźnienia grupowego  $\tau(\omega, C)$  kaskadowego połączenia filtra NOI i zaprojektowanego w ten sposób korektora fazy są pokazane na rysunku 2. Jak widać na tym rysunku, w znacznej części pasma przepustowego charakterystyka  $\tau(\omega, C)$  ma mniejsze zafalowania, czyli stanowi lepsze przybliżenie stałego opóźnienia grupowego, niż w przypadku aproksymacji równomiernie falistej w całym pasmie przepustowym.

## 5. ZAKOŃCZENIE

W pracy przedstawiono metodę projektowania korektora fazy umożliwiającą uzyskanie równomiernie falistego przebiegu funkcji błędu. W przeciwieństwie do wielu metod opisanych w literaturze, w metodzie tej warunki stabilności korektora są bezpośrednio uwzględniane w procesie projektowania. Spełnienie tych warunków

dla określonego zbioru współczynników jest mianowicie sprawdzane w trakcie wykonywania obliczeń i gdy nie są one spełnione, do funkcji celu wprowadzany jest odpowiedni człon reprezentujący karę za przekroczenie ograniczeń. Jeżeli więc dla określonych danych wejściowych otrzymamy rozwiązanie rozpatrywanego zadania, odpowiada ono korektorowi stabilnemu. W metodach opisanych w pracach [10, 22, 13] nie jest możliwe uwzględnianie warunków stabilności w trakcie wykonywania obliczeń.

Zaletą zaproponowanej metody jest również jej elastyczność. Jeżeli w części pasma przepustowego bliskiej  $\omega_{p2}$  może wystąpić nieco większe zafalowanie charakterystyki opóźnienia grupowego niż w pozostałej części pasma przepustowego, to zamiast  $n + 2$  wyrównujemy jedynie  $n + 1$  odległości  $\Delta\tau_i(C)$ , w wyniku czego otrzymujemy mniejsze zafalowania charakterystyki opóźnienia grupowego przy niższych częstotliwościach, niż w przypadku aproksymacji równomiernie falistej. Postępowanie takie nie jest możliwe w przypadku metod opisanych w literaturze. Ponadto, w przypadku zaproponowanej metody możliwe jest stosunkowo łatwe uwzględnienie, oprócz warunków stabilności, również ewentualnych innych dodatkowych warunków, zarówno liniowych, jak i nieliniowych, poprzez wprowadzenie odpowiednich ograniczeń w zadaniu optymalizacji. Zaproponowana metoda, po wprowadzeniu niewielkich modyfikacji, może zostać również wykorzystana do projektowania układów wszechprzepustowych, których charakterystyka  $\tau(\omega, C)$  aproksymuje pewną zadaną charakterystykę opóźnienia grupowego  $\tau(\omega) \neq \text{const}$ .

#### BIBLIOGRAFIA

1. S.K. Mitra, J.F. Kaiser (ed.): *Handbook for Digital Signal Processing*. John Wiley & Sons, Inc., New York 1993
2. A.V. Oppenheim, R.W. Schaffer: *Cyfrowe przetwarzanie sygnałów*. WKŁ, Warszawa 1979
3. A. Wojtkiewicz: *Elementy syntezy filtrów cyfrowych*. WNT, Warszawa 1984
4. K. Steiglitz, T.W. Parks, J.F. Kaiser: *METEOR: A constraint-based FIR filter design program*. IEEE Trans. Signal Processing, 1992, vol. 40, No. 8, pp. 1901-1909
5. J.K. Liang, R.J.P. De Figueiredo, F.C. Liu: *Design of optimal Nyquist, partial response, N-th band, and nonuniform tap spacing FIR digital filters using linear programming techniques*. IEEE Trans. Circuits and Syst., 1985, vol. CAS-32, No. 4, pp. 386-392
6. P.P. Vaidyanathan, T.Q. Nguyen: *Eigenfilters: a new approach to least-squares FIR filter design and applications including Nyquist filters*. IEEE Trans. Circuits and Syst., 1987, vol. CAS-34, No. 1, pp. 11-22
7. F. Wysocka-Schillak: *Projektowanie cyfrowego filtru SOI o liniowej fazie z uwzględnieniem wymagań dotyczących jego odpowiedzi jednostkowej*. Kwartalnik Elektroniki i Telekomunikacji, 1994, t. 40, z. 2, ss. 185-199
8. F. Wysocka-Schillak, T. Wysocki: *Design of FIR digital filters with monotone pass-band and equiripple stop-band frequency response*. Proceedings of International Symposium on Information Theory & Its Applications, Sydney, Australia, 20-24 November 1994, vol. 1, pp. 331-334
9. H. Savoji, F. Wysocka-Schillak: *Design of FIR filters with extra constraints using the modified LMS algorithm*. IEE Proceedings -- Vision, Image and Signal Processing, 1995, Vol. 142, No. 4, pp. 237-240

10. M. Lang, T.I. Laakso: *Simple and robust method for the design of allpass filters using least-squares phase error criterion*. IEEE Trans. Circuits and Syst. — II: Analog and Digital Signal Processing, 1994, vol. 41, No. 1, pp. 40-47
11. T.Q. Nguyen, T.I. Laakso, R.D. Koilpillai: *Eigenfilter approach for the design of allpass filters approximating a given phase response*. IEEE Trans. Signal Processing, 1994, vol. 42, No. 9, pp. 2257-2263
12. A.G. Deczky: *Equiripple and minimax (Chebyshev) approximations for recursive digital filters*. IEEE Trans. Acoust., Speech, Signal Processing, 1974, vol. ASSP-22, No. 2, pp. 98-112
13. M. Ikehara, M. Funaiishi, H. Kuroda: *Design of complex all-pass networks using Remez algorithm*. IEEE Trans. Circuits and Syst. — II: Analog and Digital Signal Processing, 1992, vol. 39, No. 8, pp. 549-555
14. C. Charalambous, A. Antoniou: *Equalisation of recursive digital filters*. IEE Proc. 1980, Vol. 127, Pt. G, No. 5, pp. 219-225
15. Z. Jing: *A new method for allpass digital filter design*. IEEE Trans. Acoust., Speech, Signal Processing, 1987, vol. ASSP-35, No. 11, pp. 1557-1564
16. P.A. Regalia, S.K. Mitra, P.P. Vaidyanathan: *The digital all-pass filter: A versatile signal processing building block*. Proc. IEEE, 1988, vol. 76, No. 1, pp. 19-37
17. A.T. Chottera, G. Jullien: *A linear programming approach to recursive digital filter design with linear phase*. IEEE Trans. Circuits and Syst., 1982, vol. CAS-29, No. 3, pp. 139-149
18. W. Findeisen, J. Szymanowski, A. Wierzbiński: *Teoria i metody obliczeniowe optymalizacji*. PWN, Warszawa 1980
19. T. Kręglewski, T. Rogowski, A. Ruszczyński, J. Szymanowski: *Metody optymalizacji w języku FORTRAN*, PWN, Warszawa 1984

F. WYSOCKA-SCHILLAK

## DESIGN OF IIR DELAY EQUALIZERS USING NONLINEAR OPTIMIZATION METHODS

### S u m m a r y

A method is presented for designing IIR delay equalizers using the equiripple error criterion. In the method, the considered filter design problem is transformed into an equivalent constrained minimization problem, which is solved using a penalty function shifting method. Unconstrained minimization is performed by means of a conjugate gradient method. A description of the computer program EQUA is included. Finally, a design example illustrating the application of the method is given.

*Key words:* digital filters, delay equalizers, optimization



## Sprzętowe architektury kompresji obrazu wizyjnego dla potrzeb przetwarzania w czasie rzeczywistym

KAZIMIERZ WIATR, PAWEŁ RUSSEK

*Instytut Elektroniki, Akademia Górniczo-Hutnicza w Krakowie*

*Otrzymano 1997.10.30*

*Autoryzowano 1998.03.26*

W artykule przedstawiono problemy związane z kompresją obrazów wideo w czasie rzeczywistym. W szczególności podjęto próbę określenia wymagań jakie stawiają szeroko stosowane algorytmy kompresji pod względem mocy obliczeniowych, jak również ukazania technik stosowanych obecnie w celu ich spełnienia. Rozważania dotyczą algorytmów kompresji opartych na transformacjach w dziedzinę częstotliwości, szeroko stosowanych w standardach H.261, MPEG oraz nie będącej standardem CCITT metodzie kompresji falkowej. Artykuł ukazuje możliwości związane z zastosowaniem procesorów ogólnego przeznaczenia RISC oraz DSP, wskazując na konieczność stosowania architektur dedykowanych do tych zastosowań. Podane są również przykłady realizacji takich architektur dla zadań związanych z estymacją ruchu i przekształceń typu częstotliwościowego. Zaprezentowane sprzętowe architektury kompresji obrazu będą, w przygotowywanym przez autorów projekcie, implementowane w programowalnych strukturach FPGA o bardzo dużych pojemnościach.

*Słowa kluczowe:* systemy czasu rzeczywistego, przetwarzanie obrazów, kompresja obrazów, układy programowalne

### 1. WSTĘP

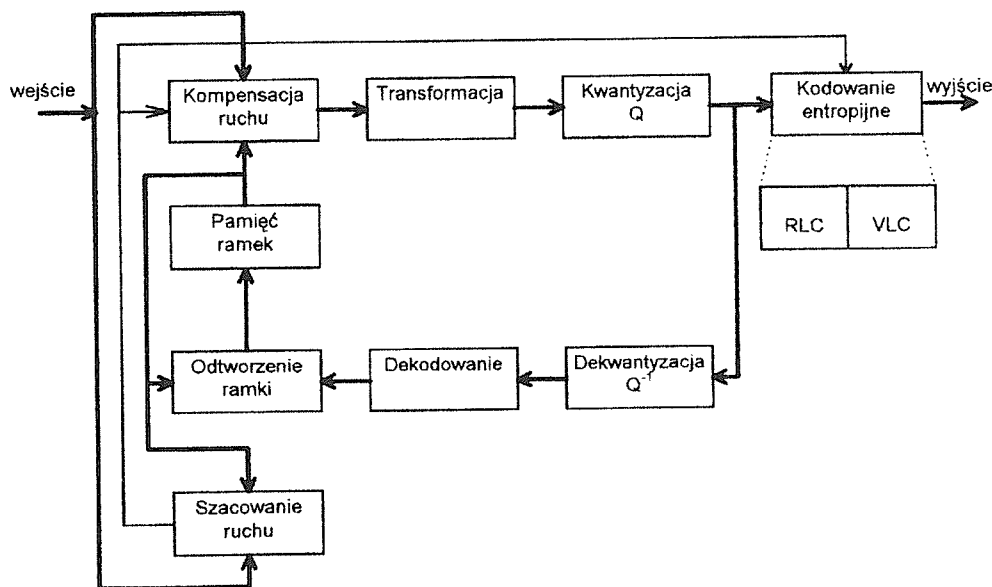
Obecny rozwój techniki cyfrowej, a w szczególności cyfrowej techniki wideo daje zagadnieniom cyfrowej kompresji obrazu szerokie pole zastosowań. Wymienić tu można telewizję cyfrową, wideokonferencje, VOD (ang. *Video On Demand*), DVD (ang. *Digital Video Disc*). Wiąże się to oczywiście z potrzebą poszukiwania tanich rozwiązań, które równocześnie zdolne byłyby sprostać wymaganiom czasowym jakie stawiają metody kompresji. Wymagania te powodują, że w celu efektywnej kompresji stosować należy układy specjalizowane o dedykowanej architekturze. Stoso-

wanie procesorów ogólnego przeznaczenia jest nieefektywne ze względu na konieczność równoległego stosowania dużej ich liczby. Masowe zastosowanie otwiera szansę na implementowanie dedykowanych architektur w technologii VLSI. Lawinowo pojawiające się ostatnio propozycje układów półprzewodnikowych do zastosowań w tej dziedzinie wydają się to potwierdzać. Stosowane architektury wykazują wiele podobieństw; możliwe jest więc ich ogólne omówienie.

## 2. UOGÓLNIONY ALGORYTM KOMPRESJI OBRAZU Z ZASTOSOWANIEM TRANSFORMACJI W DZIEDZINĘ CZĘSTOTLIWOŚCI

Prowadzone od kilku lat badania ciągle owocują propozycjami nowych algorytmów kompresji. Od najprostszych, łatwych w zastosowaniu, nie dających jednak dużych współczynników kompresji, do bardzo złożonych, które co prawda cechuje duża efektywność działania, ale ze względu na koszty implementacyjne nie nadają się obecnie do szerokiego stosowania. Metodami szeroko obecnie wykorzystywanymi praktycznie są różnego rodzaju metody oparte na transformacji informacji obrazowej w dziedzinę częstotliwości.

Metody te polegają na spostrzeżeniu, że ilość niesionej informacji zapisanej w obrazie maleje ze wzrostem rozważanej częstotliwości. Oko ludzkie wykazuje szybki spadek czułości przy rozpoznawaniu różnic jaskrawości dla większych częstotliwości. Zatem poddanie obrazu filtracji dolnoprzepustowej, w przeciwieństwie



Rys. 1. Schemat blokowy ogólnego algorytmu kompresji obrazów wideo

do filtracji górnoprzepustowej, nie powoduje takiej degradacji jego treści, że nie jest możliwe rozpoznanie zawartej w nim treści. Jeżeli więc wydzielić poszczególne składowe częstotliwościowe informacji obrazowej, to można ograniczać, lub wręcz pomijać, informacje wysokoczęstotliwościowe. Ograniczać można w ten sposób ilość koniecznych do zakodowania danych.

Rys. 1 przedstawia schemat blokowy uogólnionego układu kodera bazującego na podanej wyżej zasadzie. Wspomniana transformacja do dziedziny częstotliwości wykonywana jest w bloku TRANSFORMACJA. Następnie dane zostają poddane procesowi kwantyzacji w bloku KWANTYZACJA. Współczynniki opisujące obraz częstotliwościowo są kwantowane z różną dokładnością. Zgodnie z tym co napisano powyżej, im większej częstotliwości dany współczynnik odpowiada tym mniejsza precyzja przypisana mu w procesie kwantowania. Indeksy poziomów kwantyzacji, wygenerowane przez blok KWANTYZACJA, tworzą ciąg z dużą ilością zer. Ciąg ten jest dalej kodowany w bloku KODOWANIE ENTROPIJNE, realizującego operację dwustopniową: RLC (ang. *run-length coding*) i VLC (ang. *variable-length coding*). RLC koder kompresje strumień danych zastępując ciągi zer ich długością, VLC zastępuje symbole występujące w ciągu słowami kodowymi o zmiennej liczbie bitów zgodnie z zasadą, aby symbolom często powtarzającym się przypisać krótkie słowo kodowe, a rzadko występującym odpowiednio dłuższe.

Podczas kodowania sekwencji wideo kolejne kodowane ramki są podobne. Różnice najczęściej zaznaczają się przesunięciem całej lub części prezentowanej sceny. Dlatego opisane operacje transformacji, kwantowania, itd. wykonuje się faktycznie nie na poszczególnych obrazach, ale na różnicy powstającej z odjęcia dwóch następujących po sobie obrazów. Jeżeli odpowiednio skorygować wzajemne przesunięcia, okaże się, że energia sygnału różnicy obrazów jest znacznie mniejsza niż energia zawarta w obrazie oryginalnym. Prowadzi to do zmniejszenia liczby bitów niezbędnych do zakodowania sekwencji obrazów. Odpowiednie oszacowanie przesunięcia wykonywane jest w bloku SZACOWANIE RUCHU, a odjęcie poszczególnych ramek w bloku KOMPENSACJA RUCHU.

### 3. POTRZEBNE MOCE OBLICZENIOWE

Szanse na szerokie stosowanie mają te metody kompresji obrazu, których złożoność algorytmów, a co za tym idzie wymagana moc obliczeniowa, odpowiadają obecnemu stanowi techniki cyfrowej. Możliwości obliczeniowe stosowanych komputerów są jednym z głównych czynników branych pod uwagę przy opracowywaniu nowych standardów kompresji obrazu. Jest wiele miar złożoności algorytmów. Dla cyfrowej kompresji obrazów wideo przyjęto miarę ilości koniecznych do jej wykonania operacji procesora typu RISC: MOPS (ang. *million operations per second*) lub GOPS (ang. *giga operations per second*).

Dla przykładu obliczmy ilość operacji RISC koniecznych do obliczenia szeroko stosowanej przez algorytmy kompresji transformaty DCT bloku  $8 \times 8$ . Transformatę

taką można obliczyć ze wzoru.

$$Y = T * X * T^{-1}$$

gdzie:  $Y$  — wynikowa macierz współczynników,  $X$  — wejściowa macierz wartości pixeli obrazu,  $T$  — macierz prosta współczynników charakterystycznych transformacji DCT,  $T^{-1}$  — macierz odwrócona  $T$ .

Jeżeli przyjąć model procesora RISC, w którym operacje arytmetyczne wykonywane są na rejestrach i dysponującego operującymi na pamięci zewnętrznej instrukcjami typu ŁADUJ i ZAPAMIĘTAJ, to aby policzyć pojedynczy współczynnik macierzy powstającej z pomnożenia dwóch macierzy  $8 \times 8$ , z których współczynniki jednej są w rejestrach procesora (macierz  $T$ ), a drugiej w pamięci zewnętrznej (macierz  $X$ ) należy wykonać: 8 instrukcji typu ŁADUJ, 8 dodawań, 8 mnożeń i jedną operację ZAPAMIĘTAJ. Łącznie 25 instrukcji typu RISC dla jednego elementu macierzy. Obliczenie całej macierzy  $Y$  wymaga zatem  $2 \cdot 8 \cdot 8 \cdot 25 = 3\,200$  instrukcji.

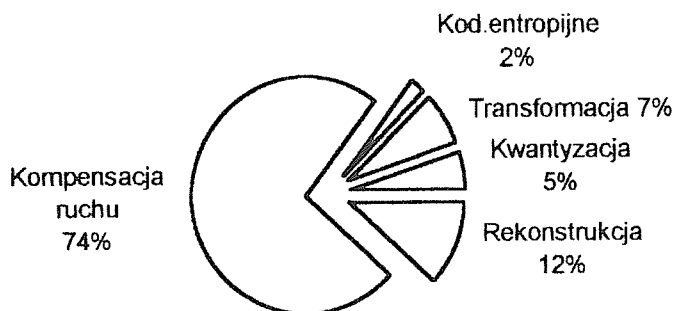
Standard formatu CIF ramki zakłada rozdzielczość  $352 \cdot 288$  punktów na ramkę o formacie 4:2:0 (Luminancja  $Y$  —  $352 \cdot 288$  punktów; Chrominancja  $Cr$  i  $Cb$  —  $176 \cdot 288$  punktów) i 30 ramek na sekundę. Oznacza to, że aby zakodować pojedynczą ramkę należy wykonać 2376 operacji DCT na bloczkach  $8 \times 8$ . Przy 30 klatkach na sekundę konieczna moc obliczeniowa dla przeprowadzenia samej tylko transformacji wynosi  $3200 \cdot 2376 \cdot 30$  (około 228 MOPS). Wartość ta może być zredukowana do 60 MOPS zakładając użycie szybkiego algorytmu transformacji DCT (FDCT).

Najbardziej pracochłonna jest operacja kompensacji ruchu. Standardy przewidują, że kompensacja, może być przeprowadzona nie tylko względem poprzednich, ale także następnych ramek (użyteczne przy zmianie sceny), bądź nawet średniej dwóch: poprzedzającej i następnej. Enkoder w procesie kompresji podejmuje decyzję: kodowanie względem której ramki jest najbardziej korzystne, a to oczywiście wymaga wykonania kilku kompensacji i wybranie najlepszej.

Obliczenia przeprowadzone dla opracowanego dla potrzeb videokonferencji standardu kompresji H.261 bazującego na ramce CIF pokazują, że aby wykonać wszystkie operacje algorytmu kompresji należy dysponować procesorem RISC o mocy około 1GOPS. Nie jest to mało biorąc pod uwagę fakt, że powszechnie stosowane obecnie procesory ogólnego przeznaczenia oferują ok. 300-600 MOPS, a standard H.261 należy do najmniej wymagających.

Zamiarem autorów jest aby algorytmy kompresji były realizowane w czasie rzeczywistym, rozumianym jako nadążanie sprzętu obliczeniowego za napływającym bezpośrednio z kamery sygnałem wizyjnym (po przetworzeniu na postać cyfrową za pośrednictwem przetwornika A/C). Zatem, aby sprostać wymaganiom stawianym przez te algorytmy należy sięgnąć do specjalizowanych rozwiązań. Rys. 2 przedstawia procentowy udział poszczególnych faz kompresji w ogólnym zapotrzebowaniu na pracę procesora. Widać na nim, które operacje należałoby objąć rozwiązaniami specjalizowanymi w pierwszej kolejności.

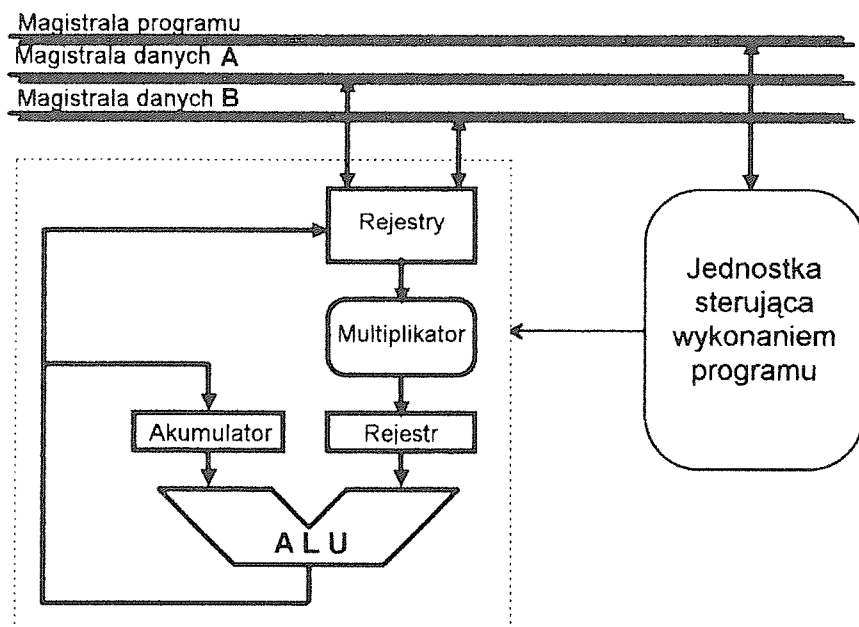




Rys. 2. Udział poszczególnych faz kompresji w ogólnym zapotrzebowaniu na pracę procesora

#### 4. PROCESORY DSP

Opracowane dla potrzeb cyfrowej obróbki sygnałów procesory DSP (ang. *Digital Signal Processor*) różnią się swoją architekturą wewnętrzną od klasycznych procesorów RISC. Podstawową operacją większości algorytmów cyfrowej obróbki sygnałów jest operacja typu:  $y = y + a * x$ . Wykonanie takiej operacji przez procesor ogólnego stosowania wymaga dwóch załadowań z pamięci, dodawania i mnożenia. Procesor DSP posiada budowę umożliwiającą zrobienie tego samego w jednej instrukcji. Rys. 3 przedstawia budowę prostego procesora typu DSP. Składa się on



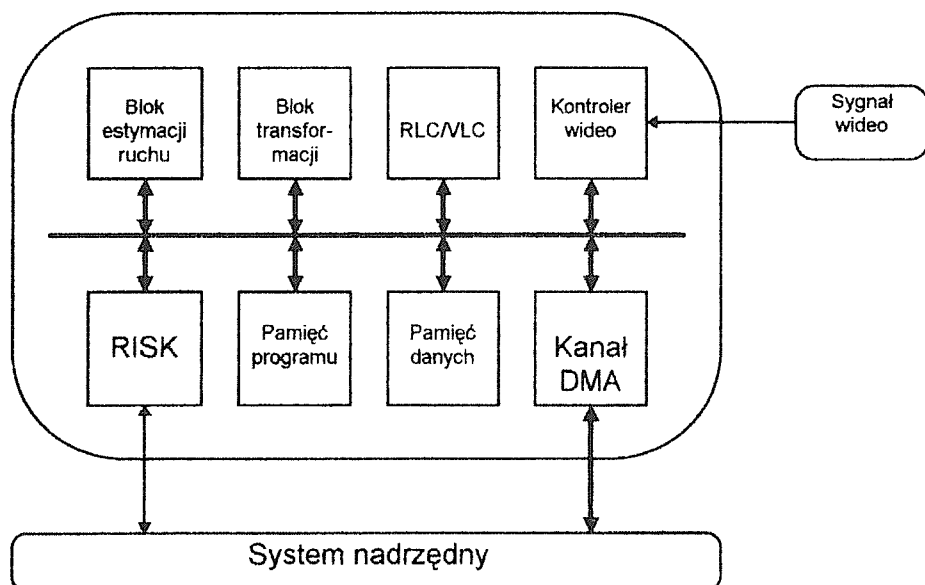
Rys. 3. Budowa rdzenia procesora DSP

z bloku rejestrów, bloku multiplikatora i jednostki arytmetyczno-logicznej. Ponadto procesor zawiera więcej niż jedną magistralę do komunikacji z zasobami zewnętrznymi. Poszczególne bloki pracują równolegle. Równocześnie zachodzi załadowanie nowych argumentów do rejestrów, mnożenie i dodawanie. Dlatego składnia assemblera procesora DSP różni się od składni tradycyjnego procesora tym, że każda instrukcja zawiera osobne polecenia dla każdej części składowej. Podobnie jak w procesorach RISC, wykonanie jednej instrukcji wymaga trzech taktów zegara procesora, ale wykonywane są jednocześnie trzy instrukcje (praca potokowa). W efekcie otrzymujemy wydajność jednej instrukcji na takt. Dodatkowo wydzielona jest magistrala programu, co umożliwia pobranie następnej instrukcji do wykonania w trakcie wykonywania przez procesor dostępu do przetwarzanych aktualnie danych. Inną ważną cechą procesorów DSP jest możliwość wielokrotnego powtarzania instrukcji bez konieczności tworzenia pętli programowych za pomocą tradycyjnych instrukcji skoku — stanowi to dodatkową możliwość zmniejszenia ilości koniecznych do wykonania algorytmu cykli procesora. Przykładowo procesor TMS320C40 firmy Texas Instruments osiąga moc obliczeniową 275 MOPS (co nie wyczerpuje potrzeb dla kompresji obrazów).

Procesory RISC oraz DSP zawdzięczają swoją szybkość i efektywność działania filozofii opartej na stwierdzeniu „im prostsze rozwiązanie tym szybsze”. Stąd także wypływają takie ich cechy jak: uproszczona lista rozkazów, sposób adresowania danych i niezbyt rozbudowana architektura.

## 5. ARCHITEKTURY DEDYKOWANE DO KOMPRESJI OBRAZU

Zastosowanie procesorów DSP do kompresji obrazów wideo pokazuje jak dedykowana do danych zastosowań architektura usprawnia i przyspiesza obliczenia. Procesory DSP zbudowano z myślą o przetwarzaniu sygnałów — szczególną formą takiego przetwarzania jest kompresja obrazu. Procesory DSP dobrze wykorzystują walory jakie daje wprowadzenie architektury MISD (ang. *Multiple Instruction — stream Single Data — stream*) czyli wykonywania wielu instrukcji równocześnie. Chcąc bardziej usprawnić proces kompresji należy dalej rozwijać filozofię DSP: po pierwsze w kierunku specjalizowania procesorów do kodowania obrazu w zakresie jednej operacji, a wręcz nawet jednego konkretnego algorytmu, po drugie w kierunku architektury MIMD (ang. *Multiple Instruction — stream Multiple Data — stream*) czyli do wykonywania wielu instrukcji i wielu danych równocześnie; w takim przypadku architektura MIMD bazowałaby na zrównoległaniu architektur MISD. Przedstawione na rys. 2 obciążenie procesora wskazuje, że celowym jest budowanie specjalnych procesorów do kompensacji ruchu i transformacji w dziedzinę częstotliwości. Układ taki nie wymagałby programu, ponieważ wykonuje tylko jeden konkretny algorytm, a wprowadzane dane przetwarzane są równolegle i kolejno na poszczególnych etapach. Zintegrowanie takich specjalizowanych modułów w jednym systemie (lub lepiej w jednym układzie VLSI) pozwala na zbudowanie układu bardzo wydajnego enkodera, którego poszczególne części składowe odpowiadają za



Rys. 4. Schemat blokowy układu do kompresji obrazu

szybkie wykonanie poszczególnych faz procesu kodowania. Rys. 4 przedstawia propozycję takiej hybrydy. Rdzeń stanowi procesor RISC lub DSP, który steruje procesem, wykonując także pomniejsze obliczenia konieczne podczas kodowania oraz podejmuje decyzje kontrolujące działanie innych bloków. Odpowiada on również za komunikację ze światem zewnętrznym — resztą systemu (zazwyczaj komputerem), którego częścią jest enkoder. Pozostałe moduły to specjalizowane w kierunku konkretnych obliczeń jednostki, którym procesor główny powierza właściwe zadania. Komunikacja pomiędzy modułami odbywa się po wielu szybkich magistralach, tak aby transfer danych nie spowalniał całego procesu.

W dalszej części artykułu omówione zostaną przykładowe rozwiązania sprzętowe dedykowane do obliczeń transformacji i estymacji ruchu.

## 6. ARYTMETYKA ROZPROSZONA

Arytmetyka rozproszona (ang. *Distributed Arithmetic*) jest metodą szybkiego obliczania iloczynu wielu sum [9]. Wielokrotne sumowanie typu:

$$S = \sum_{i=0}^N C_i * x_i$$

jest podstawą algorytmów transformacji w dziedzinę częstotliwości, niezależnie od tego, czy kompresja bazuje na DCT, FFT, czy na filtrowaniu falkowym. Przy tego typu operacjach współczynniki  $C_i$  w powyższym równaniu są stałe, a wartości

$x$  reprezentują punkty obrazu. Przy założeniu, że próbki obrazu są zapisane w kodzie U2, mamy:

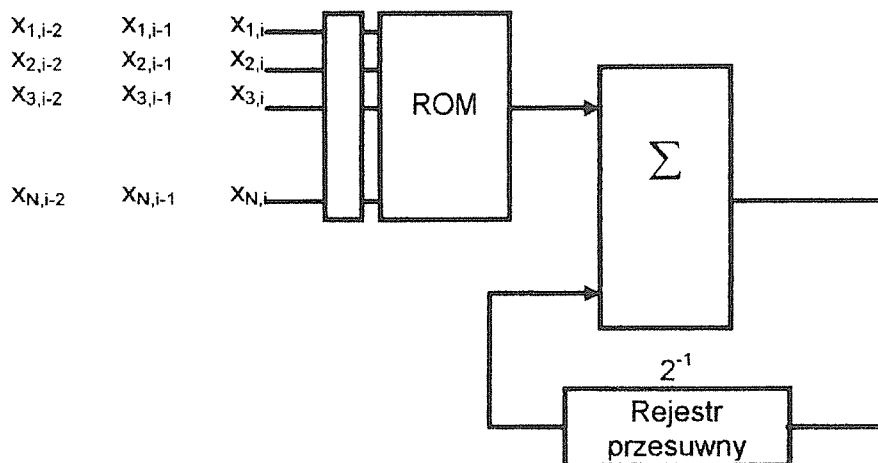
$$x = -x_M * 2^M + \sum_{j=0}^{M-1} x_j * 2^j$$

gdzie:  $x_j$  —  $j$ -ty bit wartości  $x$ .

Wyrażenie na  $S$  można więc zapisać:

$$S = \sum_{i=0}^N C_i * \left[ -x_{iM} * 2^M + \sum_{j=0}^{M-1} x_{ij} * 2^j \right] = - \left[ \sum_{i=0}^N C_i * x_{iM} \right] * 2^M + \sum_{j=0}^{M-1} \left[ \sum_{i=0}^N C_i * x_{ij} \right] * 2^j.$$

Ponieważ  $x_{ij}$  jest wartością binarną, a  $C_i$  to stałe współczynniki charakterystyczne dla danej operacji,  $\sum_{i=0}^N C_i * x_{ij}$  może przyjmować tylko  $2^N$  różnych wartości. Istnieje więc możliwość zapamiętania wszystkich wyników sumowań w pamięci adresowanej wektorem  $x$ . To rozwiązanie przyspiesza późniejsze obliczenia, ponieważ czasochłonne mnożenie sprowadza się do pojedynczego odczytu wyniku z pamięci. Konieczne dodatkowe mnożenie przez potęgę 2 to prosty przesuw wyniku w rejestrze. Reszta operacji potrzebna do obliczenia sumy  $S$  to już tylko dodawanie i odejmowanie odczytanych z pamięci i odpowiednio przesuniętych wartości. Rys. 5 przedstawia propozycję procesora bazującego na podanym sposobie. Kolejne współczynniki  $x$ , po przekształceniu na odpowiednią wartość w pamięci wyników, są dodawane, przesuwane i akumulowane w rejestrze. Trzeba zwrócić uwagę, że podane rozwiązanie zamiast przesuwac w lewo kolejne  $x$  przesuwa w prawo przechowywane w rejestrze kolejne sumy cząstkowe dodawania. Ogromną zaletą tej architektury jest jej niebywała prostota. Małym kosztem można więc powielić



Rys. 5. Moduł kalkulujący metodą arytmetyki rozproszonej

moduły przedstawione na rysunku i wykonywać operację na kolejnych  $x$  nie szeregowo, ale równoległe. Celowe jest równoległe zastosowanie nawet  $N$  takich modułów.

## 7. ALGORYTM KOMPENSACJI RUCHU

Kompensacja ruchu odgrywa kluczową rolę w procesie kompresji obrazu. Istnieje wiele algorytmów kompensacji, lecz te najbardziej efektywne są równocześnie najbardziej czasochłonne. Najbardziej popularne i najczęściej stosowane są kompromisem wydajności i złożoności obliczeniowej.

Najszerzej stosowany algorytm zakłada podział całej ramki na jednakowej wielkości bloki. Procesowi kompensacji ruchu poddawany jest każdy blok w aktualnie kodowanej ramce. Polega on na znalezieniu w ramce względem której następuje kodowanie obszaru o rozmiarach kodowanego bloku który najlepiej do niego pasuje. Najczęściej stosowanym kryterium podobieństwa dwóch bloków jest średni błąd wartości odpowiadających sobie punktów. Można to wyrazić wzorem:

$$E = \sum_{i/0}^{N-1} \sum_{j/0}^{N-1} \left| X_{(A+i+p), (B+j+q)} - Y_{(A+i), (B+j)} \right|,$$

gdzie:  $i, j$  — współrzędne punktów porównywanych obszarów,  $X, Y$  — porównywane obrazy,  $N$  — wielkość obszarów,  $A, B$  — współrzędne położenia obszaru w ramce obrazu,  $p, q$  — współrzędne wektora przesunięcia.

Za wektor przesunięcia  $V(p, q)$  przyjmowany jest ten dla którego błąd  $E$  osiąga minimum. Wielkość przeszukiwanego obszaru wyznacza nakłady czasowe przeznaczone na obliczenia. Z drugiej strony dla kodowanych sekwencji przesunięcie obrazu pomiędzy ramkami ma ograniczony charakter. Praktycznie przyjmowane wartości  $p$  i  $q$  zawierają się w przedziale  $[-7, 8]$ . Podana metoda porównywania każdego bloku z każdym w ramach przeszukiwanego obszaru daje najlepsze wyniki poszukiwania minimum błędu, ale dla 30 ramek o rozmiarze  $720 \times 480$  kodowanych w czasie 1 sekundy (standardowy sygnał telewizyjny) wymaga ok. 30 GOPS przy przeszukiwaniu obszaru  $[-15, 16]$ .

Wykonanie obliczeń dla jednego wektora przesunięcia dla bloku o rozmiarach  $16 \times 16$  wymaga wykonania operacji odejmowania, uwzględnienia wartości bezwzględnej i dodawania dla 256 pikseli ( $3 \cdot 256$  operacji). Przy przeszukiwaniu obszaru  $[-15, 16]$  musimy obliczyć 1024 lokacji ( $3 \cdot 256 \cdot 1024$ ). Ramkę  $720 \times 480$  pikseli musimy podzielić na 1350 bloków mamy więc  $3 \cdot 256 \cdot 1024 \cdot 1350$  operacji co 30 część sekundy. Razem  $3 \cdot 256 \cdot 1024 \cdot 1350 \cdot 30$ , a więc ok. 31 GOPS.

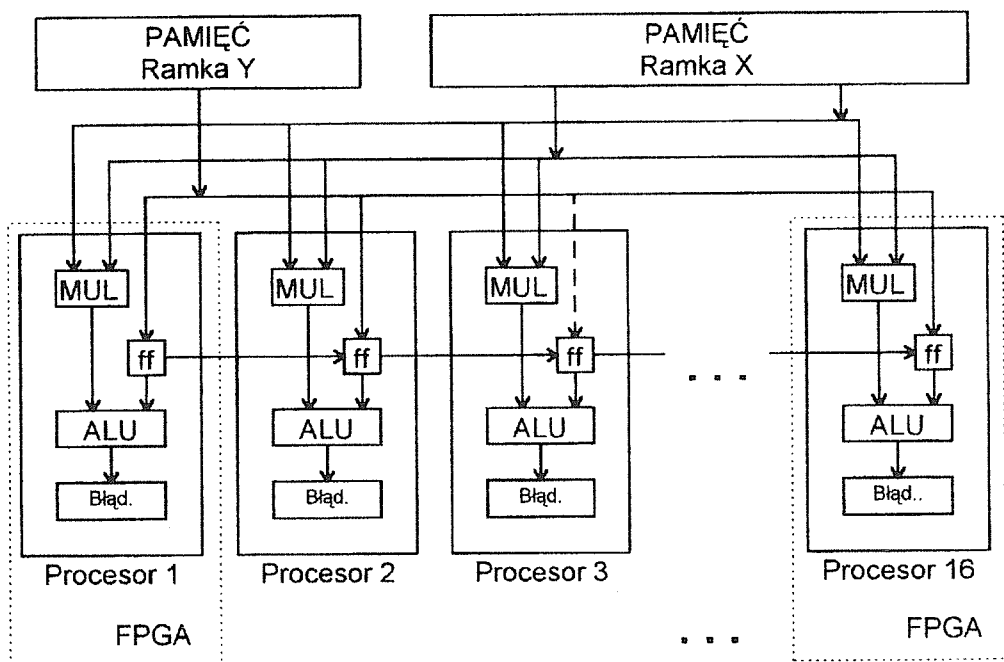
## 8. ARCHITEKTURA PROCESORA KOMPENSACJI RUCHU

Z podanego wzoru na kompensację ruchu wynika, że aby przeliczyć błąd dla poszczególnych wektorów  $V(p, q)$  sekwencyjnie, konieczne jest wielokrotne czytanie z pamięci wartości tych samych punktów. Np: dla  $i = 0, j = 0, p = 0, q = 0$  czytamy

ten sam punkt obszaru  $X$  co dla  $i = 1, j = 1, p = -1, q = -1$ . Ponadto punkty  $Y$  czytane są wielokrotnie dla różnych  $p$  i  $q$ . Powoduje to spowolnienie procesu obliczeń o czas dostępu do danych i obciążenie magistrali niepotrzebnym transferem z pamięci zewnętrznej.

Podany przykład [2] zakłada porównywanie bloków o rozmiarze  $16 \times 16$  punktów z wektorem przesunięcia  $[-7.. 8, -7.. 8]$ . Pomysł architektury zakłada równoległe obliczenia 16 wartości błędów średnich równocześnie. Raz załadowana wartość  $X_{(i+p), (j+q)}$  wykorzystana jest w 16 współbieżnie pracujących układach arytmetycznych. Układy te odpowiadają za obliczenie kolejnych wektorów przesunięć:  $[-7, q], [-6, q], [-5, q], \dots, [8, q]$ .

Ponieważ dla danego punktu  $X_{(i+p), (j+q)}$  wykonanie operacji odejmowania w ALU dla kolejnych wektorów  $V(p, q)$  wymaga kolejnych wartości punktów  $Y_{ij}$  należałoby równoległe załadować 8 kolejnych punktów  $Y$  z pamięci. Rozwiązaniem jest zastosowanie rejestru przesuwającego poprzez który wartość przesuwana się od jednego układu do następnego. W ten sposób kolejne punkty  $y_{i,j}$  trafiają w odpowiednim czasie do odpowiedniego ALU. To rozwiązanie, jak zobaczymy później, wymaga jednak wprowadzenia dodatkowej magistrali przesyłającej punkty obrazu  $X$  i wprowadzenie multiplexera na wejściu ALU, który wybiera odpowiednio dane o obrazie  $X$  (Rysunek 6). Proces obliczeń przebiega następująco: do Procesora 1 ładowana jest wartość  $X_{A+0-7, B+0-7}$  i  $Y_{A+0, B+0}$  ( $i = 0, j = 0, p = -7, q = -7$ ) i wykonane zostaje pierwsze odejmowanie dla wektora  $[-7, -7]$ , wszystkie



Rys. 6. Schemat blokowy procesora do równoległego obliczania wektorów przesunięć

[illegible]

punktu w którym wymagane jest przejście do następnej linii w bloku  $X$  tzn. po załadowaniu  $Y_{A+0, B+1}$  następującego po  $Y_{A+15, B+0}$  do pierwszego procesora. Wtedy to Procesor 1 wymaga do dalszych obliczeń elementu  $X_{A+0-7, B+1-7}$ , ale pozostałe procesory wymagają elementu  $X_{A+16-7, B+0-7}$ . Dwie magistrale przekazują procesorom obydwie potrzebne elementy. W następnym kroku również Procesor 2 zakończy obliczenia dla linii pierwszej i przejdzie do obliczeń dla linii drugiej. Wtedy Procesory 1 i 2 załadują element  $X_{A+1-7, B+1-7}$  po pierwszej magistrali, a reszta procesorów element  $X_{A+17-7, B+0-7}$  po drugiej. Kolejne procesory będą w następnych etapach zakańczać pracę w pierwszej linii i przechodzić do linii drugiej. Po wykonaniu 256 cykli ładowania, odejmowania i akumulacji błąd  $E$  dla pierwszego wektora  $[-7, -7]$  jest gotowy do odczytu, a po następnych 15 cyklach gotowa jest reszta obliczeń dla pozostałych wektorów. W tabeli 1 przedstawiono kolejne fazy pracy poszczególnych procesorów.

## 9. PODSUMOWANIE

W ostatnim okresie intensywnie rozwija się nowa dziedzina szybkiej realizacji obliczeń poprzez stosowanie struktur programowalnych FPGA, dedykowanych dla zadań obliczeniowych użytkownika (ang. *FPGA for Custom Computing Machines*) [10]. Niniejsze opracowanie mieści się w tym nurcie — związane jest bowiem z implementacją algorytmów kompresji obrazów wizyjnych w programowalnych strukturach FPGA celem zapewnienia odpowiednio wysokiej szybkości realizacji tych algorytmów.

Opracowywana architektura (w przygotowywanym projekcie) pochłonie wiele układów FPGA o największych pojemnościach. W szczególności zastosowane zostaną układy serii XC4000 firmy Xilinx, których pojemności sięgają aktualnie 62000 bramek [18]. Zostaną one jednak skonfigurowane w taki sposób, aby była możliwa ich dynamiczna rekonfiguracja w trybie normalnej pracy. Pozwoli to na zmianę algorytmów kompresji obrazu w zależności od charakteru aktualnie obserwowanych obiektów, oświetlenia i dynamiki ruchu na obserwowanej scenie. Umożliwi także zmianę parametrów i algorytmów kompresji w zależności od aktualnie wymaganej rozdzielczości, jakości obrazu itp. Wykorzystane tu zostaną wcześniejsze doświadczenia autorów w projektowaniu nowoczesnych rekonfigurowalnych architektur obliczeniowych (ang. *Reconfigurable System Architecture*) [14].

Kompresja obrazów często wiąże się z wcześniejszą poprawą ich jakości. Różnego rodzaju filtracje, realizowane w trybie czasu rzeczywistego, wymagają znacznych mocy obliczeniowych. Opracowany przez autorów system potokowego przetwarzania danych wizyjnych przez dedykowane procesory sprzętowe zapewnia wysokie parametry jakości przetwarzania przy jednoczesnym zachowaniu minimalnego czasu wnoszonego opóźnienia [15, 16, 17]. Jednocześnie umieszczony na początku architektury potokowej specjalny translator zapewnia przekodowanie obrazu z międzyliniowością na obraz kolejnoliniowy, wymagany przez większość



algorytmów kompresji obrazu. Translator wprowadza minimalne opóźnienie równe czasowi trwania jednej linii. Zastosowanie powyższych rozwiązań sprzętowych poprawi jakość przetwarzanych obrazów i przystosuje je do potrzeb systemu sprzętowej kompresji obrazu (np. kolejnoliniowość).

## BIBLIOGRAFIA

1. *Advanced Imaging*. PTN Publishing Co., New York, June 1995
2. V. Bhaskaran, K. Konstantinides: *Image and Video Compression Standards*. Kluwer Academic Publishers, 1995
3. A.L. Decega: *The Technology of Parallel Processing -- Parallel Processing Architectures and VLSI Hardware*. Prentice Hall, New Jersey 1989
4. D.L. Gall: *MPEG: A Video Compression Standard for Multimedia Applications*. Communications of the ACM/April 1991/Vol. 34. No. 4
5. K. Hwang, F. Briggs: *Computer Architecture and Parallel Processing*. McGraw-Hill 1984
6. Y. Iwata, M. Kawamata, T. Higuchi: *Design of fine VLSI array processor for real-time 2-D digital filtering*. IEEE Symposium on Circuits and Systems, Chicago 1993, p.1559-1562
7. P.P. Jonker: *Morphological Image Processing: Architecture and VLSI design*. Kluwer — Holland 1992
8. T. Łuba, M.A. Markowski, B. Zbierzchowski: *Komputerowe projektowanie układów cyfrowych w strukturach PLD*. WKŁ, Warszawa 1993
9. L. Mintzer: *FIR filters with the Xilinx FPGA*. FPGA'92 ACM/SIGDA Workshop on FPGA, p. 129-134
10. K.L. Pocek, J. Arnold: *Proceedings IEEE Syposium: FPGAs for Custom Computing Machines*. IEEE Computer Society Press, Washington-Brussels-Tokyo, 1996
11. SGS-THOMSON Microelectronics: *Image Processing. Databook*. 1991
12. R. Tadeusiewicz: *Systemy wizyjne robotów przemysłowych*. WNT, Warszawa 1992
13. Texas-Instruments: *TMS320C80 (MPV) Parallel Processor. User's Guide*. 1995
14. K. Wiatr: *Pipelined Architecture of Reconfigurable Specialised Processors for a Real-Time Image Data Pre-Processing*. Proc. of the IEEE International Conference on Signal Processing — ICSP-96, Beijing — China 14-18.10.1996, 649-652
15. K. Wiatr: *Architektura specjalizowanych procesorów potokowych do obróbki wstępnej danych wizyjnych w czasie rzeczywistym*. Kwartalnik Elektroniki i Telekomunikacji 1997, t. 43, z. 1, ss. 121-148
16. K. Wiatr: *Specialised Architecture of Dedicated Hardware Processors for a Real-Time Image Data Pre-Processing*. Proc of the EUROMICRO Int. Conf.: Real-Time Systems, IEEE Computer Press, Washington-Brussels-Tokyo 1997, p. 180
17. K. Wiatr: *Dedicated Hardware Processors for a real-Time Image Data Pre-processing Implemented in FPGA Structure*. 9th International Conference on Image Analysis and Processing — ICIAP'97, Springer Verlag, Berlin 1997, p. 69-75
18. XILINX: *Data Book*. San Jose, 1996

K. WIATR, P. RUSSEK

HARDWARE ARCHITECTURES OF IMAGE COMPRESSION  
FOR REAL-TIME DATA COMPUTATION

## S u m m a r y

The following article is concerned with real time video compression problems. Particularly it emphasis a present compression algorithms computing complexity and attempts to show techniques that are used to fulfil that complexity. Practical solutions widely implemented in MPEG or H.261 and wavelet compression codecs are presented.

The article compares performance of RISC and DSP processors in video compression and unveils a necessity of specialised architectures for this purpose. Hardware solutions for efficient motion estimation and frequency transforms are proposed. The hardware is to be implemented in the high capacity FPGA structures. The article is a part of the research supported by The Polish Science Committee.

*Key words:* real-time systems, image processing, image compression, programmable logic device.

# Zastosowanie konwersji modulacji częstotliwości na modulację amplitudy w światłowodowych systemach transmisji o dużych przepływnościach

KRZYSZTOF PERLICKI

*Instytut Telekomunikacji, Politechnika Warszawska*

*Otrzymano 1998.01.15*

*Autoryzowano 1998.08.07*

W pierwszej części artykułu przedstawiono analizę zjawiska konwersji modulacji częstotliwości na modulację amplitudy w standardowym światłowodzie jednomodowym, która powodowana jest przez dyspersję chromatyczną włókna optycznego. Część druga opisuje właściwości systemu transmisji danych o przepływności 5 Gbit/s wykorzystujący zjawisko konwersji.

**Słowa kluczowe:** telekomunikacja światłowodowa, konwersja modulacji częstotliwości na modulację amplitudy, dyspersja światłowodu.

## 1. WPROWADZENIE

Współczesne systemy światłowodowe powszechnie wykorzystują jako medium transmisyjne światłowody standardowe o wartości współczynnika dyspersji chromatycznej rzędu  $17 \text{ ps/nm} \cdot \text{km}$  dla długości fali równej  $1550 \text{ nm}$ . Taka wartość dyspersji chromatycznej powoduje, że przy stosowaniu zewnętrznej modulacji natężenia światła emitowanego przez laser maksymalna odległość, na którą można przesłać dane w kodzie NRZ z przepływnością  $10 \text{ Gbit/s}$  albo w kodzie bifazowym typu Manchester z przepływnością  $5 \text{ Gbit/s}$ , bez pogorszenia parametrów systemu, wynosi tylko  $36,7 \text{ km}$ . Przy bezpośredniej modulacji lasera, ze względu na występowanie zjawiska migotania (ang. *Chirp*), zakres odległości jest znacznie mniejszy i wynosi  $3,6 \text{ km}$  [6]. Zwiększenie zasięgu transmisji wymaga zastosowania specjalnych rozwiązań technicznych; do najpopularniejszych spośród nich należy zaliczyć: technikę *prechirpingu*, konwersję modulacji częstotliwości światła na modulację

amplitudy z wykorzystaniem metod interferometrycznych oraz fotoniczne techniki kompensacji dyspersji. Rozwiązania te zwiększają zakres transmisji do około 100 km [5]. W 1993 roku po raz pierwszy zaprezentowano transmisję światłowodową opartą na konwersji modulacji częstotliwości światła na modulację amplitudy wywołaną dyspersją chromatyczną włókna optycznego. Metoda ta pozwala na przesyłanie informacji w kodzie NRZ z przepływnością rzędu 10 Gbit/s na odległość 160 km [5, 8, 9]. Podstawową wadą tego rozwiązania jest konieczność doboru wielkości zmiany częstotliwości do długości toru światłowodowego oraz otrzymywanie na końcu światłowodu wielopoziomowego sygnału optycznego wymagającego dodatkowej obróbki w celu uzyskania dwupoziomowego sygnału elektrycznego. Realizowane jest to powszechnie przez złożony układ decyzyjny połączony z filtrem o pasmie dopasowanym do długości traktu optotelekomunikacyjnego [8].

Podstawowym celem artykułu, obok analizy samego zjawiska konwersji modulacji częstotliwości na modulację amplitudy jest przedstawienie takiego sposobu zmiany częstotliwości światła emitowanego przez laser, który pozwoli uzyskać dwupoziomowy sygnał optyczny, co w konsekwencji umożliwi uproszczenie części odbiorczej systemu.

## 2. PRZYBLIŻONA ANALIZA KONWERSJI MODULACJI CZĘSTOTLIWOŚCI NA MODULACJĘ AMPLITUDY W ŚWIATŁOWODZIE STANDARDOWYM

Przebieg z modulacją częstotliwości ma postać:

$$\Psi_{FM}(t) = A_{we} \exp \{ j(\omega_0 t + \varphi(t)) \} \quad (2.1)$$

gdzie:

$A_{we}$  — amplituda

$\omega_0$  — pulsacja nośnej

$\varphi$  — faza.

Jeżeli sygnał optyczny jest generowany przez laser sterowany prądem stałym o wartości  $I_0$  to modulację częstotliwości można uzyskać poprzez dodanie do prądu stałego małego, zmiennego w czasie prądu modulującego  $I_m$ . Całkowita wartość prądu sterującego laser wynosi wówczas:

$$I = I_0 + I_m. \quad (2.2)$$

W przypadku idealnym zmiany częstotliwości powinny być proporcjonalne do zmian prądu modulującego:

$$f_d(t) = KI_m(t). \quad (2.3)$$

Zmianę fazy określamy zależnością:

$$\frac{\varphi(t)}{2\pi} = K \int_0^t I_m(t) * H_v(t) dt, \quad (2.4)$$

przy czym:

stała  $K$  w (2.3) i (2.4) określona jest w następujący sposób:

$$K = \frac{f_{d,\max}(t)}{I_{m,\max}(t)} \quad (2.5)$$

$f_d$  — dewiacja częstotliwość

$f_{d,\max}$  — maksymalna dewiacja częstotliwość

$I_{m,\max}$  — wartość prądu modulującego przy  $f_d = f_{d,\max}$

$H_v$  — znormalizowana odpowiedź częstotliwościowa lasera.

Znormalizowana odpowiedź częstotliwościowa lasera ma następującą postać [4]:

$$H_v(j\omega) = \left| \frac{\Omega_d \omega_r^2}{4_d \Omega_c - Q_1 \omega_r^2} \right| \frac{(\Omega_d + j\Omega_v)(\Omega_c + j\Omega_v) - Q_1 \omega_r^2}{(\Omega_d + j\Omega_v)(\omega_r^2 - \Omega_v^2 + j\Omega\Omega_v)} \quad (2.6)$$

przy czym:

$$\Omega_d = \frac{1}{\tau} + (1 - \chi \tau_p^2 \omega_r^2) \tau_p \omega_r^2 \quad (2.7)$$

$$\Omega_c = \chi \tau_p \omega_r^2 \quad (2.8)$$

$$\Omega = \Omega_d + \Omega_c \quad (2.9)$$

$$\chi = \frac{\varepsilon}{A \tau_p}, \quad (2.10)$$

gdzie:

$\tau_p$  — czas życia fotonu we wnętrzu rezonansowej lasera;

$\tau$  — czas życia elektronu;

$f_d$  — częstotliwość rezonansowa;

$Q_1$  — parametr strojenia (*tuning parameter*);

$\chi$  — dekrement tłumienia (*damping parameter*);

$\varepsilon$  — współczynnik kompresji wzmocnienia (*gain compression coefficient*).

Uwzględniając zależność (2.6) i korzystając z tego, że operacji całkowania w dziedzinie czasu odpowiada pomnożenie widma przez operator  $\frac{1}{j\omega}$  oraz wiedząc, że widmo splotu dwóch przebiegów jest iloczynem widm składowych splotu można wzór (2.4) przedstawić w następującej postaci:

$$\frac{\varphi(t)}{2\pi} = K \mathfrak{F}^{-1} \left( \mathfrak{F}(I_m(t)) \cdot \frac{H_v(j\omega)}{j\omega} \right) \quad (2.11)$$

gdzie:

$\mathfrak{F}$  — transformata Fouriera;

$\mathfrak{F}^{-1}$  — odwrotna transformata Fouriera.

Otrzymaną ze wzoru (2.11) zmianę fazy w czasie  $\varphi(t)$  podstawia się do wzoru (2.1), w wyniku czego otrzymujemy przebieg fali świetlnej zmodulowanej częstotliwościowo na wyjściu lasera.

Konwersja modulacji częstotliwościowej fali świetlnej na modulację amplitudy następuje w wyniku zachodzenia zjawiska dyspersji chromatycznej światłowodu. W zakresie tzw. dyspersji anormalnej, tj. w zakresie długości fal przy których  $\beta_2 < 0$ , składowe przesyłanego sygnału o większych częstotliwościach (mniejszych długościach fali) poruszają się szybciej niż składowe o mniejszych częstotliwościach (większych długościach fali). W wyniku dyspersji chromatycznej składowe sygnału o różnych częstotliwościach rozchodzą się z różnymi prędkościami grupowymi w wyniku czego przebywają dany odcinek światłowodu w różnych czasach. Powoduje to nakładanie się lub wygaszanie przesyłanych impulsów na wyjściu światłowodu. Za opis propagacji fali światła w światłowodzie odpowiedzialny jest składnik  $\exp(-j\beta L)$ , gdzie:

$L$  — długość światłowodu

$\beta$  — stała fazowa.

Rozwijając  $\beta$  wokół pulsacji nośnej  $\omega_0$  otrzymujemy:

$$\beta = \beta_0 + \omega_d \beta_1 + \frac{1}{2} \omega_d^2 \beta_2 + \dots \quad (2.12)$$

gdzie  $\omega_d = \omega - \omega_0$ .

W szeregu tym człony trzeciego i wyższych rzędów zostały pominięte.

Składowa  $\beta_0$  związana jest z fazą fali nośnej. Pochodne stałej fazowej  $\beta_1$  i  $\beta_2$  względem pulsacji odnoszą się odpowiednio do jednostkowego czasu przejścia modu i dyspersji prędkości grupowej. Zakładając, że wykładnik składnika związanego z drugą pochodną stałej propagacji jest mały, co otrzymujemy przy założeniu, że:

$$\omega_d \ll \sqrt{\frac{2}{|\beta_2| L}}, \quad (2.13)$$

można go przybliżyć w następujący sposób:

$$\exp\left(-j\beta_2 \frac{L}{2} \omega_d^2\right) \approx 1 - j\beta_2 \frac{L}{2} \omega_d^2. \quad (2.14)$$

Dodatkowo zakładając, że ograniczamy się do analizy wpływu zjawiska dyspersji prędkości grupowej, zaniebując składowe, które odnoszą się odpowiednio do opóźnienia fazowego fali nośnej i opóźnienia czasowego sygnału modulującego, otrzymujemy nie uwzględniając strat w światłowodzie następującą postać przebiegu z modulacją częstotliwości na końcu światłowodu o długości  $L$ :

$$\Psi_{FM}(j\omega, L) = \Im(A \exp\{j\varphi(t)\}) \cdot \left(1 - j\beta_2 \frac{L}{2} \omega_d^2\right). \quad (2.15)$$

Przechodząc do dziedziny czasu otrzymujemy:

$$\Psi_{FM}(t, L) = [A \exp \{j\varphi(t)\}] * \left[ \mathfrak{F}^{-1} \left( 1 - j\beta_2 \frac{L}{2} \omega_d^2 \right) \right]. \quad (2.16)$$

Korzystamy z następujących właściwości transformaty delty Diraca [15]:

$$\mathfrak{F}^{-1}(1) = \delta(t) \quad (2.17)$$

$$\mathfrak{F}^{-1}[(j\omega)^2] = \delta''(t). \quad (2.18)$$

Na podstawie (2.17) i (2.18) odwrotna transformata Fouriera składnika związanego z drugą pochodną stałej propagacji jest równa:

$$\mathfrak{F}^{-1} \left( 1 - j\beta_2 \frac{L}{2} \omega_d^2 \right) = \delta(t) + jP\delta''(t), \quad (2.19)$$

gdzie:

$\delta$  — delta Diraca;

$$P = \frac{L \, d\tau}{2 \, d\lambda} \frac{\lambda^2}{2\pi c} = - \frac{L \, d\tau}{2 \, d\omega}, \quad (2.20)$$

przy czym:

$$\frac{d\tau}{d\omega} = \beta_2, \quad (2.21)$$

$$\frac{d\tau}{d\omega} = D \quad (2.22)$$

$D$  — współczynnik dyspersji;

$c$  — prędkość światła;

$\lambda$  — długość fali świetlnej.

Po podstawieniu (2.19) do (2.16), przebieg z modulacją częstotliwości ma następującą postać:

$$\Psi_{FM}(t, L) = [A \exp \{j\varphi(t)\}] * [\{\delta(t) + jP\delta''(t)\}]. \quad (2.23)$$

Operacja splotu danej funkcji  $f(t)$  z deltą Diraca  $\delta(t)$  jest równa [8]:

$$f(t) * \delta(t) = f(t). \quad (2.24)$$

Operacja splotu danej funkcji z wielkością  $\delta''(t)$  jest równoważna  $n$ -krotnemu różniczkowaniu funkcji  $f(t)$  [15]:

$$f(t) * \delta''(t) = (-1)^n \frac{\partial^n f(t)}{\partial t^n}. \quad (2.25)$$

Korzystając z (2.24) i (2.25) i wykonując operację splotu otrzymujemy:

$$\Psi_{FM}(t, L) = A_{we} \left( 1 + jP \frac{\partial^2}{\partial t^2} \right) \exp \{ j\varphi(t) \}. \quad (2.26)$$

Przebieg mocy optycznej ( $P_{wy}$ ) na końcu światłowodu o długości  $L$  jest równy:

$$P_{wy}(t, L) = |\Psi_{FM}(t, L)|^2, \quad (2.27)$$

czyli:

$$P_{wy}(t, L) = P_{we} \left| \left( 1 + jP \frac{\partial^2}{\partial t^2} \right) \exp \{ j\varphi(t) \} \right|^2; \quad (2.28)$$

W wyniku przekształceń otrzymujemy [7]:

$$P_{wy}(t, L) = P_{we} \left( 1 - 2P \frac{\partial^2 \varphi(t)}{\partial t^2} \right). \quad (2.29)$$

Wiedząc, że faza jest całką pulsacji po czasie równanie (2.29) można przedstawić w następującej postaci:

$$P_{wy}(t, L) = P_{we} \left( 1 - 2P \frac{\partial \omega_d(t)}{\partial t} \right). \quad (2.30)$$

Zależność (2.30) wskazuje, że kształt otrzymanych w wyniku zjawiska konwersji modulacji częstotliwości na modulację amplitudy impulsów jest pochodną częstotliwości po czasie. Poziom mocy na wyjściu światłowodu zależy od czasu trwania impulsu, maksymalnej dewiacji częstotliwości, długości światłowodu i wartości współczynnika dyspersji oraz mocy na wejściu światłowodu. Jeżeli zmiany pulsacji w czasie mają charakter przebiegu w postaci:

$$\omega_d(t) = - \frac{\omega_{d,max}}{T} t, \quad (2.31)$$

gdzie  $T$  jest czasem trwania przebiegu, a  $\omega_{d,max}$  maksymalną dewiacją częstotliwości, to uzyskana, w wyniku konwersji modulacji częstotliwości na modulację amplitudy, zmiana mocy w funkcji czasu na końcu światłowodu ma postać zbliżoną do przebiegu prostokątnego o wartości mocy proporcjonalnej do  $\frac{\omega_{d,max}}{T}$ . Wówczas zależność (2.30) ma postać:

$$P_{wy}(t, L) = P_{we} \left( 1 + 2P \frac{\omega_{d,max}}{T} \right). \quad (2.32)$$

Przy założeniu, że rozpatrujemy zjawisko konwersji dla światłowodu standardowego o współczynniku dyspersji  $D = 17 \text{ ps/nm} \cdot \text{km}$  i długości traktu światłowodowego ( $L$ ) z zakresu 100 km-200 km, to ze względu na założenie (2.13) wzór (2.30) pod względem ilościowym jest słuszny dla dewiacji pulsacji  $f_{d,max} \ll 2\pi \cdot 7,3 \text{ GHz}$  (przy  $L = 100 \text{ km}$ ) i  $f_{d,max} \ll 2\pi \cdot 3,65 \text{ GHz}$  (przy  $L = 200 \text{ km}$ ). Wyniki uzyskane z zależ-



ności (2.30) porównano z wynikami otrzymanymi z symulacji numerycznej zjawiska konwersji modulacji częstotliwości na modulację amplitudy w światłowodzie standardowym. Symulacja numeryczna została przeprowadzona przez rozwiązanie nieliniowego równania Schrödingera Metodą Dwukrokową (*Split-Step Fourier Method*) [12]. Analiza porównawcza uzyskanych wyników wskazuje, że zgodność między obydwoima rozwiązaniami wzrasta wraz ze zmniejszeniem maksymalnej wartości dewiacji pulsacji. Otrzymanie zbieżności wyników na poziomie 20% uzyskujemy dla maksymalnej dewiacji

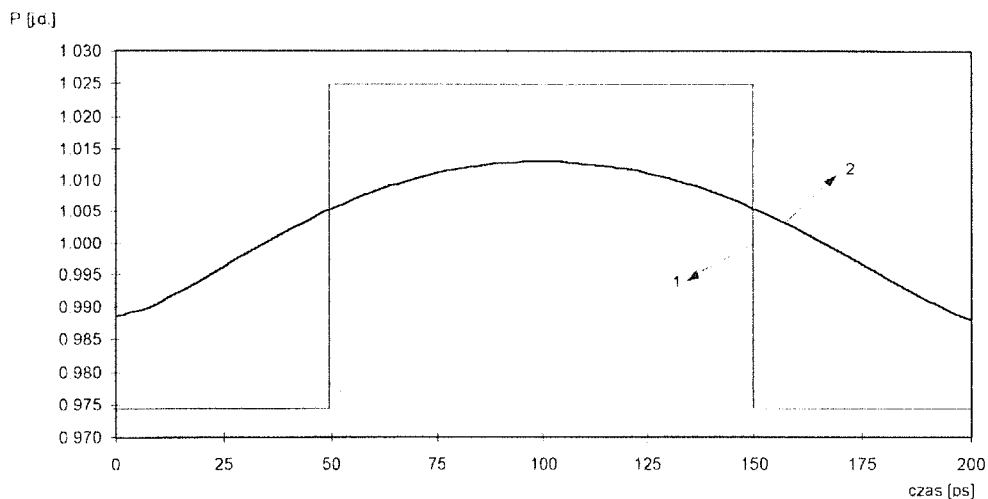
pulsacji  $\omega_{d,\max} \leq 0,05 \frac{2\pi Tc}{LD\lambda^2}$ , gdzie  $T$  — czas trwania zmiany pulsacji,  $c$  — prędkość światła,  $\lambda$  — długość fali światła,  $L$  — długość światłowodu,  $D$  — współczynnik dyspersji.

Na rys. 1 przedstawiono przebieg mocy impulsu optycznego otrzymanego z symulacji numerycznej dla zmiany częstotliwości określonej zależnością (2.31) przy

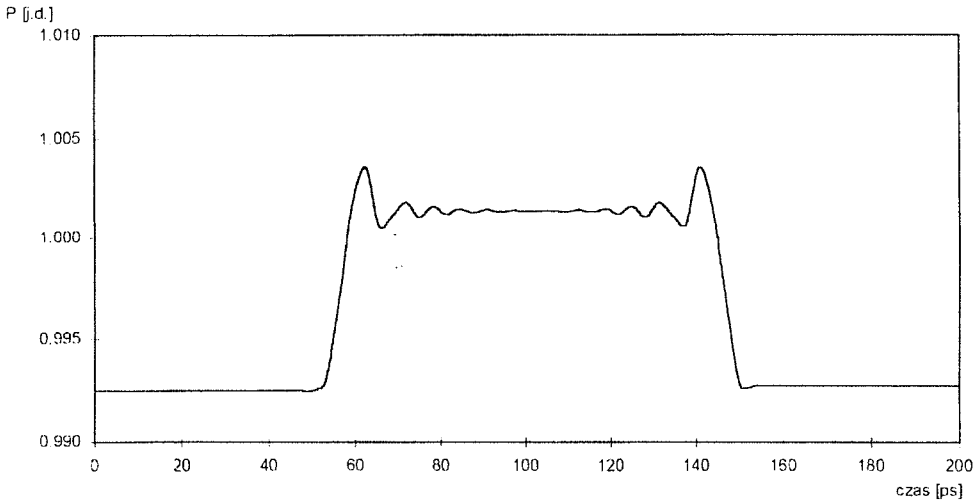
$\omega_{d,\max} = 0,05 \frac{2\pi Tc}{LD\lambda^2}$  założono, że  $L = 140$  km,  $\lambda = 1550$  nm,  $D = 17 \frac{\text{ps}}{\text{nm} \cdot \text{km}}$ ,

$T = 100$  ps. Pomimo, że wzor (2.32) można stosować jedynie w zakresie małych wartości  $\omega_{d,\max}$  to oddaje on ogólny, jakościowy charakter zjawiska konwersji modulacji częstotliwości na modulację amplitudy, kształt otrzymanych impulsów jest zbliżony do pochodnej pulsacji po czasie; jest to szczególnie wyraźne w zakresie małych dewiacji pulsacji. Na rys. 2 pokazano kształt impulsu dla

$\omega_{d,\max} = 0,001 \frac{2\pi Tc}{LD\lambda^2}$  przy zmianie pulsacji określonej wzorem (2.31). Przyjęto takie same wartości dla  $L$ ,  $\lambda$ ,  $D$  i  $T$  jak dla przykładu z rys. 1. Podsumowując można powiedzieć, że w każdym przypadku poziom mocy otrzymanego w wyniku konwer-



Rys. 1. Kształt impulsu o czasie trwania  $T = 100$  ps przy  $\omega_{d,\max} = 0,05 \frac{2\pi Tc}{LD\lambda^2}$  otrzymanego przez rozwiązanie zależności (2.32) — 1 i drogą symulacji numerycznej — 2;  $P$  — moc optyczna j.d. — jednostki dowolne



Rys. 2. Kształt impulsu o czasie trwania  $T = 100$  ps przy  $\omega_{d,\max} = 0,001 \frac{2\pi Tc}{LD\lambda^2} P$  — moc optyczna, j.d.  
— jednostki dowolne

sji impulsu optycznego na wyjściu światłowodu zależny od:

- relacji między czasem trwania impulsu a maksymalną dewiacją pulsacji;
- długości światłowodu;
- wartości współczynnika dyspersji;
- mocy na wejściu światłowodu.

### 3. OPIS ZJAWISKA KONWERSJI MODULACJI CZĘSTOTLIWOŚCI NA MODULACJĘ AMPLITUDEY Z DWUPOZIOMOWYM SYGNAŁEM WYJŚCIOWYM

Wyniki otrzymane z zależności (2.30) nie oddają dokładnego kształtu impulsów optycznych uzyskanych w wyniku konwersji modulacji częstotliwości na modulację amplitudy. Zgodność wyników otrzymanych z tej zależności są zgodne z wynikami numerycznymi tylko w zakresie bardzo małych wartości dewiacji pulsacji. Poniżej przedstawiona analiza doprowadza do takiego opisu omawianego zjawiska, który pozwala na uzyskanie wyników zgodnych z wynikami otrzymanymi drogą numeryczną w szerokim zakresie zmian maksymalnej dewiacji pulsacji.

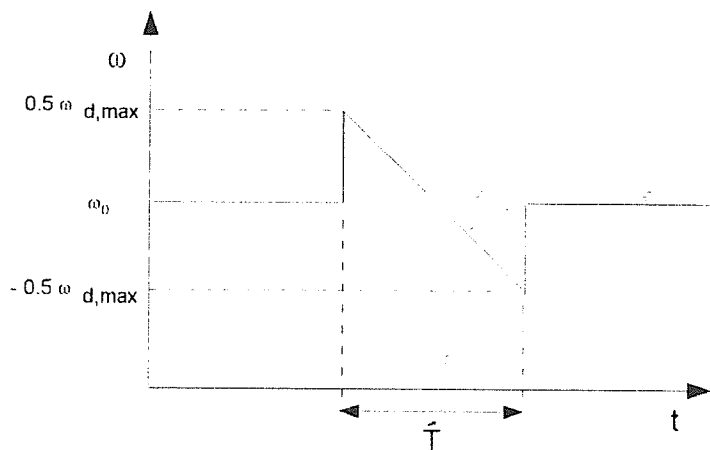
Dla poniższego opisu zjawiska konwersji przyjmujemy, że fala zmodulowana częstotliwościowo ma postać:

$$\Psi(t) = A_{we} \operatorname{Re}(\exp\{j\varphi(t)\}), \quad (3.1)$$

gdzie:

$A_{we}$  — amplituda,

$\varphi$  — faza.



Rys. 3. Zmiana pulsacji  $\omega_d$  w funkcji czasu  $t$  (oznaczenia w tekście)

Dla przypadku przedstawionego na rys. 3 zmianę fazy w czasie można opisać w następującej formie:

$$\varphi(t) = \int_0^t \omega_d(t) dt; \quad (3.2)$$

$$\omega_d(t) = -\frac{\omega_{d,\max} \cdot t}{T}. \quad (3.3)$$

Ze wzoru (3.3) i (3.4) otrzymujemy:

$$\varphi(t) = -\frac{\omega_{d,\max}}{T} \cdot \frac{t^2}{2}. \quad (3.4)$$

Przyjęto, że przebieg mocy, przy pominięciu strat, na końcu światłowodu o długości  $L$  wynosi:

$$P_{wy}(t, L) = \left| A_{we} \left( 1 + \frac{m_f}{2} \Psi_{F-A}(t, L) \right) \right|^2, \quad (3.5)$$

gdzie:  $m_f$  — wskaźnik modulacji równy  $\frac{f_{d,\max} T}{2}$  przy czym  $f_{d,\max}$  jest odległością między największą, a najmniejszą częstotliwością, a  $T$  czasem trwania bitu.  $\Psi_{F-A}$  — czynnik związany z konwersją modulacji częstotliwości na modulację amplitudy. Korzystając z transformaty Fouriera i odwrotnej transformaty Fouriera czynnik związany z konwersją można przedstawić w następującej postaci:

$$\Psi_{F-A}(t, L) = \int_{-\infty}^{\infty} \frac{d\omega}{2\pi} \exp(-j\beta L) \exp(j\omega t) \int_{-\frac{T}{2}}^{\frac{T}{2}} d\tau \cos\{\varphi(\tau)\} \exp(-j\omega \tau). \quad (3.6)$$

We wzorze (3.6) całka po  $d\tau$  jest transformatą Fouriera przebiegu wejściowego, a całka po  $d\omega$  jest odwrotną transformatą Fouriera.

Zależność (3.6) dalej przekształcamy do następującej postaci:

$$\Psi_{F-A}(t, L) = \int_{-\frac{T}{2}}^{\frac{T}{2}} d\tau \int_{-\infty}^{\infty} \frac{d\omega}{2\pi} \cos\{\varphi(\tau)\} \exp(-j\beta L) \exp(j\omega(t-\tau)). \quad (3.7)$$

Rozwijając stałą fazy  $\beta$  wokół pulsacji nośnej  $\omega$  otrzymujemy:

$$\beta = \beta_0 + (\omega - \omega_0)\beta_1 + \frac{1}{2}(\omega - \omega_0)^2\beta_2 + \dots \quad (3.8)$$

Składowa  $\beta_0$  związana jest z fazą fali nośnej. Pochodne stałej fazowej  $\beta_1$  i  $\beta_2$  względem pulsacji odnoszą się odpowiednio do jednostkowego czasu przejścia modu i dyspersji prędkości grupowej. W dalszej analizie ograniczamy się do analizy wpływu zjawiska dyspersji prędkości grupowej, pomijamy  $\beta_0 L$  i  $\beta_1 L$  odnoszące się odpowiednio do opóźnienia fazowego fali nośnej i opóźnienia czasowego sygnału modulującego. Biorąc pod uwagę zależności (3.2), (3.7), (3.8) i powyższe założenia uzyskujemy:

$$\Psi_{F-A}(t, L) = \frac{1}{2\pi} \int_{-\frac{T}{2}}^{\frac{T}{2}} d\tau \cos\left(-\frac{\omega_{d,\max}\tau^2}{T} \frac{1}{2}\right) \times \int_{-\infty}^{\infty} d\omega \exp\left(-j\omega^2 \frac{L}{2}\right) \exp(j\omega(t-\tau)). \quad (3.9)$$

Wiedząc, że  $\cos\left(-\frac{\omega_{d,\max}\tau^2}{T} \frac{1}{2}\right) = \cos\left(\frac{\omega_{d,\max}\tau^2}{T} \frac{1}{2}\right)$  i po rozpisaniu  $\cos\left(\frac{\omega_{d,\max}\tau^2}{T} \frac{1}{2}\right)$  na

następującą postać: 
$$\frac{e\left(j\frac{\omega_{d,\max}\tau^2}{T} \frac{1}{2}\right) + e\left(-j\frac{\omega_{d,\max}\tau^2}{T} \frac{1}{2}\right)}{2}$$
 i rozwiązaniu całki po  $d\omega$  z zależności (3.9) otrzymujemy:

$$\begin{aligned} \Psi_{F-A}(t, L) = & \frac{1}{4\pi} \sqrt{\frac{2\pi}{i\beta_2 L}} \exp\left(-j\frac{t^2}{2\beta_2 L}\right) \times \\ & \times \left\{ \int_{-\frac{T}{2}}^{\frac{T}{2}} d\tau \exp\left(j\frac{\tau^2}{2} \left(\frac{\omega_{d,\max}}{T} - \frac{1}{\beta_2 L}\right)\right) \exp\left(j\tau \frac{t}{\beta_2 L}\right) + \right. \\ & \left. + \int_{-\frac{T}{2}}^{\frac{T}{2}} d\tau \exp\left(j\frac{\tau^2}{2} \left(-\frac{\omega_{d,\max}}{T} - \frac{1}{\beta_2 L}\right)\right) \exp\left(j\tau \frac{t}{\beta_2 L}\right) \right\}. \end{aligned} \quad (3.10)$$

Wyniki otrzymane z rozwiązania zależności (3.10) porównano z wynikami obliczeń numerycznych dotyczących konwersji modulacji częstotliwości na modulację amplitudy dla zmiany częstotliwości pokazanej na rys. 3. Do analizy numerycznej zjawiska konwersji w światłowodzie standardowym wykorzystano wyniki otrzymane z rozwiązania numerycznego nieliniowego równania Schrödingera. Analizie pod-

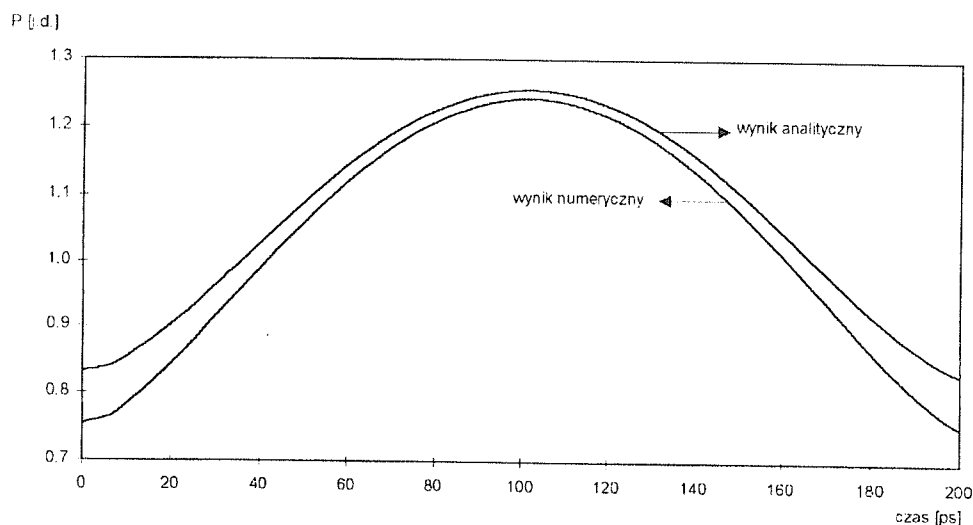
dana została konwersja zmiany pulsacji  $\omega_d$  w funkcji czasu  $t$  pokazanej na rys. 3, przy przyjęciu czasu trwania zmiany  $T = 100$  ps. Badania zostały przeprowadzone dla maksymalnej dewiacji pulsacji równej:

$$\omega_{d,\max} = k \frac{2\pi Tc}{LD\lambda^2},$$

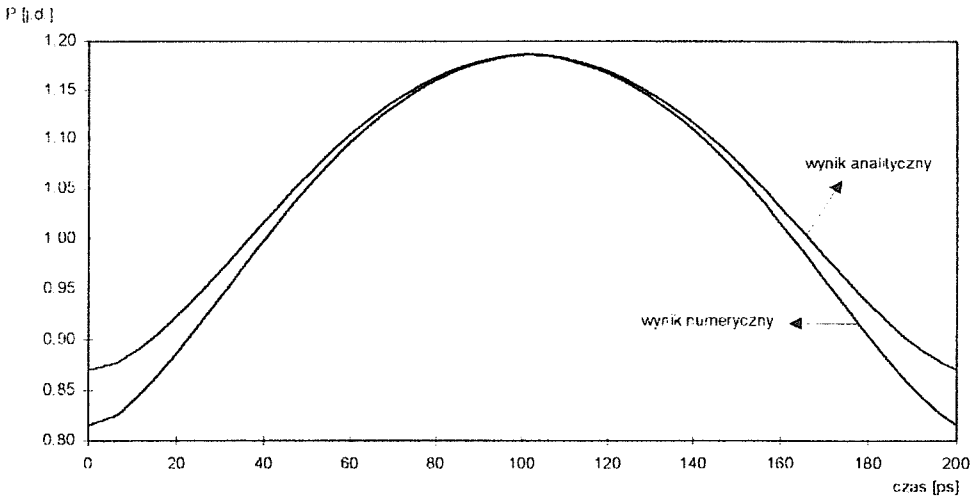
gdzie  $k = 0,25; 0,5; 0,75$  i  $1$ ,  $T$  — czas trwania zmiany pulsacji,  $c$  — prędkość światła,  $\lambda$  — długość fali światła,  $L$  — długość światłowodu,  $D$  — współczynnik dyspersji.

Przyjęto, że  $P_{ve} = 1$  [j.d.],  $T = 100$  ps,  $\lambda = 1550$  nm,  $L = 140$  km i  $D = 17 \frac{\text{ps}}{\text{nm} \cdot \text{km}}$ . Poniżej, na rysunkach 4÷7 przedstawiono kształty impulsu uzyskanego w wyniku rozwiązania zależności (3.10) i numerycznej symulacji efektu konwersji.

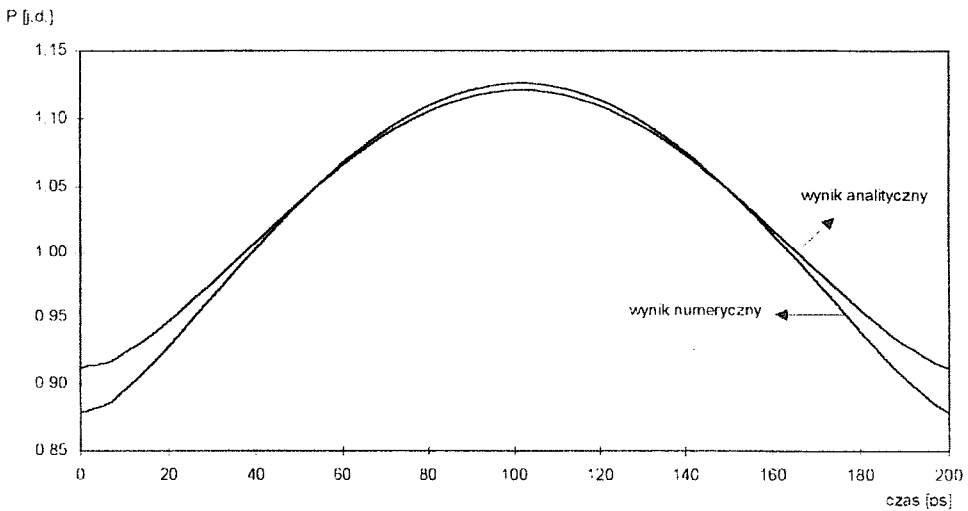
Analiza konwersji pojedynczego impulsu pozwala na stwierdzenie, że rozwiązanie równania (3.10) daje rozwiązanie porównywalne z wynikami uzyskanymi drogą numeryczną. Maksymalne różnice między wynikami danymi przez obie metody nie przekraczają 10%. Od  $\omega_{d,\max} = \frac{2\pi Tc}{LD\lambda^2}$  do  $\omega_{d,\max} = 0,75 \frac{2\pi Tc}{LD\lambda^2}$  maksymalna wartość mocy impulsu uzyskana z zależności (3.10) jest większa od rozwiązania numerycznego. Maleje ona, zbliżając się do wyniku uzyskanego numerycznie, wraz ze zmniejszeniem wartości  $\omega_{d,\max}$ . Od  $\omega_{d,\max} = 0,5 \frac{2\pi Tc}{LD\lambda^2}$  wartość maksymalna mocy



Rys. 4. Kształt impulsu przy  $\omega_{d,\max} = \frac{2\pi Tc}{LD\lambda^2}$ ; j.d. — jednostki dowolne

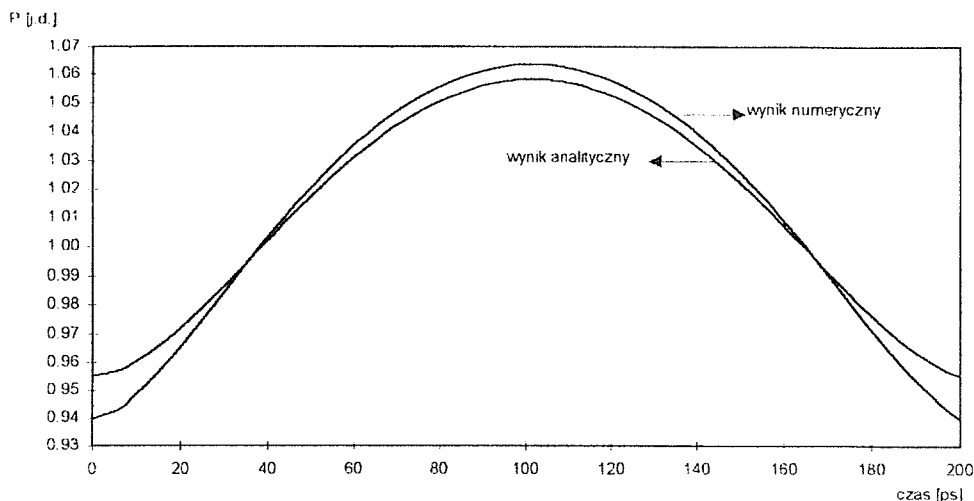


Rys. 5. Kształt impulsu przy  $\omega_{d,\max} = 0,75 \frac{2\pi Tc}{LD^2}$ ; j.d. --- jednostki dowolne



Rys. 6. Kształt impulsu przy  $\omega_{d,\max} = 0,5 \frac{2\pi Tc}{LD^2}$ ; j.d. --- jednostki dowolne

impulsu uzyskana z (3.10) jest mniejsza od wartości uzyskanej numerycznie, różnica między nimi wzrasta wraz ze zmniejszeniem wartości  $\omega_{d,\max}$ . Na podstawie przeprowadzonych badań można stwierdzić, że zależność (3.10) pozwala na dokładną analizę impulsów optycznych otrzymanych w światłowodzie w wyniku konwersji modulacji częstotliwości na modulację amplitudy, tak co do ich kształtu jak i co do ich poziomu mocy.



Rys. 7. Kształt impulsu przy  $\omega_{d,max} = 0,25 \frac{2\pi Tc}{LD\lambda^2}$ ; j.d. --- jednostki dowolne

#### 4. SYSTEM TRANSMISJI Z WYKORZYSTANIEM KONWERSJI MODULACJI CZĘSTOTLIWOŚCI NA MODULACJĘ AMPLITUDY

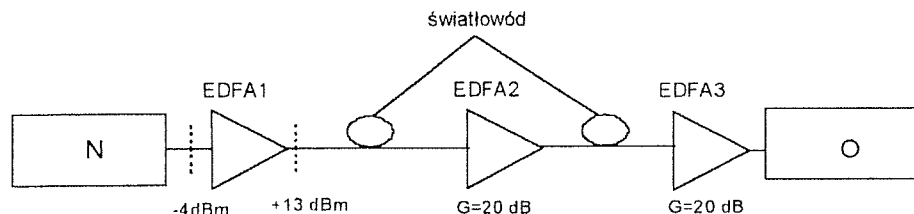
Poniżej przedstawiono metodę transmisji danych o przepływności 5 Gbit/s zakodowanych za pomocą kodu bifazowego (Manchester) opartą na konwersji modulacji częstotliwości na modulację amplitudy. Kody bifazowe, mimo swej podstawowej wady jaką jest ich większa szerokość pasma w porównaniu z innymi kodami, mają takie zalety jak: samosynchronizacja, brak składowej stałej czy też możliwość wykrywania za ich pomocą błędów w transmisji [17]. W kodzie bifazowym maksymalna długość ciągu stanów wysokich (jedynek logiczne) i stanów niskich (zera logiczne) wynosi dwa. Ułatwia to odpowiedni dobór zmian pulsacji, które reprezentują stan wysoki i stan niski; dla przepływności 5 Gbit/s najkrótszy stan wysoki i niski trwa 100 ps, a najdłuższy 200 ps. Transmisja z wykorzystaniem konwersji modulacji częstotliwości na modulację amplitudy została przebadana w systemie opisanym w podrozdziale 4.1. Zastosowanie zjawiska konwersji do transmisji danych wymaga rozwiązania problemu wyboru charakteru zmiany pulsacji w funkcji czasu, dotyczy tego podrozdział 4.2, oraz doboru poziomu maksymalnej zmiany pulsacji (podrozdział 4.3).

##### 4.1. Opis badanego systemu

Analiza transmisji z wykorzystaniem konwersji została przeprowadzona dla systemu przedstawionego na rys. 8 o następujących parametrach:

##### 1. Nadajnik

- moc wyjściowa z lasera — 4 dBm,
- moc wyjściowa ze wzmacniacza końcowego (EDFA1) + 13 dBm.



Rys. 8. Schemat badanego systemu. N — nadajnik (laser), EDFA1 — wzmacniacz końcowy nadajnika, EDFA2 — wzmacniacz umieszczony na 100 km, EDFA3 — przedwzmacniacz w odbiorniku, O — odbiornik z filtrem

## 2. Tor optotelekomunikacyjny

- wzmacnienie wzmacniacza EDFA2 równe 20 dB,
- światłowód jednomodowy standardowy o współczynniku dyspersji równej  $D = 17 \text{ ps/nm} \cdot \text{km}$  i tłumienności  $\alpha = 0,2 \text{ dB/km}$ .

## 3. Odbiornik

- wzmacnienie przedwzmacniacza w odbiorniku (EDFA3) równe 20 dB,
- wartość fotoprądu wytworzonego przez fotodiodę wynosi:

$$i(t) = \frac{\eta \lambda q}{hc} P_{wy}(t) + I_d, \quad (4.1)$$

gdzie:  $I_d$  — prąd ciemny,  $\lambda$  — długość fali światła,  $\eta$  — sprawność kwantowa detektora,  $q$  — ładunek elektronu,  $h$  — stała Plancka,  $c$  — prędkość światła.

- w odbiorniku zastosowano filtr Butterwortha o następującej charakterystyce amplitudowej [16]:

$$H_{FR}(j\omega) = \frac{1}{\sqrt{1 + \left(\frac{\omega}{\omega_R}\right)^2}}, \quad (4.2)$$

gdzie:

$\omega_R$  — pulsacja przy której następuje spadek wartości  $H_{FR}(j\omega)$  o 3 dB;

$\omega$  — pulsacja.

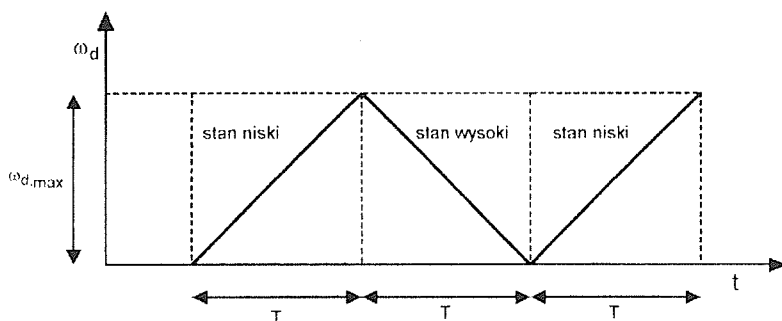
W celu określenia zachowania się systemu w funkcji długości traktu światłowodowego ( $L$ ) przyjęto, że:

- dla małych odległości tj. dla  $L$  od 0 km do 30 km w systemie występuje przedwzmacniacz EDFA3 i brak jest wzmacniaczy EDFA1 i EDFA2;
- dla średnich odległości tj. dla  $L$  od 40 km do 120 km obok przedwzmacniacza EDFA3 występuje wzmacniacz końcowy EDFA1;
- dla dużych odległości tj. dla  $L \geq 140 \text{ km}$  dodatkowo w celu skompensowania strat na setnym kilometrze umieszczono wzmacniacz przelotowy EDFA2.



#### 4.2. Zmiana pulsacji w funkcji czasu

Charakter zmiany pulsacji w funkcji czasu został uwarunkowany założeniem, że na końcu traktu światłowodowego wymagane jest otrzymanie dwupoziomowego impulsu optycznego. Analiza przedstawiona w rozdziale 2 i 3, w szczególności wzory (2.30)-(2.32) wskazują, że spełnienie tego założenia wymaga takiej zmiany pulsacji w funkcji czasu jak to przedstawiono na rys. 3. Dla zakresu tzw. dyspersji anormalnej, tj. w zakresie długości fal przy których  $\beta_2 < 0$ , składowe przesyłanego sygnału o większych częstotliwościach (mniejszych długościach fali) poruszają się szybciej niż składowe o mniejszych częstotliwościach (większych długościach fali). W związku z tym stan niski powinien być reprezentowany przez zmianę pulsacji określoną zależnością  $\omega_d(t) = \frac{\omega_{d,\max}}{T}t$ , a stan wysoki przez  $\omega_d(t) = \omega_{d,\max} - \frac{\omega_{d,\max}}{T}t$ , z odpowiednim przesunięciem w skali czasu (rys. 9).



Rys. 9. Zmiany pulsacji odpowiadające poszczególnym stanom logicznym

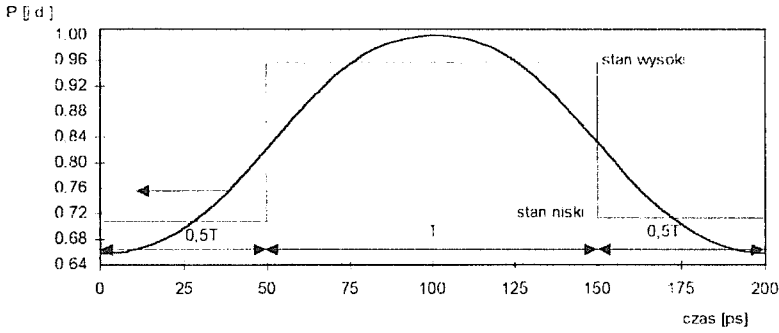
Na rys. 10÷13 przedstawiono zmianę poziomu mocy optycznej impulsów jaką uzyskano w wyniku zjawiska konwersji dla sytuacji z rys. 9 przy różnych wartościach  $\omega_{d,\max}$ . Przyjęto, że maksymalna dewiacja pulsacji jest równa:

$$\omega_{d,\max} = k \frac{2\pi Tc}{LD\lambda^2}, \quad (4.3)$$

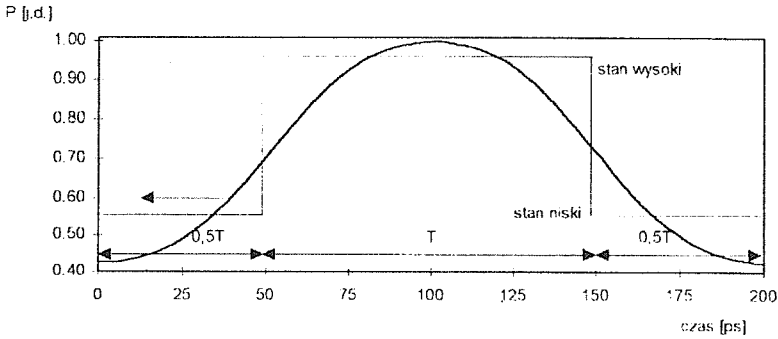
gdzie:  $k = 0,25; 0,5; 1$  i  $2$ ,  $T$  — czas trwania impulsu,  $c$  — prędkość światła,  $\lambda$  — długość fali światła,  $L$  — długość światłowodu,  $D$  — współczynnik dyspersji.

Przyjęto, że:  $T = 100$  ps,  $\lambda = 1550$  nm,  $L = 140$  km,  $D = 17 \frac{\text{ps}}{\text{nm} \cdot \text{km}}$ ;

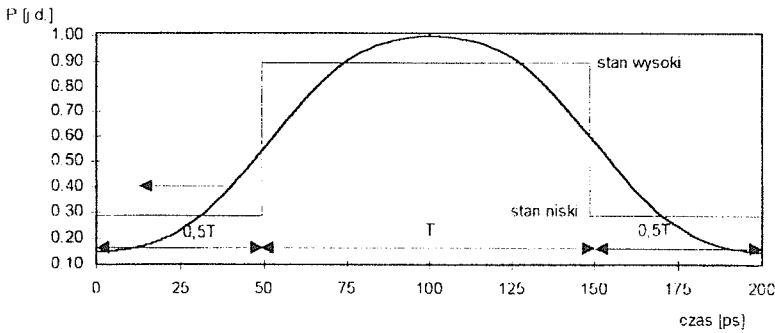
Moc optyczna impulsów otrzymanych w wyniku zjawiska konwersji dla przypadku zmian częstotliwości pokazanych na rys. 9 wzrasta wraz ze zwiększeniem maksymalnej dewiacji pulsacji  $\omega_{d,\max}$ . W szerokim zakresie zmian maksymalnej



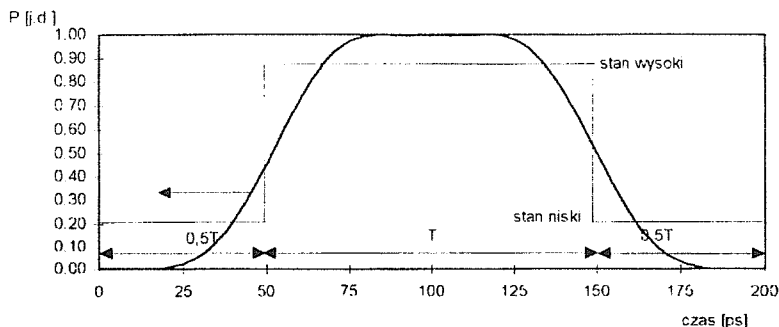
Rys. 10. Impuls optyczny otrzymany w wyniku zjawiska konwersji przy  $\omega_{d,\max} = 0,25 \frac{2\pi Tc}{LD\lambda^2}$ ; P — znormalizowana moc optyczna, j.d. — jednostki dowolne



Rys. 11. Impuls optyczny otrzymany w wyniku zjawiska konwersji przy  $\omega_{d,\max} = 0,5 \frac{2\pi Tc}{LD\lambda^2}$ ; P — znormalizowana moc optyczna, j.d. — jednostki dowolne



Rys. 12. Impuls optyczny otrzymany w wyniku zjawiska konwersji przy  $\omega_{d,\max} = \frac{2\pi Tc}{LD\lambda^2}$ ; P — znormalizowana moc optyczna, j.d. — jednostki dowolne



Rys. 13. Impuls optyczny otrzymany w wyniku zjawiska konwersji przy  $\omega_{d,\max} = 2 \frac{2\pi Tc}{LD\lambda^2}$ ; P — znormalizowana moc optyczna, j.d. — jednostki dowolne

wartości dewiacji pulsacji, tj. do  $\omega_{d,\max} = 2 \frac{2\pi Tc}{LD\lambda^2}$ , maksymalne wartości stanu wysokiego i minimalna wartość stanu niskiego przypadają w połowie ich czasu trwania. Stabilność kształtu impulsu w tak szerokim zakresie zmian  $\omega_{d,\max}$  jest prawdopodobnie spowodowana wymianą mocy między impulsami zajmującymi przedziały czasowe sąsiednich bitów, sąsiednie impulsy spełniają w pewnym sensie rolę stabilizującą. Wybierając moment próbkowania w połowie czasu trwania stanu i określając poziom progowy jako różnicę między poziomem mocy optycznej w połowie czasu trwania stanu wysokiego, a poziomem mocy optycznej w połowie czasu trwania stanu niskiego, można prawidłowo rozróżnić stan niski od stanu wysokiego dla szerokiego zakresu zmian  $\omega_{d,\max}$  (do  $k = 2$ ).

#### 4.3. Wartość maksymalnej zmiany pulsacji

Najprostszym rozwiązaniem problemu wyboru maksymalnej wartości zmiany pulsacji  $\omega_{d,\max}$  dla stanów o czasie trwania 100 ps i 200 ps byłoby przyjęcie założenia, że dla stanów o różnym czasie trwania nachylenie zmiany pulsacji w funkcji czasu jest jednakowe tzn.

$$\frac{\omega_{d,\max, 100 \text{ ps}}}{T_{100 \text{ ps}}} = \frac{\omega_{d,\max, 200 \text{ ps}}}{T_{200 \text{ ps}}}, \quad (4.4)$$

gdzie:

$\omega_{d,\max, 100 \text{ ps}}$  — maksymalna dewiacja pulsacji dla stanu o czasie  $T = 100 \text{ ps}$  ( $T_{100 \text{ ps}}$ );

$\omega_{d,\max, 200 \text{ ps}}$  — maksymalna dewiacja pulsacji dla stanu o czasie  $T = 200 \text{ ps}$  ( $T_{200 \text{ ps}}$ ).

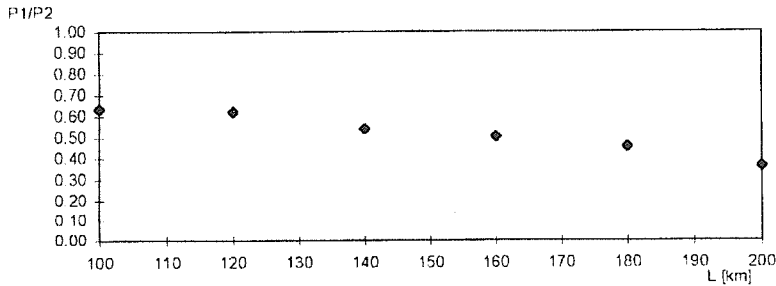
W przypadku jednakowego nachylenia zmiany pulsacji w funkcji czasu powinien zostać spełniony następujący warunek:

$$\omega_{d,\max, 200 \text{ ps}} = 2 \cdot \omega_{d,\max, 100 \text{ ps}}. \quad (4.5)$$

Poniżej, na rys. 14, pokazano stosunek mocy optycznej impulsu o czasie trwania 100 ps (P1) do mocy optycznej impulsu o czasie trwania 200 ps (P2) dla długości traktu światłowodowego  $L = 100$  km,  $L = 120$  km,  $L = 160$  km,  $L = 180$  km i  $L = 200$  km. Założono, że:

$$\omega_{d,\max,100\text{ ps}} = 0,5 \frac{2\pi Tc}{LD\lambda^2}. \quad (4.6)$$

Na rys. 14 widoczne jest, że przy założeniu (4.5) nie jesteśmy w stanie otrzymać jednakowego poziomu mocy optycznej dla impulsów o różnym czasie trwania dla danej długości toru światłowodowego ( $L$ ). Potwierdzają to badania numeryczne



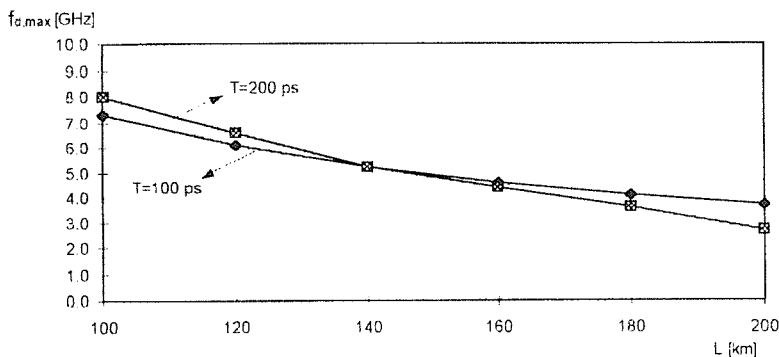
Rys. 14. Stosunek mocy optycznej impulsu o czasie trwania 100 ps do mocy optycznej impulsu o czasie trwania 200 ps w funkcji długości toru światłowodowego

przeprowadzone również dla innych wartości  $\omega_{d,\max,100\text{ ps}}$ . Do znalezienia odpowiednich par maksymalnych dewiacji pulsacji, tj. pulsacji dla  $T = 100$  ps i  $T = 200$  ps, które pozwoliłyby na uzyskanie zbliżonych poziomów mocy przyjęto na wstępie określoną wartość  $\omega_{d,\max}$  dla  $T = 100$  ps i dalej na tej podstawie poszukiwano takiej wartości zmiany pulsacji dla  $T = 200$  ps przy której otrzymamy poziom mocy równy z poziomem mocy dla krótszego stanu. Przyjęto też, że dopuszczalny błąd w różnicy mocy optycznej dla obydwu czasów trwania jest na poziomie mniejszym od 1%. Założono, że  $\omega_{d,\max}$  dla  $T = 100$  ps równa się:

$$\omega_{d,\max,100\text{ ps}} = \frac{2\pi Tc}{LD\lambda^2}, \quad (4.7)$$

gdzie  $T$ ,  $c$ ,  $L$ ,  $D$ ,  $\lambda$  zgodnie ze wzorem (4.3).

Na rys. 15 przedstawiono znalezione pary maksymalnej zmiany pulsacji dla stanów o różnym czasie trwania pozwalające na uzyskanie w wyniku zjawiska konwersji, na danej długości traktu światłowodowego ( $L$ ), jednakowego poziomu mocy optycznej. Na podstawie rys. 15 możemy dobrać odpowiednie pary wartości pulsacji dla długości toru światłowodowego z zakresu od 100 km do 200 km. Dla  $L = 160$  km maksymalna dewiacja pulsacji wynosi dla  $T = 100$  ps  $\omega_{d,\max} = 2\pi \cdot 4,6$  GHz, a dla  $T = 200$  ps  $\omega_{d,\max} = 2\pi \cdot 4,4$  GHz. Dla tych danych



Rys. 15. Pary maksymalnej zmiany pulsacji dla stanów o czasie trwania  $T = 100$  ps i  $T = 200$  ps

w funkcji długości toru światłowodowego przy założeniu  $\omega_{d,\max,100\text{ ps}} = \frac{2\pi Tc}{LD\lambda^2}$ ;  $\omega_{d,\max} = 2\pi f_{d,\max}$

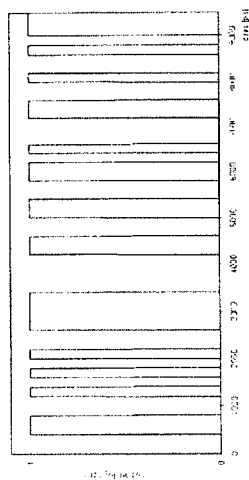
przedstawiono na rys. 16 zasadę działania systemu transmisji wykorzystującego konwersję modulacji częstotliwości na modulację amplitudy. Rys. 16d) przedstawia poziom mocy impulsów uzyskanych w wyniku efektu konwersji po 160 km. Pomimo doboru maksymalnych zmian pulsacji pod kątem otrzymania impulsów o jednakowym poziomie mocy optycznej uzyskano przesunięcie między impulsami o czasie trwania  $T = 100$  ps, a impulsami o czasie trwania  $T = 200$  ps ( $\Delta P_d$ ). Maksymalna wartość tego przesunięcia wynosi około 0,5 mW, przy wartości mocy impulsów równej 1,7 mW, co stanowi blisko 30% mocy impulsu. Taka sytuacja jak ta którą przedstawiono występuje również i dla innych długości traktu światłowodowego i odpowiadających im par wartości maksymalnej zmiany pulsacji z rys. 15. Przesunięcie między impulsami powoduje, że na wyjściu światłowodu otrzymujemy sygnał wielopoziomowy, co z kolei utrudnia wybór właściwego poziomu progowego. Przyjęto, że przesunięcie impulsów o takiej wartości jest zbyt duże, aby system ten poprawnie działał. Dalej przeprowadzono analizę systemu przy założeniu, że  $\omega_{d,\max}$  dla  $T = 100$  ps jest równa:

$$\omega_{d,\max,100\text{ ps}} = 0,5 \frac{2\pi Tc}{LD\lambda^2}, \quad (4.8)$$

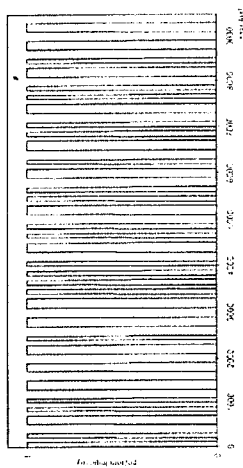
gdzie  $T$ ,  $c$ ,  $L$ ,  $D$ ,  $\lambda$  zgodnie ze wzorem (4.3).

Na rys. 17 przedstawiono znalezione pary maksymalnej zmiany pulsacji dla stanów o różnym czasie trwania pozwalające na uzyskanie w wyniku zjawiska konwersji, na długości traktu światłowodowego ( $L$ ), jednakowych poziomów mocy optycznej. Dla  $L = 160$  km maksymalna dewiacja pulsacji wynosi dla  $T = 100$  ps  $\omega_{d,\max} = 2\pi \cdot 2,3$  GHz, a dla  $T = 200$  ps  $\omega_{d,\max} = 2\pi \cdot 2,45$  GHz. Podobnie jak w poprzednim przypadku poniżej przedstawiono zasadę działania systemu transmisji wykorzystującego konwersję modulacji częstotliwości na modulację amplitudy dla  $L = 160$  km i  $\omega_{d,\max,100\text{ ps}} = 2\pi \cdot 2,3$  GHz oraz  $\omega_{d,\max,200\text{ ps}} = 2\pi \cdot 2,45$  GHz. Rys. 18d) przedstawia poziom mocy impulsów uzyskanych w wyniku efektu konwersji po

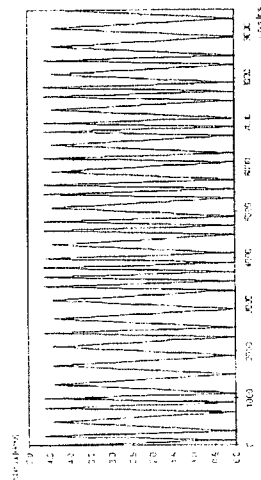
a)



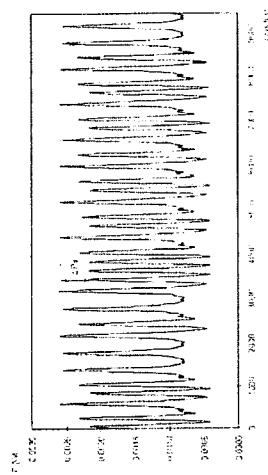
b)



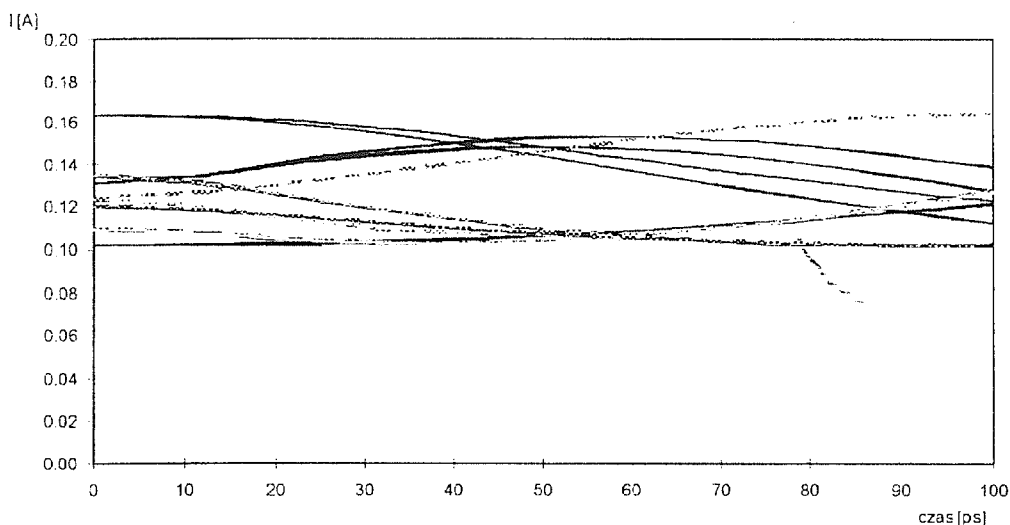
c)



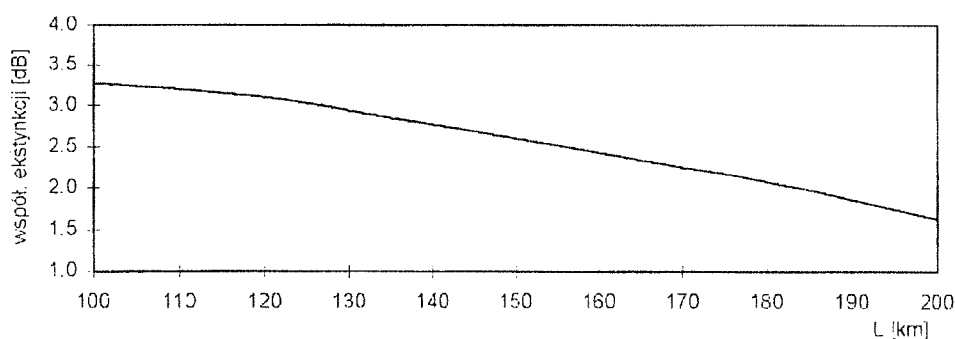
d)



Rys. 16. Transmisja danych z wykorzystaniem konwersji modulacji częstotliwości na modulację amplitudy dla  $L = 160$  km i  $\omega_{d,\max,100\text{ps}} = 2\pi \cdot 4,6$  GHz oraz  $\omega_{d,\max,200\text{ps}} = 2\pi \cdot 4,4$  GHz; a) ciąg przesyłanych danych, b) ich reprezentacja w kodzie bifazowym Manchester, c) odpowiadająca im zmiana częstotliwości, d) przebieg mocy optycznej na końcu toru światłowodowego;  $f_{d,\max}$  — maksymalna częstotliwości, P — moc optyczna



Rys. 17. Pary maksymalnej zmiany pulsacji dla stanów o czasie trwania  $T = 100$  ps i  $T = 200$  ps w funkcji długości toru światłowodowego przy założeniu  $\omega_{d,\max,100\text{ ps}} = 0,5 \frac{2\pi Tc}{LD\lambda^2}$ ,  $\omega_{d,\max} = 2\pi f_{d,\max}$

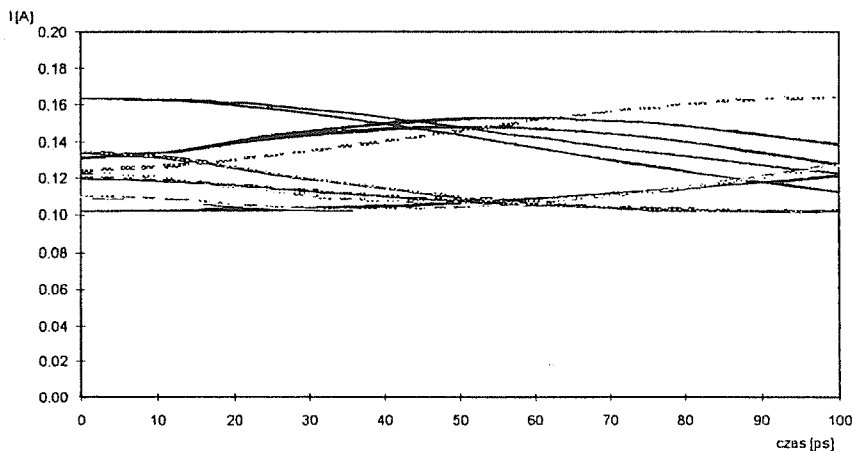


Rys. 18. Transmisja danych z wykorzystaniem konwersji modulacji częstotliwości na modulację amplitudy dla  $L = 160$  km i  $\omega_{d,\max,100\text{ ps}} = 2\pi \cdot 2,45$  GHz; a) ciąg przesyłanych danych, b) ich reprezentacja w kodzie bifazowym Manchester, c) odpowiadająca im zmiana częstotliwości, d) przebieg mocy optycznej na koncu toru światłowodowego;  $f_{d,\max}$  — maksymalna dewiacja częstotliwości,  $P$  — moc optyczna

160 km. W porównaniu z wynikami przedstawionymi na rys. 16d) przesunięcie między impulsami o czasie trwania  $T = 100$  ps, a impulsami o czasie trwania  $T = 200$  ps ( $\Delta P_d$ ) jest znacznie mniejsze. Największe przesunięcie jest na poziomie 0,15 mW przy mocy impulsu równej 0,92 mW. Wielkość ta wynosi 16% mocy impulsów, czyli jest prawie dwa razy mniejsze niż w poprzednim przypadku. Podobna sytuacja jak ta którą przedstawiono występuje również i dla innych długości traktu światłowodowego i odpowiadających im par wartości maksymalnej zmiany pulsacji z rys. 17. Na podstawie przeprowadzonej analizy relacji pomiędzy długością toru światłowodowego ( $L$ ), a wartościami dewiacji pulsacji  $\omega_{d, \max, 100 \text{ ps}}$  i  $\omega_{d, \max, 200 \text{ ps}}$  można zauważyć, że wraz ze wzrostem wartości maksymalnej dewiacji pulsacji wzrasta wartość przesunięcia ( $\Delta P_d$ ) między impulsami reprezentującymi stany o różnym czasie trwania. Wartość tego przesunięcia maleje wraz ze zmniejszeniem wartości zmiany pulsacji, a co za tym idzie również wraz ze zmniejszeniem poziomu mocy impulsów. Mamy stąd wyraźne ograniczenia doboru odpowiednich par wartości pulsacji:

- ograniczeniem „od dołu” jest otrzymanie małych poziomów mocy na wyjściu światłowodu;
- ograniczeniem „od góry” jest wzrost wartości  $\Delta P_d$ .

Przyjęto, że podstawą do dalszych badań systemu transmisji z wykorzystaniem konwersji modulacji częstotliwości na modulację amplitudy będzie przyjęcie założenia (4.8) i przyjęcie zmiany pulsacji dla poszczególnych długości toru pokazanych na rys. 17. Na rys. 19 pokazano wykres oczkowy jaki uzyskano dla systemu opisanego w podrozdziale 4.2, przy długości traktu światłowodowego  $L = 160$  km,  $f_R$  filtru Butterwortha równy 5 GHz, maksymalnej dewiacji pulsacji  $\omega_{d, \max, 100 \text{ ps}} = 2\pi \cdot 2,3$  GHz i  $\omega_{d, \max, 200 \text{ ps}} = 2\pi \cdot 2,45$  GHz i dla pseudoprzypadkowej sekwencji stanów niskich i wysokich o długości  $2^7 - 1$ .



Rys. 19. Wykres oczkowy dla  $L = 160$  km,  $f_R = 5$  GHz,  $\omega_{d, \max, 100 \text{ ps}} = 2\pi \cdot 2,3$  GHz  
i  $\omega_{d, \max, 200 \text{ ps}} = 2\pi \cdot 2,45$  GHz



## 5. WYNIKI SYMULACJI NUMERYCZNEJ

Dla długości traktu światłowodowego od 100 km do 200 km przebadano system pokazany na rys. 4 przy maksymalnej dewiacji pulsacji  $\omega_{dm \max, 100 \text{ ps}}$  i  $\omega_{d, \max, 200 \text{ ps}}$  przedstawionej na rys. 17. System przebadano pod kątem:

- wartości optycznego współczynnika ekstynkcji (EX);
- wpływu szumów przedwzmacniacza EDFA na parametry badanego systemu.

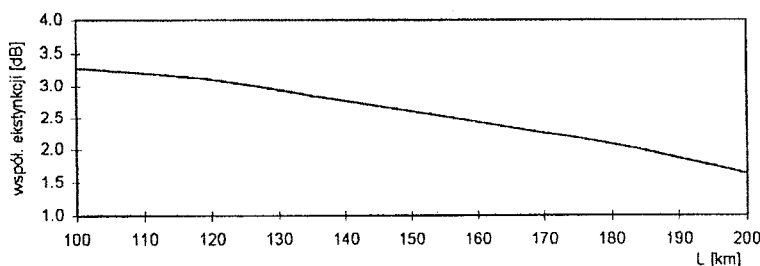
Badania zostały przeprowadzone na pseudoprzykładowej sekwencji stanów niskich i wysokich o długości  $2^7 - 1$ .

### 5.1. Współczynnik ekstynkcji

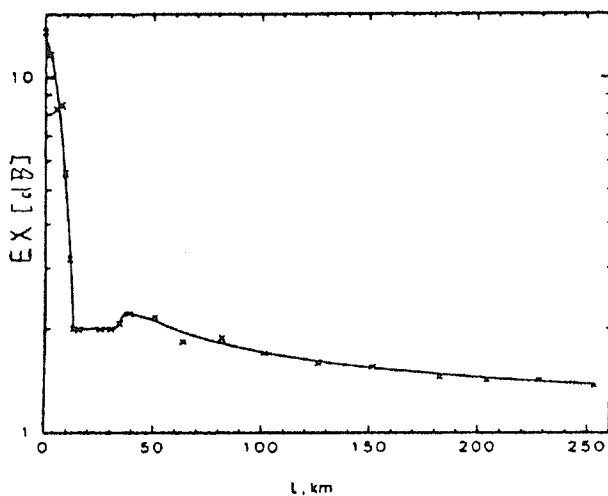
Przeprowadzono badania mające na celu określenie wartości optycznego współczynnika ekstynkcji dla transmisji danych na odległości od 100 km do 200 km. Wartość optycznego współczynnika ekstynkcji określono jako iloraz średniej wartości mocy, otrzymanej w miejscu próbkowania dla stanu wysokiego ( $\bar{P}_1$ ) i średniej wartości mocy dla stanu niskiego ( $\bar{P}_0$ ), zgodnie z zależnością:

$$\text{EX} = 10 \log \left( \frac{\bar{P}_1}{\bar{P}_0} \right). \quad (5.1)$$

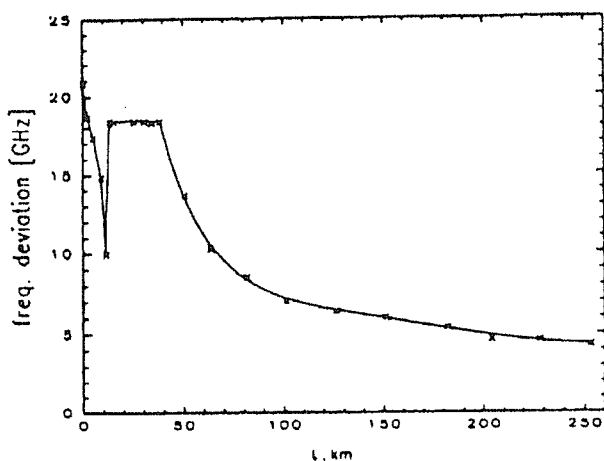
Na rys. 20 przedstawiono zmianę wartości optycznego współczynnika ekstynkcji w funkcji długości toru światłowodowego. Rys. 21 pokazuje wartość optycznego współczynnika ekstynkcji w funkcji długości traktu światłowodowego [5]. Wyniki zostały przez autorów otrzymane drogą eksperymentalną przy zmianach częstotliwości przedstawionych na rys. 22. W interesującym nas zakresie długości  $L$  tj. od 100 km do 200 km uzyskane wyniki są lepsze od przedstawionych w [5], pomimo stosowania prawie dwukrotnie mniejszej wartości maksymalnej dewiacji częstotliwości, która ma zasadnicze znaczenie dla wartości współczynnika ekstynkcji. Dla badanego systemu wartość optycznego współczynnika ekstynkcji w przedziale od 100 km do 180 km nie spada poniżej 2,0 dB; tylko dla  $L = 200$  km jest on mniejszy od 2,0 dB i wynosi 1,64 dB. Otrzymanie takich wartości współczynnika ekstynkcji wynika z zastosowania stosunkowo niewielkich wartości maksymalnej dewiacji



Rys. 20. Wartość współczynnika ekstynkcji w funkcji długości toru światłowodowego



Rys. 21. Wartość współczynnika ekstynkcji — EX w funkcji długości toru światłowodowego [5]



Rys. 22. Zmiany częstotliwości w funkcji długości toru światłowodowego [5]

pulsacji (rys.17). Wartość tego współczynnika, przy danej długości toru, wzrasta wraz ze wzrostem  $\omega_{d, \max}$ , ale zastosowanie większych zmian pulsacji ograniczone jest wystąpieniem zjawiska opisanego w podrozdziale 4.3.

## 5.2. Wpływ szumów wzmacniacza EDFA na parametry systemu

Tor odbiorczy badanego systemu (rys. 8) składa się z przedwzmacniacza EDFA3, fotodetektora i filtru. W celu określenia poziomu sygnału od szumu konieczne jest uwzględnienie wpływu szumów wzmacniacza EDFA3 na odbierany sygnał. Prąd powstający w detektorze wynosi:

$$I = I_p + I_d, \quad (5.2)$$

gdzie:  $I_d$  — prąd ciemny;

$$I_p = RGP_s, \quad (5.3)$$

gdzie  $R$  jest czułością detektora wynoszącą:

$$R = \frac{\eta q \lambda}{hc}, \quad (5.4)$$

przy czym:  $\lambda$  — długość fali światła,  $\eta$  — sprawność kwantowa detektora,  $q$  — ładunek elektronu,  $h$  — stała Plancka,  $c$  — prędkość światła,  $P_s$  — moc optyczna sygnału na wejściu wzmacniacza,  $G$  — wzmocnienie wzmacniacza.

Prąd szumów ma zerową średnią, a jego moc określa się wzorem:

$$\sigma^2 = \sigma_T^2 + \sigma_s^2 + \sigma_{\text{sig-sp}}^2 + \sigma_{\text{s-sp}}^2, \quad (5.5)$$

gdzie:

$$\text{szum termiczny: } \sigma_T^2 = (3kT/R_L) \Delta f \quad (5.6)$$

$$\text{szum śrutowy: } \sigma_s^2 = 2q[R(GP_s + P_{\text{sp}}) + I_d] \Delta f \quad (5.7)$$

$$\text{szum zdudnienia emisji spontanicznej: } \sigma_{\text{sp-sp}}^2 = 4R^2 S_{\text{sp}}^2 P_s \Delta f \quad (5.8)$$

$$\text{szum zdudnienia sygnału z emisją spontaniczną: } \sigma_{\text{sig-sp}}^2 = 4R^2 S_{\text{sp}} P_s \Delta f \quad (5.9)$$

szum zdudnienia prądu śrutowego z emisją spontaniczną:

$$\sigma_{\text{s-sp}}^2 = 4qRS_{\text{sp}} \Delta v_{\text{opt}} \Delta f \quad (5.10)$$

$k$  — stała Boltzmanna;

$T$  — temperatura;

$R_L$  — rezystancja wejściowa;

$\Delta f$  — pasmo elektryczne odbiornika;

$\Delta v_{\text{opt}}$  — pasmo optyczne szumu emisji spontanicznej;

$P_{\text{sp}}$  — moc emisji spontanicznej wzmacniacza, która jest równa:

$$P_{\text{sp}} = S_{\text{sp}} \Delta v_{\text{sp}} \quad (5.11)$$

przy czym:  $S_{\text{sp}}$  — gęstość widmowa mocy szumów emisji spontanicznej (ASE),  $\Delta v_{\text{sp}}$  — efektywne pasmo emisji spontanicznej, równe w praktyce pasmu wzmacniacza lub filtru optycznego,  $I_d$  — prąd ciemny detektora.

Przyjmując, że  $S_{\text{sp}} = (G - 1)n_{\text{sp}}h\nu$ , gdzie  $n_{\text{sp}}$  jest współczynnikiem inwersji oraz podstawiając pod  $n_{\text{sp}}$  częściej stosowany współczynnik szumu wzmacniacza  $F_n \approx 2n_{\text{sp}}$  (dla  $G \gg 1$ ); wyrażenia (5.7)-(5.10) można przedstawić w następującej postaci [13]:

$$\sigma_s^2 = 2q\eta GP_s \Delta f / (h\nu); \quad (5.12)$$

$$\sigma_{sp-sp}^2 = (q\eta GF_n)^2 \Delta v_{opt} \Delta f; \quad (5.13)$$

$$\sigma_{sig-sp}^2 = 2(q\eta G)^2 F_n P_s \Delta f / (h\nu); \quad (5.14)$$

$$\sigma_{s-sp}^2 = 4q^2 \eta GF_n \Delta v_{opt} \Delta f; \quad (5.15)$$

Porównując powyższe zależności można stwierdzić, że  $\sigma_s^2$  jest mniejsze od  $\sigma_{sig-sp}^2$  o czynnik  $\eta GF_n$  i może zostać pominięte. Składowik  $\sigma_{s-sp}^2$  jest pomijalnie mały w porównaniu z  $\sigma_{sp-sp}^2$ . Szum termiczny może zostać pominięty w porównaniu ze składnikami  $\sigma_{sig-sp}^2$  i  $\sigma_{sp-sp}^2$ . Przy takich założeniach oraz przy przyjęciu, że moc związana ze stanem niskim nie jest równa zero, prądy szumów  $\sigma_1$  i  $\sigma_0$  mogą być przybliżone zależnościami:

$$\sigma_1 = (\sigma_{sig-sp}^2 + \sigma_{sp-sp}^2)^{\frac{1}{2}} \quad (5.16)$$

$$\sigma_0 = (\sigma_{sig-sp}^2 + \sigma_{sp-sp}^2)^{\frac{1}{2}}. \quad (5.17)$$

Wprowadzamy parametr  $Q$ , będący stosunkiem różnicy średniej wartości prądu odpowiadającego stanowi wysokiemu i niskiemu do poziomu szumów [13]:

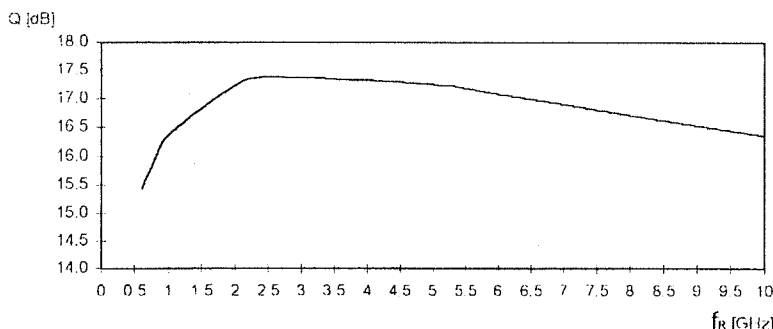
$$Q = \frac{\bar{I}_1 - \bar{I}_0}{\sigma_1 + \sigma_0}. \quad (5.18)$$

Wielkości  $\bar{I}_1$  i  $\bar{I}_0$  odpowiadają średniej wartości prądu związanej odpowiednio ze stanem wysokim i niskim. Przy obliczeniach przyjęto następujące założenia:  $\Delta v_{opt} = 125$  GHz,  $q = 1,6 \cdot 10^{-19}$  C,  $\eta = 0,9$ ,  $G = 20$  dB,  $h = 6,63 \cdot 10^{-34}$  J · s,  $c = 3,0 \cdot 10^8$  m/s,  $\lambda = 1550$  nm,  $F_n = 3$  dB,  $I_d = 10$  nA.

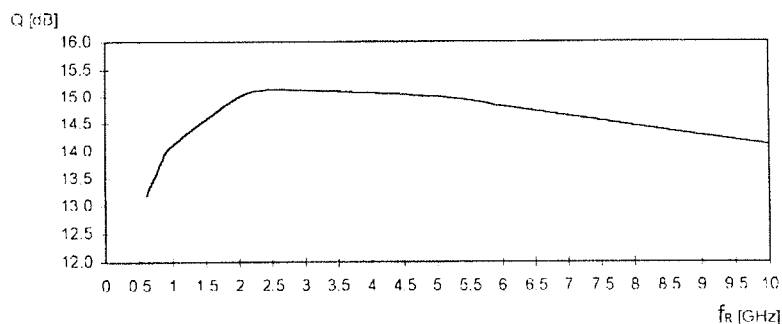
Na rys. 23-28 przedstawiono zależność parametru  $Q$  w funkcji  $f_R$  filtru dla długości toru światłowodowego równego odpowiednio 100 km, 120 km, 140 km, 160 km, 180 km i 200 km.

Zauważalny jest wzrost wartości parametru  $Q$  wraz ze wzrostem  $f_R$  filtru do wartości  $f_R = 6$  GHz.

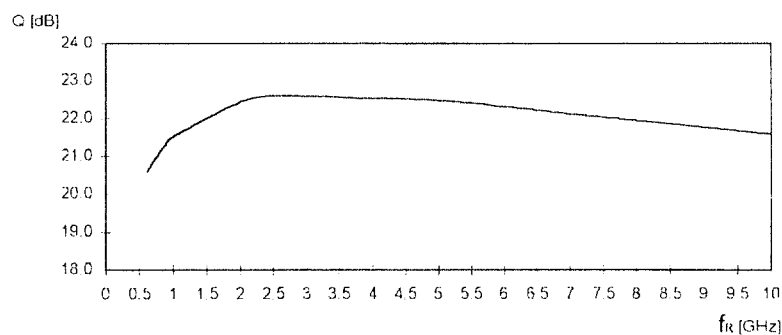
W tab. 1 przedstawiono jego wartość przy  $f_R = 5$  GHz.



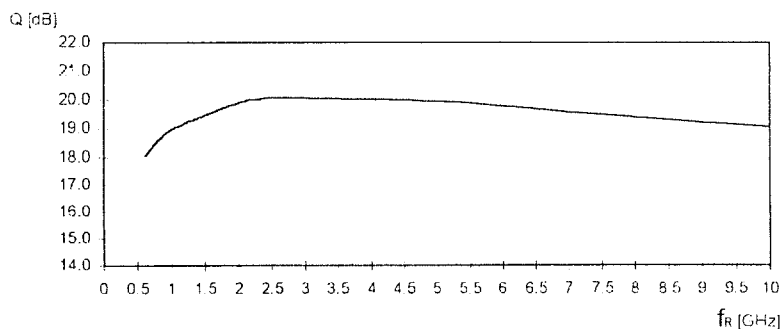
Rys. 23. Zmiana parametru  $Q$  w funkcji  $f_R$  dla długości toru światłowodowego  $L = 100$  km



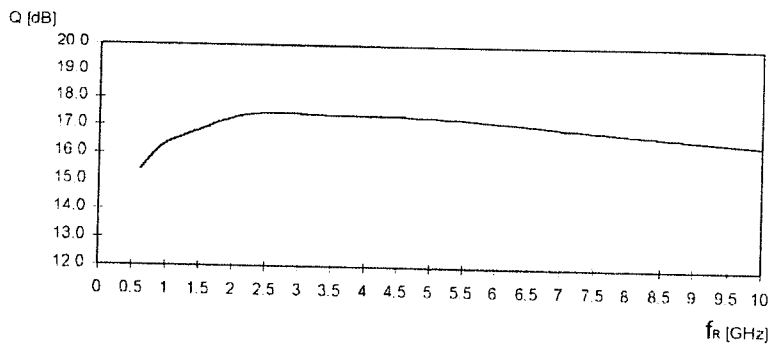
Rys. 24. Zmiana parametru  $Q$  w funkcji  $f_R$  dla długości toru światłowodowego  $L = 120$  km



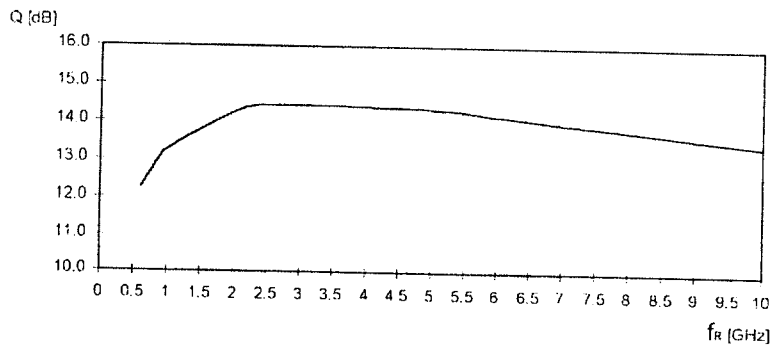
Rys. 25. Zmiana parametru  $Q$  w funkcji  $f_R$  dla długości toru światłowodowego  $L = 140$  km



Rys. 26. Zmiana parametru  $Q$  w funkcji  $f_R$  dla długości toru światłowodowego  $L = 160$  km



Rys. 27. Zmiana parametru  $Q$  w funkcji  $f_R$  dla długości toru światłowodowego  $L = 180$  km



Rys. 28. Zmiana parametru  $Q$  w funkcji  $f_R$  dla długości toru światłowodowego  $L = 200$  km

Tabela I  
Wartość parametru  $Q$  dla różnych długości toru światłowodowego przy  $f_R = 5$  GHz

| Długość traktu światłowodowego [km] | Q [dB] |
|-------------------------------------|--------|
| 100                                 | 17,39  |
| 120                                 | 15,14  |
| 140                                 | 22,59  |
| 160                                 | 20,07  |
| 170                                 | 17,46  |
| 180                                 | 14,43  |

Tabela 2  
Wartość parametru  $Q$  dla różnych długości toru  
światłowodowego przy  $f_R = 10$  GHz

| Długość traktu<br>światłowodowego [km] | $Q$ [dB] |
|----------------------------------------|----------|
| 100                                    | 16,36    |
| 120                                    | 14,10    |
| 140                                    | 21,59    |
| 160                                    | 19,07    |
| 170                                    | 16,45    |
| 180                                    | 13,45    |

Wartość parametru  $Q$  dla każdego z przypadku, w zakresie od  $f_R = 2,5$  GHz do  $f_R = 5$  GHz, zmienia się na poziomie 0,1 dB. Od  $f_R = 5$  GHz do  $f_R = 10$  GHz jego wartość maleje dochodząc przy  $f_R = 10$  GHz do wartości pokazanych w tab. 2. Spadek wartości parametru  $Q$  w zakresie od  $f_R = 5$  GHz do  $f_R = 10$  GHz wynosi tylko około 1 dB wynika ze wzrostu współczynnika ekstynkcji wraz ze wzrostem  $f_R$  filtru. Jego wpływ na  $Q$  jest w tym zakresie znacznie silniejszy niż wpływ na jego wartość szerokości pasma filtru.

Na podstawie przeprowadzonych badań można przyjąć jako optymalną szerokości pasma odbiornika dla tej transmisji wartość 5 GHz. Wielkość ta jest zgodna z szerokością pasma Nyquista, która w przypadku zastosowania do przesyłania danych kodu bifazowego Manchester powinna wynosić  $\frac{1}{2T_{\min}}$ , gdzie  $T_{\min}$  jest czasem trwania najkrótszego stanu wysokiego lub niskiego, w naszym przypadku  $T_{\min}$  równa się 100 ps.

#### BIBLIOGRAFIA

1. P. Henry: *Lightwave primer*. IEEE Journal of Quantum Electronics. 1985, vol. QE-21, p. 1862
2. E. Bocheve, E. de Carvalho: *FM-AM conversion by material dispersion in an optical fiber*. Optics Letters. 1981, vol. 6, nr 2, p. 58
3. E. Bocheve, E. de Carvalho: *Conversion of a wideband frequency-modulated signal to amplitude modulation through dispersion in an optical fiber*. Optics Letters. 1982, vol. 7, nr 3, p. 139.
4. D. Marcuse: *Computer simulation of FSK laser spectra and of FSK-to-ASK conversion*. IEEE Journal of Lightwave Technology. 1990. vol. 8, nr 7, p. 1110
5. B. Wedding, B. Franz, B. Juninger: *10-Gb/s optical transmission up to 253 km via standard single-mode fiber using the method of dispersion-supported transmission*. IEEE Journal of Lightwave Technology. 1994, vol. 12, nr 10, p. 1720
6. S. Fujita, M. Kitamura, T. Torikai: *10 Gbit/s, 100 km optical fibre transmission experiment using high-speed MQW DFB-LD and back-illuminated GaInAs APD*. Electronics Letters. 1989, vol. 25, p. 702

7. K. Petermann: *FM-AM noise conversion in dispersion single-mode fibre transmission lines*. Electronics Letters. 1990, vol. 26, nr 25, p. 2097
8. B. Wedding: *New method for optical transmission beyond dispersion limit*. Electronics Letters. 1992, vol. 28, nr 14, p. 1298
9. B. Wedding, B. Franz: *Unregenerated optical transmission at 10 Gbit/s via 204 km of standard singlemode fibre using a directly modulated laser diode*. Electronics Letters. 1993, vol. 29, nr 4, p. 402
10. C. Kurtzke, A. Gnauck: *Operating principle of in-line amplified dispersion-supported transmission*. Electronics Letters. 1993, vol. 29, nr 22, p. 1969
11. C. Kurtzke, S. Kindt, K. Petermann: *Impact of residual amplitude modulation on the performance of dispersion-mediated nonlinearity-enhanced transmission*. Electronics Letters. 1994, vol. 30, nr 12, pp. 988
12. G.P. Agrawal: *Nonlinear fiber optics*. Academic Press Inc., New York 1989, pp. 374-381
13. G.P. Agrawal: *Fiber-optical communication systems*. John Wiley & Sons, Inc., New York 1992, p. 44-48
14. A. Majewski: *Nieliniowa optyka światłowodowa*. Wydawnictwa Politechniki Warszawskiej, Warszawa 1993, ss. 74-80
15. J. Osowski: *Zarys rachunku operatorowego*. WNT, Warszawa 1981, p. 266-268
16. M.S. Rodan: *Systemy telekomunikacji analogowej i cyfrowej*. WNT, Warszawa 1983, ss. 80-83
17. H.B. Killen: *Transmisja cyfrowa w systemach światłowodowych i satelitarnych*. WKiŁ, Warszawa 1992, ss. 42-49

#### K. PERLICKI

#### S u m m a r y

The conversion of frequency modulated signal to amplitude modulation is described. A novel modulation scheme for dispersion supported transmission is proposed.

A new analysis of conversion of a frequency-modulated signal to amplitude modulation through dispersion in an optical fiber is presented. In this article is compared results of this analysis and results numerical simulation phenomena of FM-AM conversion. The propagation of optical field through a fiber is described by nonlinear Schrödinger equation, which is solved numerically by the split-step Fourier method.

The computer simulations are carried out for a system with the following characteristics. The transmitted optical signal is modulated by an augmented PRBS of length  $2^7 - 1$ . The fiber dispersion is 17 ps/nm · km, the loss is 0,2 dB/km. Optical in-line amplifier is inserted at 100 km. The gain of the amplifier is chosen to compensation of the fiber loss. A first order Butterworth filter at the receiver is used. Binary frequency modulation of transmitter laser a linear frequency modulation along with the Manchester coding is used. It gives optical binary signal at the receiver with greater values of excitation coefficient as compared to the previous research at the region of 100 km  $< L < 200$  km.

*Key words:* optical communication, frequency modulation to amplitude modulation conversion, dispersion supported transmission.



## SPIS TREŚCI

|                                                                                                                                                      |     |
|------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| W. Czajewski: Zbiory przybliżone w rozpoznawaniu znaków                                                                                              | 143 |
| Z. Drzycimski, J. Majewski, Z. Zakrzewski, I. Jaworski: Własności korelacyjne i widmowe niestacjonarnych sygnałów zmodulowanych                      | 155 |
| D. Kania: Pojemność kodowa programowalnych transkoderów                                                                                              | 193 |
| A. Pfizner, B. Jucht-Suchta, A. Świt: Półempiryczne warunki brzegowe w złączu p–n dla symulacji przyrządów półprzewodnikowych                        | 205 |
| L. Tomawski: Zmienna pojemność z użyciem konweytorów prądowych sterowana przez rezystor                                                              | 219 |
| F. Wysocka-Schillak: Projektowanie cyfrowego korektora fazy NOI przy użyciu metod optymalizacji nieliniowej                                          | 231 |
| K. Wiatr, P. Russek: Sprzętowe architektury kompresji obrazu wizyjnego dla potrzeb przetwarzania w czasie rzeczywistym                               | 245 |
| K. Perlicki: Zastosowanie konwersji modulacji częstotliwości na modulację amplitudy w światłowodowych systemach transmisji o dużych przepływnościach | 259 |

## CONTENTS

|                                                                                                                                |     |
|--------------------------------------------------------------------------------------------------------------------------------|-----|
| W. Czajewski: A new approach to OCR                                                                                            | 143 |
| Z. Drzycimski, J. Majewski, Z. Zakrzewski, I. Jaworski: Correlation and spectral properties of nonstationary modulated signals | 155 |
| D. Kania: Coding capacity of programmable transcoder                                                                           | 193 |
| A. Pfizner, B. Jucht-Suchta, A. Świt: Semi empirical boundary conditions at p–n junctions for device simulation                | 205 |
| L. Tomawski: Varying capacitances with the use of current conveyors controlled by a resistor                                   | 219 |
| F. Wysocka-Schillak: Design of IIR delay equalizers using nonlinear optimization methods                                       | 231 |
| K. Wiatr, P. Russek: Hardware architectures of image compression for a real time data computation                              | 245 |
| K. Perlicki: High capacity optical transmission using FM–AM supported transmission                                             | 259 |

# INFORMACJE DLA AUTORÓW

Redakcja przyjmuje do publikowania prace oryginalne, przeglądowe i monograficzne wchodzące w zakres szeroko pojętej elektroniki. Ponieważ KWARTALNIK ELEKTRONIKI i TELEKOMUNIKACJI jest czasopismem Komitetu Elektroniki i Telekomunikacji Polskiej Akademii Nauk, w związku z tym na jego łamach znajdują się prace naukowe dotyczące podstaw teoretycznych i zastosowań z zakresu elektroniki, telekomunikacji, mikroelektroniki, optoelektroniki, radiotechniki i elektroniki medycznej.

Artykuły powinny charakteryzować oryginalne ujęcie zagadnienia, własna klasyfikacja, krytyczna ocena (teorii lub metod), omówienie aktualnego stanu, lub postępu danej gałęzi techniki oraz omówienie perspektyw rozwojowych.

Artykuły publikowane w innych czasopismach nie mogą być kierowane do druku w Kwartalniku Elektroniki i Telekomunikacji w drugiej kolejności zgłoszenia.

Objętość artykułu nie powinna przekraczać 30 stron po około 1800 znaków na stronie.

**Wymagania podstawowe:** Artykuły należy nadsyłać w maszynopisie pisanym jednostronnie lub na wyraźnym czarno-białym wydruku komputerowym, w 2 egzemplarzach, w języku polskim lub angielskim wybranym przez autora. Tekst artykułu musi być poprzedzony tytułem pracy, imieniem i nazwiskiem Autora wraz z podaniem miejsca jego pracy. Wszystkie strony muszą mieć numerację ciągłą.

**Sposób pisania tekstu:** Tekst powinien być pisany bez używania wyróżnień, a w szczególności nie dopuszcza się spacjiowania, podkreślania i pisania tekstu dużymi literami z wyjątkiem wyrazów, które umownie pisze się dużymi literami (np. FORTRAN). Proponowane wyróżnienia Autor może zaznaczyć w maszynopisie zwykłym ołówkiem za pomocą przyjętych znaków adiustacyjnych) np. podkreślenie linią przerywaną oznacza spaciowanie (rozstrzelenie), podkreślenie linią ciągłą — pogrubienie, podkreślenie wężykiem — kursywa. Tekst powinien być napisany z podwójnym odstępem między wierszami, tytuły i podtytuły małymi literami. Marginesy z każdej strony powinny mieć około 35 mm. Przy podziale pracy na rozdziały i podrozdziały cyfrowe ich oznaczenia nie powinny być większe niż III stopnia (np. 4.1.1).

**Sposób pisania tablic:** Tablice powinny być pisane na oddzielnych stronach. Tytuły rubryk pionowych i poziomych powinny być napisane małymi literami z podwójnym odstępem między wierszami. Przypisy (notki) dotyczące tablic należy pisać bezpośrednio pod tablicą. Tablice należy numerować kolejno liczbami arabskimi, u góry każdej tablicy podać tytuł. Tablice umieścić na końcu maszynopisu. Przyjmowane są tablice algorytmów i programy na wydrukach komputerowych. W tym przypadku zachowany jest ich oryginalny układ.

**Sposób pisania wzorów matematycznych:** Rozmieszczenie znaków, cyfr, liter i odstępów powinno być zbliżone do rozmieszczenia elementów druku. Wskaźniki i wykładniki potęg powinny być napisane wyraźnie i być prawidłowo obniżone lub podwyższone w stosunku do linii wiersza podstawowego. Znaki nad literami i cyframi, całkami i in. symbolami (strzałki, linie, kropki, daszki) powinny być umieszczone dokładnie nad tymi elementami, do których się odnoszą. Numery wzorów cyframi arabskimi powinny być kolejne i umieszczone w nawiasach okrągłych z prawej strony. Nazwy jednostek, symbole literowe i graficzne powinny być zgodne z wytycznymi IEE (International Electrotechnical Commission) oraz ISO (International Organization of Standardization).

**Powołania:** Powołania na publikacje powinny być umieszczone na ostatnich stronach tekstu pod tytułem „Bibliografia”, opatrzone numeracją kolejną bez nawiasów. Numeracja ta powinna być zgodna z odnośnikami w treści artykułu. Przykłady opisu publikacji:

- periodycznej: F. Valdoni: A new millimetre wave satellite, E.T.T. 1990, vol. 2, nr 5, p. 553
- nieperiodycznej: K. Andersen: A resource allocation framework. XVI International Symposium, Stockholm (Sweden), May 1991, paper A 2.4
- książki: Y.P. Tvidis: Operation and modelling of the MOS transistors. Mc Graw-Hill, New York 1987, p. 141 — 148

**Materiał ilustracyjny:** Rysunki powinny być wykonane wyraźnie, na papierze gładkim, lub milimetrowym w formacie nie mniejszym niż  $9 \times 12$  cm. Mogą być także w postaci wydruku komputerowego. Fotografie lub diapozytywy przyjmowane są raczej czarno-białe w formacie nie przekraczającym  $10 \times 15$  cm. Na marginesie każdego rysunku i na odwrocie fotografii powinno być napisane ołówkiem imię i nazwisko Autora oraz skrót tytułu artykułu, do którego są przeznaczone. Spis podpisów pod rysunki i fotografie powinien być umieszczony na oddzielnej stronie.

**Streszczenia:** Do każdego artykułu musi być dołączone obszerne — do 60 wierszy — streszczenie z tytułem artykułu. Streszczenia (Summary) muszą być w języku polskim i angielskim. Streszczenie powinno wyjaśniać główny cel pracy, wskazywać korzyści i ograniczenia, możliwe zastosowania i zalecenia dla dalszego rozwoju danej gałęzi techniki. Pod streszczeniami powinny być podane słowa kluczowe.

Autorowi przysługuje bezpłatnie 20 odbitek artykułu. Dodatkowe egzemplarze odbitek, lub cały zeszyt Autor może zamówić u wydawcy na własny koszt.

Autora obowiązuje korekta autorska, którą powinien wykonać w ciągu 3 dni od daty otrzymania tekstu z Redakcji i zwrócić osobiście, lub listownie pod adresem Redakcji. Korekta powinna być naniesiona na przekazanych Autorowi szpaltach na marginesach ew. na osobnym arkuszu w przypadku uzupełnień tekstu większych niż dwa wiersze. W przypadku nie zwrócenia korekty w terminie, korektę przeprowadza Redakcja Techniczna Wydawcy.

Redakcja prosi Autorów o powiadomienie ją o zmianie miejsca pracy i adresu prywatnego.

# INFORMATION FOR THE AUTHORS OF K.E.T.

The KWARTALNIK ELEKTRONIKI I TELEKOMUNIKACJI (K.E.T.) — ELECTRONICS AND TELECOMMUNICATION QUARTERLY is a journal of the Committee for Electronics and Telecommunications of the Polish Academy of Sciences and is a direct continuation for 44 years of the quarterly ROZPRAWY ELEKTROTECHNICZE, that has been published also under the auspices of the Polish Academy of Sciences.

Its pages contain scientific articles concerned with theory and applications of electronics, telecommunication, microelectronics, optoelectronics, radioengineering and medical electronics.

Accordingly, its pages contain scientific articles, that are original works both contributory and concerned with comparative analyses of methods or systems synthetic reviews of a given field, state-of-art discussions of a given field, and presentations of developments in a given technology.

Regular papers are accepted for publication in Polish or English, but title, the extensive summary and figure caption must be bilingual. Bilingual summary must show the significance of the results shown in the paper, emphasizing limitations and advantages, possible applications and recommendations for further works.

K.E.T. is meant for specialists and students with above subjects.

## FORMAL PREPARATION OF MANUSCRIPT

### Texts

Texts to be submitted to K.E.T. may be typed on one side of the paper only with double spacing between the lines and a 4 cm margin or using of wordprocessing with text formatters based on MS DOS, APPLE/MAC such as Microsoft Word. In any case the text formatter must be specified on the diskette. Each text should be preceded by: a sheet with paper title, name (in full) and surname of the Author/s, their Affiliations and relevant addresses.

All text pages should be numbered. Regular papers should have a length of no more than 30 pages.

### Printing types and symbols

For letteral symbols, units and graphical symbols, the recommendations of IEE (International Electrotechnical Commission) and ISO (International Organisation of Standardisation) must be followed.

- In case of manuscripts conventionally typed (i.e. without wordprocessing), letters to be printed in italic should be underlined; letters to be printed in bold should be double underlined.
  - Formulae and footnotes should be numbered in accordance with the quotation order followed in the text.
  - The serial number of each formula should be written at right side, between round brackets.
  - The serial number of footnotes should be written between round brackets and upper.
  - Mathematical details, ancillary to the main discussion of the paper may be included in one or more appendixes, that are printed after conclusions and before references.
- In manuscript containing lot of symbols, a glossary may be included after introduction.

### References

References are usually gathered at the end of the text and numbered according to the order of quotation in the text. The serial numbers should be written without square brackets (i.e. 1., 2....)

### Illustrations

The originals of illustrations may be either drawings or photographs (black or colour). Each copy of paper must include two complete set of illustrations.

All rights reserved. No part of the publication may be reproduced, stored in a retrieval system or transmitted, in any form or by any means: electronic, photocopying, recording or otherwise, without the prior permission of the copyright owner — Relaction K.E.T.

**Subscription for external subscribers:**

The promotional subscription price in 1997 is S 90 including postage for institutions. Subscriptions should be sent to the publisher, Polish Scientific Publishers PWN Ltd, Journal Division, Miodowa 10, 00-250 Warsaw, POLAND, fax (48) (22) 826 09 50, (48) (22) 826 71 63, with a copy by fast to Electronics and Telecommunications Quarterly 00-665 Warsaw, ul. Nowowiejska 15/19, p. 470. Subscription is occupied after showing a cheque or transfer documents. Our bank account is as follows: Bank account: PBK VIII O/Warszawa nr 11101037-1052-2700-1-43. At subscriber's request this journal will be air mailed at additional postage to 50% gross price to European countries and 65% overseas.

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

Ruch S.A. fulfils foreign customers orders, starting from any issue in the calendar year: tel. (48) (22) 620 120 39; fax (48) (22) 620 17 62.

# Kwartalnik Elektroniki i Telekomunikacji

## WARUNKI PRENUMERATY

**Wpłaty na prenumeratę** przyjmują jednostki kolportażowe RUCH S.A., i wszystkie urzędy pocztowe na terenie całego kraju, właściwe dla miejsca zamieszkania lub siedziby prenumeratora oraz doręczyciele w miejscowościach, gdzie dostęp do urzędu jest utrudniony.

Od osób lub instytucji, zamieszkałych lub mieszkających się w miejscowościach, w których nie ma jednostek kolportażowych RUCH, wpłaty należy wносить do RUCH S.A. Oddział Krajowej Dystrybucji Prasy, 00-958 Warszawa, ul. Towarowa 28. Konto: PBK S.A. XIII Oddział Warszawa nr 11101053-16551-2700-1-67. RUCH S.A. i POCZTA POLSKA zapewniają dostawę pod wskazanym adresem pocztą zwykłą w ramach opłaconej prenumeraty.

**Terminy przyjmowania wpłat na prenumeratę krajową:** RUCH S.A. do 5 XII na I kw., do 5 III na II kw., do 5 VI na III kw. i do 5 IX na IV kw. POCZTA POLSKA do 25 XI na I kw., do 25 na II kw., do 25 V na III kw. i do 25 na IV kw.

**Na zagranicę prenumeratę przyjmuje** RUCH S.A. Oddział Krajowej Dystrybucji Prasy, 00-958 Warszawa, ul. Towarowa 28. Konto: PBK S.A. XII Oddział Warszawa nr 11101053-16551-2700-1-67. Dostawa odbywa się pocztą zwykłą w ramach opłaconej prenumeraty, z wyjątkiem zlecenia dostawy pocztą lotniczą, której koszt w pełni pokrywa zleceniodawca.

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

**Terminy wpłat:** do 20 XI na I kw., do 20 II na II kw., do 20 V na III kw. i do 20 VIII na IV kw.

RUCH S.A. fulfils foreign customers' orders, starting from any issue in the calendar year: tel. (48)(22) 620 10 19; fax (48)(22) 620 10 39.

**Prenumeratę przyjmuje Zakład Kolportażu Wydawnictwa SIGMA-NOT**, 00-716 Warszawa, ul. Bartycka 20, skr. poczt. 1004, termin wpłat do 5 XII roku poprzedzającego.

Konto PBK S.A. III Oddział Warszawa nr 11101024-1573-2720-3-28

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 49-30-86 lub 40-00-21 wew. 249, 295, 299.

**Bieżące numery** można nabyć w **Księgarni Wydawnictwa Naukowego PWN**, ul. Miodowa 10, 00-251 Warszawa, tel. (22) 695 40 26. Również można je nabyć, a także zamówić (przesyłka za zaliczeniem pocztowym) w Księgarni Ośrodka Rozpowszechniania Wydawnictw Naukowych PAN, ul. Twarda 51/55, 00-818 Warszawa, tel. (22) 697 88 35.

**All journals edited by PWN are available through:**

Foreign Trade Enterprise  
ARS POLONA  
Krakowskie Przedmieście 7,  
00-068 Warszawa, Poland  
fax (48) (22) 826 86 73

or

Polish Scientific Publishers PWN Ltd  
10 Miodowa Str.  
00-251 Warszawa, Poland  
fax (48) (22) 826 09 50  
fax (48) (22) 695 42 88

W roku 1997

prenumerata roczna w kraju

34,00 zł

prenumerata roczna

90 \$ USA