

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

KWARTALNIK
ELEKTRONIKI I TELEKOMUNIKACJI

ELECTRONICS AND
TELECOMMUNICATIONS
QUARTERLY

TOM 45 — ZESZYT 1



WYDAWNICTWO NAUKOWE PWN
WARSZAWA 1999

RADA REDAKCYJNA

Przewodniczący

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

Członkowie

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

REDAKCJA

Redaktor Naczelny

prof. dr hab. inż. WIESŁAW WOLIŃSKI

Zastępca Redaktora Naczelnego

doc. dr inż. KRYSTYN PLEWKO

Sekretarz Odpowiedzialny

mgr KRYSTYNA LELAKOWSKA

ADRES REDAKCJI

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

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

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

WYDAWNICTWO NAUKOWE PWN
Warszawa, ul. Miodowa 10

Ark. wyd. 9,0 Ark. druk. 7,5

Podpisano do druku w kwietniu 1999 r.

Papier offsetowy kl. III 70 g. B-1

Druk ukończono w maju 1999 r.

Skład: *Amel*, Warszawa

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

Szanowni Autorzy

„Kwartalnik Elektroniki i Telekomunikacji” — Electronics and Telecommunications Quarterly jest kontynuatorem tradycji powstałego 45 lat temu kwartalnika p.t. „Rozprawy Elektrotechniczne”.

Kwartalnik jest czasopismem Komitetu Elektroniki i Telekomunikacji Polskiej Akademii Nauk. Wydawany jest przez Wydawnictwo Naukowe PWN SA w Warszawie. Kwartalnik jest czasopismem naukowym, na którego łamach są publikowane artykuły i komunikaty prezentujące wyniki oryginalnych prac teoretycznych i doświadczalnych, a także przeglądowych. Związane są one z szeroko rozumianymi dziedzinami współczesnej elektroniki, telekomunikacji, mikroelektroniki, optoelektroniki, radiotechniki i elektroniki medycznej.

Autorami publikacji są wybitni naukowcy, znani specjaliści o wieloletnim doświadczeniu, a także młodzi badacze — głównie doktoranci.

Artykuły charakteryzują się oryginalnym ujęciem zagadnienia, interesującymi wynikami badań, krytyczną oceną teorii lub metod, omówieniem aktualnego stanu, lub postępu danej gałęzi techniki oraz omówieniem perspektyw rozwojowych. Sposób pisania matematycznej części artykułów zgodny jest z wytycznymi IEE (International Electronics Commission) oraz ISO (International Organization of Standardization).

Wszystkie publikowane w Kwartalniku artykuły są recenzowane przez znanych krajowych specjalistów, co zapewnia że publikacje te są uznawane jako autorski dorobek naukowy. Opublikowanie w kwartalniku wyników prac naukowych zrealizowanych w ramach „GRANTów” Komitetu Badań Naukowych spełnia więc jeden z wymogów stawianych tym pracom.

Czasopismo dociera do wszystkich zajmujących się elektroniką i telekomunikacją krajowych ośrodków naukowych oraz technicznych, a także szeregu instytucji zagranicznych. Jest ponadto prenumerowane przez liczne grono specjalistów i biblioteki.

Każdy Autor otrzymuje bezpłatnie 20 egzemplarzy nadbitek swojego artykułu, co ułatwia przesłanie go do indywidualnie wybranych przez Autora osób i instytucji w kraju, lub za granicą. Ułatwia to dodatkowo fakt, że w Kwartalniku są publikowane artykuły w języku angielskim.

Nadesłane do redakcji artykuły są publikowane w terminie około pół roku, w przypadku sprawnej współpracy Autora z redakcją. Wytyczne dla Autorów dotyczące formy publikacji są zamieszczone w zeszytach Kwartalnika. Można je także otrzymać w siedzibie redakcji.

Artykuły można dostarczać osobiście, lub pocztą pod adresem zamieszczanym na stronie redakcyjnej w każdym zeszycie.

Redakcja

Wp

Wpływ zmian parametrów obiektu opisanego transmitancją

$$G(s) = \frac{k}{(Ts + 1)^n}$$

JANUSZ HALAWA

Instytut Cybernetyki Technicznej, Politechnika Wrocławska

Otrzymano 1998.11.04

Autoryzowano 1999.03.24

W pracy przedstawiono wpływ zmian rzędu n i stałej czasowej modelu T opisanego transmitancją

$$G_n(s) = \frac{k}{(Ts + 1)^n} \quad n = 2, 3, 4, \dots$$

na funkcjonały

$$Q_{1n} = \int_0^{\infty} [(k - y_{n1}(t))] dt$$

i

$$Q = \sqrt{\int_0^{\infty} [y_n^k(t) - y_n(t)]^2 dt},$$

gdzie $y_n(t)$ i $y_n^k(t)$ są odpowiedziami modeli na wymuszenie $x(t) = 1$. Transmitancjami tymi opisuje się szereg obiektów przemysłowych. W wypadku występowania zmian ich parametrów w toku eksploatacji można stosować do sterowania układy adaptacyjne. Układy te mogą wykorzystać do sterowania zmiany przyjętego kryterium np. Q_{n1} lub Q .

Słowa kluczowe: modele matematyczne transmitancji, obiekty energetyczne, funkcje ciągłe.

Szereg obiektów energetycznych można opisać transmitancją

$$G_n(s) = \frac{y_n(s)}{x(s)} = \frac{k}{(Ts + 1)^n}, \quad (1)$$

lub transmitancją

$$G_{n0}(s) = \frac{y_{n0}(s)}{x(s)} = \frac{k}{(Ts + 1)^n} e^{st_0}. \quad (2)$$

W niektórych obiektach (na przykład w kotłach energetycznych) zauważa się zmianę współczynników T lub n albo obu tych współczynników łącznie w czasie eksploatacji. Dokonując identyfikacji współczynników transmitancji (1) na początku okresu eksploatacji możemy tak uzyskaną transmitancję traktować jako wyjściowy model matematyczny. Zdejmując zaś tę odpowiedź po pewnym czasie eksploatacji można ocenić zmiany parametrów obiektu i wygenerować sygnał do przestrojenia nastaw regulatorów lub np. automatycznego czyszczenia elementów obiektu. W warunkach przemysłowych pomiary obciążone są błędem wynikającym z zakłóceń. W pracy rozpatrzono modele bez zakłóceń (np. po ich odfiltrowaniu). W przestrzeni funkcji ciągłych można zdefiniować odległość dwóch funkcji ciągłych jako

$$Q = \sqrt{\int_0^{\infty} [(y_n^*(t) - y_n(t))]^2 dt}. \quad (3)$$

W pracy tej, we wzorze (3) i (4) $y_n(t)$ oraz y_n^* przyjmuje się jako odpowiedzi przyjętych modeli uproszczonych postaci (1) na wymuszenie skokowe $x(t) = 1(t)$.

1 kryterium

$$Q_{1n} = \int_0^{\infty} [(k - y_n(t))] dt \quad (4)$$

W wypadku transmitancji (1) otrzymuje się

$$\int_0^{\infty} [(k - y_n(t))] dt = knT,$$

gdzie $y_n(t)$ jest odpowiedzią modelu n -tego rzędu na wymuszenie skokowe $x(t) = 1(t)$ a $\lim_{t \rightarrow \infty} y_n(t) = k$.

Wynika stąd, że zmiana obiektu o jeden rząd powoduje zmianę kryterium o wartość:

$$Q_{1,n-1} - Q_{1n} = \int_0^{\infty} (y_{1,n-1} - y_{1,n}) dt = kT. \quad (5)$$

Zakładając więc model postaci (1) można otrzymać tę samą wartość kryterium Q_1 dla modelu np.

$$G_3(s) = \frac{y_3(s, 2)}{x(s)} = \frac{1}{(2s + 1)^3}$$

i dla modelu

$$(2) \quad G_2(s) = \frac{y_2(s, 3)}{x(s)} = \frac{1}{(3s + 1)^2}.$$

waża się w czasie początku wyjściowy eksploatacji estrojenia u. W wa- zakłóceń. zestrzeni

Uwaga ta jest istotna przy ustalaniu algorytmu identyfikacji obiektów opisanych transmitancjami typu (1) dla regulatorów samonastrajalnych sterujących na podstawie zmian przyjętego kryterium sterowania. Jak wynika z (4) nie można powiedzieć czy zmianę kryterium Q_1 spowodowała zmiana stałej czasowej T obiektu, czy zmiana rzędu obiektu. Zakładając stałość jednego parametru (T lub n) można ocenić zmiany drugiego parametru.

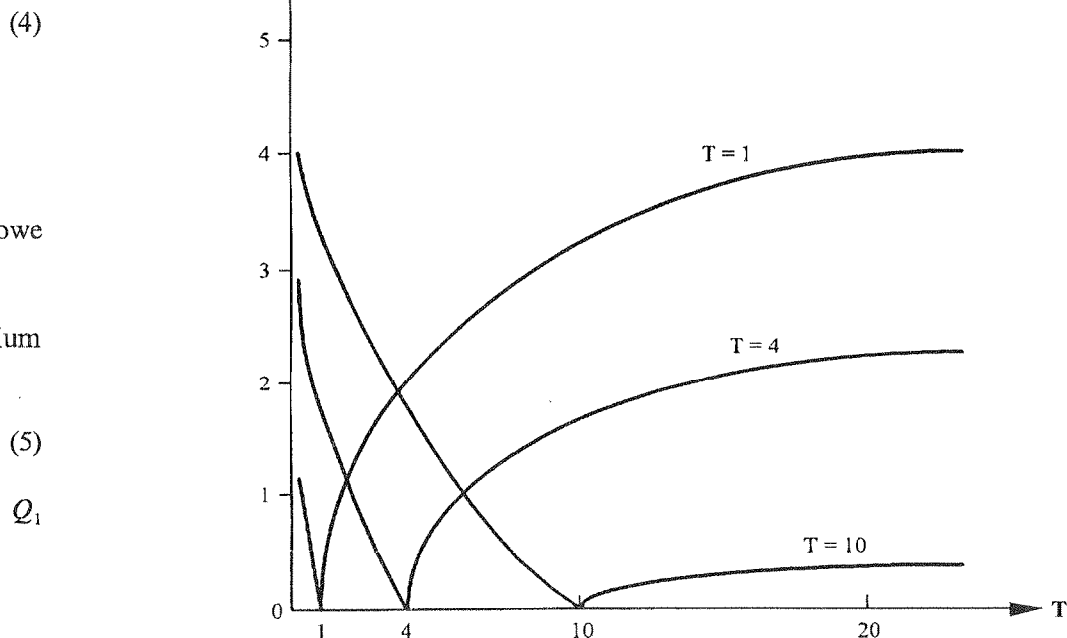
2 kryterium

$$Q = \sqrt{\int_0^{\infty} [(y_n^*(t) - y_n(t))]^2 dt}.$$

Rozpatrzone zostaną zmiany wartości kryterium Q powstałe w wyniku zmian wartości stałej czasowej T . Przykładowo, przyjęto obiekt postaci

$$(3) \quad G(s) = \frac{y_3(s)}{x(s)} = \frac{1}{(Ts + 1)^3}. \quad (6)$$

powiedzi
= 1(t).



erium Q_1

Rys. 1. Wpływ zmian parametru T modelu (6) na wartości kryterium Q

Funkcja $y_n^*(t)$ jest odpowiedzią skokową modelu tego samego rzędu dla różnej wartości stałej czasowej T . Na rysunku 1 przedstawione są wykresy zmian kryterium Q dla obiektów o trzech stałych czasowych $T = 1$, $T = 4$ i $T = 10$.

Jak wynika z tych wykresów te same wartości kryterium można uzyskać dla różnych wartości stałej T . Należy więc rozpatrywać zmiany T w dwóch przedziałach $[0, T]$ i $[T, \infty]$. Lepszą identyfikację zmian stałej czasowej T uzyskuje się dla małych T i w przedziale $(0, T)$.

Wniosek: Podane w pracy uwagi mogą być wykorzystane przy układaniu algorytmów dla regulatorów samonastawialnych sterujących obiektami, których modele matematyczne można opisać transmitancją (1).

BIBLIOGRAFIA

1. H. G ó r e c k i: *Analiza i synteza układów regulacji z opóźnieniem*. WNT, Warszawa, 1971
2. A. N i e d e r l i ń s k i: *Systemy komputerowe automatyki przemysłowej*. WNT, Warszawa, 1985
3. A. N i e d e r l i ń s k i, J. M o ś c i ń s k i, Z. O g o n o w s k i: *Regulacja adaptacyjna*. WNT, Warszawa 1995
4. L. T r y b u s: *Regulatory wielofunkcyjne*. WNT, Warszawa, 1992

J. HALAWA

THE INFLUENCE OF CHANGE OF THE PARAMETERS OF THE OBJECT DESCRIBED BY TRANSFER FUNCTION

$$G(s) = \frac{k}{(Ts + 1)^n}$$

S u m m a r y

In this paper the influence of change of the parameters T and n of the transfer function $G(s) = \frac{k}{(Ts + 1)^n}$ to the two function criteria is considered. The transfer function described a series objects. The result may be used in system with self-tuning controllers.

Key words: transfer function, controllers, mathematical models.

a różnej
ryteriumskać dla
działach
a małychalgoryt-
modele71
wa, 1985
ina. WNT,function
d a series

Komputerowo wspomagany dobór nastaw regulatorów dla obiektów bez wyrównania

JANUSZ HALAWA

Instytut Cybernetyki Technicznej, Politechnika Wrocławska

Otrzymano 1998.11.17

Autoryzowano 1999.03.24

W pracy przedstawiono metodę doboru nastaw regulatorów typu PI i PID współpracujących z obiektem bez wyrównania opisanym transmitancją $G_0(s) = \frac{k_0}{s}$. Podano wykresy odpowiedzi zamkniętego układu sterowania na funkcję skokową otrzymane w wyniku symulacji komputerowej. Drogą wspomagania komputerowego ustalono zależności pozwalające na wyznaczenie nastaw regulatorów.

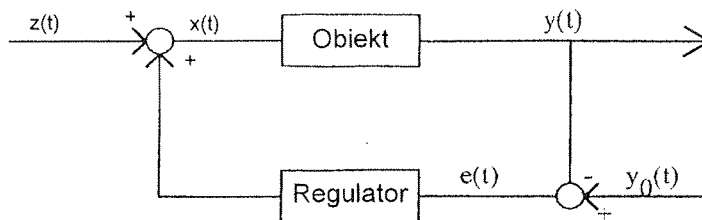
Słowa kluczowe: regulatory PI, PID, układy sterowania, symulacja komputerowa.

W literaturze omawiającej dobór nastaw regulatorów najwięcej uwagi poświęcono układom z obiektami inercyjnymi z wyrównaniem. Materiał dotyczący nastaw regulatorów w wypadku obiektów bez wyrównania jest znacznie skromniejszy. W monografii [1] podane są wzory służące do wyznaczania nastaw regulatorów typu P, PI i PID pracujących w układzie zamkniętym przedstawionym na rysunku 1 w wypadku obiektu opisanego transmitancją

$$G_0(s) = \frac{1}{Ts} e^{-sto}.$$

W poniższej pracy rozpatrzona zostanie dynamika układu zamkniętego zawierającego obiekt opisany transmitancją

$$G_0(s) = \frac{k_0}{s} \quad (1)$$



Rys. 1. Schemat blokowy zamkniętego układu regulacji

z regulatorem PI i PID o transmitancji

$$G_r(s) = k_p \left(1 + \frac{1}{T_i s} + T_w s \right). \quad (2)$$

Przedstawione zostaną odpowiedzi skokowe układu zamkniętego na zmiany wartości zadanej, przy zerowej wartości zakłócenia. Sygnały wejściowe przedstawione na rysunku 1 przyjmują więc wartości $x(1) = 1(t)$ i $z(t) = 0$. Jest to pewien wyidealizowany obiekt, gdyż rzeczywiste obiekty lepiej opisują transmitancje

$$G_{01}(s) = \frac{k}{s} e^{-sto}, \quad (3)$$

$$G_{02}(s) = \frac{k}{s(s+a)^n}, \quad n = 1, 2, 3, \quad (4)$$

$$G_{03} = \frac{k}{s(s+a)^n} e^{-sto}. \quad (5)$$

Regulator (2) pisany transmitancją jest również regulatorem o idealnej części różniczkującej.

Układ zamknięty zawierający obiekt (1) i regulator (2) opisany jest transmitancją

$$G_z(s) = \frac{y(s)}{y_0(s)} = \frac{\frac{k_0 T_w}{k_0 k_p T_w + 1} s^2 + \frac{k_0 k_p}{k_0 k_p T_w + 1} s + \frac{k_0 k_p}{(k_0 k_p T_w + 1) T_i}}{s^2 + \frac{k_0 k_p}{k_0 k_p T_w + 1} s + \frac{k_0 k_p}{(k_0 k_p T_w + 1) T_i}}. \quad (6)$$

Przyjmując $T_w = 0$ w (6) otrzymuje się transmitancję zamkniętego układu regulacji z obiektem (1) i regulatorem PI. Wówczas

$$G_z(s) = \frac{y(s)}{y_0(s)} = \frac{k_0 k_p s + \frac{k_0 k_p}{T_i}}{s^2 + k_0 k_p s + \frac{k_0 k_p}{T_i}}. \quad (7)$$

Gdyby
odpowi
nego tra

bądź z
nikami

W p
wykresy
przedsta



Rys. 2. W

Z p
o wyk
czasie c
Aby ot
(9) $\xi =$
Wyl
rys. 4.
Z p

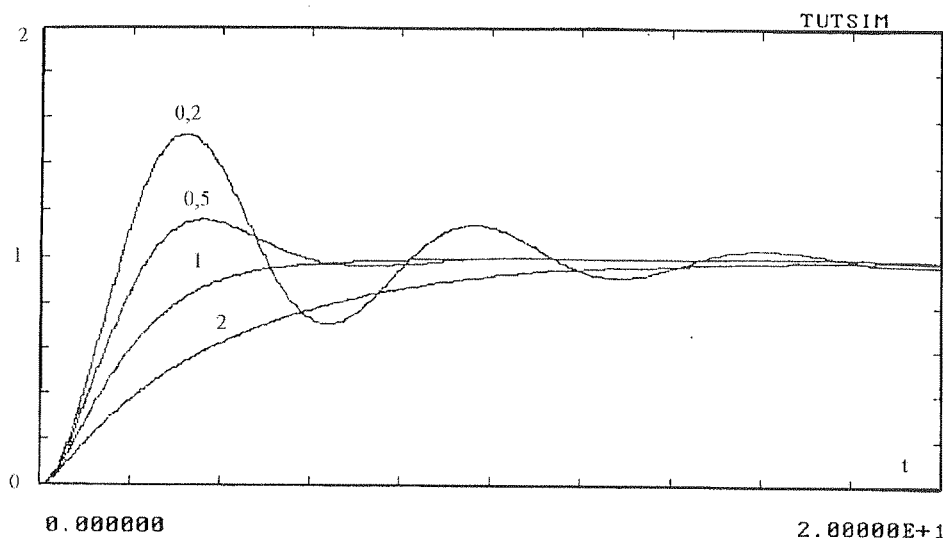
Gdyby w transmitancji (7) pominąć współczynnik $k_0 k_p s$ występujący w jej liczniku, odpowiedzi układu zamkniętego byłyby identyczne z odpowiedziami układu opisanego transmitancją

$$G_2(s) = \frac{\omega_0^2}{s^2 + 2\xi\omega_0 s + \omega_0^2}, \quad (8)$$

bądź z rozwiązaniami liniowego równania różniczkowego ze stałymi współczynnikami postaci

$$y^{(2)} + 2\xi\omega_0 y^{(1)} + \omega_0^2 y = \omega_0^2 x(t). \quad (9)$$

W podręcznikach poświęconych podstawom fizyki i automatyki podawane są wykresy odpowiedzi równania (9) na skokową funkcję wymuszającą. Na rysunku 2 przedstawione są przykładowo w/w wykresy dla $\omega_0 = 1$ i $\xi = 0,2, \xi = 0,5, \xi = 1, \xi = 2$.



Rys. 2. Wykresy odpowiedzi równania różniczkowego (9) dla $\omega_0 = 1$ i $\xi = 0,2, 0,5, 1, 2$ na skokową funkcję wymuszającą $x(t) = 1$

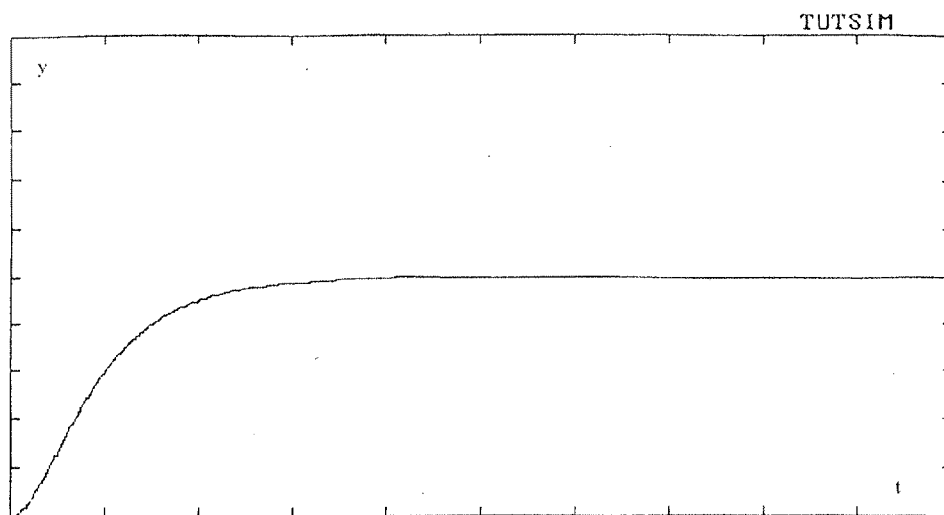
Z punktu widzenia sterowania automatycznego interesowała by nas odpowiedź o wykresie przedstawionym na rysunku 3. Czyli odpowiedź bez oscylacji i krótkim czasie osiągnięcia wartości ustalonej y_u .

Aby otrzymać odpowiedź przedstawioną na rysunku 3 należy przyjąć w równaniu (9) $\xi = 1$ i ω_0 możliwie duże.

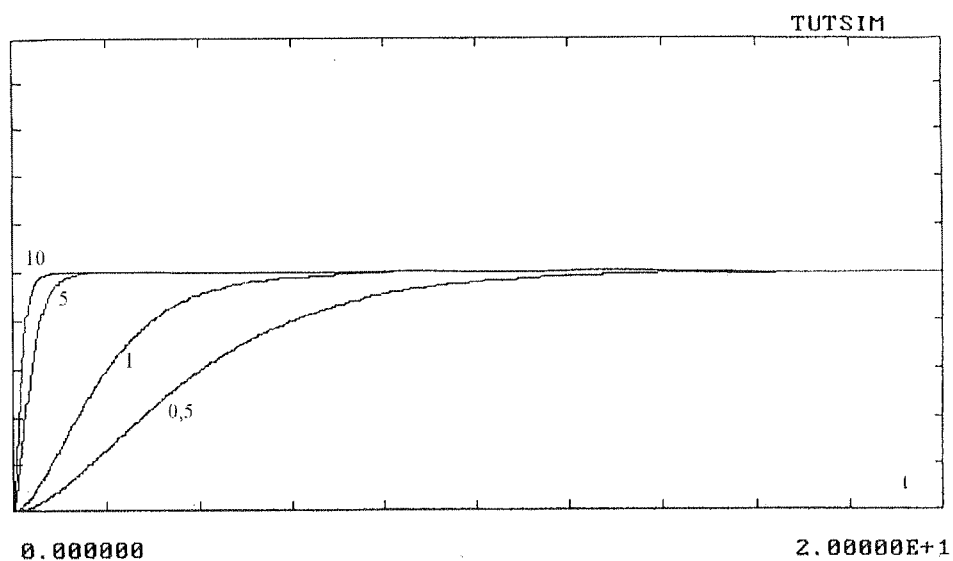
Wykresy odpowiedzi równania (9) $\xi = 1$ i $\omega_0 = 0,5, 1, 5, 10$ przedstawione są na rys. 4.

Z porównania mianowników transmitancji (7) i (8) otrzymuje się

$$s^2 + k_0 k_p s + \frac{k_0 k_p}{T_i} = s^2 + 2\xi\omega_0 s + \omega_0^2,$$



Rys. 3. Oczekiwany wykres odpowiedzi układu zamkniętego

Rys. 4. Odpowiedzi równania (9) dla $\xi = 1$ i $\omega_0 = 0,5, 1, 5, 10$

a stąd

$$k_0 k_p = 2\xi \omega_0 \quad (10)$$

$$\frac{k_0 k_p}{T_i} = \omega_0^2. \quad (11)$$

Z równ

W regu
działac

Maksym

Należy
transm
Na
cji z reg
(13) i (
Jak
20%. I



0.0

Ry

Z równań tych można wyznaczyć wartości nastaw

$$k_p = \frac{2\xi\omega_0}{k_0}, \quad (12)$$

$$T_i = \frac{2\xi}{\omega_0}. \quad (13)$$

W regulatorach przemysłowych zarówno k_p jak T_i zawiera się w pewnych przedziałach. Jeśli więc

$$k_p \leq k_{p\max} \quad \text{stąd} \quad \omega_0^2 \leq \frac{k_0 k_{p\max}}{T_i}.$$

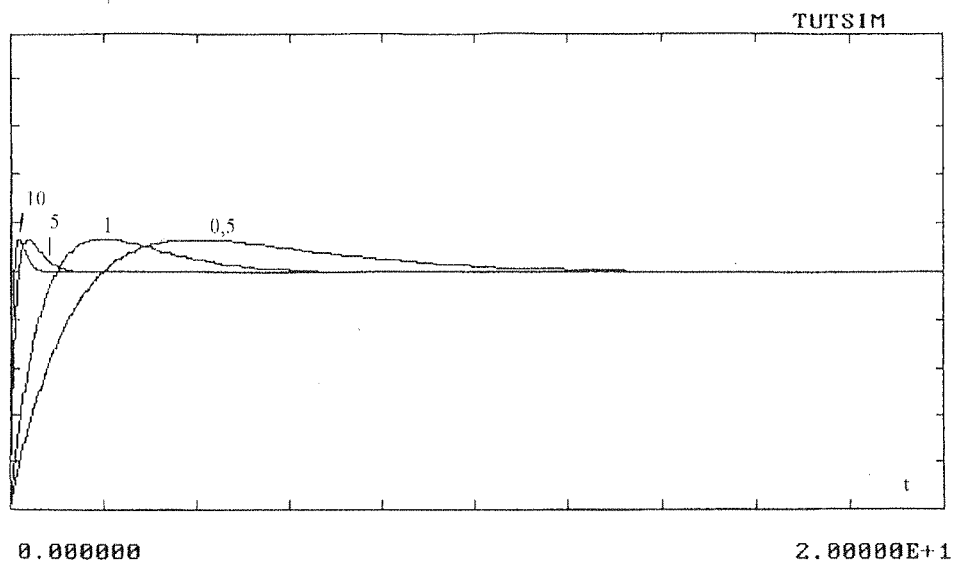
Maksymalna możliwa do przyjęcia pulsacja w tym wypadku wynosi

$$\omega_0^2 = \frac{k_0 \cdot k_{p\max}}{T_{i\min}}.$$

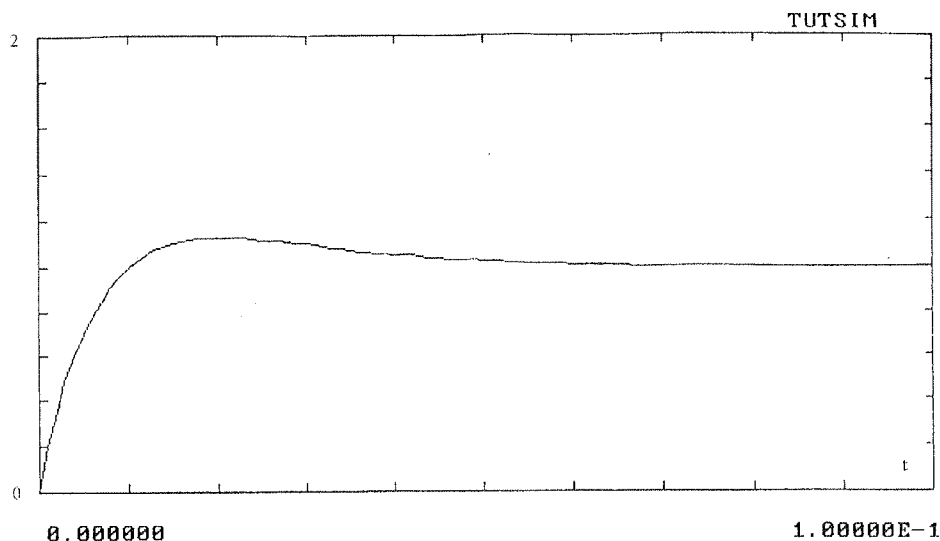
Należy jednak uwzględnić wpływ współczynnika $k_0 k_p s$ występującego w liczniku transmitancji na odpowiedź układu opisanego tą transmitancją.

Na rysunku 5 przedstawione są wykresy odpowiedzi zamkniętego układu regulacji z regulatorem PI opisanego transmitancją (7) dla nastaw wyliczonych ze wzorów (13) i (14) po przyjęciu wartości $\omega_0 = 0,5, 1, 5, 10$ oraz $\xi = 1$.

Jak wynika z rysunku 5 i 10 uzyskuje się odpowiedzi z przeregulowaniem około 20%. Im większe przyjęte ω_0 , tym krótszy czas ustalania się odpowiedzi.



Rys. 5. Wykresy odpowiedzi zamkniętego układu regulacji dla $\omega_0 = 0,5, 1, 5, 10$ oraz $\xi = 1$



Rys. 6. Wykres odpowiedzi zamkniętego układu regulacji dla $\omega_0 = 100$ i $\xi = 1$

Na rysunku 6 przedstawiono wykres odpowiedzi zamkniętego układu regulacji dla $\omega_0 = 100$ i $\xi = 1$, czyli dla nastaw $k_p = 200$ i $T_i = 0,02$. Wykres ten byłby niewidoczny na rysunku 5.

Jak wynika z rys. 6 przeregulowanie nie uległo zmianie, natomiast znacznie skrócił się czas ustalania się odpowiedzi. Z badań komputerowych wynika, że lepsze nastawy uzyskuje się przyjmując ξ we wzorach (13) i (14) większe od 1, np. 4.

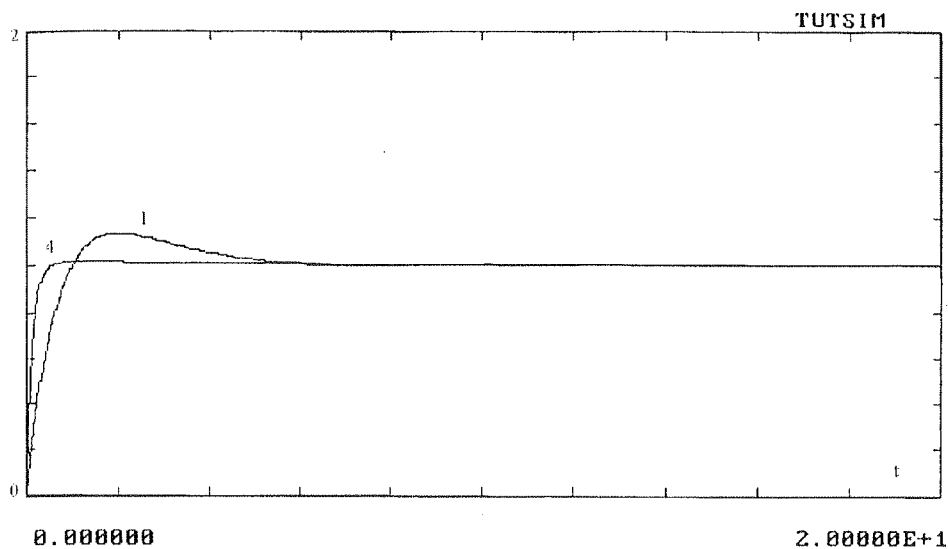
Przykład. Przyjmijmy $k_0 = 1$, $\omega_0 = 1$, $\xi = 4$. Wówczas otrzymujemy $k_p = 8$, $T_i = 8$. Na rysunku 7 przedstawiono wykresy odpowiedzi zamkniętego układu regulacji z obiektem opisanym transmitancją (1) i z regulatorem PI o nastawach $k_p = 8$ i $T_i = 8$, a także dla nastaw $k_p = 2$ i $T_i = 2$ wyznaczonych dla $\xi = 1$.

Jak wynika z rysunku 7 przyjęcie większego ξ od 1 jest celowe.

Przedstawioną metodę nie można stosować wprost do doboru nastaw regulatorów typu PID. Na podstawie transmitancji (6) mamy dwa równania i trzy niewiadome k_p , T_i i T_w . Proponuje się przyjąć jakąś wartość T_w w zależności od T_i np. $T_w = 0,1 T_i$. Załóżmy, że dobór nastaw został dokonany dla regulatora PI a stałą T_w przyjęto jako stałą $T_w = 0,1 T_i$. Po wyliczeniu wartości nastaw należy zaobserwować wykresy odpowiedzi układu zamkniętego korzystając z symulacji komputerowej.

Rozpatrzmy przykład doboru nastaw regulatora PID dla obiektu opisanego transmitancją (1), w której $k_0 = 1$. Przyjęto $\omega_0 = 1$. Wówczas dla regulatora PI wyliczamy $k_p = 2$ i $T_i = 2$. Przyjmijmy $T_w = 0,1 T_i$. W takim wypadku transmitancja (6) przyjmuje postać

$$G_z(s) = \frac{y(s)}{y_0(s)} = \frac{0,1429s^2 + 1,4286s + 0,7143}{s^2 + 1,4286s + 0,7143}. \quad (15)$$

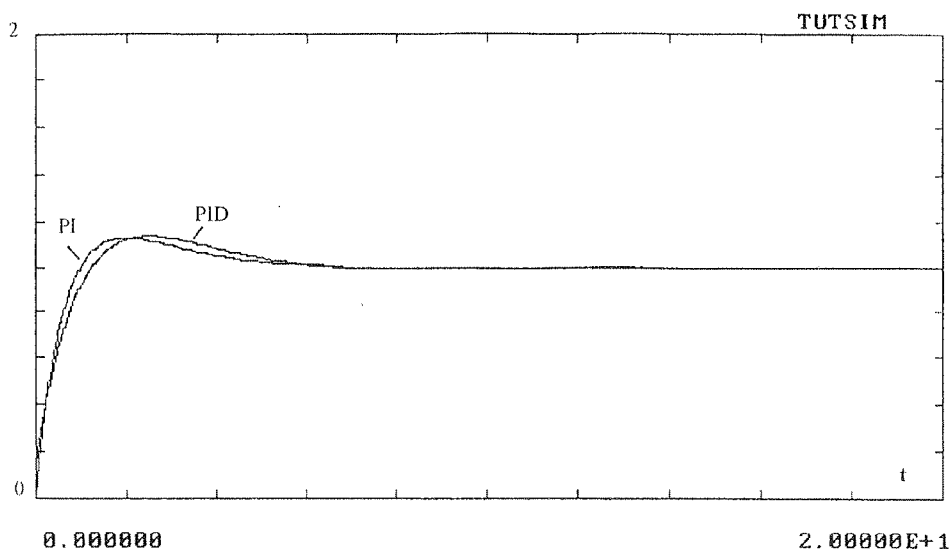


Rys. 7. Odpowiedzi zamkniętych układów regulacji dla $k_p = 2$, $T_i = 2$ ($\xi = 1$) oraz $k_p = 8$, $T_i = 8$ ($\xi = 4$)

Na rysunku 8 przedstawione są odpowiedzi na skok jednostkowy układu opisanego transmitancją (7) wyliczonej dla $k_p = 2$ i $T_i = 2$, czyli

$$G_{z_{PI}}(s) = \frac{2s + 1}{s^2 + 2s + 1} \quad (16)$$

i transmitancją (15) opisującą układ z regulatorem PID.

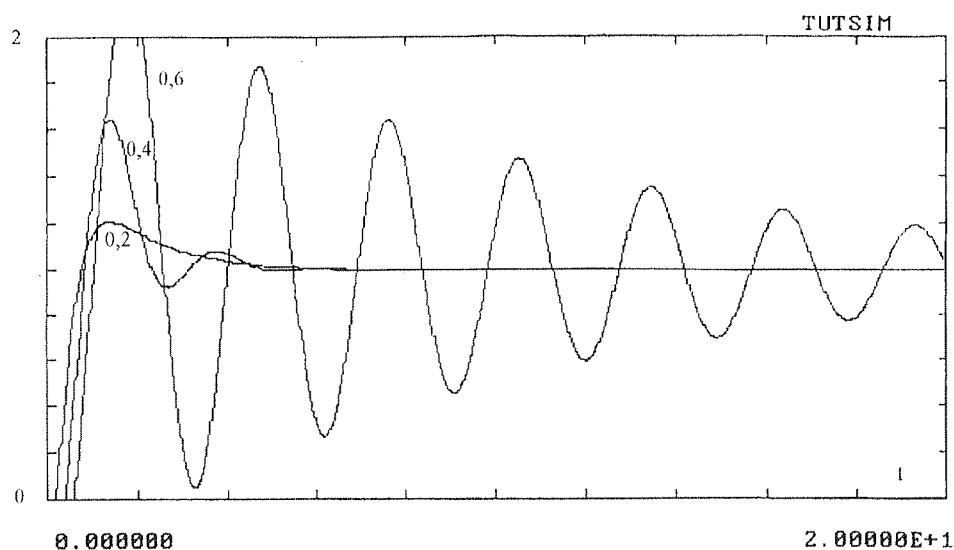


Rys. 8. Odpowiedzi skokowe układów opisanych transmitancjami (15) i (16)

Nastawy regulatora PID wyznaczone w podany wyżej sposób mogą być wykorzystane do wyznaczenia lepszych nastaw którąś z metod optymalizacji. Można je wykorzystać jako punkty początkowe.

Interesującym jest jak wpływa obecność członu opóźniającego o transmitancji $G_{0p}(s) = e^{-st_0}$ w transmitancji obiektu.

Przykład. Obiekt jest opisany wówczas transmitancją (3). Dla obiektu o transmitancji $G_0(s) = \frac{1}{s}$ i $\omega_0 = 1$, $\xi = 1$ obliczone nastawy regulatora PI wynoszą $k_p = 2$, $T_i = 2$. W wypadku pojawienia się w obiekcie opóźnień $t_0 = 0,2$ 0,4 i 0,6 wykresy odpowiedzi układu zamkniętego przedstawione są na rysunku 9.



Rys. 9. Odpowiedź zamkniętego układu regulacji dla $k_0 = 1$, $k_p = 2$, $T_i = 2$ i $t_0 = 0,2, 0,4, 0,6$

Jak wynika z rysunku 9 proponowana metoda doboru nastaw regulatorów może być stosowana w wypadku małych opóźnień. W tych wypadkach wzmocnienie k_p nie może być zbyt duże (np. wyliczone dla $\xi = 1$). Występowanie członu opóźniającego obniża wartość dopuszczalnego wzmocnienia regulatora ze względu na stabilność zamkniętego układu regulacji.

Wniosek: Nastawy regulatora PI obliczane są z prostej zależności (12) i (13) po przyjęciu maksymalnie dopuszczalnych wartości ω_0 i $\xi > 1$ np. 3 lub 4.

BIBLIOGRAFIA

1. J. Pułaczewski: *Dobór nastaw regulatorów przemysłowych*. WNT, Warszawa 1966

Key words: Regulators PI, PID, control circuits computer simulations.

ć wyko-
Można je

smitan-
k_p = 2,
wykresy

nsmitan-
k_p = 2,
wykresy

0E+1
4, 0,6

ów może
nie k_p nie
niającego
tabilność
(13) po

Decimation as an effective method of finding sets of m-sequences and other groups of pseudorandom codes

ANDRZEJ MAZUR

*Zakład Teletransmisji
Instytut Telekomunikacji, Politechnika Warszawska*

Received 1998.10.04

Authorized 1999.02.10

The pseudo-random codes are widely used in spread-spectrum applications where a nearly random signal is needed for spreading the data bandwidth. In the last forty years many new classes of spreading codes have been introduced however still the family of m-sequences have a key position in the theory of pseudo-random sequences. That is why the effective algorithm for finding the irreducible polynomials corresponding to each sequence belonging to the set of m-sequences of given period p is still needed. This paper presents a solution of the described problem. In practise the presented method can be easily translated into a computer program and is an alternative to the sieve method shown in [1]. Decimation can be used to define new classes of codes, which is also presented in this paper.

Key words: pseudorandom codes, polynomials, computer programs.

1. INTRODUCTION

Every maximal-length linear feedback shift register sequence (m-sequence) $\{a_n\} = \{a_1, a_2, a_3, \dots, a_p\}$ is associated with the mod-2 linear recurrence [1], [2], which has a form

$$a_n = c_1 a_{n-1} + c_2 a_{n-2} + \dots + c_r a_{n-r} \quad (1)$$

where $c_1, c_2, c_3, \dots, c_r$ are coefficient symbols which do not depend on n and can have the value belonging to the set $c_i \in \{0, 1\}$.

Having the equation (1) which defines a particular sequence $\{a_n\}$ and according to Fig. 1 it is possible to build the shift register with linear feedback unit which can generate this sequence.

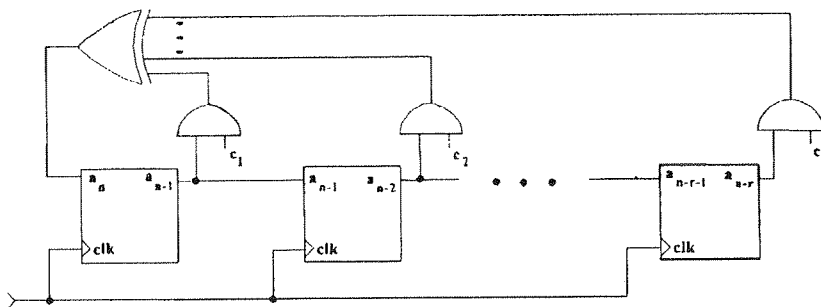


Fig. 1. Shift register with linear feedback

The second equation which fully describes the sequence $\{a_n\}$ and the shift register which produced it, is the characteristic polynomial [1], [2] defined as

$$f(x) = 1 - \sum_{n=1}^{\infty} c_n x^n. \quad (2)$$

Every m-sequence has a length

$$p = 2^r - 1, \quad (3)$$

where r is the number of tubes in shift register producing this sequence. There are only

$$\lambda_2(r) = \frac{\varphi(2^r - 1)}{r} \quad (4)$$

m-sequences generated in r -tubes shift register. Every m-sequence follows the randomness postulates given in [1], [3].

Let $\{a_n\} = \{a_1, a_2, \dots, a_p\}$ be an m-sequence of period p . According to (1) the coefficient symbols which fully describe the feedback unit of shift register generator which produces $\{a_n\}$ can be found by solving the system of equations mod-2 created as follows

$$\begin{cases} a_k^i = c_1 a_{i-1} + c_2 a_{i-2} + \dots + c_r a_{i-r} \\ a_{i-1} = c_1 a_{i-2} + c_2 a_{i-3} + \dots + c_r a_{i-1-r} \\ \dots \\ a_{i-r+1} = c_1 a_{i-r} + c_2 a_{i-r-1} + \dots + c_r a_{i-2r+1} \end{cases} \quad (5)$$

where $i \geq (2r + 1)$.

2. THE COSET DECOMPOSITION

It can be proven that the set of integers from 1 to $p-1$, which no factors in common with p , where p is defined by (3), form a multiplicative Abelian group, with

respect to the operation: multiplication mod- p . This group can be decomposed into cosets according to the pattern:

1. Create a subgroup of numbers with r elements such as

$$C_0 = \langle 2^0, 2^1, 2^2, \dots, 2^{r-1} \rangle \quad \text{where } r = \log_2(p+1).$$

2. Find the smallest integer k less than p , which has not appeared in C_0 and is relatively prime to p .
3. The coset C_i where i belongs to the set of integers from 1 to $\lambda_2(r) - 1$ is created by the multiplication mod- p of every element from C_{i-1} by integer k .
Cosets created in this way extended by C_0 are called proper cosets. There are

$$Y(p) = \left[\frac{1}{r} \sum_{d/r} \varphi(r) * 2^{r_d} \right] - 1 \quad (6)$$

all kinds of cosets modulo- p (proper and improper) [1]. That is why the $[Y(p) - \lambda_2(r)]$ improper cosets must be created in the next steps.

4. Find the smallest integer k , which is less than p and not relatively prime to p .
5. The coset C_i , where i is bigger than $\lambda_2(r) - 1$, is created by the multiplication mod- p of every element from C_0 by integer k . It must be noticed that if some of the results of this operation appeared in the previously created cosets, then the new one must replace multiplier k .
6. If the created cosets (proper and improper) consist of all integers belonging to the set $\langle 1..(p-1) \rangle$ then the next lacking improper cosets must be created by filling them with all zeros.

Example 1. Make the decomposition into cosets of multiplicative groups for $p_1 = 31$ and $p_2 = 63$. For $p_1 = 31$.

3. BASE SEQUENCE DECIMATION

According to [4] decimation can be defined as follows:

Definition 2. If q denotes a positive integer then the operation of choosing every q^{th} element of sequence $\{a_n\}$ is called a decimation by q of sequence $\{a_n\}$, and will be denoted by $\{a_n[q]\}$.

The properties of the sequence $\{a_n[q]\}$, which is a product of decimation, depend on q (see Fig. 2).

Let $\{a_n[q_1]\}$ and $\{a_n[q_2]\}$ be products of decimation of m -sequence $\{a_n\}$ of period p . If the set of integers from 1 to $p-1$ is decomposed into a group of cosets while q_1 and q_2 belong to the same proper coset then the sequences $\{a_n[q_1]\}$ and $\{a_n[q_2]\}$ are different by no more than a phase shift. On the other hand if the numbers q_1 and q_2 belong to different proper cosets then $a_n[q_1]$ and $a_n[q_2]$ are different m -sequences of length p . The product of decomposition $\{a_n[q]\}$, where q is an integer which belongs to one of improper cosets, does not belong to the set of m -sequences of period p . According to (4) there are only $\lambda_2(r)$ m -sequences of given length p . That is why

Table 1

Coset decomposition for $p_1 = 31$

Number of coset	Multiplier	Type of coset	$p_1 = 31$				
			1	2	4	8	16
C_0	x	proper	1	2	4	8	16
C_1	3	proper	3	6	12	24	17
C_2	3	proper	9	18	5	10	20
C_3	3	proper	27	23	15	30	29
C_4	3	proper	19	7	14	28	25
C_5	3	proper	26	21	11	22	13
C_6	0	improper	0	0	0	0	0

Table 2

Coset decomposition for $p_2 = 63$

Number of coset	Multiplier	Type of coset	$p_2 = 63$					
			1	2	4	8	16	32
C_0	x	proper	1	2	4	8	16	32
C_1	5	proper	5	10	20	40	17	34
C_2	5	proper	25	50	37	11	22	44
C_3	5	proper	62	61	59	55	47	31
C_4	5	proper	58	53	43	23	46	29
C_5	5	proper	38	13	26	52	41	19
C_6	3	improper	3	6	12	24	48	33
C_7	7	improper	7	14	28	56	49	35
C_8	9	improper	9	18	36	9	18	36
C_9	15	improper	15	30	60	57	51	39
C_{10}	21	improper	21	42	21	42	21	42
C_{11}	27	improper	27	54	45	27	54	45
C_{12}	0	improper	0	0	0	0	0	0
C_{13}	0	improper	0	0	0	0	0	0
C_{14}	0	improper	0	0	0	0	0	0

there are also only $\lambda_2(r)$ proper cosets into which the set of integers from 1 to $p-1$ can be decomposed.

Example 2. For the base sequence $\{a_n\}$ described by polynomial $f(x) = 1 + x + x^3$ find products of its decimation for $q = 2, 3, 4, 6$.

$\{a_n\} = \{a_1, a_2, a_3, a_4, \dots, a_7\} = \{1, 1, 1, 0, 1, 0, 0\}$

able 1

16
17
20
29
25
13
0

2

32
34
44
31
29
19
33
35
36
39
42
45
0
0
0

1 to $p-1$ $+x+x^3$

then

$$\{a[2]\} = \{a_2, a_4, a_6, a_8, a_{10}, a_{12}, a_{14}\} = a_2, a_4, a_6, a_1, a_3, a_5, a_7 \pmod{7^1}$$

$$\{a[2]\} = \{1, 0, 0, 1, 1, 1, 0\}$$

according to (4) it is possible to write

$$\begin{cases} a_8 = a_7 c_1 + a_6 c_2 + a_5 c_3 \\ a_7 = a_6 c_1 + a_5 c_2 + a_4 c_3 \\ a_6 = a_5 c_1 + a_4 c_2 + a_3 c_3 \end{cases} \Rightarrow \begin{cases} 1 = c_2 + c_3 \\ 0 = c_1 + c_2 + c_3 \\ 1 = c_1 + c_2 \end{cases}$$

then

$$c_1 = c_3 = 1, c_2 = 0 \Rightarrow \{a[2]\} : f_2(x) = 1 + x + x^3$$

$$\{a[3]\} = \{a_3, a_6, a_9, a_{12}, a_{15}, a_{18}, a_{21}\} = (a_3, a_6, a_2, a_5, a_1, a_4, a_7)$$

$$\{a[3]\} = (1, 0, 1, 1, 1, 0, 0) \Rightarrow f_3(x) = 1 + x^2 + x^3$$

$$\{a[4]\} = (0, 1, 1, 1, 0, 1, 0) \Rightarrow f_4(x) = 1 + x + x^3$$

$$\{a[6]\} = (0, 1, 0, 1, 1, 1, 0) \Rightarrow f_6(x) = 1 + x^2 + x^3.$$

It is easy to notice that for $q \in C_0 = \langle 1, 2, 4 \rangle$ the products of $\{a_n\}$'s decimation are phase shifts of $\{a_n\}$, and for $q \in C_1 = \langle 3, 5, 6 \rangle$ the products are new sequences different from $\{a_n\}$. The set of created sequences different from $\{a_n\}$ extended by the base sequence $\{a_n\}$ (where $\{a_n\}$ is a sequence of maximal-length), consists of all maximal-length sequences possible for the given p (where p is a period of $\{a_n\}$). (Fig. 2, Fig. 3).

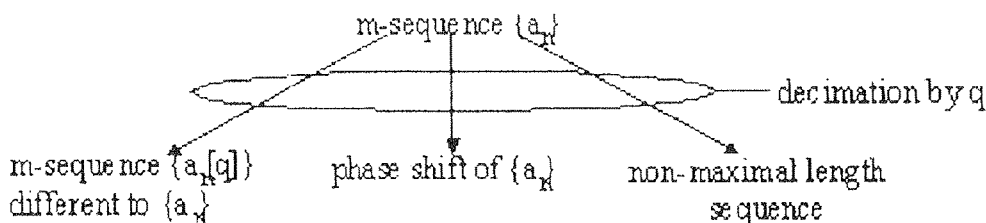
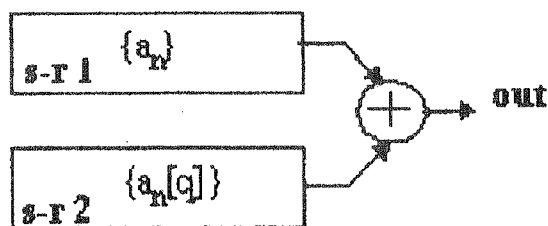
Fig. 2. Possible products of $\{a_n\}$ decimation

Fig. 3. The structure of shift register generator which can produce the set of sequences defined by (7) or (9)

As presented in [4] decimation can also be used for finding the classes of spreading codes different from the set of m-sequences. In this case the integer q has to follow different conditions compared with the decimation, which produces a new m-sequence.

¹ $a_{(k \cdot p) \bmod p} = a_p$ for every integer k

For example

If $q = 2^{\lfloor \frac{r+2}{2} \rfloor} + 1$ and $r \neq 0 \pmod{4^2}$ then $\{a_n\}$ and $\{a_n[q]\}$ are called a preferred pair.

A set of Gold sequences of period p , consisting of $p + 2$ sequences, can be defined [4] as

$$G = \langle \{a_n\}, \{a_n[q]\}, \{a_n\} \oplus \{a_n[q]\}, \{a_n\} \oplus T\{a_n[q]\}, \dots, \{a_n\} \oplus T^{p-1}\{a_n[q]\} \rangle, \quad (7)$$

where $\{a_n\}, \{a_n[q]\}$ is the preferred pair.

$T^i\{a_n\}$ denotes an i -chips phase shift of $\{a_n\}$.

For every pair of sequences belonging to the set G the peak value of the even cross-correlation function can be written as

$$\theta_{\max}(\{a_n\}, \{a_n[q]\}) = 1 + 2^{\lfloor \frac{r+2}{2} \rfloor}. \quad (8)$$

Let r be even and $q = 2^{\frac{r}{2}} + 1$ then the product of decimation by q of $\{a_n\}$ is a sequence of length $2^{\frac{r}{2}} - 1$ ($\{a_n\}$ is a sequence of period p).

The set of sequences defined as

$$K_s \langle \{a_n\}, \{a_n\} \oplus \{a_n[q]\}, \{a_n\} \oplus T\{a_n[q]\}, \dots, \{a_n\} \oplus T^{2^{\frac{r}{2}}-2}\{a_n[q]\} \rangle \quad (9)$$

is called the small set of Kasami sequences.

For every pair of sequences belonging to the set K_s the peak value of the even cross-correlation function can be written as

$$\theta_{\max}(\{a_n\}, \{a_n[q]\}) = 1 + 2^{\frac{r}{2}}. \quad (10)$$

Example 2. Find the small set of Kasami sequences for $\{a_n\}$ described by the polynomial $1 + x + x^4$

In this case

$$\{a_n\} = \{0, 1, 1, 1, 1, 0, 1, 0, 1, 1, 0, 0, 1, 0, 0\}$$

and

$$\{a_n[5]\} = \{1, 1, 0\}, T\{a_n[5]\} = \{0, 1, 1\}, T^2\{a_n[5]\} = \{1, 0, 1\},$$

then the set K_s can be written as

$$K_s = \langle \{0, 1, 1, 1, 1, 0, 1, 0, 1, 1, 0, 0\}, \{1, 0, 0, 0, 0, 1, 1, 0, 1, 0, 0, 1\},$$

$$\{0, 0, 0, 1, 0, 1, 1, 1, 0, 1, 1, 1\}, \{1, 1, 0, 0, 1, 1, 0, 0, 0, 0, 1, 0\}, \{0, 0, 1, 1\} \rangle$$

Sequences belonging to the set defined by (7) and (9) can be generated in shift register generator with structure shown in Fig. 3.

² Notice that decimation other than presented here can also lead to preferred three-valued cross-correlation functions [4]. That is why some authors define Gold sequences as the set which is a product of any preferred pair (also those which are a product of decimation different to the presented one).

As
set of
difficu
(only
other
like G

$\{a_n\}$
 a_i
 $\{a_n[q]\}$
 c_i
 C_0
 C_1
 $f(x)$
 G
 K_s
mod-
 p
 r
 $T^i\{a_n$
 $[x]$
 $Y(p)$
 $\theta_{\max}$
 $\lambda_2(x)$

$\varphi(x)$

$\sum_{d|r}$

Table 3
Values of functions $\lambda_2(r)$; $Y(p)$ and $\varphi(n)$ which can be helpful to
analyse the example

r	p	$\lambda_2(r)$	$Y(p)$	$\varphi(n)$
1	1	1	1	1
2	3	1	2	1
3	7	2	3	2
4	15	2	5	2
5	31	6	7	4
6	63	6	15	2

CONCLUSION

As shown in Example 2 the decimation could be used to find easily the set of m-sequences in case one m-sequence from this set is given. it is not difficult to fulfil such condition, because the tables of irreducible polynomials (only a few members from each set) are presented in many books [3]. Among other applications of the presented methods is creating new classes of codes like Gold, Kasami or Dual-BCH [4].

NOMENCLATURE

$\{a_n\}$	m-sequence
a_i	particular bits of m-sequence $\{a_n\}$
$\{a_n[q]\}$	product of decimation by q of $\{a_n\}$
c_i	i-th symbol coefficient in linear recurrence formula
C_0	subgroup of numbers with r elements such as $C_0 = \langle 2^0, \dots, 2^{r-1} \rangle$
C_i	i-th coset which is a product of decomposition of Abelian group
$f(x)$	characteristic polynomial of given m-sequence
G	set of Gold codes of given period
K_s	small set of Kasami codes of given period
mod-i	operation modulo i
p	period of pseudo-random sequence
r	number of tubes of shift register
$T^i\{a_n\}$	i-chip shift of sequence $\{a_n\}$
$[x]$	integer part of the real number x
$Y(p)$	total number of proper and improper cosets for given p
$\theta_{\max}$	maximum value of even cross-correlation function
$\lambda_2(x)$	the number of polynomials mod-2 of degree r which have a maximum exponent
$\varphi(x)$	Euler φ -function
$\sum_{d r}$	the sum extended over all positive divisors d of r

SELECTED BIBLIOGRAPHY

1. S.W. Golomb: *Shift Register Sequences*. Aegean Park Press, 1982
2. R.L. Picholtz: *Theory of Spread-Spectrum Communications-A Tutorial*. IEEE Transcation on Communications, 1982, vol. Com-30, pp. 855—884
3. R.C. Dixon: *Spread Spectrum Systems with Commerical Applications*. John Willey & Sons INC, 1992
4. D. Sarwate, M.B. Pursley: *Crosscorelation Properties of Pseudorandom and Related Sequences*. Proceedings of the IEEE, vol. 68, no. 5, May 1980

A. MAZUR

DECYMACJA JAKO EFEKTYWNA METODA KONSTRUKCJI ZBIOROW M-SEKWENCJI
ORAZ INNYCH GRUP KODÓW PSEUDOŁOSOWYCH

Streszczenie

Jako jeden z głównych obszarów zastosowań sekwencji pseudolosowych wymienić należy systemy z celowo rozproszonym widmem (ang. *Spread Spectrum Systems*). Z uwagi na coraz większe zainteresowanie rozwiązaniami bazującymi na tej technologii warto zastanowić się nad sposobem wyboru ciągów rozpraszających, które stanowią czynnik decydujący o jakości projektowanego systemu. W artykule przedstawiono podstawy teoretyczne metody decymacji, która zdaniem autora stanowi znakomite narzędzie do wyszukiwania zbiorów kodów pseudolosowych o zadanym okresie. Omawiana metoda polega na odpowiednim wyborze (ze stałym ściśle określonym skokiem q) elementów znanej m -sekwencji co może prowadzić w zależności od przyjętego skoku do utworzenia innej m -sekwencji posiadającej ten sam okres co sekwencja decymowana bądź sekwencji pseudolosowej, która może posłużyć do budowy nowej klasy kodów cechujących się np: lepszymi w stosunku do m -sekwencji własnościami korelacyjnymi (kody Golda, Kasaniego). Produkt decymacji jest jednoznacznie określony przez wartość przyjętego skoku q , co sprawia, że w niniejszej pracy zagadnieniu wyboru wartości parametru q poświęcono szczególną uwagę.

Słowa kluczowe: metody decymacji, kody pseudolosowe, kody Golda, kody Kasaniego, programy komputerowe.

Szeregowanie wiadomości w rozproszonych systemach czasu rzeczywistego wykorzystujących magistrale miejscowe

WENCJI

JAN WEREWKA

Katedra Automatyki, Akademia Górniczo-Hutnicza, Kraków

SŁAWOMIR ŻABA

*Instytut Elektrotechniki i Elektroniki Przemysłowej, Politechnika Krakowska**Otrzymano 1998.11.20**Autoryzowano 1999.03.10*

ży systemy
zaintereso-
oru ciągów
W artykule
znakomite
na metoda
n-sekwencji
dającej ten
do budowy
relacyjnymi
przyjętego
poświęcono

programy

Jednym z popularnych rozwiązań stosowanych w obszarze systemów sterujących jest użycie rozproszonego systemu komputerowego. Poszczególne węzły takiego systemu są często łączone poprzez magistralę miejscową (*fieldbus*). Istotnym zadaniem staje się wówczas sprawdzenie w takim systemie dochowania ograniczeń czasu rzeczywistego, w szczególności dla przesyłanych wiadomości poprzez sieć. W ostatnim okresie nastąpił rozwój analitycznych metod szeregowania i sprawdzania spełnienia warunków czasu rzeczywistego dla środowiska scentralizowanego systemów operacyjnych czasu rzeczywistego, a jednocześnie podejmowane są działania mające na celu modyfikację tych metod, aby zastosować je do badania systemów rozproszonych czasu rzeczywistego, a w szczególności do szeregowania wiadomości przesyłanych poprzez sieć.

W pierwszej części artykułu dokonano porównania pomiędzy szeregowaniem zadań a szeregowaniem wiadomości. Następnie przedyskutowano prosty model systemu rozproszonego, na podstawie którego został przyjęty model strumienia wiadomości. Jako metodę szeregowania i sprawdzenia spełnienia warunków czasu rzeczywistego wybrano metodę Generalised Rate Monotonic Scheduling (GRMS). Przedstawiono zastosowanie tej metody dla badania spełnienia warunków czasu rzeczywistego wiadomości przesyłanych poprzez magistrale miejscowe stosujące metodę odpytań i metodę przekazywania żetonu w warstwie dostępu do medium fizycznego. Powyższe rozważania zobrazowane zostały przykładem obliczeniowym dla magistrali PROFIBUS DP.

Słowa kluczowe: rozproszony system komputerowy, system czasu rzeczywistego, magistrala miejscowa, szeregowanie, metoda GRMS

1. WSTĘP

W komputerowych systemach sterowania występuje tendencja do rozproszenia sprzętu i oprogramowania tych systemów. Obecnie w systemach produkcyjnych występuje kierunek zastępowania tradycyjnych, zcentralizowanych systemów sterowania i akwizycji danych przez systemy rozproszone realizujące przesyłanie danych poprzez magistrale miejscowe [9, 13]. Magistrale miejscowe tym różnią się od zwykłych sieci komputerowych, że w większości przypadków gwarantują deterministyczny czas odpowiedzi, co umożliwia spełnienie warunków czasu rzeczywistego (*RT* — *Real Time*). Dla systemów operacyjnych czasu rzeczywistego z globalnym programem szeregującym zadania, rozwiniętych zostało wiele analitycznych metod szeregowania i sprawdzania dotrzymania ograniczeń czasu rzeczywistego. Jedną z klasycznych metod, należącą do tej grupy, jest metoda Generalised Rate Monotonic Scheduling (GRMS). Metoda GRMS zapewnia, że jeżeli wykorzystanie systemu znajduje się poniżej pewnego poziomu i stosowany jest odpowiedni algorytm przydziału priorytetów, to wszystkie zadania dotrzymają swoich ograniczeń czasowych. Metoda ta może być także stosowana do badania rozproszonych systemów czasu rzeczywistego (*Distributed Real Time Systems*), opartych o magistrale miejscowe, jeżeli tylko każdy węzeł posiada deterministyczny dostęp do warstwy fizycznej sieci. W przypadku magistral: PROFIBUS, CAN, InterBus-S i Modbus, warstwa sterowania dostępem do medium fizycznego tych sieci udostępnia mechanizm umożliwiający węzłom spełnienie tego warunku. Dlatego też możliwe jest zastosowanie metody GRMS do badania dotrzymania ograniczeń czasu rzeczywistego dla systemów rozproszonych bazujących na wyżej wymienionych magistralach miejscowych.

2. WPLYW PROTOKOŁU SIECIOWEGO NA SZEREGOWANIE WIADOMOŚCI

Mimo znaczących różnic pomiędzy szeregowaniem zadań w procesorze i wiadomości na magistrali, istnieją także pewne analogie:

- Planowanie i szeregowanie zadań na procesorze oznacza uszeregowanie ich w odpowiedniej kolejności do uruchamiania. Podobnie, możliwe jest rozpatrzenie szeregowania wiadomości na magistrali jako odpowiedniego uszeregowania ich do wysłania na medium fizyczne.
- Dwa lub więcej zadań nie może być jednocześnie wykonywanych przez jeden procesor w tym samym czasie; podobnie, dwie lub więcej wiadomości nie mogą być jednocześnie transmitowane przez medium (nie rozważamy sieci wielokanałowych).
- Zadanie ma swój czas wykonania zależny od długości kodu. Podobnie wiadomości posiadają swój czas transmisji zależny od długości wiadomości (ilości bitów).

Chcąc rozważać zagadnienie szeregowania wiadomości w odniesieniu do szeregowania zadań należy przede wszystkim rozpatrzyć wpływ poszczególnych warstw protokołu sieciowego modelu OSI/ISO [15] na zagadnienie szeregowania:

⇒ *Wa*
bita
⇒ *Poa*
la
Poa
nac
szer
⇒ *Poa*
dzia
inic
wy
wię
dog
kac
kni
dla
⇒ *Wa*
siec
segr
⇒ *Wa*
dan
war
mie
nik
fun
wia
kol
nier
prz
siec
⇒ *Wa*
kon
totr
⇒ *Wa*
wid
⇒ *Wa*
umi
Jak
wy są
wiadom
otrzym

- ⇒ *Warstwa fizyczna* — dotyczy sposobu przesyłania danych w postaci strumienia bitów. Z punktu widzenia zagadnienia szeregowania warstwa ta jest nieistotna.
- ⇒ *Podwarstwa sterowania dostępem do medium* (Medium Access Control) — określa w jaki sposób jest sprawowana kontrola nad ośrodkiem transmisyjnym. Podaje metodę dostępu do sieci tj. na jakich zasadach stacja uzyskuje prawo do nadawania danych i w jaki sposób kontroluje prace sieci. Z punktu widzenia szeregowania wiadomości jest to warstwa najważniejsza.
- ⇒ *Podwarstwa sterowania łączem logicznym* (Logical Link Control) — uniezależnia działanie wyższych warstw sieci od technicznych szczegółów budowy sieci. LLC inicjalizuje wymianę danych sterujących, organizuje wymianę danych użytkowych, prowadzi kontrolę i korektę błędów transmisji. Działania te odpowiadają więc mechanizmom odporności na błędy (*fault tolerance*) jak np. technika 'watch dog' dla zadań. Główna różnica jest taka, że w przypadku protokołów komunikacyjnych stosowanie mechanizmów odporności na błędy jest w zasadzie nieuniknione, podczas gdy w przypadku procesora taka technika używana jest tylko dla aplikacji wymagających wysokiej niezawodności.
- ⇒ *Warstwa sieciowa* — zajmuje się przekazywaniem danych przez sieć lub kilka sieci od nadawcy do adresata. Magistrale miejscowe składają się z jednego segmentu i nie mają z reguły tej warstwy (z wyjątkiem sieci LonWorks).
- ⇒ *Warstwa transportowa* — zadaniem tej warstwy jest bezbłędne przekazywanie danych między połączonymi systemami. W pewnym sensie odpowiada to funkcji warstwy LLC, z tym że podwarstwa zapewnia komunikację bezpośrednio pomiędzy dwoma węzłami, podczas gdy warstwa transportowa pomiędzy użytkownikami końcowymi, którymi może być wiele aplikacji jednocześnie. Kolejną funkcją tej warstwy jest segmentacja danych polegająca na dzieleniu długich wiadomości na mniejsze bloki i zapewnienie przesłania ich w odpowiedniej kolejności, bez zagubienia ani powielenia. Funkcja ta z punktu widzenia zagadnienia szeregowania odpowiada ograniczaniu czasu wykonywania zadania w celu przydzielenia procesora innym zadaniom, jednakże, podobnie jak warstwa sieciowa, z reguły nie występuje w magistralach miejscowych.
- ⇒ *Warstwa sesji* — zapewnia połączenie logiczne, czyli sesję pomiędzy dwoma konwersującymi systemami. Z punktu widzenia szeregowania jej rola jest nieistotna.
- ⇒ *Warstwa prezentacji* — zajmuje się interpretacją przesyłanych danych. Z punktu widzenia szeregowania jej rola jest także nieistotna.
- ⇒ *Warstwa aplikacji* — istotna o tyle dla zagadnienia szeregowania, że tutaj są umieszczone zadania użytkownika.

Jak zostało pokazane, z punktu widzenia zagadnienia szeregowania jedne warstwy są bardziej, a inne mniej ważne. Najistotniejszą rolę w procesie szeregowania wiadomości odgrywa podwarstwa MAC, określającą w jaki sposób i kiedy węzeł otrzymuje dostęp do dzielonego zasobu jakim jest medium fizyczne sieci.

3. SZEREGOWANIE ZADAŃ A SZEREGOWANIE WIADOŚCI

W celu zastosowania metod szeregowania zadań do metod szeregowania wiadomości należy dokonać analizy porównawczej obu metod. Poniżej przedstawiono najistotniejsze aspekty z zakresu szeregowania zadań i ich odniesienia do szeregowania wiadomości [1]:

• *Właściwości czasowe*

Istniejące algorytmy planowania i szeregowania zadań uwzględniają następujące parametry czasowe:

- ⇒ czas gotowości (czas, przed którym uruchomienie zadania jest niemożliwe),
- ⇒ czas wykonania (czas potrzebny do wykonania się zadania bez wyłączenia),
- ⇒ maksymalny czas wykonywania zadania (maksymalny czas jaki może upłynąć od wystartowania zadania do jego zakończenia),
- ⇒ ograniczenie czasowe (czas, przed którym zadanie musi zostać zakończone).

Wymienione parametry mogą być przypisane praktycznie bez żadnych zmian do wiadomości. Wiadomości mają także swój czas gotowości (po sformatowaniu ramki), czas przesyłania (odpowiadający czasowi wykonywania), maksymalny czas przesyłania oraz ograniczenie czasowe.

• *Priorytety*

Jednym ze sposobów szeregowania zadań jest przypisanie im priorytetów (statycznie lub dynamicznie). W przypadku statycznego przydziału priorytetów każde zadanie ma przypisany stały priorytet, który nie zmienia się przy kolejnych wznowieniach wykonywania zadania algorytmy z dynamicznym przydziałem priorytetów pozwalają na zmianę priorytetu zadania tzn.: w kolejnych wznowieniach wykonywania zadanie może mieć różne priorytety. Zadania są następnie szeregowane zgodnie z przypisanymi wcześniej priorytetami. Szeregowanie wiadomości również może być określone poprzez przydzielenie im statycznych lub dynamicznych priorytetów. Przydzielenie priorytetów jest realizowane z reguły nie przez użytkownika, a przez algorytm szeregujący na podstawie właściwości (z reguły czasowych) zadań lub wiadomości.

• *Wywłaszczanie*

Wywłaszczanie zadań jest operacją polegającą na zawieszeniu uruchomionego zadania, a następnie ponownym jego uruchomieniu od punktu zatrzymania. Najpowszechniejszą przyczyną wywłaszczania zadań jest potrzeba uaktywnienia zadania o wyższym priorytecie. Generalnie, algorytmy szeregowania z wywłaszczaniem zadań są lepsze, jakkolwiek przy wykorzystywaniu dzielonych zasobów mogą występować takie zjawiska jak: odwrócenie priorytetu lub zakleszczenie.

OŚCI

nia wiado-
dstawiono
szeregow-

astępujące

liwe),
szczenia),
płynąć od

zione).

zmian do
natowaniu
nalny czas

ów (staty-
ów każde
wznowie-
riorytetów
ykonywa-
ne zgodnie
może być
riorytetów.
a, a przez
zadań lub

omionego
ania. Naj-
ia zadania
szczaniem
ów mogą

W przypadku szeregowania wiadomości zdawać się może, że należy stosować tylko algorytmy bez wyłuszczania, bowiem wiadomość, która zaczęła być transmitowana, nie może zostać wyłuszczona przez inną.

Jednakże w przypadku magistral miejscowych można w zasadzie stosować algorytmy z wyłuszczaniem. Wynika to z charakteru transmisji danych w magistralach miejscowych polegającego na ciągłym przesyłaniu krótkich wiadomości. Oznacza to, że czasy przesyłania wiadomości są małe w stosunku do ograniczeń czasowych wiadomości. Przykładowo, dla magistrali CAN przesłanie 1 bajtu danych (plus 44 bity narzutu protokołu), przy prędkości 500 kbit/s wynosi 104 μ s. Ponieważ typowe ograniczenia czasowe są rzędu ms, to pomimo braku wyłuszczania wiadomości, możemy z powodzeniem stosować algorytmy z szeregowaniem z wyłuszczaniem, gdyż czasy transmisji są bardzo krótkie w stosunku do ograniczeń czasowych.

- *Związki następstwa i poprzedzania*

Dotychczas rozważane były ograniczenia pojedynczych zadań. W systemach rozproszonych mamy do czynienia z oddziaływaniem zadań w celu osiągnięcia nadrzędnego celu. Należy przez to rozumieć, że zadania w systemie rozproszonym nie są niezależne i muszą być wykonywane w odpowiednim porządku, dzielić krytyczne zasoby, itp. Zatem istnieją silne zależności poprzedzania i następstwa zadań w dostępie do dzielonych zasobów.

W przypadku szeregowania wiadomości takie zależności także występują, jak chociażby czekanie na potwierdzenie wysłanej wiadomości w celu podjęcia dalszego działania.

- *Rozproszenie zasobów*

Zadanie, aby zostało wykonane, wymaga dostępu do odpowiednich zasobów (procesor, pamięć, porty WE/WY, itp.). Podobnie, aby przesłać wiadomość wykorzystuje się interfejs sieciowy, pamięć, medium fizyczne. Wykorzystywanie wielu zasobów, w celu uzyskania pożądanego efektu, może prowadzić do wystąpienia zakleszczenia zadań.

Takie sytuacje mogą wystąpić także w przypadku sieci, a dokładnie dla sieci wielokanałowych. Na przykład, jeżeli węzeł chce przesłać wiadomość przez dwa kanały jednocześnie, może dojść do zakleszczenia, jeżeli inny węzeł próbuje wykonać tą samą operację rezerwując kanały w odwrotnej kolejności.

- *Szeregowanie jednoprocessorowe i wieloprocessorowe*

System RT może bazować na jednym procesorze lub być oparty o architekturę wieloprocessorową i wówczas zadanie może być wykonywane przez jeden z procesorów. Podobnie sytuacja wygląda w przypadku systemów sieciowych: może on być złożony z jednego medium fizycznego lub z kilku redundantnych mediów i wówczas wiadomość może być transmitowana przez którekolwiek z nich.

• Szeregowanie lokalne i szeregowanie globalne

System RT może być złożony z jednego węzła lub z kilku wzajemnie połączonych węzłów (węzeł może stanowić dowolne urządzenie przyłączone do sieci jak np.: PLC, mikrokomputer, inteligentne WE/WY).

W związku z powyższym podziałem struktury systemów RT, algorytmy szeregowania zadań można podzielić na lokalne i globalne. Lokalne szeregowanie dotyczy tylko jednego węzła. Jeżeli węzły komunikują się wzajemnie w celu zaplanowania kolejności wykonywania zadań, szeregowanie staje się globalne — program szeregujący musi zdecydować kiedy i gdzie ma zostać wykonane zadanie (w szeregowaniu lokalnym decyduje się tylko kiedy zadanie ma zostać wykonane).

• Szeregowanie statyczne i dynamiczne

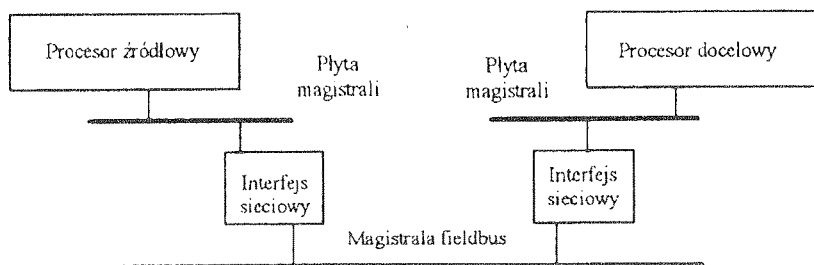
Szeregowanie statyczne jest wykonywane przed uruchomieniem systemu, natomiast metody dynamiczne szeregują zadania w trakcie działania systemu. Wybór metody szeregowania zależy od zestawu zadań w systemie: dla zadań okresowych dobrze nadają się algorytmy szeregowania statycznego, a do zadań aperiodycznych algorytmy dynamiczne.

W przypadku szeregowania wiadomości, wybór metody będzie zależał także od typu wiadomości występujących w systemie (okresowe, sporadyczne).

4. MODEL SYSTEMU ROZPROSZONEGO

Niech rozważany system rozproszony składa się z węzłów przyłączonych do magistrali, a każdy węzeł może zawierać: procesor, moduł WE/WY, płytę magistrali (*backplane*), interfejs sieciowy (rys. 1).

Zadania wykonywane przez system polegają na wysłaniu wiadomości z procesora źródłowego do docelowego, który wykonuje odpowiednie działanie na podstawie danych zawartych w otrzymanej wiadomości. Na tak sformułowane zadanie nałożone jest ograniczenie czasowe (*end-to-end deadline*), obejmujące przedział czasu od rozpoczęcia wysyłania wiadomości z procesora źródłowego do momentu zakoń-



Rys. 1. Prosty model systemu rozproszonego

czenia j
temie s
powyż
⇒ okre
⇒ po z
inter
⇒ inter
su s
⇒ inter
do p
⇒ wiad
Dziś
proszon
akcji. O
zsumow
porówn
Jeże
zasoby,
można
przez m
rali na
okresow
kół siec
(rys. 2).

pm₁...pm_n



Elemen

- PM o mess
- PM
- tp_i je
- ts_j o
- spora

łączonych
np.: PLC,

szeregowania
dotyczy
planowania
n szeregu-
regowaniu

mu, nato-
u. Wybór
kresowych
odycznych

ze od typu

onych do
magistrali

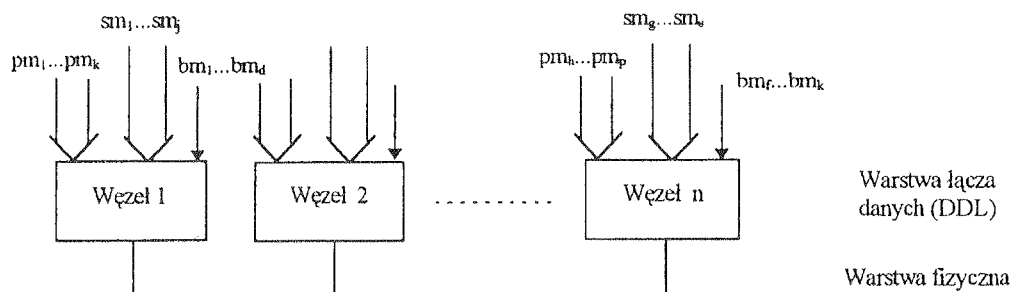
z procesor-
podstawie
ie nałożo-
czasu od
tu zakoń-

czenia jej przetwarzania na procesorze docelowym. Niech wszystkie zasoby w systemie stosują szeregowania ze statycznym przydziałem priorytetu. Sformułowany powyżej problem można więc przedstawić jako następującą sekwencję akcji:

- ⇒ określone zadanie wykonuje się na procesorze w węźle źródłowym,
- ⇒ po zakończeniu działania zadanie wysyła wiadomość poprzez płytę magistrali do interfejsu sieciowego,
- ⇒ interfejs sieciowy w węźle źródłowym wysyła wiadomość poprzez sieć do interfejsu sieciowego węzła docelowego,
- ⇒ interfejs sieciowy węzła docelowego przesyła wiadomość poprzez płytę magistrali do procesora docelowego,
- ⇒ wiadomość jest przetwarzana na procesorze docelowym.

Dzięki takiemu schematowi postępowania w przyjętym modelu systemu rozproszonego, można dla każdego z zasobów określić czas wykonywania odpowiedniej akcji. Obliczanie czasu reakcji systemu (*end-to-end response*) będzie więc polegać na zsumowaniu czasów reakcji poszczególnych zasobów. Wartość ta będzie następnie porównana z ograniczeniem czasowym zadania (*end-to-end deadline*).

Jeżeli zostanie określony najgorszy czas wykonywania akcji przez poszczególne zasoby, to sumując te czasy i odejmując je od ograniczenia czasowego zadania, można określić ograniczenie czasowe dla wiadomości, które mają być przesłane przez magistralę. Przy opracowywaniu metod szeregowania wiadomości na magistrali należy zdefiniować model wiadomości. Przyjęty model obejmuje wiadomości okresowe, sporadyczne oraz wiadomości w tle (bez ograniczeń czasowych), a protokół sieciowy reprezentowany jest (na podstawie rozdziału 2) przez warstwę DDL (rys. 2).



Rys. 2. Model strumienia wiadomości i protokołu sieciowego

Elementy modelu [7]:

- PM (*periodic message*) oznacza wektor p wiadomości okresowych, SM (*sporadic message*) oznacza wektor s wiadomości sporadycznych, czyli:
 $PM = \{pm_1, pm_2, \dots, pm_k, \dots, pm_p\}$, $SM = \{sm_1, sm_2, \dots, sm_s\}$
- tp_i jest okresem występowania wiadomości okresowej pm_i , $i = 1, 2, \dots, p$,
- ts_j oznacza minimalny czas pomiędzy kolejnymi pojawieniami się wiadomości sporadycznej sm_j , $j = 1, 2, \dots, s$,

- dp_i (ds_j) oznaczają graniczne przedziały czasowe dla zrealizowania przesłania wiadomości pm_i (sm_j). Jeżeli wiadomość pm_i (sm_j) została przekazana w chwili t do wysłania, to ograniczenie czasowe wiadomości pm_i wynosi $t + dp_i$ i $t + ds_j$ dla wiadomości sm_j . Zakłada się spełnienie warunku $dp_i \leq tp_i$ dla wiadomości okresowych oraz $ds_j \leq ts_j$ dla wiadomości sporadycznych.
- przez cp_i (cs_j) oznaczamy przedział czasu, jaki potrzebny jest do nadawania wiadomości pm_i (sm_j). Wartość cp_i (cs_j) możemy wyznaczyć z zależności:

$$cp_i = \frac{1}{\alpha}(bp_i + \beta) \quad (1)$$

$$cs_j = \frac{1}{\alpha}(bs_j + \beta)$$

gdzie: α — prędkość transmisji w [bajt/s],

β — dane nadmiarowe, czyli te elementy, które są dołączane do transmitowanej wiadomości przez warstwę łącza danych i warstwę fizyczną, np.: pole adresowe, pole kontrolne, itp,

bp_i — długość wiadomości pm_i [bajt],

bs_j — długość wiadomości sm_j [bajt].

Oczywiście przedział czasu nadawania wiadomości powinien być mniejszy od jej ograniczenia czasowego, co zapisujemy $cp_i < dp_i$ ($cs_j < ds_j$).

- BM (background message) oznacza wektor k wiadomości w tle, czyli:

$$BM = \{bm_1, bm_2, \dots, bm_p, \dots, bm_k\}.$$

- wiadomości są niezależne od siebie.

Przedstawiony model wiadomości posłuży do opracowania metod szeregowania wiadomości na magistrali.

5. ZASTOSOWANIE METODY GRMS DO SZEREGOWANIA ZADAŃ

Do najbardziej popularnych metod szeregowania zadań ze statycznym priorytetem należy metoda GRMS [3, 4, 14]. Algorytm szeregowania zadań w metodzie GRMS działa według następujących reguł:

- każdemu zadaniu zostaje przydzielony priorytet,
 - priorytety są przydzielane do zadań zgodnie z zasadą, że priorytet jest odwrotnie proporcjonalny do okresu występowania zadania, czyli zadanie z krótszym okresem występowania otrzymuje wyższy priorytet (stąd wywodzi się nazwa metody),
 - zadanie o wyższym prioryecie może wywłaszczyć zadanie o prioryecie niższym.
- W literaturze podawane są dwa najistotniejsze twierdzenia dotyczące metody GRMS.

Twierdzenie

Niech t_i oznacza czas występowania zadania i . Jeżeli $t_i < t_{i+1}$, to zadanie i ma wyższy priorytet niż zadanie $i+1$.
Jeżeli $t_i < t_{i+1}$, to zadanie i ma wyższy priorytet niż zadanie $i+1$.
i spełnia

to są spełnione warunki Gran. Oznacza to, że mniejsze wartości t_i wykorzystywane do szeregowania zadań. Jeżeli $t_i < t_{i+1}$, to zadanie i ma wyższy priorytet niż zadanie $i+1$.
przypadki twierdzenia.

Twierdzenie

Niech τ_i oznacza czas występowania zadania i . Jeżeli $\tau_i < \tau_{i+1}$, to zadanie i ma wyższy priorytet niż zadanie $i+1$.
Innymi słowy, jeżeli $\tau_i < \tau_{i+1}$, to zadanie i ma wyższy priorytet niż zadanie $i+1$.
ograniczone

gdzie: $\bar{x}$

Wartość

reprezentacja

Twierdzenie 1 GRMS (*warunek wystarczający*) [5]

Niech $\tau = \{\tau_1, \dots, \tau_n\}$ będzie zbiorem n niezależnych okresowych zadań uszeregowanych według malejącego priorytetu, posiadających czas wykonania c_i oraz okres występowania t_i taki, że ograniczenie czasowe d_i równe jest okresowi występowania zadania ($d_i = t_i$). Zadanie τ_i poprzedza zadanie τ_{i+1} , gdy zadanie τ_i ma wyższy priorytet niż zadanie τ_{i+1} . Dla metody GRMS τ_i poprzedza τ_{i+1} gdy $t_i < t_{i+1}$.

Jeżeli dla tego zbioru zadań stosowany jest algorytm szeregowania GRMS i spełniony jest warunek:

$$(1) \quad \sum_{i=1}^n \frac{c_i}{t_i} \leq n(2^{1/2} - 1) \quad (2)$$

to są spełnione ograniczenia czasowe dla wszystkich zadań z tego zbioru.

Graniczne wykorzystanie zasobów $n(2^{1/n} - 1)$ dla dużych n dąży do $\ln 2 = 0,69$. Oznacza to, że dla każdej liczby zadań, dla których wykorzystanie zasobów jest mniejsze od 0,69, będą spełnione wymagania RT. W czasie, gdy zasoby nie są wykorzystane przez zadania z ograniczeniami czasowymi, mogą być wykonywane pozostałe zadania.

Jeżeli dla pewnego zbioru zadań nie jest spełniony warunek podany w twierdzeniu 1, to nie oznacza to jeszcze, że ograniczenia RT nie są spełnione. W takim przypadku dla sprawdzenia spełnienia ograniczeń czasowych należy posłużyć się twierdzeniem, podającym warunek konieczny i wystarczający.

Twierdzenie 2 GRMS (*warunek konieczny i wystarczający*) [5]

Niech $\tau = \{\tau_1, \dots, \tau_n\}$ jest zbiorem niezależnych zadań okresowych, dla których ograniczenia czasowe są mniejsze lub równe okresowi występowania. Jeżeli zadanie τ_i dotrzyma pierwszego ograniczenia czasowego, gdy w tym samym czasie wystartowały zadania o wyższym priorytecie, to dotrzyma ono swoich ograniczeń czasowych również w przyszłości.

Innymi słowy, dla zadania τ_i (zadania od τ_1 do τ_{i-1} mają wyższy priorytet) z okresem występowania t_i , ograniczeniem czasowym $d_i \leq t_i$ i czasem wykonania c_i zakładając, że w czasie $t = 0$ wystartowały zadania o wyższym priorytecie, będą spełnione ograniczenia czasowe jeżeli istnieje taki czas t , że:

$$w(i, t) = c_1 \left\lceil \frac{t}{t_1} \right\rceil + \dots + \left\lceil \frac{t}{t_i} \right\rceil = \sum_{j=1}^i c_j \left\lceil \frac{t}{t_j} \right\rceil = t \leq d_i \quad (3)$$

gdzie: $\lceil x \rceil$ oznacza najmniejszą liczbą całkowitą większą lub równą x .

Wartość $\left\lceil \frac{t}{t_i} \right\rceil$ reprezentuje liczbą aktywacji zadania τ_i w przedziale czasu t , a $c_i \left\lceil \frac{t}{t_i} \right\rceil$ reprezentuje żądany przez to zadanie odcinek czasu w przedziale czasu t . Innymi

słowy, poszukujemy takiego przedziału czasu t , w którym wykona się zadanie τ_i (oraz oczywiście wszystkie kolejne aktywacje zadań o priorytetach wyższych od τ_i), a jednocześnie wartość t będzie mniejsza lub równa ograniczeniu czasowemu zadania τ_i .

Wzór (3) określany jest mianem *sprawydenia czasu zakończenia* (*completion time test*). Jeżeli w systemie występuje blokowanie zadań, to maksymalny czas tego blokowania należy dodać do czasu wykonywania zadania.

Należy jednakże zaznaczyć, że dla wielu rzeczywistych systemów trudno jest opisać zadania tak, jak to podano w założeniach twierdzenia GRMS (wszystkie są cykliczne i mają ściśle zdefiniowane parametry c , t i d). Odstępstwa dotyczą z reguły dwóch aspektów: nie wszystkie zadania są cykliczne, a czas wykonywania zadań nie jest stały. Należy przyjąć wówczas następujące założenia:

- Dla zadań sporadycznych określany jest minimalny czas występowania zadania. Zadania sprowadzane są do grupy zadań okresowych,
- Jeżeli dla pewnych zadań nie jest niemożliwe wyznaczenie tego czasu, należy stworzyć dodatkowe zadanie w systemie do obsługi tychże zadań, tzw. serwer zadań sporadycznych (*sporadic server*). Serwer sporadyczny otrzymuje najwyższy priorytet poprzez uczynienie jego okresu występowania najkrótszym w systemie, a maksymalny czas wykonywania powinien wynosić taką wartość, przy której cały system jest nadal szeregowalny.
- Dla każdego zadania należy określić średni czas wykonywania i maksymalny czas wykonywania.
- Zadaniom o ostrych ograniczeniach czasowych należy przydzielić wyższe priorytety niż zadaniom o ograniczeniach łagodnych.

Opierając się na powyższych założeniach należy tak zaprojektować system, aby był on szeregowalny dla wszystkich zadań dla ich średnich czasów wykonywania oraz istniało uszeregowanie dla zadań o ostrych ograniczeniach czasowych dla ich maksymalnych czasów wykonywania.

Taki sposób postępowania odpowiada analizie najgorszego przypadku dla zadań o ostrych ograniczeniach czasowych. Zaprojektowany w ten sposób system w razie przeciążenia będzie stopniowo zaniedbywał wykonywanie zadań mniej istotnych, wykonując wszystkie zadania najważniejsze.

6. PODSTAWOWE MODELE MAGISTRAL MIEJSCOWYCH

Najistotniejszą rolę w procesie szeregowania wiadomości odgrywa podwarstwa MAC, określająca w jaki sposób i kiedy węzeł otrzymuje dostęp do medium fizycznego. W celu dokonania analizy metod szeregowania wiadomości należy więc dokonać przeglądu rozwiązań stosowanych w magistralach miejscowych. Poniżej rozważono magistrale, które umożliwiają deterministyczny dostęp do medium fizycznego [10, 11, 12].

• Metod

Metod
względ
periody
przekazu

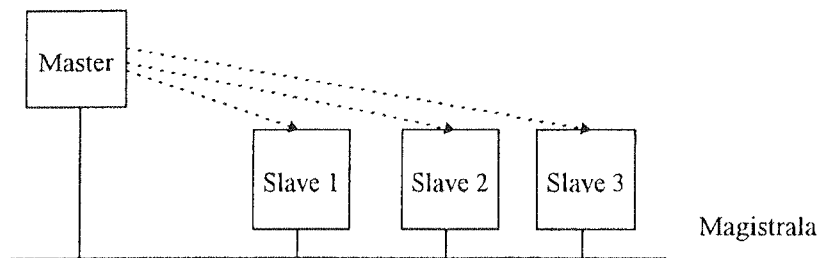
Wię
a praca
idealny
komunik
rozwiąza
komunik
lub koszt
zuje się
liwiają tr
nosć syst
tej metod

• Token

W si
punktu.
(token p
Jeżeli
poprzez
gdyż cza
wylczon
przesłani
kołu pol
wszystki
N jest li
globalny

• Metoda odpytań (*polling*)

Metoda odpytań jest popularna w systemach z magistralami miejscowymi, ze względu na determinizm i prostotę. W tej metodzie wydzielona jest stacja master okresowo odpytująca stacje slave poprzez wysyłanie odpowiednich wiadomości, przekazując im w ten sposób zgodę na transmisję w sieci (rys. 3).



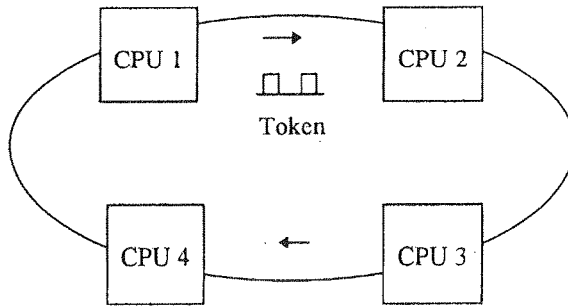
Rys. 3. Magistrala z protokołem odpytań

Większość oprogramowania protokołu jest przechowywana w stacji master, a praca związana z komunikacją ze stacją slave jest minimalna. Protokół ten jest idealny w przypadku centralnej akwizycji danych, dla których nie jest wymagana komunikacja typu peer-to-peer i nie są stosowane globalne priorytety. Wadą tego rozwiązania jest, to, że wystąpienie błędu w węźle master powoduje przerwanie komunikacji, natomiast użycie drugiej stacji typu master jest często niedopuszczalne lub kosztowne. Metoda odpytań zajmuje sporą część pasma transmisyjnego i wykazuje się słabą efektywnością. Niektóre warianty rozwiązań tego protokołu umożliwiają transmisję danych pomiędzy stacjami slave poprzez stację master. Niezawodność systemu zwiększa się poprzez użycie większej liczby węzłów master. Odmiana tej metody jest stosowana przez sieć PROFIBUS DP i Modbus.

• Token Ring

W sieci Token-Ring węzły są połączone w pierścień stosując łącza punkt-do-punktu. Specjalny sygnał znacznika (*token*) jest przekazywany od węzła do węzła (*token passing*) wzdłuż pierścienia (rys. 4).

Jeżeli węzeł ma coś do wysłania, to zatrzymuje znacznik, przesyła wiadomość poprzez pierścień, a następnie oddaje znacznik. Protokół ten jest deterministyczny, gdyż czas analizowania na znacznik dla najgorszego przypadku może zostać łatwo wyliczony. Protokół jest efektywny przy dużym obciążeniu, gdyż wydatek czasu na przesłanie pustego znacznika jest zminimalizowany. Zwykle implementacja protokołu polega na opóźnieniu o 1 bit w każdym węźle tak, że znacznik może odwiedzić wszystkie węzły w $N + T$ jednostkach czasu przeznaczonych na transmisję 1 bitu. N jest liczbą stacji w sieci, a T jest długością (liczbą bitów) wiadomości. Priorytet globalny jest ustalany poprzez zmianę pola priorytetu w znaczniku. Pole to oznacza,

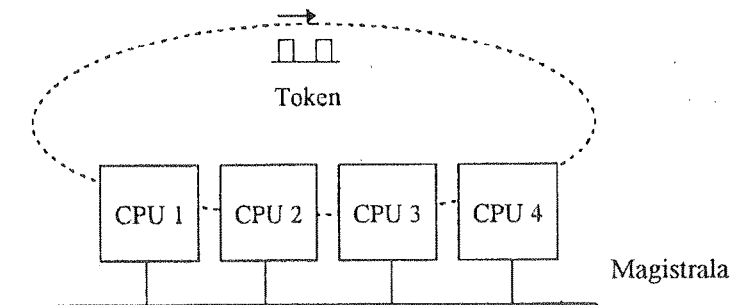


Rys. 4. Przekazywanie znacznika w sieci Token Ring

że tylko węzły o wyższym priorytecie niż to podano w polu priorytetu mogą przesyłać dane. Inicjalizacja wiadomości znacznika, wykrycie przypadkowo zduplikowanego lub utraconego znacznika wpływa na złożoność i koszt protokołu. Jednym z większych problemów w tej metodzie jest możliwość przerwy w okablowaniu lub awaria jednego z węzłów, powodujące zanik komunikacji w sieci. Rozwiązaniem tego problemu jest sprzętowe omijanie węzłów lub dodatkowy pierścień, co jednak powoduje wzrost kosztów. Sieć pierścieniowa jest typu punkt-do-punktu, a zatem dobrze nadaje się do rozwiązań światłowodowych. Przykładami sieci pierścieniowych są InterBus-S i Sercos.

• Token Bus

Działanie Token Bus przypomina Token Ring, ale znacznik jest przekazywany z węzła do węzła poprzez pierścień wirtualny (rys. 5). Węzły posiadające znacznik mają dostęp do sieci. Token Bus wydajnie pracuje przy dużych obciążeniach sieci, z dużym stopniem determinizmu. Token Bus przesyła wiadomość do wszystkich węzłów, a nie bit po bicie, jak to było w Token Ring. Minimalny czas obiegu znacznika wynosi teraz $N \cdot T$, a nie $N + T$ jak to było dla Token Ring (nie występuje



Rys. 5. Przekazywanie znacznika w protokole Token Bus

równole
nieprak
jednego
usuwan
identyf
dobrze
wiele m
Automa
resource

W d
oparty
i oznac
token p
metody
time) w
stępow
nika dla
odpytar
czas, tzn
i wyłącz

gdzie: u

N_i — z

Syst

dlatego

celu sto

inne w

$c_{i0} = T$

kowego

występo

najwyżs

W d

jeżeli r

stosowa

równoległość w połączeniach). Zasada globalnego priorytetu dla wiadomości jest niepraktyczna. W odróżnieniu od Token Ring przerwa w okablowaniu lub awaria jednego z węzłów nie musi unieruchomić całej sieci. W przypadku dodawania lub usuwania węzła stosowana jest długa procedura rekonfiguracyjna, polegająca na identyfikacji przez każdy węzeł swoich sąsiadów. Ponieważ topologie magistralowe dobrze pasują dla systemów produkcyjnych, dlatego metoda ta jest stosowana przez wiele magistral miejscowych: PROFIBUS FMS, P-NET MAP (*Manufacturing Automation Protocol*), a także została zaadaptowana przez ARCnet (*Attached resource Computer Network*).

7. ZASTOSOWANIE METODY GRMS DO ANALIZY SYSTEMÓW ROZPROSZONYCH

W celu zastosowania metody GRMS do analizy systemów rozproszonych opartych na magistralach miejscowych, wprowadzone zostaną dodatkowe pojęcia i oznaczenia. Oznaczmy czas obiegu znacznika (dla systemu stosującego metodę token passing) lub czas cyklu odpytania węzłów slave przez układ master (dla metody odpytań) jako docelowy czas obiegu znacznika $TTRT$ (*target token rotation time*) w ten sposób, że $TTRT \leq t_{\min}$, gdzie $t_{\min}$ jest najkrótszym okresem występowania zadania w systemie. Oznaczamy w_T jako czas na przekazywanie znacznika dla metody token passing, lub czas na zapytanie przez stację master dla metody odpytań (rozważamy przypadek, gdy stacja master tylko zbiera dane). Pozostały czas, tzn. $TTRT - w_T$ rozdzielamy do poszczególnych węzłów. Czas, w którym węzeł i wyłącznie dysponuje magistralą oznaczony jest przez h_i i obliczany jest wg wzoru:

$$h_i = \frac{u_i}{u} (TTRT - w_T) \quad (4)$$

gdzie: $u = u_1 + \dots + u_n$, a u_i jest stopniem wykorzystania sieci przez stację i $u_i = \sum_{j \in N_i} \frac{c_j}{t_j}$
 N_i — zbiór wiadomości należących do stacji s_i .

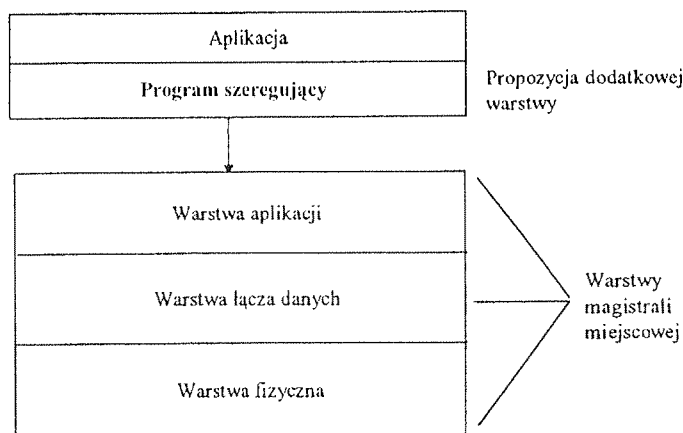
System będzie spełniał wymagania RT, jeżeli każdy węzeł będzie je spełniał, dlatego analiza systemu zostanie przeprowadzona dla każdego węzła osobno. W tym celu stosowany jest model zachowania węzła. W węźle i obciążenie magistrali przez inne węzły możemy przedstawić jako dodatkowe zadanie τ_{i0} o czasie trwania $c_{i0} = TTRT - h_i$ oraz o okresie występowania $t_{i0} = TTRT$. Czas trwania dodatkowego zadania c_{i0} obrazuje zajęcie magistrali przez inne węzły, natomiast okres występowania zadania t_{i0} jest najkrótszy, bo wynosi $TTRT$, co gwarantuje zadaniu najwyższy priorytet.

W dalszej części opracowania indeksy i przy τ_{i0} , c_{i0} , t_{i0} mogą zostać opuszczone, jeżeli rozważania będą dotyczyć ustalonego węzła. W takim przypadku będą stosowane oznaczenia τ_0 , c_0 , t_0 .

Program szeregujący

Stosowanie algorytmu GRMS (lub jakiegokolwiek innego algorytmu szeregowania z priorytetami) jest możliwe tylko wtedy, jeżeli potrafimy zapewnić, iż w danym momencie wysyłana jest wiadomość o największym priorytecie. Systemy sieciowe z reguły dysponują kolejkowaniem wiadomości typu FIFO (*First In First Out*) i dlatego nie zawsze spełniają wymagania algorytmów priorytetowych.

W celu umożliwienia realizacji szeregowania priorytetowego należy pomiędzy warstwą aplikacji, która wytwarza dane do wysłania, a warstwą systemu sieciowego zaimplementować dodatkowe oprogramowanie — program szeregujący (rys. 6). Oprogramowanie to będzie przyjmować od aplikacji dane do wysłania, nadawać priorytety tym danym (według określonego algorytmu) oraz przekazywać je do bufora systemu sieciowego. Dane o aktualnie największym priorytecie będą przekazywane w momencie sygnalizacji przez system sieciowy gotowości do wysłania następnych danych.



Rys. 6. Umieszczenie programu szeregującego

Do zalet takiego usytuowania programu szeregującego możemy zaliczyć:

- brak ingerencji w warstwy protokołu sieciowego,
- dane są selekcionowane przed przekazaniem ich do bufora systemu sieciowego (istnieje możliwość wyłączenia danych o mniejszym priorytecie przez dane o większym priorytecie, czego nie można dokonać po przesłaniu ich do warstwy aplikacji magistrali miejscowej),
- w razie konieczności można zmienić program szeregujący na inny.

Rozważyć należy dwie implementacje omawianego oprogramowania:

- I. Program na bieżąco oblicza czas jaki pozostał do dyspozycji węzła. Jeżeli okaże się, że dane, które ma aktualnie przekazać do bufora sieciowego nie zostaną wysłane, to dane te nie zostaną przekazane do bufora.
- II. Program przekazuje dane do bufora systemu sieciowego, jeżeli tylko bufor ten zostanie opróżniony. W tym przypadku, jeżeli wiadomość nie zostanie wysłana

z po
sieci
pierw
Róż.
algorytm
Po drug
wytworz
wysłane
w przyp
W y
węzłowi
wysłanie
w chw
zostanie
opóźnie
wliczają
odpowi

Blokow

W s
uwzglę
prioryt
przez d
dowoln
szemu c
Dla
stępują

Dochow

Dot
za pom

gdzie: [

szeregowania
w danym
y sieciowe
(First Out)

z warstwą
implemento-
nowanie to
ym danym
sieciowego.
sygnalizacji

z powodu braku wystarczającej ilości czasu, pozostaje ona w buforze systemu sieciowego. Gdy węzeł odzyska prawo transmisji zostaje ona wysłana jako pierwsza.

Różnica pomiędzy tymi dwoma implementacjami jest zasadnicza. Implementacja algorytmu wg pkt. I zapewni, że czas przyznany węzłowi nie zostanie przekroczony. Po drugie, w czasie, gdy węzeł nie wysyłał żadnych wiadomości, aplikacje mogły wytworzyć nowe dane do wysłania o wyższym priorytecie, które powinny być wysłane przed danymi obecnymi w buforze systemu sieciowego. Stanie się tak tylko w przypadku implementacji algorytmu wg pkt. I.

W wielu protokołach sieciowych dopuszcza się przekroczenie przyznanego węzłowi czasu nadawania. Przykładowo, w magistrali PROFIBUS FMS [8] wysłanie wiadomości oznacza rozpoczęcie transakcji, która zakończy się dopiero w chwili otrzymania odpowiedzi. Jeżeli w czasie wykonywania transakcji zostanie przekroczony przyznany węzłowi czas do jego dyspozycji, to pomimo opóźnienia węzeł kontynuuje wykonanie transakcji aż do jej pełnego zakończenia, wliczając w to ewentualne powtarzanie komunikatów, które pozostały bez odpowiedzi.

Blokowanie wiadomości

W systemach sieciowych nie ma wywłaszczania wiadomości dlatego musimy uwzględnić czas blokowania bt . Blokowanie występuje wtedy, gdy dane o niższym priorytecie przesłane do bufora systemu sieciowego nie mogą być wywłaszczone przez dane o wyższym priorytecie. Ponieważ przez magistralę może być przesyłana dowolna wiadomość, dlatego też czas blokowania wiadomości jest równy najdłuższemu czasowi trwania transmisji wiadomości o niższym priorytecie.

Dla każdego zadania $\tau_i \in \tau$ czas blokowania bt_i wyznaczany jest według następującego wzoru:

czyć:

$$bt_i = \begin{cases} 0 & \text{dla } i = n \\ bt_n & \text{dla } i = n - 1 \\ \max \{bt_{i+1}, \dots, bt_n\} & \text{dla } i < n - 1 \end{cases}$$

sieciowego
przez dane
do warstwy

Dochowanie ograniczeń RT

Dotrzymanie ograniczeń czasowych jest sprawdzane dla każdego węzła z osobna za pomocą zależności:

żeeli okaże
nie zostaną

$$w(i-1, t) = c_0 \left\lceil \frac{t}{t_0} \right\rceil + c_1 \left\lceil \frac{t}{t_1} \right\rceil + \dots + c_{i-1} \left\lceil \frac{t}{t_{i-1}} \right\rceil = \sum_{j=0}^{i-1} c_j \left\lceil \frac{t}{t_j} \right\rceil = t \quad (5)$$

o bufor ten
nie wysłana

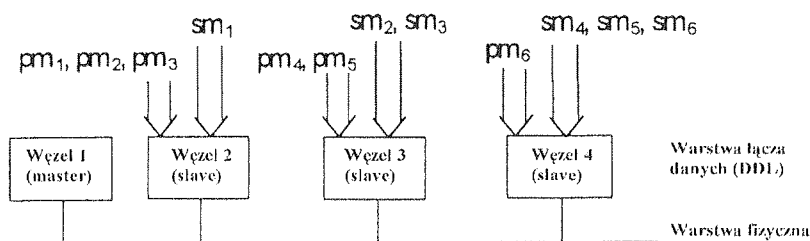
$$tz_i = w(i-1, t) + c_i \leq d_i$$

gdzie: $[x]$ oznacza najmniejszą liczbę całkowitą większą od x .

Jak łatwo zauważyć, wzór (5) uległ modyfikacji w stosunku do wzoru (3). Wynika to z faktu, iż w systemie nie ma wywłaszczania wiadomości i nie trzeba uwzględniać wykonywania zadania o wyższym priorytecie w trakcie wykonywania zadania τ_i . W pierwszej fazie poszukiwana jest taka wartość czasu t , w którym wykonają się wszystkie zadania o priorytetach wyższych od zadania τ_i , czyli zadania $\tau_0 \dots \tau_{i-1}$ (wraz z kolejnymi aktywacjami). Jeżeli istnieje taka wartość t , to poszukiwany czas zakończenia wykonania t_{z_i} dla zadania τ_i otrzymywany jest przez dodanie do wartości t czasu trwania zadania τ_i . Dokonywane to jest dlatego, gdyż nie ma możliwości wywłaszczenia zadania τ_i po jego uruchomieniu. W praktyce wyznaczenie t_{z_i} ze wzoru (5) odbywa się przy użyciu algorytmu iteracyjnego, którego postać zależy od rodzaju magistrali miejscowej.

8. PRZYKŁAD OBLICZENIOWY DLA PROFIBUS DP

Rozważony zostaje system akwizycji danych z magistralą PROFIBUS DP [2, 6, 8]. Dla zadanego modelu strumienia wiadomości przyjęto system złożony z czterech węzłów. Sposób przydziału wiadomości okresowych PM i sporadycznych SM ilustruje rys. 7. Wiadomości ze stacji 2, 3 i 4 mają być przekazywane do stacji centralnej (węzeł 1).



Rys. 7. Przykładowy model wiadomości

Przyjęto następujące parametry wysyłanych wiadomości:

- okres występowania wiadomości peridycznych TP = {1.6, 3.2, 6.4, 1.6, 3.2, 3.2} [ms],
- ograniczenia czasowe wiadomości periodycznych DP = TP
- minimalny okres występowania wiadomości sporadycznych TS = {4, 4, 4, 4, 4, 4} [ms]
- ograniczenia czasowe wiadomości sporadycznych DS = TS,
- długość wiadomości: BP = {3, 5, 8, 4, 2, 4} [bajt]
BS = {9, 5, 3, 5, 5, 8} [bajt].

Przyjęto szybkość transmisji równą 1 [Mbit/s]. Postać ramki zapytania stacji master i ramki odpowiedzi, w której przesyłane są dane, przedstawiono poniżej.

Ramka zapytania

SYN	SD1	DA	SA	FC	FCS	ED
-----	-----	----	----	----	-----	----

Ramka odpowiedzi

SYN	SD2	LE	LEr	DA	SA	FC	DATA _ UNIT	FCS	ED
-----	-----	----	-----	----	----	----	-------------	-----	----

Poszczególne pola mają następujące znaczenie:

SYN	bity synchronizacji, przynajmniej 33 bity
SD1	znacznik początku, kod = 10H, ramka pytania stacji master
SD2	znacznik początku, kod = 68H, ramka przesyłanej zmiennej
LE	długość, 4 do 249 bajtów
LEr	długość zwrotna
DA	adres przeznaczenia
SA	adres źródłowy
FC	znak sterujący
DATA _ UNIT	jednostka danych, długość 1 do 246 bajtów
FCS	suma kontrolna
ED	znacznik końca

Każdy znak składa się z 11 bitów: 8 bitów danych i 3 bitów technicznych

Rozwiązanie:

Czas przesyłania wiadomości będzie obliczany ze wzoru (1), przy czym:

$$\alpha = 1000000 \text{ [bit/sek]}$$

$$\beta = \text{SYN} + \text{SD2} + \text{LE} + \text{LEr} + \text{DA} + \text{SA} + \text{FC} + \text{FCS} + \text{ED} .$$

Czasy przesyłania wyznaczane są więc ze wzorów:

$$cp_i = \frac{\text{SYN} + (bp_i + \text{SD2} + \text{LE} + \text{LEr} + \text{DA} + \text{SA} + \text{FC} + \text{FCS} + \text{ED}) * 11}{1000000} [\text{ms}]$$

$$cs_i = \frac{\text{SYN} + (bs_i + \text{SD2} + \text{LE} + \text{LEr} + \text{DA} + \text{SA} + \text{FC} + \text{FCS} + \text{ED}) * 11}{1000000} [\text{ms}]$$

1.6, 3.2,

$$cp_1 = \frac{33 + 11 * 11}{1000000} = 0,154 [\text{ms}]$$

$$cs_1 = 0.220 \text{ ms}$$

4, 4, 4,

$$cp_2 = 0.176 [\text{ms}]$$

$$cs_2 = 0.176 [\text{ms}]$$

$$cp_3 = 0.209 [\text{ms}]$$

$$cs_3 = 0.154 [\text{ms}]$$

$$cp_4 = 0.165 [\text{ms}]$$

$$cs_4 = 0.176 [\text{ms}]$$

$$cp_5 = 0.143 [\text{ms}]$$

$$cs_5 = 0.176 [\text{ms}]$$

$$cp_6 = 0.165 [\text{ms}]$$

$$cs_6 = 0.209 [\text{ms}]$$

cji master

Czas trwania ramki zapytania stacji master cm wyznaczany jest ze wzoru:

$$cm = \frac{SYN + (SDI + DA + SA + FC + FCS + ED) * 11}{1000000} = \frac{33 + 6 * 11}{1000000} = 0,99 \text{ [ms]},$$

czyli straty czasowe na zapytania przez stację master trzech stacji slave wyniosą $w_T = 3 * 0,99 = 0,297 \text{ [ms]}$. Zapytanie stacji master nie zawiera użytecznej informacji, ale jest niezbędne do prawidłowego funkcjonowania systemu.

Czas cyklu odpytań $TTRT$ dla wszystkich stacji w systemie musi być mniejszy niż $t_{\min} = 1,6 \text{ [ms]}$, gdyż $t_{\min} = \min \{TP\} = 1,6 \text{ [ms]}$. Czas $TTRT$ jest parametrem ustalonym dla danego rozproszonego systemu sterowania. W przykładzie arbitralnie przyjęto, że $TTRT = 1,5 \text{ [ms]}$. Dla zadanego przydziału wiadomości do węzłów, wyznaczany jest stopień wykorzystania magistrali przez poszczególne węzły ze wzorów:

$$u_2 = \frac{cp_1}{tp_1} + \frac{cp_2}{tp_2} + \frac{cp_3}{tp_3} + \frac{cs_1}{ts_1} = 0,239$$

$$u_3 = \frac{cp_4}{tp_4} + \frac{cp_5}{tp_5} + \frac{cs_2}{ts_2} + \frac{cs_3}{ts_3} = 0,230$$

$$u_4 = \frac{cp_6}{tp_6} + \frac{cs_4}{ts_4} + \frac{cs_5}{ts_5} + \frac{cs_6}{ts_6} = 0,192$$

Czasy zajętości magistrali przez poszczególne węzły wynoszą odpowiednio:

$$h_2 = \frac{u_2}{u_2 + u_3 + u_4} (TTRT - w_T) = 0,435 \text{ [ms]}$$

$$h_3 = 0,418 \text{ [ms]}$$

$$h_4 = 0,350 \text{ [ms]}$$

PROFIBUS DP sotsuje metodę odpytań [2, 7]. W przykładzie rozważony jest tylko system akwizycji danych, a nie sterowania. Oznacza to, że dane będą przesyłane tylko ze stacji slave do stacji master. Sytuację tę ilustruje rys. 8. Analiza czasowa dla węzła 1 nie jest rozpatrywana, gdyż jest to stacja master, której zadaniem jest zebranie wiadomości od pozostałych trzech węzłów i nie są na nią nałożone ograniczenia RT.

Do zbioru zadań w węzłach 2, 3, 4 należy dołączyć zadania τ_{i0} , τ_{i0} , τ_{i0} reprezentujące obieg znacznika. Parametry czasowe tych zadań przedstawia tablica 1.

Dotrzymanie ograniczeń czasowych jest sprawdzane za pomocą zależności (5) w każdym węźle.

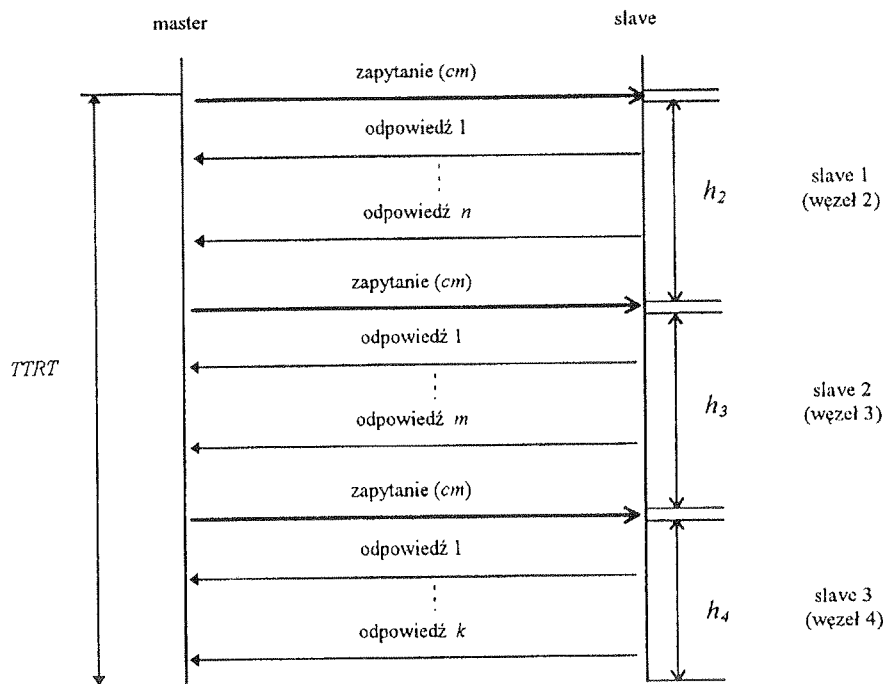
TTT

Numer v
2
3
4

Sprawd

W celu
każdego
dla węz
okresu

Cza
trwania
zadanie
wynosz



Rys. 8. Organizacja przesyłania danych w systemie

Tablica 1

Parametry czasowe zadania τ_{i0} dla poszczególnych węzłów

Numer węzła i	Czas trwania zadania τ_{i0}	Okres występowania zadania τ_{i0}
2	$c_{20} = TTRT - h_2 = 1.5 - 0.435 = 1.065$ ms	$t_{20} = TTRT = 1.5$ ms
3	$c_{30} = TTRT - h_3 = 1.5 - 0.418 = 1.082$ ms	$t_{30} = TTRT = 1.5$ ms
4	$c_{40} = TTRT - h_4 = 1.5 - 0.350 = 1.150$ ms	$t_{40} = TTRT = 1.5$ ms

Sprawdzenie spełnienia ograniczeń czasowych dla węzła 2.

W celu przeprowadzenia analizy czasowej wygodne jest poszeregowanie zadań dla każdego z węzłów z osobną i wprowadzenie oddzielnych oznaczeń. Tabela 2 zawiera dla węzła 2 opis poszeregowanych zadań wg malejącego priorytetu (rosnącego okresu występowania zadania).

Czasy blokowania tb_1 , tb_2 i tb_3 odpowiednio dla zadań τ_1 , τ_2 i τ_3 są równe czasowi trwania najdłuższego zadania w systemie o niższym priorytecie. W tym przypadku zadanie wysłania wiadomości sporadycznej sm_1 ma najdłuższy czas wykonania wynoszący 0.220 [ms] oraz priorytet niższy od zadań τ_1 i τ_2 . Dlatego też czasy

Tablica 2

Parametry czasowe poszeregowanych zadań w węźle 2

Nazwa zadania τ_i		Okres t_i [ms]	Czas trwania c_i [ms]	Ograniczenie czasowe d_i [ms]
τ_0	zadanie τ_{20}	1.5	1.065	1.5
τ_1	wysłanie wiadomości pm_1	1.6	0.154	1.6
τ_2	wysłanie wiadomości pm_2	3.2	0.176	3.2
τ_3	wysłanie wiadomości sm_1	4	0.220	4
τ_4	wysłanie wiadomości pm_3	6.4	0.209	6.4

blokowania tb_1 i tb_2 będą wynosić 0.220 [ms]. Natomiast czas blokowania tb_3 będzie równy czasowi trwania zadania τ_4 , czyli $tb_3 = 0.209$ [ms].
Symbol td_i oznacza dopuszczalny (maksymalny) czas do wykorzystania przez węzeł, którego nie wolno przekroczyć węzłowi podczas wysyłania wiadomości przy i -tym zapytaniu przez stację master.

Sprawdzenie warunków dla zadania τ_1 (wysłanie wiadomości pm_1).

$tz_0 = c_0 = 1.065$ ms
Przyjęcie $tz_0 = c_0$ oznacza, że analiza jest wykonywana dla najgorszego przypadku dla zadania. Tzn. zostało założone, że nastąpiła nie tylko aktywacja wszystkich zadań o wyższym priorytecie, ale też, że węzeł mógł utracić prawo do transmisji wiadomości.
 $tz_1 = tz_0 + c_1 = 1.065 + 0.154 = 1.219$ ms $< td_1 = TTRT = 1.5$ ms,
Uwzględniając czas blokowania zachodzi:
 $tz'_1 = tz_1 + tb_1 = 1.219 + 0.220 = 1.439 < d_1$
Dla zadania τ_1 będą spełnione ograniczenia czasowe.

Sprawdzenie warunków dla zadania τ_2 (wysłanie wiadomości pm_2).

Przy wyznaczaniu zależności czasowych dla zadania τ_2 należy uwzględnić występowanie zadania τ_1 , które ma wyższy priorytet:
 $tz_0 = c_0 + c_1 = 1.065 + 0.154 = 1.219$ ms $< td_1$,
 $tz_1 = tz_0 + c_2 = 1.219 + 0.176 = 1.395$ ms $< td_1$.
Uwzględniając czas blokowania zachodzi:
 $tz'_1 = tz_1 + tb_2 = 1.395 + 0.220 = 1.615$ ms $< d_2$
Zatem dla zadania τ_2 będą również spełnione ograniczenia czasowe.
W podobny sposób dokonywane jest sprawdzenie ograniczeń czasu rzeczywistego dla pozostałych zadań.

Sprawdzenie
 $tz_0 = c_0 +$
 $tz_1 = tz_0 +$
Przesyła
maksyma
iż progra
 $tz_0 = 1.39$
 $tz_1 = tz_0 -$
 $tz_2 = tz_1 -$
 $tz_3 = tz_2 -$
Uwzględni
 $tz'_3 = tz_3 -$
Także dla
rys. 9.

Zadanie τ_3

Sprawdzenie
 $tz_0 = 2 * c_0$
 $tz_1 = tz_0 +$
Aktywowa
 $tz_1 = tz_0 +$
 $tz_2 = tz_1 +$
 $tz_3 = tz_2 +$
 $tz_4 = tz_3 +$

Tablica 2

e czasowe d_i [ms]
0.5
0.6
0.2
0.4
0.4

Sprawdzenie warunków dla zadania τ_3 (wysłanie wiadomości sm_1).

$$tz_0 = c_0 + c_1 + c_2 = 1.065 + 0.154 + 0.176 = 1.395 \text{ ms} < td_1$$

$$tz_1 = tz_0 + c_3 = 1.395 + 0.220 = 1.615 \text{ ms} > td_1.$$

Przesyłanie wiadomości sm_1 o czasie trwania c_3 spowodowałoby przekroczenie maksymalnego czasu td_1 , jaki pozostaje do wyłącznej dyspozycji węzła. Oznacza to, iż program szeregujący nie przekaże tej wiadomości do wysłania. W chwili $tz_0 = 1.395$ ms aktywowane zostanie zatem zadanie τ_0 .

$$tz_1 = tz_0 + c_0 = 1.395 + 1.065 = 2.46 \text{ ms} < td_2 = tz_0 + TTRT = 2.895 \text{ ms},$$

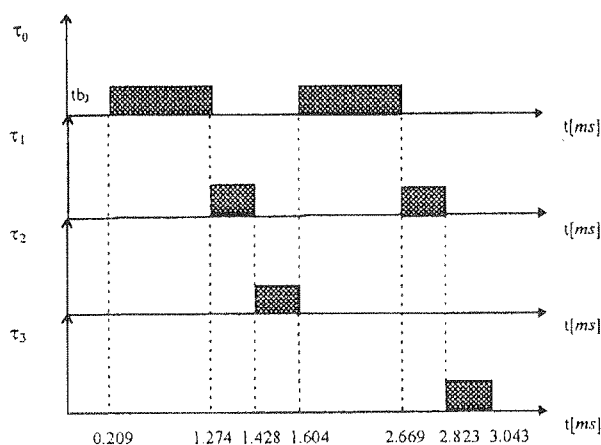
$$tz_2 = tz_1 + c_1 = 2.614 \text{ ms} < td_2,$$

$$tz_3 = tz_2 + c_3 = 2.614 + 0.220 = 2.834 \text{ ms} < td_2.$$

Uwzględniając czas blokowania zachodzi:

$$tz'_3 = tz_3 + tb_3 = 2.834 + 0.209 = 3.043 \text{ ms} < d_3$$

Także dla zadania τ_3 będą spełnione ograniczenia czasowe. Sytuację tę ilustruje rys. 9.

Rys. 9. Diagramy czasowe dla zadania τ_3

Zadanie τ_3 , zgodnie z wyliczeniami, zostanie zakończone po czasie 3.043 ms.

Sprawdzenie warunków dla zadania τ_4 (wysłanie wiadomości pm_3).

$$tz_0 = 2 * c_0 + 2 * c_1 + c_2 + c_3 = 2.834 \text{ ms} < td_2$$

$$tz_1 = tz_0 + c_4 = 2.834 + 0.209 + 3.043 \text{ ms} > td_2.$$

Aktywowane zostanie zadanie τ_0 , co zostanie uwzględnione w obliczeniach:

$$tz_1 = tz_0 + c_0 = 2.895 + 1.065 = 3.899 \text{ ms} < td_3 = tz_0 + TTRT = 4.334 \text{ ms},$$

$$tz_2 = tz_1 + c_1 = 4.053 \text{ ms} < td_3$$

$$tz_3 = tz_2 + c_2 = 4.229 \text{ ms} < td_3$$

$$tz_4 = tz_3 + c_3 = 4.449 \text{ ms} > td_3.$$

Znowu zostanie aktywowane zadanie τ_0 , co zostanie uwzględnione w obliczeniach:

$$tz_4 = tz_3 + c_0 = 5.294 \text{ ms}$$

$$tz_5 = tz_4 + c_1 = 5.448 \text{ ms} < td_4 = tz_3 + TTRT = 5.729 \text{ ms}$$

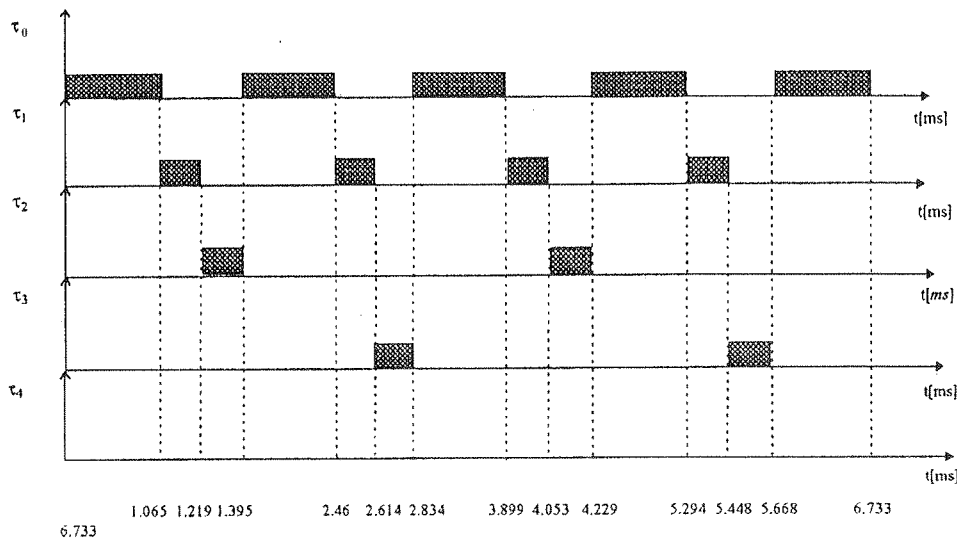
$$tz_6 = tz_5 + c_3 = 5.668 \text{ ms} < td_4$$

$$tz_7 = tz_6 + c_4 = 5.887 > td_4.$$

Znowu zostało aktywowane zadanie τ_0 :

$$tz_7 = tz_6 + c_0 = 6.733 \text{ ms} > dp_3$$

Dla zadania τ_4 nie będą dotrzymane ograniczenia czasowe (rys. 10).



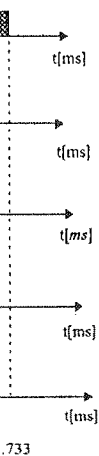
Rys. 10. Diagramy czasowe dla zadania τ_4

Dla węzłów 3 i 4 można wykazać, postępując w podobny sposób, że wykonywane w nich zadania dotrzymają ograniczeń czasowych.

9. ALGORYTM SPRAWDZANIA WARUNKÓW RT DLA POJEDYNCZEGO ZADANIA

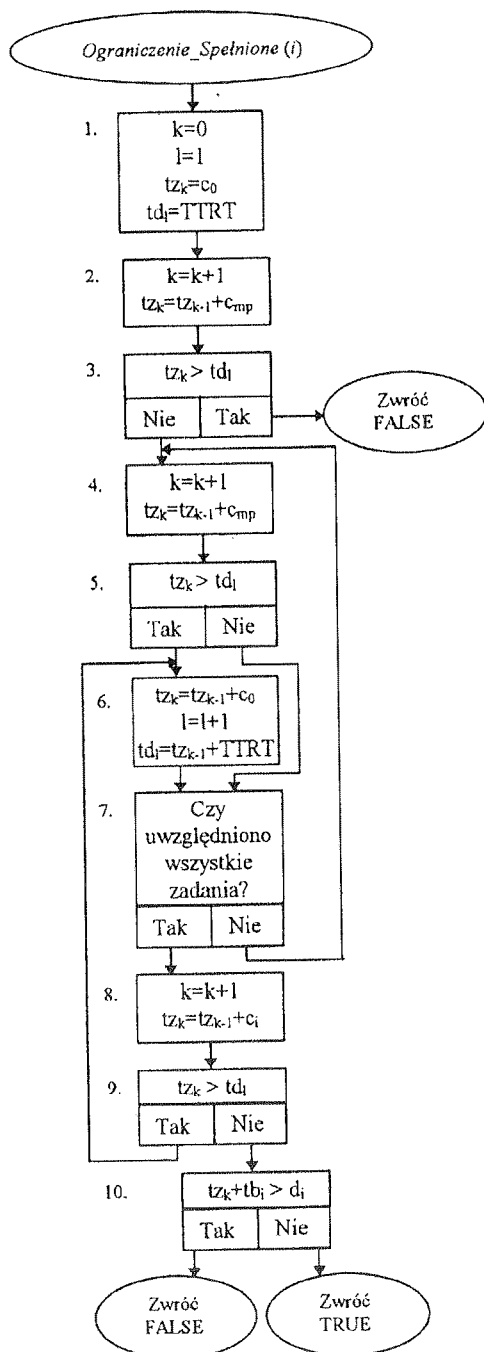
Na podstawie przykładu z rozdziału 8 widać, że istnieje ogólny sposób postępowania do sprawdzenia spełnienia ograniczeń czasu rzeczywistego dla dowolnego zadania. Sposób wyznaczenia spełnienia tych ograniczeń można zapisać w formie algorytmu przedstawionego przy użyciu schematu blokowego (rys. 11). Algorytm przedstawiono w postaci funkcji o nazwie *Ograniczenie Spełnione* (i) z jednym parametrem i , który jest indeksem zadania, dla którego wykonywana jest analiza. Funkcja zwraca wartość *TRUE*, jeżeli ograniczenia RT są spełnione, w przeciwnym przypadku zwraca wartość *FALSE*.

liczeniach:



ykonywane

sposób po-
dla dowol-
na zapisać
o (rys. 11).
pełnione (i)
ywana jest
spełnione,



Rys. 11. Schemat algorytmu do wyznaczania spełnienia ograniczeń RT

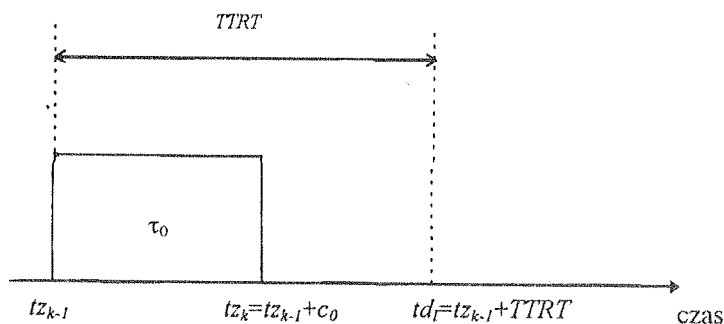
W algorytmie stosowane są następujące oznaczenia:

- $\tau = \{\tau_1, \dots, \tau_i\}$ — zbiór zadań posortowanych od najwyższego priorytetu do najniższego,
- c_{mp} — czas trwania zadania o aktualnie najwyższym priorytecie z pośród zadań gotowych do wykonania,
- k — kolejny numer iteracji przy obliczaniu tz ,
- tz_k — czas zakończenia wykonywania zadań,
- l — kolejny numer zapytania przez stację master węzła, w którym znajduje się zadanie τ_i ,
- td_l — oznacza dopuszczalny (maksymalny) czas do wykorzystania przez węzeł przy l -tym zapytaniu przez stację master. Inaczej, oznacza to chwilę czasową, której nie wolno przekroczyć węzłowi podczas wysyłania przez niego wiadomości przy l -tym zapytaniu przez stację master,
- tb_i — czas blokowania zadania τ_i .

W celu wyznaczenia spełnienia ograniczeń czasowych w algorytmie musi być wykonanych szereg iteracji (rys. 11). Wyjaśnione zostaną operacje wykonywane w wybranych blokach schematu.

Jeżeli $tz_1 > td_1 = TTRT$ (blok 3 schematu), to sprawdzanie ulega zakończeniu, ponieważ system jest nieszegregowalny. Przydzielony czas do dyspozycji przez węzeł jest zbyt krótki, aby wykonać zadanie τ_1 . Stąd wynika wniosek, że jeżeli czas trwania któregośkolwiek zadania w węzle będzie dłuższy niż czas przydzielony do wyłącznej dyspozycji przez ten węzeł, to system będzie nieszegregowalny.

Jeżeli dodanie czasu c_{mp} (blok 4 schematu) spowoduje przekroczenie maksymalnego czasu, jaki był do dyspozycji stacji (blok 5 schematu), to oznacza to, że zadanie nie zostanie wykonane, a wtedy zadanie τ_0 , obrazujące zajęcie magistrali przez inne węzły, zostanie uaktywnione w czasie odpowiadającym wartości tz_{k-1} . Zatem węzeł odzyska ponownie prawo do wysyłania swoich wiadomości w chwili czasowej $tz_k = tz_{k-1} + c_0$. Z powyższego wynika, że chwila czasowa, która nie będzie mogła być przekroczona, wyniesie $td_l = tz_{k-1} + TTRT$. Sytuację tę ilustruje rys. 12.



Rys. 12. Zależności czasowe przy aktywacji zadania τ_0

W blo
zadania c
żadnego :

Przy p
tego, waż
sprawdze
czasowy
włączone
następują

- najważ
widzen
ności je
- mimo z
wania v
szerego
- metoda
czasow
protoko
ny, mo
nia wia

W op
warunków
rzystany v
nych syste

Praca
systemów

1. C. C a
distribut
2. W. B o
Automa
3. M. K I
Analysis
4. M.H. K
Industria
5. L. S h a
A Frame
p. 68 — 8

W bloku 7 schematu sprawdzany jest warunek, czy uwzględniono wszystkie zadania o priorytetach wyższych od zadania τ_i i czy nie nastąpiła kolejna aktywacja żadnego z nich.

10. WNIOSKI

Przy projektowaniu rozproszonych systemów komputerowych czasu rzeczywistego, ważną rolę odgrywa analiza czasowa takich systemów. Analiza ta pozwala na sprawdzenie prawidłowości działania systemów w sensie spełnienia ograniczeń czasowych. Oczywiście metody takie stanowią wycinek badań, które powinny być włączone w fazę projektowania. Z omówionych w artykule zagadnień wynikają następujące spostrzeżenia:

- najważniejszą warstwą protokołu sieciowego magistral miejscowych, z punktu widzenia zagadnienia szeregowania wiadomości, jest warstwa DDL, a w szczególności jej podwarstwa MAC, sterująca dostępem do medium,
- mimo znaczących różnic dotyczących szeregowania zadań procesorem i szeregowania wiadomości na magistrali, istnieje wiele podobieństw i dlatego algorytmy szeregowania zadań mogą być użyte do szeregowania wiadomości,
- metoda GRMS należy do klasycznych metod sprawdzania spełnienia ograniczeń czasowych i została opracowana dla globalnego programu szeregującego. Jeżeli protokół dostępu do medium fizycznego magistrali miejscowej jest deterministyczny, możliwe staje się użycie tej metody w systemie rozproszonym, do szeregowania wiadomości na magistrali miejscowej.

W opracowaniu przedstawiono także algorytm do sprawdzania dochowania warunków czasu rzeczywistego dla dowolnego zadania. Algorytm ten będzie wykorzystany w projektowanym oprogramowaniu przeznaczonym do analizy rozproszonych systemów czasu rzeczywistego bazujących na magistralach miejscowych.

Praca została w ramach grantu KBN Nr 8T11C 039 14: Analiza i projektowanie systemów komputerowych czasu rzeczywistego o różnym stopniu rozproszenia.

BIBLIOGRAFIA

1. C. Cardeira, Z. Mammari: *A schedulability analysis of task and network traffic in distributed real-time systems*. Measurement, Vol. 15, 1995, pp. 71–83
2. W. Boroń, J. Wyłupek: PROFIBUS — standard otwartych sieci przemysłowych. Pomiary, Automatyka i Kontrola. 1996, Nr 11, pp. 313–317
3. M. Klein et al: *A Practitioner's Handbook for Real-Time Analysis: Guide to Rate-Monotonic Analysis for Real-Time Systems*. Kluwer Academic Publishers, Boston, July 1993
4. M.H. Klein, L.P. Lehoczky, R. Rajkumar: *Rate-Monotonic Analysis for Real-Time Industrial Computing*. Computer January 1994, pp. 24–33
5. L. Sha, R. Rajkumar, S. Sathaye: *Generalised Rate-Monotonic Scheduling Theory: A Framework for Developing Real Time Systems*. Proc. of the IEEE. Vol. 82, No. 1, Jan. 1994, p. 68–82

6. A. Syrczyński: *Przemysłowa sieć lokalna PROFIBUS*. Biuletyn PIAP, Nr 3/185, Warszawa 1996, p. 3–41
7. F. Vasques, G. Juanolet: *Fieldbus MAC Mechanism for Hard Real Time Data Communication Support*. CNRS LAAS, LAAS Report 93396, 1993
8. M. i Volz: *PROFIBUS. Technische Druckschrift. PROFIBUS Nutzerorganisation*. Oktober 1995
9. J. Werewka, A. Drwał, S. Żaba: *Zastosowanie magistral miejscowych w sterowaniu maszyn i napędów elektrycznych*. II Konferencja Naukowo-Techniczna. Zastosowanie Komputerów w Elektrotechnice, Poznań 7–9 kwiecień 1997
10. J. Werewka, S. Żaba, A. Drwał: *Protokoły dostępu i charakterystyki czasowe magistral miejscowych*. Pomiary Automatyka Robotyka — Miesięcznik Naukowo-Techniczny. Wrzesień 1997
11. J. Werewka, A. Drwał, S. Żaba: *Protokoły dostępu do medium fizycznego dla wybranych sieci przemysłowych*. AUTOMATION'97 — Konferencja Naukowo-Techniczna, Automatyzacja — Nowości i Perspektywy, 5–7 marzec 1997, PIAP, Tom 2, p. 355–366, Warszawa 1997
12. J. Werewka, A. Drwał, S. Żaba: *Zagadnienie badania charakterystyk czasowych i przepustowości magistral miejscowych* IV Konferencja Systemy Czasu Rzeczywistego, 22–25 wrzesień 1997 Szklarska Poręba, p. 325–336
13. J. Werewka, A. Drwał, S. Żaba: *Zastosowanie magistral miejscowych w sterowaniu i akwizycji danych w systemach produkcyjnych*. Wydawnictwo AGH „Automatyka”, Półrocznik Tom I 1997, p. 99–110
14. J. Zalewski: *What Every Engineer Needs to Know about Rate-Monotonic Scheduling: A Tutorial Advanced Multimicroprocessor Bus Architectures*. Ed. Janusz Zalewski. IEEE Comp. Soc. Press 1995, p. 321–335
15. T. Zydorowicz: *PC i sieci komputerowe*. Wydawnictwo PLJ, Warszawa 1993
16. S. Żaba: *Zastosowanie metody GRMS do badania magistral miejscowych w systemach czasu rzeczywistego*. I Krajowa Konferencja „Metody i systemy komputerowe w badaniach naukowych i projektowaniu inżynierskim”, 25–26 listopad 1997, Kraków, p. 677–684

J. WEREWKA, S. ŻABA

MESSAGE SCHEDULING IN DISTRIBUTED REAL TIME SYSTEM BASED ON FIELDBUS

Summary

The popular solution for control systems is to use distributed computer system. The computers in such a system are often connected by fieldbus network. Investigation of real time behaviour of distributed computer control systems based on fieldbus is an important task. Recently, analytic methods for test compliance with real time requirements have been developed for centralised real time operating systems. At the same time attempts to modify these methods have been undertaken, so as to use them to examine distributed real time systems, especially for scheduling the messages transmitted through the network. In the first part of the paper a comparison is made between task and message scheduling. The model of distributed computer system for control is discussed. The message scheduling basing on GRMS method is discussed for two types of network protocols used in fieldbuses: polling and token passing. The message blocking is also considered in the model. An example illustrates method for checking of RT-constraint fulfilment for PROFIBUS DP fieldbus.

Key words: distributed control system, real time system, fieldbus network, scheduling, GRMS method.

Projektowanie struktur sieci magistral miejscowych dla rozproszonych systemów sterowania

ANDRZEJ DRWAŁ

Instytut Elektrotechniki i Elektroniki Przemysłowej, Politechnika Krakowska

JAN WEREWKA

Katedra Automatyki, Akademia Górniczo-Hutnicza, Kraków

Otrzymano 1998.11.20

Autoryzowano 1999.03.15

W rozproszonych systemach sterowania zachodzi wymiana informacji pomiędzy urządzeniami poprzez przemysłową sieć komputerową noszącą nazwę magistrali miejscowej (fieldbus). W skład rozproszonego systemu sterowania wchodzi stacje sterujące i zbierające dane. W artykule został rozwinięty model obrazujący przepływ danych pomiędzy stacjami. Podczas projektowania rozproszonych systemów sterowania czasu rzeczywistego należy sprawdzić dotrzymanie warunków czasu rzeczywistego RT. Dotrzymanie warunków RT zależy od wielu czynników, jednym z nich jest struktura sieci. W artykule zostanie przeprowadzona analiza wpływu struktury sieci na dotrzymanie warunków RT w systemie. Dotrzymanie warunków RT zostanie sprawdzone przy użyciu metody GRMS. Przedstawiony został przykład obrazujący wpływ struktury sieci na dochowanie warunków RT w systemie wymiany danych. Poruszone zostało również zagadnienie wyboru optymalnej struktury sieci.

Słowa kluczowe: Systemy czasu rzeczywistego, sterowanie rozproszone, magistrale miejscowe, optymalizacja struktury sieci.

1. WSTĘP

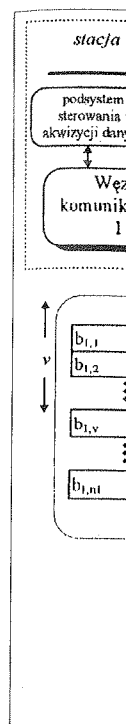
W systemach produkcyjnych, dla zapewnienia sprawnego sterowania produkcją, niezbędna jest szybka i pewna komunikacja [10]. Przesyłanie danych pomiędzy poszczególnymi ogniwami systemu produkcyjnego odbywa się przy użyciu magistral miejscowych. Właściwe zaprojektowanie systemu rozproszonego opartego na magistralach miejscowych wymaga przeprowadzenia dokładnej analizy wymagań stawia-

nych przez system oraz możliwości ich spełnienia dla założonych parametrów magistral [1]. Projektowanie systemu rozproszonego powinno być wykonane w oparciu o model formalny. Model ten może zostać użyty dla wyznaczenia takiej struktury systemu, w której będą spełnione narzucone ograniczenia czasowe. Sposób sprawdzenia dochowania ograniczeń RT czasu rzeczywistego (Real Time — RT) zależy od metod szeregowania wiadomości. Jedną z bardziej znanych metod szeregowania jest GRMS, przyznająca zadaniu o najkrótszym okresie występowania największy priorytet. Metoda ta może być zastosowana do szeregowania wiadomości. Zakładamy, że system składa się ze stacji podłączonych do wspólnej magistrali. Rozwiązanie takie przy dużej liczbie stacji lub intensywnej wymianie danych może prowadzić do naruszenia ograniczeń czasowych. W przypadku, gdy w zaprojektowanym systemie nie są spełnione warunki RT, należy rozważyć możliwość zmiany struktury systemu poprzez podział na podsystemy. W takim przypadku każdy podsystem będzie posiadał swoją magistralę miejscową, co zwykle powoduje odciążenie pozostałej części systemu. Spośród ogromnej liczby różnych rozwiązań wybierane są te, które spełniają ograniczenia czasowe. W przypadku, gdy istnieje wiele rozwiązań podziału systemu na podsystemy spełniające ograniczenia RT, należy dokonać wyboru najlepszej struktury w sensie przyjętego kryterium.

Artykuł będzie próbą potwierdzenia tezy, że podział sieci transmitujących dane na podsieci poprawia zachowanie systemu spełnienia warunków.

1. MODEL STRUKTURY SYSTEMU CZASU RZECZYWISTEGO BAZUJĄCEGO NA MAGISTRALACH MIEJSCOWYCH

Rozpatrzmy system sterowania składający się z $NK = \{1..m\}$ stacji. Każda ze stacji składa się z: podsystemu sterowania oraz węzła komunikacyjnego, do którego przychodzą wiadomości z innych stacji oraz z którego wysyłane są wiadomości do innych stacji (rys. 1). Oznaczmy przez $\mathbf{B}$ tablicę długości zapisu danych w systemie wyrażoną w bajtach. Każda ze stacji $l \in NK$ może posiadać maksymalnie n_l wartości różnych typów danych generowanych w tej stacji wykorzystywanych przez inne podsystemy. Element b_{lv} tablicy $\mathbf{B}$ będzie oznaczał długość danej v tworzonej na stacji l . Odznaczmy przez $\mathbf{R}$ tablicę długości przedziałów czasowych w milisekundach, która podaje, co jaki czas uaktualniane są dane na stacjach. Element r_{lv} podaje w milisekundach, co jaki odstęp czasu uaktualniane są dane v na stacji l . Czas uaktualniania danych może zależeć między innymi od szybkości przetworników pobierających dane oraz od oprogramowania wykonywanego na stacji. Podsystem sterowania i akwizycji danych generuje dane wykorzystywane lokalnie na stacji lub przeznaczone do wykorzystania przez inne stacje. Generacja danych polega na odczycie danych z czujników wchodzących w skład podsystemu lub wytworzeniu danych w wyniku obliczeń opartych o bieżący stan systemu. Stacja może być uważana za stację nadawczą, gdy dane z tej stacji są przekazywane do innych stacji systemu.

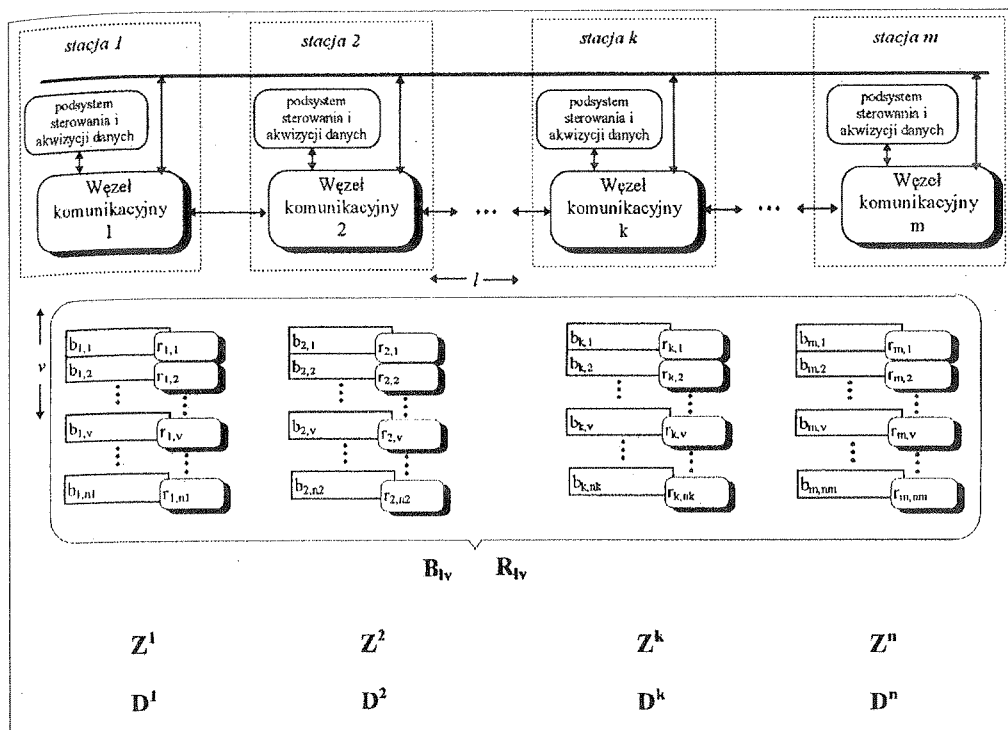


Każda cza, przyj wymaga z tablicę bin z innych s

$$z_{lv}^k =$$

Zatem realizacji s jeżeli przed nia w celu ograniczen D^k składaj czasowe dl przez stacj nych na k

parametrów
wykonane
czenia takiej
ia czasowe.
(Real Time
iej znanych
ym okresie
do szerego-
czonych do
intensywnej
. W przypa-
RT, należy
podsystemy.
miejscową,
omnej liczby
we. W przy-
spełniające
przyjętego
ających dane



Rys. 1. Model rozproszonego systemu sterowania

Każda z wymienionych stacji może być również traktowana jako stacja odbiorcza, przyjmująca dane z innych stacji, gdyż wyznaczenie sterowania na stacji wymaga znajomości wartości danych z innych podsystemów. Oznaczmy przez Z^k tablicę binarną dla stacji odbiorczej k , która określa zapotrzebowanie na wartości z innych stacji odbiorczych. Element z_{lv}^k tablicy Z^k definiowany jest jako:

$$z_{lv}^k = \begin{cases} 1 & \text{— gdy stacja } k \text{ dla wyznaczenia sterowania potrzebuje wartość} \\ & \text{danej } v \text{ ze stacji } l, \\ 0 & \text{— w przeciwnym przypadku.} \end{cases}$$

Zatem element tablicy Z^k określa zapotrzebowanie na dane, wymagane do realizacji sterowania. Zakłada się, że sterowanie może być poprawnie wyznaczone jeżeli przedział czasu od momentu powstawania wartości do momentu jej przekazania w celu rozpoczęcia wyznaczania sterowania nie będzie większy niż zadane ograniczenie czasowe. Ograniczenia czasowe dla stacji k będą opisane przez tablicę D^k składającą się z elementów d_{lv}^k . Element d_{lv}^k podaje w milisekundach ograniczenie czasowe dla pozyskiwania wartości wielkości v , powstającej na stacji l , a wymaganej przez stację k . Dla ułatwienia dalszego zapisu przyjęto, że liczba danych generowanych na każdej stacji wynosi n , gdzie n oznacza największą spośród liczb: $n_1, \dots, n_i$;

$n = \max \{n_1, \dots, n_i\}$. Stosowane w ten sposób dane znormalizowane posłużą do dogodnego zapisu formalnego modelu. Ewentualne aplikacje komputerowe, w celu przyspieszenia obliczeń, mogą dokonywać konwersji danych do postaci przydatnej do bardziej efektywnych obliczeń. W wyniku normalizacji, tablice C , R , Z^k , D^k mogą zostać przedstawione jako prostokątne $m \times n$ powstałe poprzez wypełnienie brakujących pól zerami.

Opisany powyżej system generowania i transmisji danych może zostać przedstawiony w postaci czwórki:

$$ST = \langle Z, D, R, B \rangle.$$

Przyjęto następujące oznaczenia i nazewnictwo:

$NK = \{1, 2, \dots, m\}$ — zbiór numerów węzłów nadawczo odbiorczych, a indeksy $k, l \in NK$,

$NW = \{1, 2, \dots, n\}$ — zbiór numerów danych generowanych w węźle, przy czym indeksy $v, w \in NW$.

Elementy czwórki mają następujące znaczenie:

Z^k jest tablicą zapotrzebowań informacyjnych węzła k , przy czym element tablicy przyjmuje wartości:

$$z_{lv}^k = \begin{cases} 1 & \text{— gdy węzeł } k \text{ wymaga danej } v \text{ z węzła } l, \\ 0 & \text{— w przeciwnym przypadku.} \end{cases}$$

D^k jest tablicą ograniczeń czasowych węzła k , natomiast element tablicy D^k definiowany jest jako:

$$d_{lv}^k = \begin{cases} x & \text{— gdy węzeł } k \text{ wymaga danej } v \text{ z węzła } l, \\ 0 & \text{— w przeciwnym przypadku.} \end{cases}$$

x — wartość liczbową oznaczającą ograniczenie czasowe pozyskiwania danej v z węzła l

R jest tablicą przedziałów czasowych uaktualniania danych z elementami o wartościach:

$$r_{lv} = \begin{cases} x & \text{— } x > 0 \text{ gdy istnieje dana } v \text{ w węźle } l, \\ 0 & \text{— w przeciwnym przypadku.} \end{cases}$$

x — wartość liczbową oznaczającą okres co ile dana v z węzła l jest uaktualniana.

B jest tablicą długości danych, której elementy posiadają wartości:

$$b_{lv} = \begin{cases} x & \text{— } x > 0 \text{ gdy w węźle } l \text{ występuje dana } v, \\ 0 & \text{— w przeciwnym przypadku.} \end{cases}$$

x — wartość liczbową oznaczającą długość wiadomości v w węźle l .

W po
znaczenie
(2 bajty).
Zanim m
ność zde

- czy wy
 $z_{lv}^k > 0$
- czy dl
 k istni

System

będą do
będzie p
kazywan

Spełnieni

- szybko
- obciąż

Rozważ

którego b

rzeczywi

W dalsz

Wartość

Podstaw

dane z po

czasowe.

ogranicz

okres cza

poprzez

wych wi

G — tab

g_l

Element

Każda z

podanym

nym i pr

W podsystemie sterowania mogą być tworzone dane, które mogą mieć różne znaczenie i długości. Przykładowymi wielkościami mogą być: ciśnienie zbiornika (2 bajty), poziom przekroczonego (1 bit), parametry regulatora PID (3×2 bajty). Zanim model zostanie użyty do dalszego projektowania, należy sprawdzić poprawność zdefiniowanego modelu poprzez określenie:

- czy wymagane wielkości istnieją w podsystemach, tzn. jeżeli dla stacji k zachodzi $z_{lv}^k > 0$, to czy $b_{lv} > 0$, to znaczy czy dana ta jest generowana na stacji l ,
- czy dla węzła k , $z_{kv}^k = 0$ to znaczy nie zachodzi sytuacja taka, że dla węzła k istnieje zapotrzebowanie informacyjne na dane tylko od niego samego.

2. ANALIZA WYMAGAŃ CZASU RZECZYWISTEGO

System sterowania będzie działał poprawnie, jeżeli dane tworzone na danej stacji będą dostępne wystarczająco szybko na innych stacjach. Udostępnienie danych będzie polegało na przekazywaniu wiadomości w środowisku transmisyjnym. Przekazywane wiadomości będą zawierały dane generowane na stacji.

Spełnienie ograniczeń czasowych zależy między innymi od:

- szybkości transmisji systemu,
- obciążenia transmisyjnego.

Rozważone będzie zadanie określenia struktury systemu komunikacyjnego, dla którego będą spełnione narzucone ograniczenia czasowe. Spełnienie ograniczeń czasu rzeczywistego zależy również od struktury przesyłanych danych w wiadomościach. W dalszych rozważaniach założono, że wiadomość zawiera tylko jedną wartość. Wartość taka może mieć różną długość zapisu tak jak to definiuje tablica B .

Podstawowym problemem występującym w systemie jest wyznaczenie jak często dane z poszczególnych stacji powinny być wysyłane aby spełnione były ograniczenia czasowe. Dlatego w modelu należy wyznaczyć dodatkowo tablicę minimalnych ograniczeń czasowych przekazywania danych G . Element g_{lv} tej tablicy oznacza okres czasu co jaki dana v z węzła l zostanie wysłana. Tablica ta zostaje utworzona poprzez wybór najmniejszego ograniczenia spośród wszystkich ograniczeń czasowych wiadomości v , ze wszystkich węzłów systemu.

G — tablica ograniczeń czasowych wysyłanych wiadomości

$$g_{lv} = \begin{cases} x & x > 0 \text{ — najmniejsze ograniczenie czasowe wiadomości } v \\ & \text{w systemie (jeśli wiadomość } v \text{ istnieje),} \\ 0 & \text{— jeżeli dana nie istnieje.} \end{cases}$$

Element g_{lv} tablicy G wyznaczamy z zależności:

$$\forall z_{lv}^k \neq 0; \Rightarrow g_{lv} = \min_{k \in NK} d_{lv}^k.$$

Każda ze stacji powinna wysyłać dane z okresem nie większym od przedziału czasu podanym w tablicy G . Ze względu na występujące opóźnienia w systemie transmisyjnym i przekazywania danych należy wyznaczyć okresy co jaki czas dane mają być

wysłane. Załóżmy, że każda stacja ma wysyłać daną v w jednakowych odstępach czasowych co t_{lv} . Należy określić wartości t_{lv} dla, których będą spełnione ograniczenia czasowe.

T — tablica przedziałów czasowych wysyłanych danych o elementach:

$$t_{lv} = \begin{cases} x & x > 0 \text{ — odstęp czasu co ile wysyłana jest dana } v \\ & \text{węzle } l \text{ (o ile istnieje),} \\ 0 & \text{— jeżeli dana nie istnieje,} \end{cases}$$

W projektowanym systemie powinien zachodzić warunek: $t_{lv} \leq g_{lv}$ czyli wiadomość v na stacji l musi być wysyłana z okresem mniejszym lub równym niż ograniczenie czasowe.

System generowania i transmisji danych w czasie rzeczywistym można określić też jako: $STRT = \langle Z, D, R, B, G, T \rangle$. W systemie $STRT$ należy wyznaczyć takie T , dla którego będą spełnione ograniczenia czasowe.

Dyskusja

Wysyłanie wiadomości z odstępem czasu g_{lv} leży na granicy poprawności działania systemu. Każde nawet najmniejsze opóźnienie czasowe spowoduje, że nie będą spełnione ograniczenia czasowe. Dlatego w wielu rozważaniach praktycznych przyjmuje się pewien współczynnik tolerancji $t_{lv} = a g_{lv}$ (gdzie $a < 1$), który ma zapewnić poprawne działanie systemu. Przyjęcie wartości $a = 0.1$ może w wielu przypadkach dawać zadowalające wyniki. Jednak przyjęcie arbitralnie pewnej wartości współczynnika tolerancji, nawet małej, może nie zapewnić spełnienia ograniczeń czasowych. Przyjęcie zbyt małej wartości dla współczynnika a może powodować zbyt duży wzrost kosztów systemu.

W systemach czasu rzeczywistego często stosowane są ostre ograniczenia czasowe. Przy tych ograniczeniach odpowiedź systemu ma być dokonana wewnątrz zadanego przedziału czasu. Warunek ten musi być spełniony, by zagwarantować poprawne działanie systemu [6, 8, 10].

Rozwiniętych zostało szereg metod szeregowania zadań, które można użyć do szeregowania wiadomości. Najbardziej popularną metodą statycznego szeregowania jest metoda GRMS (Generalised Rate Monotonic Scheduling) [3, 4, 5, 6, 7, 9] a najbardziej popularną metodą szeregowania dynamicznego jest metoda EDF (Earliest Deadline First) [2]. W metodzie GRMS priorytety do zadań są przypisane według zasady: im krótszy okres występowania zadania tym wyższy priorytet. Metoda EDF polega na dynamicznym przydzielaniu priorytetu. Im zadanie jest bliżej swojego ograniczenia czasowego tym większy uzyskuje priorytet. Metoda EDF może być użyta zarówno do zadań okresowych jak i nieokresowych. W dalszych rozważaniach, dla sprawdzenia dotrzymania warunków RT skoncentrowano się na metodzie GRMS.

Do sprawdzenia warunków RT stosowane są twierdzenia GRMS, z których poniżej przytoczono dwa. W twierdzeniach tych stosowane są następujące oznaczenia: liczba n — zbiór zadań okresowych, $\tau = \{\tau_1, \tau_2, \dots, \tau_n\}$ zbiór zadań okresowych,

$C = \{c_1,$
 $T = \{t_1,$
Oznacze
mi z uwa
daną v
 c_{lv} oznac
/ wielkoś

Twierdzenie

Niech
posiadają
ogranicze
szeregowa

to spełnio
Graniczn
 n dąży d
większe o
jest mniej
zadań nie
spełnione
poniższym
wania za

Twierdzenie (completi

Dla z
występow
że w cza
ogranicze

gdzie: $[x]$

- Algory
z następuj
1. Zakład
2. Oblicza
3. Spraw
4. Jeżeli t
5. Jeżeli m

h odstępach
e ogranicze-

$C = \{c_1, c_2, \dots, c_n\}$ czasy wykonania zadań,
 $T = \{t_1, t_2, \dots, t_n\}$ okresy występowania zadań.

Oznaczenia w twierdzeniach są uproszczone w porównaniu z wcześniejszymi zapisami z uwagi na to, że rozważania dotyczą pojedynczej stacji. Indeks „ lv ” oznaczający daną v na stacji l , został uproszczony do postaci „ v ”. Na przykład w nazwie c_{lv} oznaczającej czas wykonywania zadania v na stacji l , po opuszczeniu indeksu l wielkość c_v oznacza czas wykonywania zadania v .

czyli wiado-
równym niż

określić też
takie T , dla

Twierdzenie 1. Warunek wystarczający [3, 7]

Niech $\tau = \{\tau_1, \tau_2, \dots, \tau_n\}$ będzie zbiorem n niezależnych zadań okresowych posiadających czas wykonania c_i oraz okres zadadnia występowania t_i taki, że ograniczenie czasowe $d_i = t_i$. Jeżeli dla tego zbioru zadań stosowany jest algorytm szeregowania GRMS i spełniony jest warunek:

$$\sum_{i=1}^n \frac{c_i}{t_i} \leq n(2^{1/2} - 1)$$

ści działania
ędą spełnione
je się pewien
wne działanie
adawalające
rancji, nawet
alej wartości
mu.

zenia czaso-
a wewnątrz
warantować

żna użyć do
szeregowania
5, 6, 7, 9]
etoda EDF
ą przypisane
ytet. Metoda
dziej swojego
że być użyta
żaniach, dla
zie GRMS.
S, z których
ące oznacze-

to spełnione są ograniczenia czasowe dla wszystkich zadań tego zbioru.

Graniczne wykorzystanie (wartość graniczna) zasobów $(n(2^{1/2} - 1))$ dla dużych n dąży do $\ln 2 = 0.69$. Dla mniejszej liczby zadań graniczne wykorzystanie jest większe od $\ln 2$. Z tego wynika, że jeżeli dla dowolnego zbioru zadań wykorzystanie jest mniejsze od 0.69, to są spełnione ograniczenia czasowe. Jeżeli dla pewnego zbioru zadań nie jest spełniony warunek podany w twierdzeniu 1, to nie oznacza że nie są spełnione ograniczenia RT. W przypadku niespełnienia warunku można posłużyć się poniższym twierdzeniem, które może być stosowane dla każdego algorytmu szeregowania zadań ze statycznym przydziałem priorytetów:

Twierdzenie 2. Warunek konieczny i wystarczający — sprawdzenie czasu zakończenia (completion time test) [3, 7]

Dla zadania τ_i (zadania od τ_1 do τ_{i-1} mają wyższy priorytet) z okresem występowania t_i , ograniczeniem czasowym $d_i \leq t_i$ i czasem wykonania c_i , zakładając, że w czasie $t = 0$ wystartowały zadania o wyższym priorytecie, będą spełnione ograniczenia czasowe jeżeli istnieje taki czas zakończenia t , że:

$$w(i, t) = c_1 \left\lceil \frac{t}{t_1} \right\rceil + \dots + c_n \left\lceil \frac{t}{t_1} \right\rceil = \sum_{j=1}^i c_j \left\lceil \frac{t}{t_j} \right\rceil = t \leq d_i$$

gdzie: $[x]$ oznacza najmniejszą liczbę całkowitą większą lub równą x .

Algorytm wyznaczania wartości funkcji w (analiza dla zadania τ_n) składa się z następujących kroków:

1. Zakłada się $t_0 = c_1 + \dots + c_n$.
2. Oblicza się $t_1 = w(n, t_0)$.
3. Sprawdza się czy $t_1 = t_0$.
4. Jeżeli tak, to należy sprawdzić jeszcze czy $t_1 \leq d_n$ i zakończyć sprawdzanie.
5. Jeżeli nie, to należy obliczyć $t_2 = w(n, t_1)$ i sprawdzić czy $t_2 = t_1$.

6. Działania te powtarzać do uzyskania $t_i = w(n, t_{i-1})$ lub $t_i > d_n$ (przekroczenie ograniczeń czasowych).

3. ZAGADNIENIE PODZIAŁU NA PODSIECI

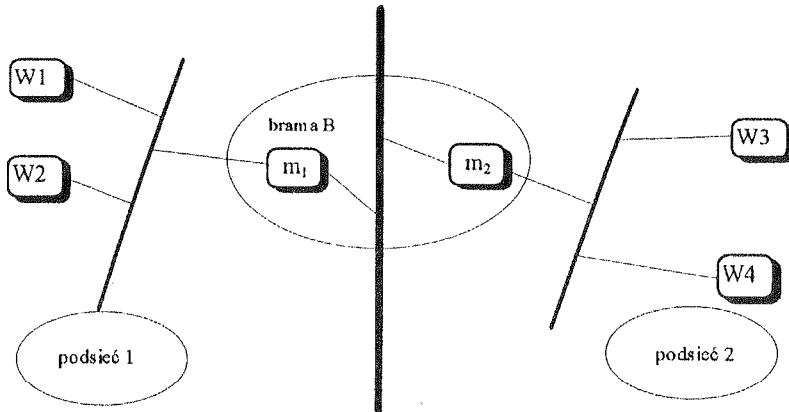
Może okazać się, że ograniczenia czasowe dla systemu nie mogą zostać dochowane. W takim przypadku należy rozważyć dwie zmiany projektu systemu, który mogą pozwolić na dochowanie ograniczeń czasowych:

- podział na podsieci (powoduje zwiększenie wydajności systemu transmisyjnego),
- grupowanie wiadomości (powoduje zmniejszenie obciążenia systemu komunikacyjnego, a przez to zmniejszenie opóźnień czasowych przesyłanych wiadomości).

W opracowaniu rozpatrywany jest w tym artykule tylko podział na podsieci, który ma doprowadzić do polepszenia zachowania się systemu w sensie spełnienia ograniczeń RT.

Założenia:

- podsieci są połączone do wspólnej magistrali, poprzez którą realizowana jest komunikacja, każdy węzeł należy do jednej, i tylko do jednej podsieci,
- ilość węzłów w podsieci może być zawarta w przedziale: $[1..m]$,
- ilość podsieci jest liczbą ze zbioru $[1..m]$.



Rys. 2. Przykład podziału na dwie podsieci

Przyjmujemy, że podsieci $1...s$ są przyłączone do wspólnej magistrali za pośrednictwem dodatkowych węzłów. Dla każdej podsieci i zostaje utworzony jeden dodatkowy węzeł pośredniczący (brama) m_i . Węzły $m_1...m_s$ tworzą dodatkową podsieć (bram), która będzie oznaczona przez zbiór bram $S_b = \{m_1, m_2, ..., m_s\}$. Oznaczmy przez S zbiór podsieci: $S = \{S_b, S_1, S_2, ..., S_s\}$.

Tablica

Dla przyk

węzły 1 i 2

Dla do

tablica zap

Z^{3i} , przy c

$$Z^{3i}_{lv} =$$

Rozwa

Wówczas

działem na

$STRTD =$

Oznac

Problem d

takiego zb

ograniczen

wania tabl

wiadomośc

węzeł k.

Warunki, k

Węzeł w po

1. Dla każ

2. Jeżeli w

Warunki, k

m_i dla podsi

Oznaczn

podobne zn

wań Z^{3i} ele

Tablica przynależności do podsieci P posiada elementy p_{ki} przy czym:

$$p_{ki} = \begin{cases} 1 & \text{gdy węzeł } k \text{ należy do podsieci } S_i, \\ 0 & \text{w przeciwnym przypadku.} \end{cases}$$

Dla przykładu z rys. 2 tablica podsieci P ma postać:

$$P = \begin{vmatrix} 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{vmatrix}$$

węzły 1 i 2 należą do podsieci pierwszej a węzły 3 i 4 do drugiej.

Dla dodatkowego węzła m_i będącego bramą do podsieci S_i zostanie określona tablica zapotrzebowań

Z^{3i} , przy czym element tablicy przyjmuje wartości:

$$z^{3i}_{lv} = \begin{cases} 1 & \text{— gdy jakikolwiek węzeł } k \notin S_i \text{ wymaga danej } v \text{ z węzła } l \in S_i, \\ 0 & \text{— w przeciwnym przypadku.} \end{cases}$$

Rozważany system STRT zostanie poszerzony o pojęcie: $S = \{S_b, S_1, S_2, \dots, S_s\}$. Wówczas system generowania i transmisji danych w czasie rzeczywistym z podziałem na podsieci można określić teraz jako:

$STRTD = \langle Z, D, R, B, G, T, S \rangle$.

Oznaczmy przez $\hat{S}$ zbiór wszystkich dopuszczalnych zbiorów S . Problem doboru struktury rozpatrywanego systemu będzie polegał na wyznaczeniu takiego zbioru S i tablicy R przy zadanych B, T, Z, D, G , dla których spełnione są ograniczenia czasu rzeczywistego. W poniższych rozważaniach, dotyczących definiowania tablic przyjęto rozgłoszeniowy sposób komunikacji pomiędzy węzłami tzn. wiadomość wysłana przez węzeł k jest dostępna w całej podsieci do której należy węzeł k .

Warunki, które musi spełniać tablica przynależności do podsieci S_i

Węzeł w powinien należeć do jednej i tylko do jednej podsieci:

1. Dla każdego węzła k istnieje podsieć S_i do której on należy:

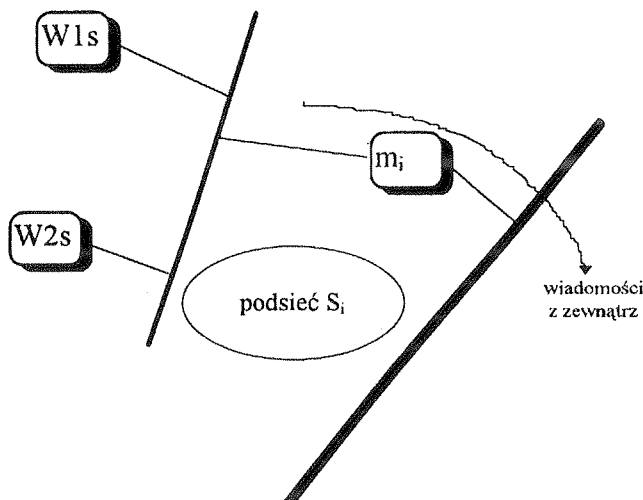
$$\forall k \exists i; \rightarrow p_{k,i} = 1$$

2. Jeżeli węzeł k należy do podsieci S_i to nie może należeć do innej podsieci:

$$\forall k; p_{k,i} = 1 \Leftrightarrow \forall S_j \in \{S/S_i\}, S \rightarrow p_{k,j} = 0$$

Warunki, które muszą być spełnione w tablicy zapotrzebowań węzła pośredniczącego m_i dla podsieci S_i :

Oznaczmy przez Z^{3i} tablicę zapotrzebowań węzła pośredniczącego m_i , która ma podobne znaczenie jak Z^i dla pozostałych węzłów (rys. 3). W tablicy zapotrzebowań Z^{3i} element $z^{3i}_{lv} = 1$, gdy dowolny węzeł nie należący do podsieci S_i wymaga



Rys. 3. Węzeł pośredniczący m_i obrazuje zapotrzebowanie zewnętrznych węzłów na wiadomości z podsieci S_i

przesłania wiadomości v z węzła l należącego do podsieci S_i . W przeciwnym przypadku $z_{lv}^{ji} = 0$:

1. Jeżeli węzeł l należy do podsieci S_i to w tablicy zapotrzebowań węzła pośredniczącego na pozycji z_{lv}^{ji} będzie 1, gdy węzeł k nie należący do podsieci S_i wymaga przesłania wiadomości z węzła l :

$$\forall v; \forall l; p_{li} = 1 \Rightarrow z_{lv}^{ji} \Leftrightarrow \exists k, k \in NK, k \neq l; p_{ki} = 0; \Rightarrow z_{lv}^k = 1$$

2. W pozostałych przypadkach element z_{lv}^{ji} ma wartość 0:

$$\forall l; \forall v; p_{li} = 0 \Rightarrow z_{lv}^{ji} = 0$$

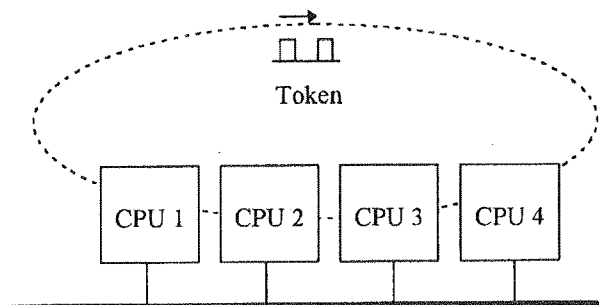
4. WYZNACZENIE WARUNKÓW RT DLA PODSIECI

Zastosowanie metody GRMS do badania systemów rozproszonych opartych o magistralę miejscowe wymaga zdefiniowania modelu tych magistral. Większość magistral bazuje na metodzie krążącego znacznika lub metodzie odpytań. W dalszych rozważaniach TTRT (*target token rotation time*) będzie traktowany jako czas obiegu znacznika (dla systemu stosującego metodę token passing) lub czas cyklu odpytania węzłów slave przez układ master (dla metody odpytań) w rozpatrywanej podsieci. Działanie magistrali stosujących metodę token passing polega na przekazywaniu znacznika z węzła do węzła poprzez pierścień wirtualny. Węzeł posiadający znacznik ma dostęp do sieci i może wysłać wiadomość w przydzielonym czasie będącym częścią całego czasu TTRT (rys. 4). W magistrali token passing wiadomość przesyłana jest do wszystkich węzłów. Ewentualna przerwa w okablowaniu lub awaria jednego z węzłów nie musi unieruchomić całej sieci.

TTRT m
występow
Przez w_T
passing),
Pozostały
sieci nast

gdzie $u_i =$
 $u = u_1 +$
podsieci
zaznaczyć
przynależ
niczący m
do podsie
modelowa
TTRT —
ważanych
mniejsze
obejmuje:
długości v
pracuje m
zresztą w
podsiecia
twierden

Dobór
który skła

Rys. 4. Przekazywanie znacznika w protokole *token passing*

TTRT musi być przedziałem czasu mniejszym lub równym najkrótszemu okresowi występowania zadania w podsieci, gdyż nie byłyby spełnione ograniczenia czasowe. Przez w_T oznaczono czas potrzebny na przekazywanie znacznika (dla metody token passing), lub czas na zapytanie przez stację master (dla metody odpytań) w podsieci. Pozostały czas, tzn. $TTRT - w_T$ rozdzielony jest do poszczególnych węzłów w podsieci następująco:

$$h_i = \frac{u_i}{u} (TTRT - w_T),$$

gdzie $u_i = c_i/t_k^i$ jest stopniem wykorzystania sieci przez stację i , przy czym zachodzi $u = u_1 + \dots + u_n$. Czas h_k^i oznacza przedział czasu dostępu stacji należącej do podsieci do magistrali w jednym obiegu znacznika lub cyklu odpytań. Należy zaznaczyć, że przynależność węzłów do podsieci jest określona poprzez tablicę przynależności do podsieci P . Dodatkowo do podsieci S_i dodany jest węzeł pośredniczący m_i . Tak więc do danej podsieci należą węzły określone tablicą przynależności do podsieci P oraz węzeł m_i . W danym węźle obciążenie magistrali w podsieci będzie modelowane przez dodatkowe zadanie o czasie trwania $TTRT - h_i$ oraz o okresie $TTRT$ — najmniejszym w podsieci, a więc o największym priorytecie. Dla rozważanych magistral przyjmujemy, że ograniczenia czasowe zadań (wiadomości) są mniejsze lub równe ich okresom występowania $d_i \leq t_p$, a czas trwania zadania c_k^i obejmuje: czas formatowania ramki, czas nadawania ramki obliczony na podstawie długości wiadomości (ilości bitów danych oraz narzut protokołu) i szybkości z jaką pracuje magistrala, czas opóźnienia związany z czasem propagacji (jak było to już zresztą wspomniane). Do analizy dotrzymania warunków RT w poszczególnych podsieciach wchodzących w skład systemu zostaną zastosowane, poznane już twierdzenia (Twierdzenie 1 oraz Twierdzenie 2) dotyczące metody GRMS.

5. PRZYKŁAD DOBORU STRUKTURY SIECI

Dobór odpowiedniej struktury sieci powinien przebiegać według algorytmu, który składa się z następujących kroków:

1. Przedstawienie systemu sieciowego w postaci modelu systemu czasu rzeczywistego opisującego ilość danych transmitowanych pomiędzy węzłami nadawczymi i odbiorczymi oraz przyjęcie warunków czasowych dla systemu.
2. Sprawdzenie dotrzymania warunków RT w oparciu o **Twierdzenie 1** (warunek wystarczający) oraz **Twierdzenie 2** (warunek konieczny i wystarczający — *completion time test*) na bazie metody GRMS przystosowanej do badania rozproszonych systemów sieciowych.
3. Sprawdzenie czy istnieje zadanie dla dowolnej stacji, dla którego nie są dotrzymane warunki RT. Jeżeli takie zadanie nie istnieje to warunki RT są spełnione i wykonanie algorytmu zostaje zakończone na tym etapie. W przeciwnym przypadku należy przejść do następnych kroków algorytmu w których zostanie zbadane, czy podział systemu na podsieci umożliwi spełnienie warunków RT.
4. Przekształcenie przyjętego modelu systemu czasu rzeczywistego na model uwzględniający podział na podsieci.
5. Sprawdzenie dotrzymania warunków RT w podsieciach w oparciu o **Twierdzenie 1** oraz **Twierdzenie 2** na bazie metody GRMS przystosowanej do badania rozproszonych systemów sieciowych.
6. Sprawdzenie czy na jakiegokolwiek stacji (w dowolnej podsieci) istnieje zadanie, dla którego nie są dotrzymane warunki RT. Jeżeli takie zadanie istnieje to warunki RT nie są spełnione w całym systemie i należy powtórnie przejść do punktu 4 algorytmu. Jeśli zostaną wyczerpane wszystkie kombinacje podziałów oznacza to, że w systemie nie będą dotrzymane warunki RT. Jeśli istnieje podział, dla którego we wszystkich węzłach zostały dotrzymane warunki RT to przyjmuje się je jako rozwiązanie dopuszczalne.
7. Wyznaczenie pozostałych parametrów systemu, np. dopuszczalnej przepustowości.
8. Ewentualna analiza innych dopuszczalnych rozwiązań (spełniających wymagania RT) w celu wybrania rozwiązania optymalnego uwzględniających inne kryteria.

$$Z^1 = \begin{vmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \\ 0 & 1 & 0 & 0 \end{vmatrix}, \quad Z^2 = \begin{vmatrix} 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 1 & 1 \\ 1 & 0 & 0 & 0 \end{vmatrix}, \quad Z^3 = \begin{vmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \end{vmatrix}, \quad Z^4 = \begin{vmatrix} 1 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \end{vmatrix}.$$

Przyjmujemy, że okresy uaktualniania danych w węzłach równe są ich ograniczeniom czasowym $t^k = d^k$. Okresy mogą być przedstawione np. w milisekundach [ms]. Tablice ograniczeń czasowych poszczególnych węzłów przedstawiono poniżej:

 $D^1 =$ $D^4 =$

Tablica

Trak
poszczeg
C czasó
jako b_{ij}
o długoś
a długoś
przyjęto
z wysła
postać:

Na p
znaczon

$$D^1 = \begin{vmatrix} 0 & 100 & 0 & 0 \\ 0 & 0 & 200 & 0 \\ 0 & 200 & 0 & 200 \\ 0 & 500 & 0 & 0 \end{vmatrix}, \quad D^2 = \begin{vmatrix} 0 & 0 & 0 & 0 \\ 100 & 0 & 0 & 0 \\ 100 & 0 & 300 & 200 \\ 700 & 0 & 0 & 0 \end{vmatrix}, \quad D^3 = \begin{vmatrix} 0 & 0 & 0 & 250 \\ 0 & 0 & 0 & 350 \\ 100 & 0 & 0 & 200 \\ 0 & 500 & 0 & 0 \end{vmatrix},$$

$$D^4 = \begin{vmatrix} 100 & 0 & 100 & 0 \\ 0 & 0 & 200 & 0 \\ 0 & 0 & 300 & 0 \\ 0 & 500 & 0 & 0 \end{vmatrix}.$$

Tablica długości danych B , wyrażonych w bajtach jest zadana następująco:

$$B = \begin{vmatrix} 10 & 10 & 10 & 20 \\ 10 & 0 & 10 & 20 \\ 10 & 10 & 10 & 10 \\ 40 & 40 & 0 & 0 \end{vmatrix}.$$

Traktując przesyłane wiadomości jako zadania, należy określić czas trwania poszczególnych zadań. W tym celu, na podstawie tablicy B wyznaczmy tablicę C czasów trwania zadań (przesłań wiadomości). Jeśli element tablicy B oznaczmy jako b_{ij} to odpowiadający element tablicy C oznacza czas przesłania danych o długości b_{ij} . W przykładzie przyjęto liniową zależność pomiędzy czasem transmisji a długością danych wyrażoną wzorem: $c_{ij} = (b_{ij} + \beta)/\alpha$. Dla omawianego przykładu przyjęto, że przepustowość $\alpha = 1$ [bajt/ms], a ilość dodatkowych bajtów związanych z wysłaniem pojedynczej wiadomości wynosi $\beta = 0$. Wówczas tablica C przyjmie postać:

$$C = \begin{vmatrix} 10 & 10 & 10 & 20 \\ 10 & 0 & 10 & 20 \\ 10 & 10 & 10 & 10 \\ 40 & 40 & 0 & 0 \end{vmatrix}.$$

Na podstawie tablicy ograniczeń czasowych D w poszczególnych węzłach wyznaczona została tablica ograniczeń czasowych G zgodnie z wcześniejszą definicją.

$$G = \begin{vmatrix} 100 & 100 & 100 & 250 \\ 100 & 0 & 200 & 350 \\ 100 & 200 & 300 & 200 \\ 700 & 500 & 0 & 0 \end{vmatrix}.$$

Przyjęto, że tablica T okresów wysyłania wiadomości równa jest G .

Sprawdzenie warunków RT dla całego systemu bez podziału na podsieci:

Wszystkie węzły są przyłączone do wspólnej magistrali (rys. 5). Do zbioru zadań w każdym węźle zostanie dodane zadanie τ_0 obrazujące zajęcie magistrali przez pozostałe węzły. Określony zostanie u_i — stopień wykorzystania sieci przez węzeł i , a następnie rozdzielony czas obiegu tokena na poszczególne węzły.

$$u_1 = 10/100 + 10/100 + 10/100 + 40/700 = 0.357$$

$$u_2 = 10/100 + 10/200 + 40/500 = 0.23$$

$$u_3 = 10/100 + 10/200 + 10/300 = 0.18$$

$$u_4 = 20/250 + 20/350 + 10/200 = 0.19$$

Sumaryczne wykorzystanie magistrali jest w tym przypadku wynosi:

$$u = u_1 + u_2 + u_3 + u_4 = 0.957$$

W przykładzie przyjęto, że dalsza analiza zostanie przeprowadzona dla systemu pracującego na zasadzie token-passing, a czas potrzebny na przekazywanie tokena wynosi $w_T = 10$, oraz czas obiegu tokena $TTRT = 100$.

Wyznaczone zostaną czasy przydzielenia znacznika (tokena) h_i dla węzła i :

$$h_k^i = (u_i/u) * (TTRT - w_T)$$

$$h_1 = (0.357/0.957) * 90 = 33.5 \text{ przyjmujemy} = 33$$

$$h_2 = (0.23/0.957) * 90 = 21.6 \text{ przyjmujemy} = 22$$

$$h_3 = (0.18/0.957) * 90 = 16.9 \text{ przyjmujemy} = 17$$

$$h_4 = (0.19/0.957) * 90 = 17.8 \text{ przyjmujemy} = 18$$

Suma czasów przydzielenia tokena poszczególnym węzłom oraz czas $TTRT$ powinny być mniejsze od najkrótszego okresu występowania zadania w systemie. Obliczone wartości czasu przydzielenia tokena do węzła zaokrągla się do wartości całkowitych (w górę jeśli część po przecinku jest większa niż 0.5) w celu uproszczenia obliczeń.

Przyjęto założenie, że jeżeli węzeł dysponuje nawet minimalną rezerwą czasową to musi wystartować nowe zadanie i zadanie zostanie w pełni wykonane bez przerw. Nie uwzględnia się również czasów blokowania zadań.

W celu dokonania analizy spełnienia ograniczeń czasowych dla węzła 4 zostanie utworzona tabela, przy czym kolejność wiadomości w tabeli jest określona okresem ich uaktualniania. Wiadomość pierwsza obrazuje zajęcie magistrali przez pozostałe węzły. Pozostałe wiadomości 2, 3, 4 odpowiadają wiadomościom 1, 2, 3 węzła 4.

Nr zad.
1
2
3
4

Warunk

5 i 4) (t
Z tablic
dotrzym
czasu za

Dla zad

$$t_{z0} = c_1 -$$

$$t_{z1} = w ($$

$$t_{z0} = t_{z1}$$

Dla zad

$$t_{z0} = c_1 -$$

$$t_{z1} = w ($$

$$t_{z2} = w ($$

$$t_{z2} = t_{z1}$$

Dla zad

$$t_{z0} = c_1 -$$

$$t_{z1} = w ($$

$$t_{z2} = w ($$

$$t_{z3} = w ($$

$$t_{z4} = w ($$

$$t_{z4} > t_4 \Rightarrow$$

niczenie

Z roz
magistral
możliwoś

Tablica 1

Zbiór zadań w węźle czwartym

Numer zadania i		Okres t_i [ms]	Czas trwania zadania c_i [ms]	$\sum_{i=1}^n \frac{c_i}{t_i}$	$i^* \left[2^{\frac{1}{i}} - 1 \right]$
1	τ_0	90	72	0.8	1
2	τ_1	200	10	0.9	0.828
3	τ_2	250	20	0.98	0.779
4	τ_3	350	20	1.037	0.757

biory zadań
strali przez
zez węzeł i ,

Warunki RT zostaną na pewno dotrzymane jeżeli $i^* \left[2^{\frac{1}{i}} - 1 \right] \geq \sum_{i=1}^n \frac{c_i}{t_i}$ (kolumny 5 i 4) (tab. 1).

Z tablicy widać, że dla zadań 2, 3 i 4 metoda ta nie rozstrzyga czy zostaną dotrzymane warunki RT. Należy skorzystać z twierdzenia drugiego (*sprawdzenie czasu zakończenia*)

Dla zadania 2:

$$t_{z0} = c_1 + c_2 = 72 + 10 = 82$$

dla systemu
anie tokena

$$t_{z1} = w(2, t_{z0}) = 1^* c_1 + 1^* c_2 = 72 + 10 = 92$$

$$t_{z0} = t_{z1} < t_{z2} \Rightarrow \text{zostaną dotrzymane warunki RT}$$

a i:

Dla zadania 3:

$$t_{z0} = c_1 + c_2 + c_3 = 72 + 10 + 20 = 102$$

$$t_{z1} = w(3, t_{z0}) = 2^* c_1 + 2^* c_2 + 1^* c_3 = 144 + 20 + 10 = 174$$

$$t_{z2} = w(3, t_{z1}) = 2^* c_1 + 2^* c_2 + 1^* c_3 = 144 + 20 + 10 = 174$$

$$t_{z2} = t_{z1} < t_{z3} \Rightarrow \text{zostaną dotrzymane warunki RT}$$

Dla zadania 4:

$$t_{z0} = c_1 + c_2 + c_3 + c_4 = 72 + 10 + 20 + 20 = 122$$

czas TTRT
w systemie.
do wartości
proszczenia

$$t_{z1} = w(4, t_{z0}) = 2^* c_1 + 2^* c_2 + c_3 + c_4 + c_5 = 144 + 20 + 20 + 20 = 204$$

$$t_{z2} = w(4, t_{z1}) = 3^* c_1 + 3^* c_2 + c_3 + c_4 = 216 + 30 + 20 + 20 = 286$$

$$t_{z3} = w(4, t_{z2}) = 3^* c_1 + 3^* c_2 + 2^* c_3 + c_4 = 216 + 30 + 40 + 20 = 306$$

czasową to
bez przerw.

$$t_{z4} = w(4, t_{z3}) = 4^* c_1 + 4^* c_2 + 2^* c_3 + c_4 + c_5 = 288 + 40 + 40 + 20 = 388$$

$t_{z4} > t_4 \Rightarrow$ nie zostaną dotrzymane warunki RT ponieważ zostało przekroczone ograniczenie czasowe zadania 4.

a 4 zostanie
na okresem
z pozostałe
3 węzła 4.

Z rozważań wynika, że sieć składająca się z 4 węzłów podłączonych do wspólnej magistrali nie spełnia warunków RT w węźle 4 dla zadania 4! Należy więc rozważyć możliwość podziału na podsieci tak, aby były spełnione warunki RT.

Sprawdzenie dochowania warunków RT przy podziale na podsieci

Proponowany jest następujący podział na dwie podsieci. Do podsieci pierwszej będą należały węzły 1 i 2, a do drugiej węzły 3 i 4.

Tablica przynależności węzłów do podsieci P będzie postaci:

$$P = \begin{vmatrix} 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{vmatrix}.$$

Odpowiada to sytuacji przedstawionej na rys. 2. Traktując węzeł m_1 jako węzeł odbierający — tablica zapotrzebowań tego węzła obrazuje wiadomości wymagane przez węzły 3 i 4 od węzłów należących do podsieci 1.

$$\hat{z}^1 = \begin{vmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{vmatrix}.$$

Sprawdzenie warunków RT dla podsieci 2:

Stopnie wykorzystania sieci przez węzły 3 i 5 wynoszą:

$$u_3 = 10/100 + 10/200 + 10/300 = 0.18$$

$$u_4 = 20/250 + 20/350 + 20/200 = 0.19.$$

Dla węzła pośredniczącego stopień wykorzystania sieci u_{m1} wynosi:

$$u_{m1} = 10/100 + 10/100 + 40/500 = 0.28$$

$$u = u_1 + u_2 + u_{m1} = 0.65.$$

Czas obiegu tokena zostanie rozdzielony na poszczególne węzły; przy czym h_i oznacza czas przydzielony dla węzła i :

$$h_i = (u_i/u) * (TTRT - w_T)$$

$$h_3 = (0.18/0.65) * 90 = 24.9 \text{ przyjmujemy } = 25$$

$$h_4 = (0.19/0.65) * 90 = 26.3 \text{ przyjmujemy } = 26$$

$$h_{m2} = (0.28/0.65) * 90 = 38.8 \text{ przyjmujemy } = 39.$$

Odpowiadający tej sytuacji zbiór wiadomości wysyłanych z węzła trzeciego zapisano w tablicy 2.

Tablica 2

Zbiór zadań w węźle trzecim po dokonaniu podziału na podsieci

Numer zadania i		Okres t_i [ms]	Czas trwania zadania c_i [ms]	$\sum_{i=1}^n \frac{c_i}{t_i}$	$i^* \left[2^{\frac{1}{i}} - 1 \right]$
1	τ_0	90	65	0.722	1
2	τ_1	100	10	0.822	0.828
3	τ_2	200	10	0.872	0.779
4	τ_3	300	10	1.905	0.757

Dla zad
kończen

Dla zad

$$t_{z0} = c_1$$

$$t_{z1} = w$$

$$t_{z0} = t_{z1}$$

Dla zad

$$t_{z0} = c_1$$

$$t_{z1} = w$$

$$t_{z2} = w$$

$$t_{z3} = w$$

$$t_{z3} = t_{z2}$$

Kolejny

tym. Zbi

Numer zadania	
1	
2	
3	
4	

Dla zad
zakończe

Dla zad

$$t_{z0} = c_1 +$$

$$t_{z1} = w(3$$

$$t_{z2} = w(3$$

$$t_{z2} = t_{z1} <$$

Dla zad

$$t_{z0} = c_1 +$$

$$t_{z1} = w(4,$$

$$t_{z2} = w(4,$$

$$t_{z2} = t_{z1} <$$

wszej będą

Dla zadań 3, 4 należy skorzystać z twierdzenia drugiego (*sprawdzenie czasu zakończenia*)

Dla zadania 3:

$$t_{z0} = c_1 + c_2 + c_3 = 65 + 10 + 10 = 85$$

$$t_{z1} = w(3, t_{z0}) = 1 * c_1 + 1 * c_2 + 1 * c_3 = 65 + 10 + 10 = 85$$

$$t_{z0} = t_{z1} < t_3 \Rightarrow \text{zostaną dotrzymane warunki RT}$$

jako węzeł
wymagane

Dla zadania 4:

$$t_{z0} = c_1 + c_2 + c_3 + c_4 = 65 + 10 + 10 + 10 = 95$$

$$t_{z1} = w(4, t_{z0}) = 2 * c_1 + c_2 + c_3 + c_4 = 130 + 10 + 10 + 10 = 160$$

$$t_{z2} = w(4, t_{z1}) = 2 * c_1 + 2 * c_2 + c_3 + c_4 = 130 + 20 + 10 + 10 = 170$$

$$t_{z3} = w(4, t_{z2}) = 2 * c_1 + 2 * c_2 + c_3 + c_4 = 130 + 20 + 10 + 10 = 170$$

$$t_{z3} = t_{z2} < t_4 \Rightarrow \text{zostaną dotrzymane warunki RT w węźle 3.}$$

Kolejnym krokiem będzie sprawdzenie dotrzymania ograniczeń RT w węźle czwartym. Zbiór wiadomości wysyłanych z węzła zapisano w tablicy 3.

Tablica 3

Zbiór zadań w węźle czwartym po dokonaniu podziału na podsieci

Numer zadania i		Okres t_i [ms]	Czas trwania zadania c_i [ms]	$\sum_{i=1}^n \frac{c_i}{t_i}$	$i * \left[2^{\frac{1}{i}} - 1 \right]$
1	τ_0	90	64	0.711	1
2	τ_1	200	10	0.761	0.828
3	τ_2	250	20	0.841	0.779
4	τ_3	350	20	1.898	0.757

zy czym h_i

Dla zadań 3 i 4 należy skorzystać z twierdzenia drugiego (*sprawdzenie czasu zakończenia*)

Dla zadania 3:

$$t_{z0} = c_1 + c_2 + c_3 = 64 + 10 + 20 = 94$$

$$t_{z1} = w(3, t_{z0}) = 2 * c_1 + 1 * c_2 + 1 * c_3 = 128 + 10 + 20 = 158$$

$$t_{z2} = w(3, t_{z1}) = 2 * c_1 + 1 * c_2 + 1 * c_3 = 128 + 10 + 20 = 158$$

$$t_{z2} = t_{z1} < t_3 \Rightarrow \text{zostaną dotrzymane warunki RT}$$

go zapisano

Tablica 2

$\left[2^{\frac{1}{i}} - 1 \right]$
1
0.828
0.779
0.757

Dla zadania 4:

$$t_{z0} = c_1 + c_2 + c_3 + c_4 = 64 + 10 + 20 = 114$$

$$t_{z1} = w(4, t_{z0}) = 2 * c_1 + 1 * c_2 + 1 * c_3 + c_4 = 128 + 10 + 20 + 20 = 178$$

$$t_{z2} = w(4, t_{z1}) = 2 * c_1 + 1 * c_2 + 1 * c_3 + c_4 = 128 + 10 + 20 + 20 = 178$$

$$t_{z2} = t_{z1} < t_4 \Rightarrow \text{zostaną dotrzymane warunki RT w węźle 4.}$$

Zbiór wiadomości wysyłanych z węzła pośredniczącego m_2 jest zapisany w tablicy 4. Dla tego węzła należy przeprowadzić podobne rozważania jak dla węzłów: trzeciego i czwartego.

Tablica 4

Zbiór zadań w węźle pośredniczącym podsieci drugiej

Numer zadania i		Okres t_i [ms]	Czas trwania zadania c_i [ms]	$\sum_{i=1}^n \frac{c_i}{t_i}$	$i^* \left[2^{\frac{1}{i}} - 1 \right]$
1	τ_0	90	67	0.744	1
2	τ_1	100	10	0.844	0.828
3	τ_2	100	10	0.944	0.779
4	τ_3	300	10	1.977	0.757

Dla zadań 2, 3 i 4 należy skorzystać z twierdzenia drugiego (*sprawdzenie czasu zakończenia*).

Przeprowadzając dokładną analizę okazuje się że dla tych zadań będą dotrzymane warunki RT.

Z powyższych rozważań wynika, że w podsieci 1 i 2 zostaną dotrzymane warunki RT.

Z tego wynika, że w całym systemie warunki RT zostaną dotrzymane.

Sprawdzenie obciążenia dla całego systemu (bez podziału na podsieci):

W rozpatrywanym przykładzie założymy, że Δ oznacza narzut komunikacyjny protokołu na wysłanie jednej wiadomości i wynosi 10 bajtów. Założymy, że maksymalna przepustowość sieci wynosi: $q_{\max} = 1.7$ [bajt/ms].



Rys. 5. Wszystkie węzły przyłączone są do jednej magistrali

Obciążenie pochodzące od węzła wysyłającego k wynosi: $q_k = \sum_{i=1}^n \frac{b_i + \Delta}{g_i}$, gdzie n jest ilością zadań w węźle k .

obciążenie od węzłów od 1 do 4:

$$q_1 = \frac{10+10}{100} + \frac{10+10}{100} + \frac{10+10}{100} + \frac{40+10}{700} = 0.67$$

$$q_2 = \frac{10+10}{200} + \frac{10+10}{200} + \frac{40+10}{500} = 0.4$$

$$q_3 = \frac{10+10}{100}$$

$$q_4 = \frac{20+10}{250}$$

Obciążenie

się z zale

$$q_c = q_1 +$$

Co stano

Z powyżs

Sprawdzenie

Od węzł

$$q_{m1} = \frac{10+10}{200}$$

$$q_{m2} = \frac{10+10}{100}$$

Obciążenie

$$q_{s1} = q_1 +$$

Co stanow

$$q_{s2} = q_3 +$$

Co stanow

Obciążenie

$$q_m = q_{m1} +$$

Co stanow

Warunki

również s

Dla pr

okresów i

się, że pr

stacje pod

spowodow

z podsieci

6. PRO

Proble

i polega n

nie jest op

any w tab-
la węzłów:

Tablica 4

$\left[2^i - 1 \right]$
1
0.828
0.779
0.757

lenie czasu

dotrzymane

dotrzymane

unikacyjny
ny, że ma-

gdzie n jest

$$q_3 = \frac{10+10}{100} + \frac{10+10}{200} + \frac{10+10}{300} = 0.36$$

$$q_4 = \frac{20+10}{250} + \frac{20+10}{350} + \frac{10+10}{200} = 0.3$$

Obciążenie sieci przy założeniu, że nie dokonujemy podziału na podsieci wyznacza się z zależności:

$$q_c = q_1 + q_2 + q_3 + q_4 = 0.67 + 0.4 + 0.36 + 0.3 = 1.73$$

Co stanowi 101.7[%] obciążenia maksymalnego $q_{\max}$.

Z powyższego widać, że została przekroczona maksymalna przepustowość dla całej sieci!

Sprawdzenie obciążenia w podsieciach:

Od węzłów pośredniczących:

$$q_{m1} = \frac{10+10}{200} + \frac{10+10}{300} + \frac{10+10}{200} = 0.26$$

$$q_{m2} = \frac{10+10}{100} + \frac{10+10}{200} + \frac{40+10}{500} = 0.4.$$

Obciążenie w podsieciach:

$$q_{s1} = q_1 + q_2 + q_{m1} = 0.67 + 0.4 + 0.26 = 1.33.$$

Co stanowi 78.2[%] obciążenia maksymalnego $q_{\max}$.

$$q_{s2} = q_3 + q_4 + q_{m2} = 0.36 + 0.3 + 0.4 = 1.06.$$

Co stanowi 62.4[%] obciążenia maksymalnego $q_{\max}$.

Obciążenie w podsieci węzłów pośredniczących:

$$q_m = q_{m1} + q_{m2} = 0.26 + 0.4 = 0.66.$$

Co stanowi 38.8[%] obciążenia maksymalnego $q_{\max}$.

Warunki przepustowości zostały dotrzymane dla każdej podsieci, zatem zostały również spełnione dla całej sieci!

Dla przyjętych danych systemu dotyczących: długości danych na stacjach oraz okresów ich uaktualniania i ograniczeń czasowych na pozyskiwanie danych okazało się, że przekroczona została przyjęta maksymalna przepustowość jeśli wszystkie stacje podłączone są do wspólnej magistrali. Podział struktury sieci na dwie podsieci spowodował, że obciążenie w podsieciach uległo zmniejszeniu i w efekcie w żadnej z podsieci maksymalna przepustowość nie została przekroczona.

6. PROBLEM OPTYMALNEJ STRUKTURY SYSTEMU SIECIOWEGO

Problem wyznaczenia optymalnej struktury systemu sieciowego jest złożony i polega na znalezieniu takiego rozwiązania, które spełnia warunki RT a jednocześnie jest optymalnym — w sensie przyjętego kryterium.

Załóżmy system złożony z m węzłów. Dla tego systemu istnieje możliwość utworzenia od 1 do m podsieci, przy czym:

- utworzenie jednej podsieci oznacza, że jest to pojedynczy podsystem, w którym spełnione są ograniczenia RT,
- utworzenie m podsieci oznacza przyłączenie wszystkich węzłów do wspólnej magistrali,
- utworzenie l podsieci (l większe od 1 i mniejsze od m) oznacza, że węzły zostają przydzielone do poszczególnych podsieci według określonego kryterium, a w każdej z podsieci należy sprawdzić dotrzymanie warunków RT.

W strukturze złożonej z m węzłów utworzenie dwu podsieci może przebiegać na wiele sposobów. Ilość ta zależy od tego czy m jest parzyste czy nie.

Dla m parzystego:

Spośród m węzłów można utworzyć dwie podsieci na $\sum_{i=1}^{\frac{m}{2}} \binom{m}{i}$ sposobów;

Dla m nieparzystego:

Spośród m węzłów można utworzyć dwie podsieci na $\sum_{i=1}^{\frac{m-1}{2}} \binom{m}{i}$ sposobów;

$$\text{gdzie } \binom{m}{i} = \frac{(m)!}{(m-i)! (i)!}.$$

Jak widać, już w przypadku podziału systemu na dwie podsieci, problem wyboru optymalnej struktury może dotyczyć ogromnej liczby możliwych rozwiązań. Celowe w związku z tym wydaje się wcześniejsze zastosowanie algorytmów heurystycznych dla ograniczenia dopuszczalnej liczby rozwiązań.

Dla wstępnego ograniczenia zbioru rozwiązań może zostać przyjęte założenie: węzeł powinien być przydzielony do takiej podsieci, aby zachodziła wymiana informacji pomiędzy nim a pozostałymi węzłami w podsieci. Zakładając, że węzeł l należy do podsieci s , węzeł k będzie należał do podsieci s jeżeli:

$$\forall k \in NK, p_{k,s} = 1 \Leftrightarrow \exists v; z_{l,v}^k = 1; \exists l; p_{l,s} = 1; \exists v; z_{k,v}^l = 1; v \in NW.$$

Może istnieć wiele rozwiązań spełniających wymagania RT. Spośród tych rozwiązań należy wybrać rozwiązanie, które:

- daje największą tolerancję spełnienia warunków RT

Jako wskaźniki jakości mogą posłużyć funkcje podane dla zbioru zadań w jednym węźle:

$$J = \sum_{i=1}^n (d_i - t_i)^2 \text{ — wskaźnik bezwzględny}$$

po zadaniach $i = 1..n$

$$J = \sum_{i=1}^n \left(\frac{d_i - t_i}{d_i} \right)^2 \text{ — wskaźnik względny}$$

po zadaniach $i = 1..n$

• prowa

$J = m$

gdzie:

s — d

Przy

rale miej

badaniu

Opracow

model jał

systemie

podsieci.

się wów

podsieci

jest polep

warunkó

doboru s

zastosow

• węzeł

zapotr

stępuje

Innymi s

z węzła i

• do pod

• do pod

najkró

• do pod

Celowe b

rozwiąz

ze wzglę

Użycie ał

rozwiąz

1. A. D r

pustowa

22—25

2. K. J e

Sporadi

utworze-

w którym

magistrali,

ły zostają

, a w każ-

biegać na

m wyboru

ń. Celowe

ystycznych

założenie:

wymiana

że węzeł l

.

śród tych

a w jednym

- prowadzi do maksymalizacji przepustowości systemu sieciowego,

$$J = \max_{s \in S} q(s)$$

gdzie: $q(s)$ jest przepustowością magistrali dla danego podziału sieci na podsieci,
 s — dany podział na podsieci, S — zbiór wszystkich podziałów na podsieci.

WNIOSKI

Przy projektowaniu systemów sterowania czasu rzeczywistego opartych o magistrale miejscowe należy rozważyć szereg zagadnień optymalnego doboru magistral. Przy badaniu właściwości systemów magistralowych wygodnie jest skorzystać z modeli. Opracowanie modelu wymaga dokładnej analizy działania systemu aby powstały model jak najwierniej oddawał sytuacje zachodzące w systemie. Jeśli w projektowanym systemie nie są spełnione warunki RT należy rozważyć możliwość podziału systemu na podsieci. Podział na podsieci wydaje się również celowy ze względu na to, że uzyskuje się wówczas prostszą, bardziej czytelną strukturę systemu. Przedstawiony model podsieci oraz oparty na nim przykład obliczeniowy pozwalają zauważyć, że możliwe jest polepszenie własności systemu (na przykład przepustowości czy też dotrzymania warunków RT) poprzez taki podział. Powstaje wówczas problem odpowiedniego doboru struktury systemu. W celu ograniczenia liczby możliwych rozwiązań należy zastosować kryteria przydziału węzłów do podsieci. Kryterium takim może być:

- węzeł ik może należeć do podsieci, do której należy węzeł iz , jeśli w tablicy zapotrzebowań Z_{ik} węzła ik , w kolumnie iz (odpowiadającej węzłowi iz), występuje co najmniej jedna jedynka.

Innymi słowy, jeśli węzeł ik wymaga przesłania co najmniej jednej wiadomości z węzła iz ,

- do podsieci należą węzły dla których występuje największa wymiana informacyjna,
- do podsieci należą węzły dla których okresy uaktualniania wiadomości (T_i) są najkrótsze,
- do podsieci należą te węzły w których występują najdłuższe wiadomości.

Celowe będzie rozważanie wielu różnych dopuszczalnych rozwiązań, dla uzyskania rozwiązania optymalnego. Zagadnienie optymalizacyjne jest jednak bardzo złożone ze względu na konieczność przeanalizowania wielkiej liczby wariantów rozwiązań. Użycie algorytmów heurystycznych pozwoli na ograniczenie liczby dopuszczalnych rozwiązań struktur systemu.

BIBLIOGRAFIA

1. A. Drwał, J. Werewka, S. Żaba: *Zagadnienie badania charakterystyk czasowych i przepustowości magistral miejscowych*. IV Konferencja Systemy Czasu Rzeczywistego, Szklarska Poręba, 22–25 IX 1997
2. K. Jeffaly, D. Stanat, C. Martel: *On Non Preemptive Scheduling of Periodic and Sporadic Tasks*. Proc. of IEEE Real Time Systems Symposium 1991

3. M. Klein et al: *A Practitioner's Handbook for Real-Time Analysis: Guide to Rate-Monotonic Analysis for Real-Time Systems*. Kluwer Academic Publishers, Boston, July 1993
4. M.H. Klein, L.P. Lehoczky, R. Rajkumar: *Rate-Monotonic Analysis for Real-Time Industrial Computing*. Computer January 1994, pp. 24–33
5. L. Liu, J. Layland: *Scheduling algorithms for multiprograming in a hard real-time environment*. J. Assoc. Comput. Mach. Vol. 20, No. 1, pp. 46–61, January 1973
6. N. Malcolm, W. Zhao: *The Timed-Token Protocol for Real-Time Communications*. Computer, January 1994, pp. 35–41
7. L. Sha, R. Rajkumar, S. Sathaye: *Generalised Rate-Monotonic Scheduling Theory: A Framework for Developing Real Time Systems*. Proc. of the IEEE. Vol. 82, No. 1, Jan. 1994, pp. 68–82
8. F. Vasques, G. Juanolle: *Fieldbus MAC Mechanism for Hard Real Time Data Communication Support*. CNRS LAAS, LAAS Report 93396, 1993, p. 9
9. J. Werewka, A. Drwal, S. Żaba: *Zastosowanie magistral miejscowych w sterowaniu maszyn i napędów elektrycznych*. II KN-T nt. Zastosowanie Komputerów w Elektrotechnice. Poznań/Kiekrz 1997
10. J. Zalewski: *What Every Engineer Needs to Know about Rate-Monotonic Scheduling: A Tutorial*. *Advanced Multimicroprocessor Bus Architectures*. Ed. Janusz Zalewski, IEEE Comp. Soc. Press 1995, pp. 321–335

A. DRWAL, J. WEREWKA

DESIGN OF A FIELDBUS NETWORK STRUCTURE FOR DISTRIBUTED CONTROL SYSTEMS

S u m m a r y

A distributed control system consists of a set of control and data acquisition stations that exchange data over an industry computer network known as a fieldbus. In the paper a model is developed which describes information flow from and to stations. In design of a distributed control systems RT (Real Time) constraints are to be considered. The fulfilment of real time constraints depends on many factors, one of them is a network structure. In the paper the influence of a network structure on RT — behaviour is analysed. The RT — conditions are calculated basing on GRMS theory. Some examples illustrates influence of a network structure on meeting of RT conditions. Optimisation problem for determining the best suitable structure is discussed.

Key words: Real time system, distributed control, fieldbus, network structure optimization.

Praca została wykonana w ramach grantu KBN nr 8T11C 039 14 pt: „Analiza i projektowanie systemów komputerowych czasu rzeczywistego o różnym stopniu rozproszenia”.

Coding

TH
connec
article
uneven
used to
compar
be used
Th
system

Key wor

The set
programmat
methods of
enhancemen

Curtis pr

v (

If the column
can be decon
Different
are presented
preclude their

Coding Capacity of PAL-based Programmable Transcoder with Uneven Number Terms per Output

DARIUSZ KANIA

Instytut Elektroniki, Politechnika Śląska, Gliwice

Received 1998.11.08

Authorized 1999.03.17

The problem of implementation digital circuits in programmable devices is strictly connected with decomposition. Programmable circuits have limited internal resources. The article presents an effective coding strategy for programmable PAL-based transcoders with uneven number terms per output cell. The proposed effective word coding algorithm was used to evaluate the coding capacity of popular programmable structures. The results were compared with those obtained by binary and Gray word coding method. The method can be used, in conjunction with other algorithms to code states in sequential automata.

The word coding algorithms presented in this paper are implemented in Decomp system which has been developed at Silesian Technical University.

Key words: decomposition, coding, Programmable Logic Devices

1. INTRODUCTION

The set of combinational functions often cannot be made to fit into any single programmable devices or logic block of CPLD, FPGA devices. The classical methods of functional decomposition are based on Ashenhurst theory and Curtis enhancement [2, 6].

Curtis proved the fundamental theorem:

$$v(X_2/X_1) \leq 2^p \Leftrightarrow f(X_2, X_1) = F[q_1(X_1), g_2(X_1), \dots, g_p(X_1), X_2] \quad (1)$$

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 equation 1.

Different algorithms of functional decomposition often based on Curtis theory are presented in [3–5, 7, 9, 12–18]. The complexity of the propose algorithms often preclude their practical application. We can observe that, after all decomposition,

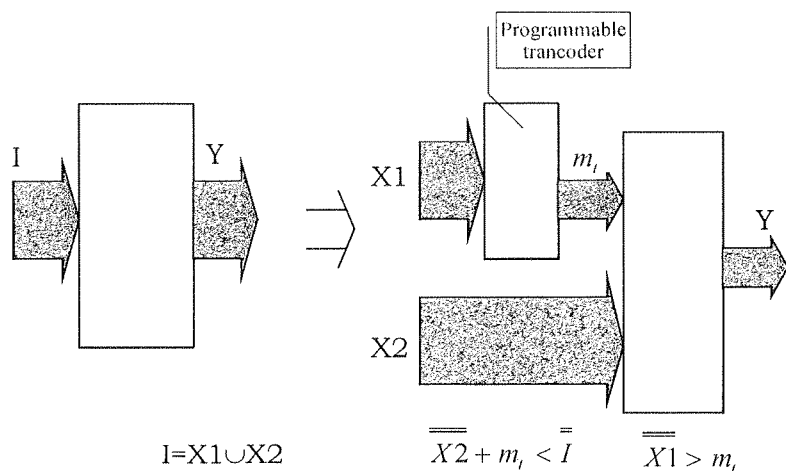


Fig. 1 The principle of partitioning the inputs based on input coding

the big circuit is broken into smaller subcircuits. These subcircuits are a specific transcoders. The simplest method of reducing the number inputs required is input coding [4, 10]. This form of decomposition is presented in Fig. 1.

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. The greater part of existing PLDs are using PAL base structures. Among those are the popular GAL, MACH, MAX devices by AMD and Altera accordingly [1]. A basic PAL structure has a strictly determined number of product terms connected to each individual output. This specific property of PAL structures introduces certain limitations. For example, if, at the stage of input partitioning, the need arises to implement a programmable transcoder, then the individual input vectors have to be coded accordingly [10].

There exist structures in which the number of product terms is not equal for each output. This is the case with the popular 18V10, 22V10, 26V12, 29M16, 32VX10 and the GLB (Global Logic Block) of pLSI circuits by Lattice. The problem of using uneven numbers of terms per output also arises during the synthesis of programmable transcoder.

It was proved [11] for PAL-based transcoders with a uniform number of terms per output that binary coding, due to its uneven use of terms, is not an efficient coding technique. In a previous paper [11] I proposed an efficient word coding algorithm for PAL-based transcoder with a uniform number of terms per output.

What is a coding capacity of programmable transcoder with uneven number terms per output?

The present paper answers to this question and compares the result with other coding method (binary, Gray) for popular GAL devices.

The efficient word coding algorithm and others 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. THEORETHICAL BACKGROUND

To discuss different word coding method for PAL-based transcoders with uneven number terms per output cell, let us first introduce some necessary terminology.

Definition 1:

Let $k = [k_1, k_2, \dots, k_m]$ be the characteristic vector of a PAL-based circuit. The value of k_i is a decimal number equal to the count of terms connected to the i -th output cell.

Definition 2:

Let $Q_j = [q_1, q_2, \dots, q_m]$ be the j -th code vector. The element q_i assumes a value of 0 or 1, and if $q_i = 1$ then term connected to the i -th output cell is used.

Definition 3:

The code matrix $Q = \begin{bmatrix} Q_1 \\ Q_2 \\ \vdots \\ Q_s \end{bmatrix}_{s \times m}$ is an $s \times m$ matrix whose rows are the various

code vectors $Q_j = [q_1, q_2, \dots, q_m]$, $1 \leq j \leq s$.

Let the code matrix $Q = \begin{bmatrix} q_{11} & q_{12} & \dots & q_{1m} \\ q_{21} & q_{22} & \dots & q_{2m} \\ \vdots & \vdots & \ddots & \vdots \\ q_{s1} & q_{s2} & \dots & q_{sm} \end{bmatrix}$, and the code vector $k = [k_1, k_2, \dots, k_m]$.

We shall say that the code matrix Q can be implemented in a circuit based on a PAL structure described by the characteristic vector k if the implementability condition is met, i.e.

$$\forall \quad 1 \leq i \leq m \quad \sum_{j=1}^s q_{ji} \leq k_i. \quad (2)$$

Definition 4:

The value $C = s$ will be called the code capacity of a PAL-based structure with the characteristic vector $k = [k_1, k_2, \dots, k_m]$, where s is the maximum number of words for using code method, which satisfies the implementability condition.

3. BINARY WORD CODING METHOD

The optimal allocation of binary code word bits to output cells is very important for PAL-based transcoders with uneven number terms per output. In the worst case

$$C_{bin\min} = 2k_{\min} + 1, \quad (3)$$

where

$$k_{\min} = \min(k_1, k_2, \dots, k_m). \quad (4)$$

If $k = [k_1, k_2, \dots, k_m]$ is the ordered characteristic vector $k_1 \leq k_2 \leq \dots \leq k_m$ and $b = b_{m-1} \dots b_1 b_0$ is the binary code number, then optimal allocation assigns the b_i -th bit to the output cell which corresponds to the $k_{(m-i)}$ -th element of the characteristic vector

$$b_i \leftrightarrow k_{(m-i)} \quad 0 \leq i \leq m-1. \quad (5)$$

Under such allocation (5) the number of possible combinations is $C_{bin\max}$

$$C_{bin\max} = \min_j \left(2^j + \left\lceil \frac{k_{(m-j)}}{2^j} \right\rceil 2^j + k_{(m-j)} \right) \quad j = 0, 1, \dots, [\lg_2(2k_{\max} + 1)], \quad (6)$$

where

j is bit number

$k_{(m-j)}$ element of the characteristic vector

$k_{\min} = \max(k_1, k_2, \dots, k_m) = k_m$

Result of optimal binary coding for 26V12 transcoder is presented in Fig. 2.

4. GRAY WORD CODING METHOD

If $k = [k_1, k_2, \dots, k_m]$ is the characteristic vector and $g = g_{m-1} \dots g_1 g_0$ is the Gray code number, then the worst allocation assigns the g_i -th bit to the output cell which corresponds to the $k_{\min}$ -th element of the characteristic vector, where

$$i = [\lg_2 k_{\min}]; \quad k_{\min} = \min(k_1, k_2, \dots, k_m). \quad (7)$$

In this case

$$C_{Gray\min} = 2^{[\lg_2(k_{\min})]} + k_{\min}. \quad (8)$$

ecture with
er of words

important
worst case

(3)

(4)

$\leq k_m$ and
as the b_i -th
characteristic

(5)

+ 1)], (6)

a Fig. 2.

s the Gray
cell which

(7)

(8)

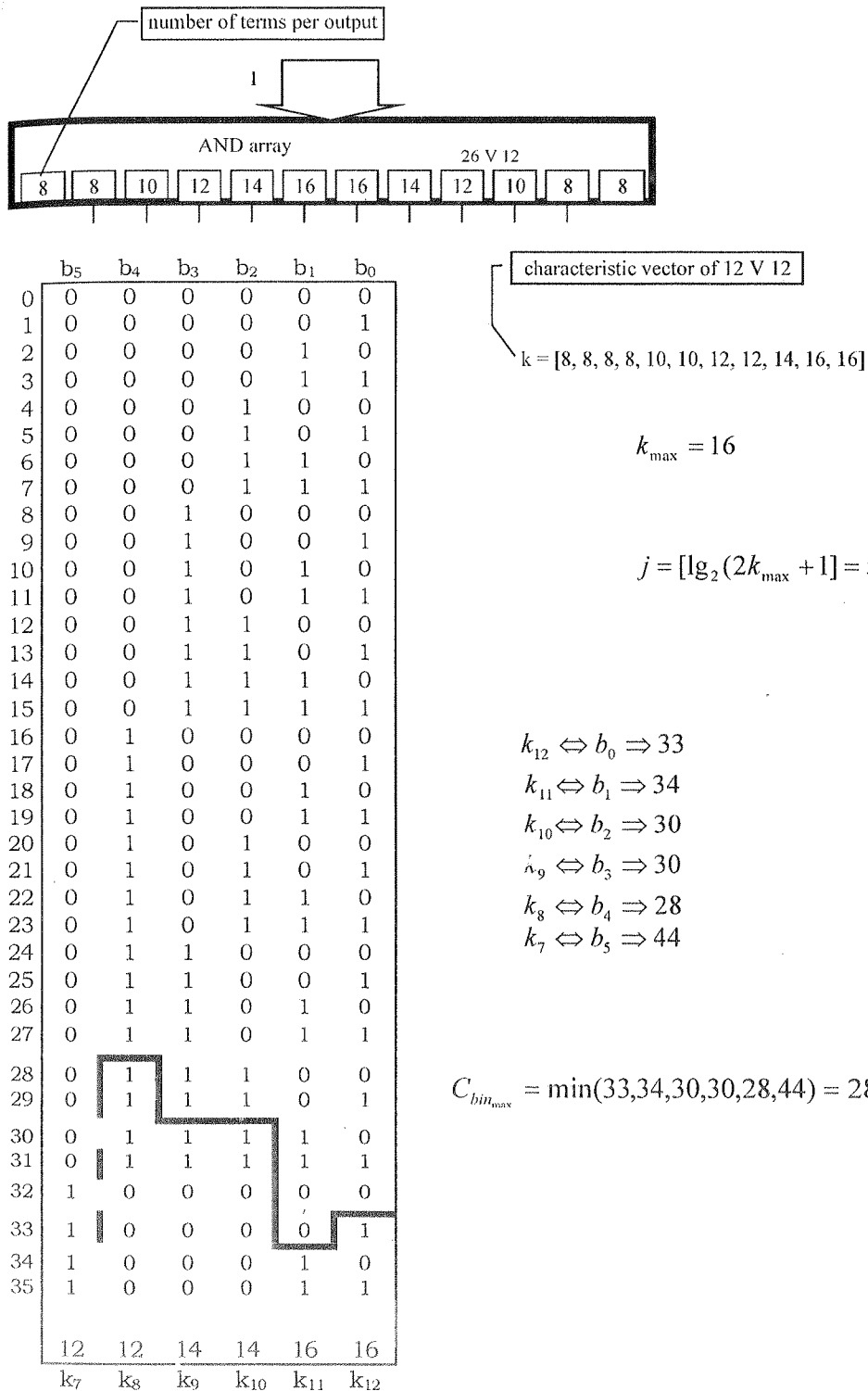


Figure 2. Result of optimal binary coding for 26V12 transcoder

If k
optimal

• if $k_{\max}$

• if $k_{\max}$

Result o

Effici
of differ
tability c

• ordering
• determ
• uniform
• supplier
• supplier

The p
character
($2 \leq i \leq n$)

The p
number fo

The un
terms con
of the pro
1 output, 2
 $k = [k_1, k_2]$
output, the

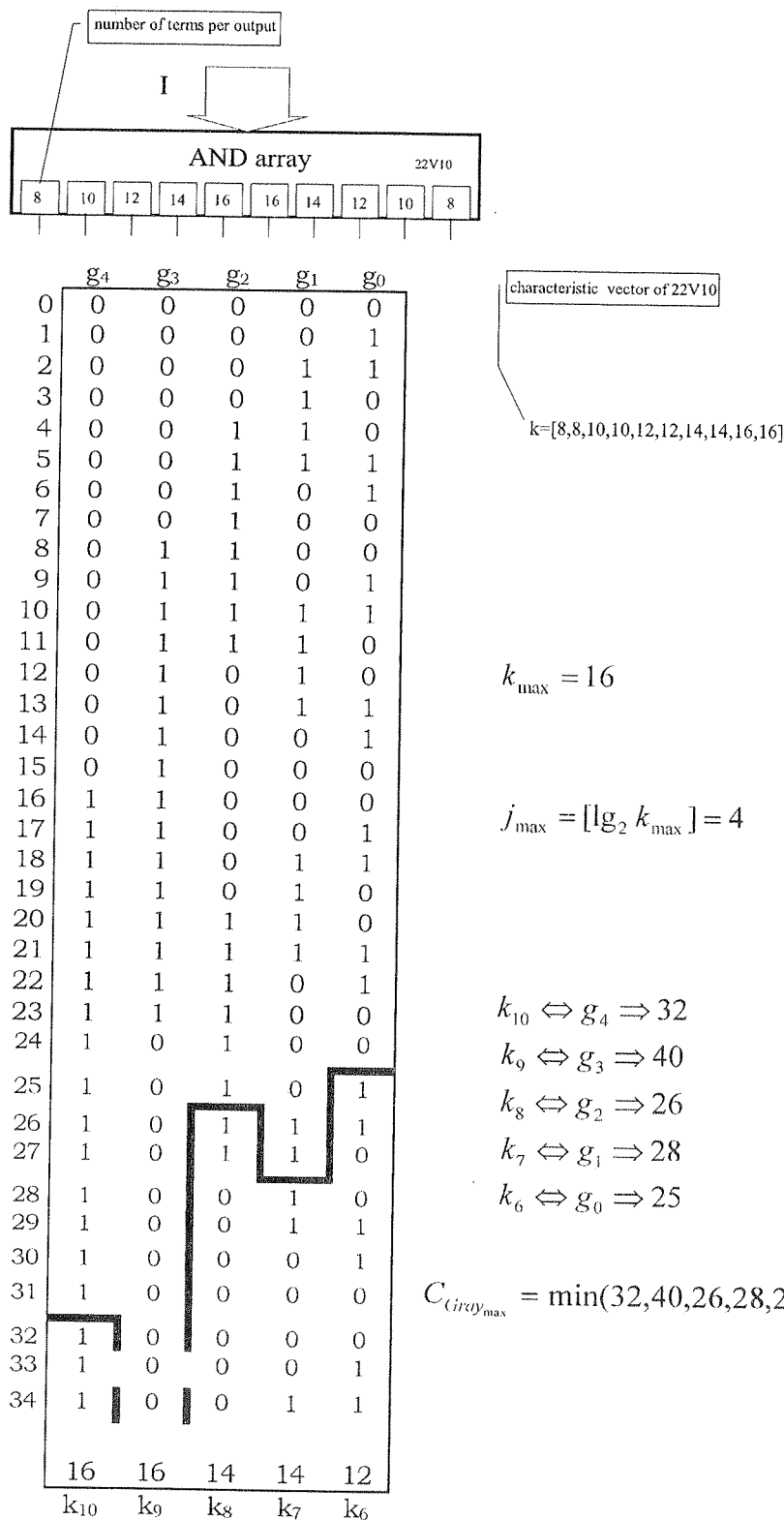


Figure 3. Result of optimal Gray coding for 22V10 transcoder

If $k = [k_1, k_2, \dots, k_m]$ is the ordered characteristic vector $k_1 \leq k_2 \leq \dots \leq k_m$, then optimal allocation is following

- if $k_{\max} = \max(k_1, k_2, \dots, k_m) = 2^i$ where $i \in N$ then

$$k_{m-i} \leftrightarrow g_{j_{\max}-i}; \quad i = 0, 1, \dots, j_{\max}; \quad j_{\max} = \lg_2 k_{\max}$$

- if $k_{\max} = \max(k_1, k_2, \dots, k_m) \neq 2^i$ where $i \in N$ then

$$k_{m-i} \leftrightarrow g_{(j_{\max}-1)-i}; \quad i = 0, 1, \dots, j_{\max} - 1$$

$$k_{m-j_{\max}} \leftrightarrow g_{j_{\max}}; \quad j_{\max} = [\lg_2 k_{\max}]$$

Result of optimal Gray coding for 22V10 structure is presented in Fig. 3.

5. AN EFFICIENT WORD CODING ALGORITHM

Efficient coding in PAL-based circuits consists in creating a maximum number of different code vectors in such a way that the code matrix meets the implementability condition (2). The whole coding process consists of several procedures

- ordering the characteristic vector
- determining the scope of uniform coding
- uniform coding
- supplementary coding
- supplementary expansion.

5.1. Ordering the characteristic vector

The procedure consists in ordering the coefficients k_i , ($1 \leq i \leq m$) of the characteristic vector to the form $k = [k_1, k_2, \dots, k_m]$, where $k_{i-1} \leq k_i \leq k_{i+1}$, ($2 \leq i \leq m-1$).

5.2. Determining the scope of uniform coding

The procedure consists in determining the coefficients p , where p is the greatest number for which the condition (9) is met.

$$\sum_{i=0}^{p-1} C_{m-1}^i \leq k_{\min}. \quad (9)$$

5.3. Uniform coding

The uniform coding procedure consists in creating words in such a way that the terms connected to the individual output cells are used uniformly. At successive stages of the procedure, groups of code words are created, in which every word activates 1 output, 2 outputs etc. If we assume a PAL-based transcoder described by the vector $k = [k_1, k_2, \dots, k_m]$, then the first group contains m different words activating one output, the second group contains $\binom{m}{2}$ different words activating two outputs etc.

This process can continue as long as the number of terms per output does not exceed $k_{\min} = \min(k_1, k_2, \dots, k_m)$. The individual code groups use $1, \binom{m-1}{1}, \binom{m-1}{2}, \binom{m-1}{3}$, etc terms connected to each output cell. Generally, p code groups can be coded under the uniform coding paradigm, where p is the greatest number for which the condition (9) is met.

5.4. Supplementary Coding

The supplementary coding procedure consists in creating code words which will activate i outputs, where $i > p$. The words are created one by one according to the following rules:

- I. The word being created must be different from any words created previously,
- II. The word being created must activate i outputs such that the sum of indexes of the individual output cells is maximum.
- III. From code words meeting condition II code word is chosen for which sum of $i-1$ indexes of output cells is maximum.
- IV. Verify if the set consisting of the code matrix Q (which contains the code words already created) and the new code word will meet the implementability condition for the given PAL structure described by its characteristic vector. If it does, then include the new code word in the matrix Q ; otherwise look for other words meeting the conditions I-IV, discarding those already included into the matrix Q and those which have already been considered and found only to satisfy the conditions I-III.

5.5. Supplementary expansion

Essentially, supplementary expansion means creating additional code words at the expense of those already created. Let us assume that the last code group to be created was " j of m " (each word in this group activates j of m outputs). The supplementary expansion procedure consists of two parts. First, expansion occurs within the most recently created code group " j of m ". The second step involves retrogressive expansion, where additional code words are created in the " j of m " group at the expense of words in the " i of m " group, where $(2 \leq i \leq j-1)$.

I shall now illustrate my algorithm by using it to determine the coding capacity of the 18V10 transcoder. The characteristic vector of this structure, when ordered, becomes $k = [8, 8, 8, 8, 8, 8, 8, 8, 10, 10]$. Knowing that $m = 10$ and $k_{\min} = 8$, we can use equation 9 to determine the scope of uniform coding. The maximum value of p for which $\sum_{i=0}^{p-1} \binom{9}{i} \leq 8$ is $p = 1$. Consequently, uniform coding will only generate the first group "1 of 10", of words activating one output each. At the next stage, i.e.

$Q =$

000
000
000
000
000
000
000
001
010
100
000
000
000
000
000
001
010
100
000
000
000
000
001
010
100
000
000
000
001
000
000
001
000
001
001
1100

$k' = 4400$

Figure 4. Th

ut does not
, $\binom{m-1}{2}$,
ups can be
r for which

which will
ding to the
reviously,
indexes of

a sum of i-1
code words
y condition
does, then
ther words
the matrix
satisfy the

e words at
group to be
puts). The
sion occurs
ep involves
e "j of m"
1).

ng capacity
en ordered,
= 8, we can
m value of

ly generate

xt stage, i.e.

Q'	0000000001	Uniform coding	0000000001	Q''
	0000000010		0000000010	
	0000000100		0000000100	
	0000001000		0000001000	
	0000010000		0000010000	
	0000100000		0000100000	
	0001000000		0001000000	
	0010000000		0010000000	
	0100000000		0100000000	
	1000000000		1000000000	
	0000000011		0000000011	
	0000000101		0000000101	
	0000001001		0000001001	
	0000010001		0000010001	
	0000100001		0000100001	
	0001000001		0001000001	
	0010000001		0010000001	
	0100000001		0100000001	
	1000000001		1000000001	
	0000000110		0000000110	
	0000001010	0000001010		
	0000010010	0000010010		
	0000100010	0000100010		
	0001000010	0001000010		
	0010000010	0010000010		
	0100000010	0100000010		
	1000000010	1000000010		
	0000001100	1000000100	expansion	
	0000010100	1000001000	expansion	
	0000100100	0100010000		
	0001000100	0100000100		
	0000011000	0000100100		
	0000101000	0001000100		
	0001001000	0010000100		
	0000011000	0000011000		
	0000101000	0100100000	expansion	
	0001001000	0100001000		
	0010001000	0001001000		
	0000110000	0010001000		
	0001010000	1000100000	expansion	
0010010000	1000010000			
0001100000	0001010000			
0010100000	0010010000			
0011000000	0001100000			
1100000000	0010100000			
	0011000000			
	1100000000			

k'= 4400000000

k''= 0000000000

Figure 4. The result of the coding with ($C_{\max}$) and without ($C_{u\&s}$) supplementary expansion procedure

supplementary coding, words are created in the code group "2 of 10". The first code vector which meets the conditions I–IV is [0000000011]; the next one — because of condition III — is [0000000101] etc. The final result of uniform and supplementary coding is the matrix Q' (Fig. 4).

Since no other code words exist that meet the conditions I–IV, the expansion procedure has to be used to create more words in the "2 of 10" group.

Let k' be a vector whose elements k'_i , ($i = 1, 2, \dots, m$) assume decimal values equal to the number of unused terms connected to the i -th output cell. After uniform and supplementary coding, $k' = [4, 4, 0, 0, 0, 0, 0, 0, 0]$. The expansion procedure, in this case, is limited to the "2 of 10" group. Some previously created words in the matrix Q' are withdrawn, and each of them is replaced with two new words, using the unused terms connected to the first and second output cell Q'' . The result of the effective coding algorithm with ($C_{\max}$) and without ($C_{u\&s}$) supplementary expansion procedure is presented in figure 4.

After the expansion procedure $k_1 \leq k_2 \leq \dots \leq k_m$, and no more code words can be created. The total coding capacity of the 18V10 transcoder is $C_{\max} = 48$. Coding capacity without supplementary expansion procedure is $C_{u\&s} = 44$.

6. EXPERIMENTAL RESULTS

The algorithm presented above can be used to determine the maximal coding capacity $C_{\max}$ of popular PAL-based transcoder with uneven numbers of terms connected to output cells. The results are presented in Table 1 and compared with the number of code words which can be obtained by others coding.

Coding capacity of selected PAL-based structures

Table 1

	$C_{bin\ min}$	$C_{bin\ max}$	$C_{Gray\ min}$	$C_{Gray\ max}$	$C_{u\&s}$	$C_{\max}$
18V10	17	20	16	18	44	48
22V10	17	28	16	25	61	61
26V12	17	28	16	25	72	72
29M16	17	28	16	25	96	97
32VX10	17	28	16	25	61	61

7. CONCLUSIONS

In the present paper the algorithm was described for PAL-based transcoders with uneven 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 programmable output type.

The efficient word coding algorithm presented above was implemented in a prototype program Decomp assisting the decomposition of combinational circuits.

In this a
some po
connecte
a somew
In such s
about th
library,
applied t
unused.

With
sequenti
MACH,

1. AMD,
2. R.L. A
Sympos
3. B. B a
Europe
4. M. B
Compa
5. S.D. B
Arrays.
6. H.A. C
Jersey,
7. F. D r
for Un
Brussel
8. D. K
with D
Network
9. D. K a
on Prog
10. D. K a
selected
11. D. K a
kacji, 1
12. T. Ł u
synthes
Kraków
13. M. R a
logic syn
ce on P
14. T. S a
drecht, 1
15. T. S a
Optimiz

In this article, the coding algorithm was used to determine the coding capacity of some popular PAL-based programmable structures with uneven number of terms connected to output cells. In many decomposition algorithms we are faced with a somewhat different problem: a given, precise number of words need to be coded. In such situations, the proposed algorithm can be used to prepare a library of data about the coding capacities of selected structure components. By analyzing this library, a suitable component can be chosen. Afterwards, my algorithm can be applied to create the required number of code words, leaving the other output cells unused.

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

REFERENCES

1. AMD, Altera, Lattice Data Book
2. R.L. Ashenhurst: *The decomposition of switching functions*. Proceedings of an International Symposium on the Theory of Switching, April 1957
3. B. Babba, M. Crastes, G. Saucier: *Input driven synthesis on PLDs and PGAs*. The European Conference on Design Automation, Brussels (Belgium), March 1992
4. M. Bolton: *Digital Systems Design with Programmable Logic*. Addison-Wesley Publishing Company, 1990, p.133-140
5. 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
6. H.A. Curtis: *The Design of switching Circuits*. D.van Nostrand Company, Inc., Princeton, New Jersey, Toronto, New York, 1962
7. 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
8. D. Kania, E. Hryniewicz, K. Pucher: *Coding Capacity of PAL-based Devices with Different Number Terms per Output*, XXI National Conference Circuits Theory and Electronic Networks, Kiekrz 1998, vol. 1, pp. 203-208
9. D. Kania: *Complex Decomposition of Multiple-output Functions*. Proc. International Conference on Programmable Devices and Systems, Ostrava 1996
10. D. Kania: *Table decomposition methods of digital circuits. Implementation of these circuits in selected PLD structures*. PhD Thesis, Silesian Technical University, Gliwice 1995
11. D. Kania: *Coding Capacity of Programmable Transcoder*, Kwartalnik Elektroniki i Telekomunikacji, 1998, vol. 44, z. 2, ss. 193-204
12. T. Łuba, H. Selveraj, A. Kraśniewski: *A new approach to FPGA-based logic synthesis*. Workshop on Design Methodologies for Microelectronics and Signal Processing, Gliwice-Kraków, 1993
13. M. Rawski, M. Nowicka, P. Tomaszewicz, T. Łuba: *Decomposition-based logic synthesis and its application in FPGA-oriented technology mapping*. Proc. International Conference on Programmable Devices and Systems, Ostrava 1996
14. T. Sasa o: *Logic Synthesis and Optimization*, Kluwer Academic Publishers, Boston/London/Dordrecht, 1993
15. T. Sasa o: *FPGA Design by Generalized Functional Decomposition, in Logic Synthesis and Optimization*, Kluwer Academic Publishers, Boston/London/Dordrecht, 1993

Table 1

$C_{\max}$
48
61
72
97
61

16. H. Selveraj, T. Łuba, M. Nowicka, B. Bignall: *Multiple-valued decomposition and its applications in data compression and technology mapping*. ICCIMA'97
17. F.A.M. Volf, *A bottom-up approach to multiple-level synthesis for look-up table based FPGAs*, CIP-Data Library Technische Iniversiteit Eindhoven 1997
18. 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 OPARTYCH NA
STRUKTURZE TYPU PAL Z RÓŻNĄ LICZBĄ TERMÓW
DOŁĄCZONYCH DO POSZCZEGÓLNYCH WYJŚĆ

Streszczenie

Realizacja układów cyfrowych wykorzystująca struktury programowalne jest ściśle związana z problemem dekompozycji. Struktury programowalne posiadają ograniczone zasoby wewnętrzne. Artykuł przedstawia efektywną metodę kodowania słów przeznaczoną do realizacji programowalnego transkodera opartego na strukturze typu PAL z różną liczbą termów dołączonych do poszczególnych wyjść. Zaproponowany, efektywny algorytm kodowania słów został wykorzystany do wyznaczenia pojemności kodowej popularnych programowalnych struktur. Wyniki zostały porównane z kodowaniem binarnym i kodowaniem kodem Gray'a. Zaprezentowana metoda wraz z innymi algorytmami może być wykorzystana do kodowania stanów automatów sekwencyjnych. Algorytm kodowania słów zaprezentowany w niniejszym artykule został zaimplementowany w systemie Decomp opracowanym w Politechnice Śląskiej.

Słowa kluczowe: dekompozycja, kodowanie, układy logiki programowalnej (PAL).

Wp:
trw

pier
weg
mał
neg
elek
Zas
no

Słow

Rury
nym, osi
1000 Gs.
znacznie
argonow
w rurze
wodnego
chłodzen
pracują
chłodzon
chłodzen

*) Pra

Wpływ pola magnetycznego pierścieniowych magnesów trwałych na wzrost mocy jonowego lasera argonowego

JERZY KĘSIK, PIOTR WARDA, WIESŁAW WOLIŃSKI

Instytut Mikroelektroniki i Optoelektroniki, Politechnika Warszawska

Otrzymano 1998.80.12

Autoryzowano 1999.02.15

Zastosowanie quasi-jednorodnego pola magnetycznego wytworzonego przez zespół pierścieniowych magnesów trwałych spowodowało znaczny wzrost mocy wyjściowej jonowego lasera argonowego. Korzystny wpływ pola magnesów jest silniejszy, zwłaszcza dla małych prądów wyładowania od jednorodnego, powszechnie stosowanego pola magnetycznego wytworzonego przez cewkę. Przyczyną tego zjawiska jest silny wzrost temperatury elektronowej plazmy wyładowania w obszarach objętych radialnym polem magnesów. Zastosowanie magnesów trwałych w laserach argonowych chłodzonych powietrzem powinno radykalnie zwiększyć ich moc wyjściową^{*)}.

Słowa kluczowe: Laser argonowy, magnesy trwałe.

1. WSTĘP

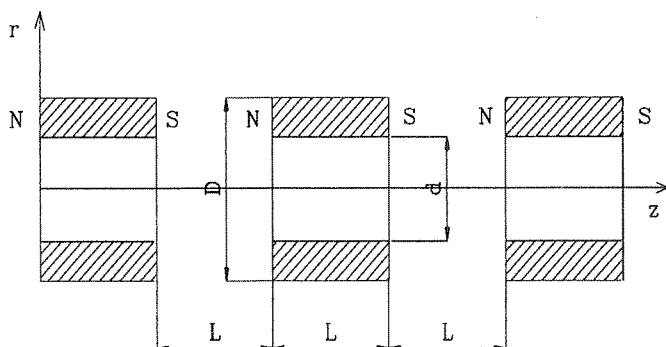
Rury wyładowcze jonowych laserów argonowych pracują zwykle w jednorodnym, osiowym polu magnetycznym wytworzonym przez cewkę o indukcji około 1000 Gs. Pole takie, poprzez zwiększenie koncentracji plazmy na osi wyładowania znacznie podnosi moc wyjściową i sprawność energetyczną tych laserów. W laserach argonowych średniej i dużej mocy (1–25 W), ze względu na dużą moc wydzielaną w rurze laserowej i cewce konieczne jest stosowanie intensywnego chłodzenia wodnego. W laserach małej mocy (10–500 mW) możliwe jest natomiast stosowanie chłodzenia powietrznego. Istotnym ograniczeniem mocy tych laserów jest fakt, że pracują one bez pola magnetycznego. Wynika to z trudności wykonania cewki chłodzonej powietrzem, wewnątrz której muszą zostać umieszczone duże radiatory chłodzenia powietrznego rury laserowej.

^{*)} Praca została wykonana w ramach projektu badawczego KBN nr 8T11B06512

Celem niniejszej pracy jest próba zastąpienia cewki zespołem magnesów trwałych o rozkładzie pola zbliżonym do wytwarzanego przez cewkę, umożliwiającą zastosowanie chłodzenia powietrznej rury laserowej i zbadanie wpływu pola magnetycznego tych magnesów na moc wyjściową lasera.

2. ROZKŁAD POLA MAGNETYCZNEGO ZESPOŁU MAGNESÓW TRWAŁYCH

Wytworzenie długiego, jednorodnego pola magnetycznego przez pojedynczy pierścieniowy magnes trwały jest praktycznie niemożliwe. Pole to, dla magnesów długości zbliżonej do długości kapilary wyładowczych rur laserowych zanika wewnątrz magnesu. Możliwe jest natomiast wytworzenie quasi-jednorodnego, osiowego pola magnetycznego przez zespół pierścieniowych magnesów trwałych, złożony z kilku magnesów, rozsuniętych między sobą o odstęp zbliżony do ich długości (rys. 1):



Rys. 1. Przykładowa konfiguracja zespołu trzech magnesów pierścieniowych dla uzyskania quasi-jednorodnego pola magnetycznego.

Rozkład indukcji B_z osiowego pola magnetycznego wzdłuż osi z i odległości r od tej osi, wytworzonego przez magnes trwały o kształcie walca średnicy D , długości L i pozostałości magnetycznej B_m przedstawia się zależnością [1]:

$$B_z(z, r, D) = \frac{B_m \cdot D}{8 \cdot \pi} \cdot \int_{-z}^{L-z} \int_0^{2\pi} \frac{[D^2/4 + r^2 + D \cdot r \cdot \cos \varphi]^{0.5}}{[z^2 + D^2/4 + r^2 + D \cdot r \cdot \cos \varphi]^{1.5}} \cdot d\varphi \cdot dz, \quad (1)$$

zaś składową radialną indukcji B_r opisuje zależność:

$$B_r(z, r, D) = \frac{B_m \cdot D}{8 \cdot \pi} \cdot \int_{-z}^{L-z} \int_0^{2\pi} \frac{[z \cdot \cos \varphi]}{[z^2 + D^2/4 + r^2 + D \cdot r \cdot \cos \varphi]^{1.5}} \cdot d\varphi \cdot dz. \quad (2)$$

Rys. 2. P
trwałych

Roz
D i śr
średnic

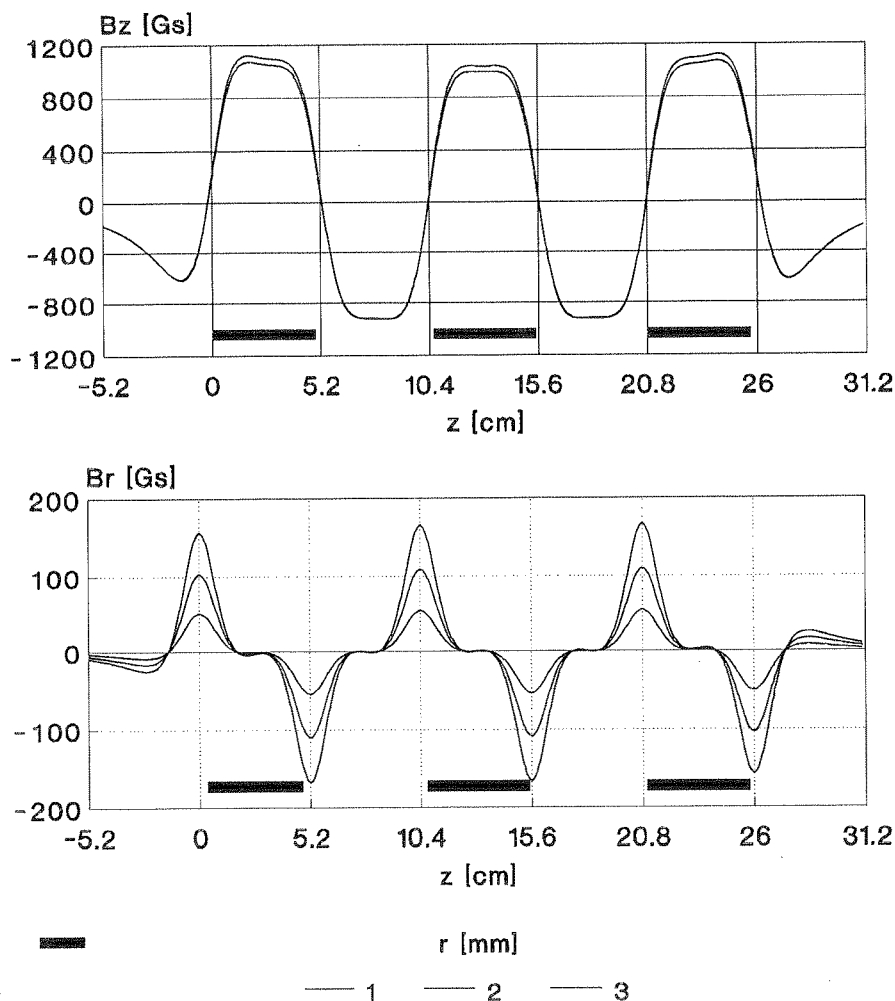
Ind
magnes

trwałych
ących za-
ola mag-

objedynczy
magnesów
nika we-
osiowe-
złożony
długości

a quasi-jed-

kości r od
długości



Rys. 2. Rozkład indukcji osiowego pola magnetycznego B_z , wytworzonego przez zespół trzech magnesów trwałych wzdłuż osi lasera z dla różnych odległości r od osi z . Obliczenia wykonano dla magnesów o rozmiarach $D = 48$ mm, $d = 34$ mm, $L = 52.5$ mm i $B_m = 10$ kGs

- (1) Rozkład pola magnetycznego magnesu pierścieniowego średnicy zewnętrznej D i średnicy wewnętrznej d stanowi superpozycję pola magnesu cylindrycznego średnicy D i pola magnesu cylindrycznego średnicy d :

$$B_z(z, r, D, d) = B_z(z, r, D) - B_z(z, r, d) \quad (3)$$

- (2) $B_r(z, r, D, d) = B_r(z, r, D) - B_r(z, r, d) \quad (4)$

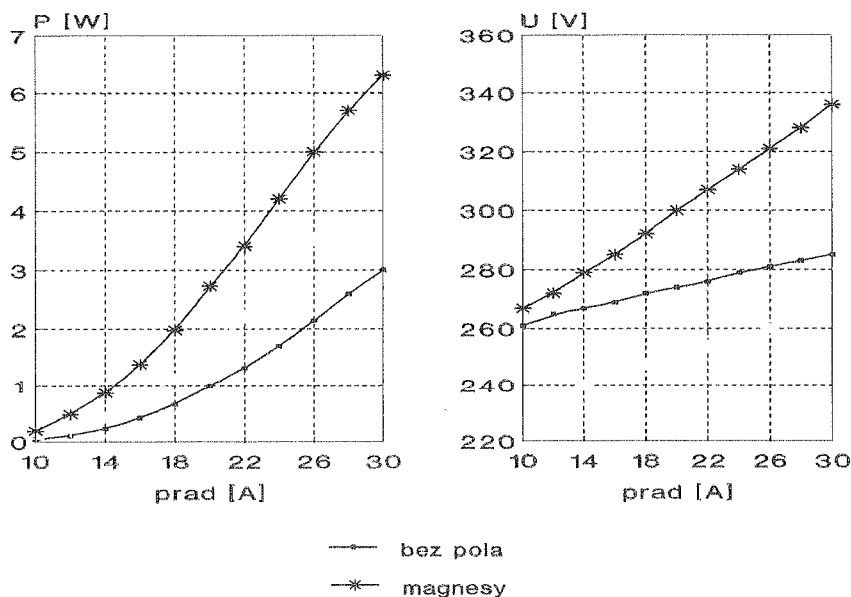
Indukcja B_z osiowego pola magnetycznego pochodzącego z takiej konfiguracji magnesów składa się z pięciu odcinków quasi-jednorodnego pola o długości równej

w przybliżeniu długości magnesu, a na krańcach magnesów gwałtownie zmienia znak na przeciwny. Indukcja B_z w niewielkim stopniu zależy od odległości od osi magnesu. Takie quasi-jednorodne pole magnetyczne, w porównaniu z jednorodnym polem wytworzonym przez cewkę, posiada niekorzystną cechę istnienia obszarów o zerowym natężeniu pola osiowego. Natomiast indukcja radialna B_r w tych obszarach osiąga wartości maksymalne, rosnące w przybliżeniu liniowo z odległością r od osi z .

3. PRZEBIEG PRAC DOŚWIADCZALNYCH

W eksperymentach użyto rurę laserową firmy Spectra Physics mod. 071, wykonaną z ceramiki berylowej. Rura laserowa wymagała zastosowania chłodzenia wodnego. Długość ośrodka aktywnego wynosiła 400 mm, zaś promień kapilary wyladowczej $R_k = 1.2$ mm.

Przeprowadzone doświadczenia polegały na porównaniu podstawowych parametrów lasera (moc wyjściowa i napięcie na rurze laserowej w funkcji parametrów pobudzenia i napełnienia) pracującego w klasycznym, jednorodnym polu magnetycznym cewki z parametrami lasera z magnesami trwałymi. Optymalną ze względu na moc wyjściową okazała się konfiguracja zespołu złożonego z 4-ch magnesów o długości L około 40 mm. Dokonano również pomiarów parametrów lasera pracującego bez pola magnetycznego. Praca taka jest charakterystyczna dla



Rys. 3. Zależność mocy wyjściowej lasera P i napięcia na rurze laserowej U od natężenia prądu wyladowania dla lasera pracującego na wszystkich liniach widmowych. Laser pracował przy optymalnym ciśnieniu argonu $p = 850$ mTr, i optymalnej indukcji pola cewki $B = 900$ Gs

laserów
lasera
przeds

Z p
trwały
nej cew
tem na
widocz
pól ma

Rys. 4. 2
wartości

Z p
wylado
Podobr
poprze
sowany

Pod
widmow
matu B

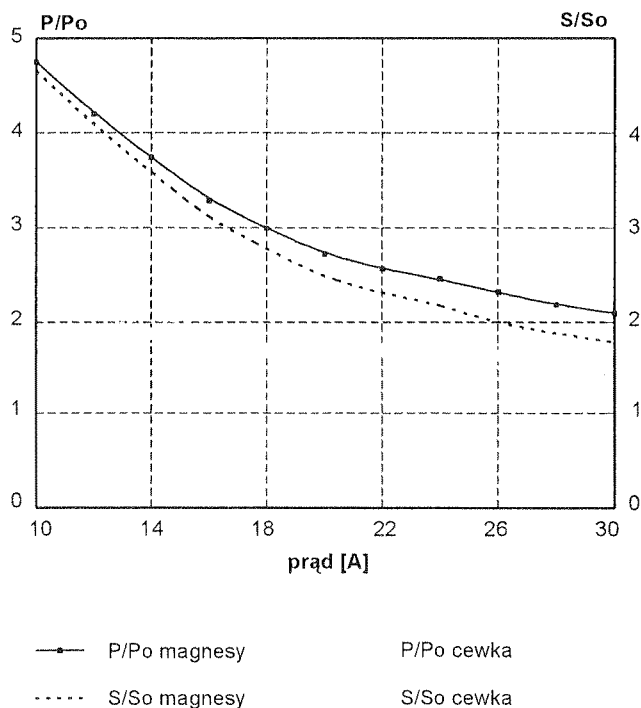
zmienia
od osi
rodnym
oszarów
w tych
ęgiłością

wyko-
odzenia
capilary

a para-
metrów
u mag-
alną ze
z 4-ch
metrów
zna dla

laserów argonowych chłodzonych powietrzem. Wyniki badań porównawczych dla lasera pracującego jednocześnie na wszystkich liniach widmowych (458–514 nm.) przedstawiono na rys. 3.

Z przeprowadzonych pomiarów wynika, że zastosowanie zespołu 4-ch magnesów trwałych znacznie podnosi moc wyjściową w porównaniu z zastosowaniem klasycznej cewki. Ten wzrost mocy okupiony jest jednak niekorzystnym, ok. 20% przyrostem napięcia zmniejszającym sprawność energetyczną lasera. Zjawisko to lepiej jest widoczne na rys. 4, na którym przedstawiono moc lasera dla obu porównywanych pól magnetycznych, zredukowaną do jej wartości bez pola magnetycznego.

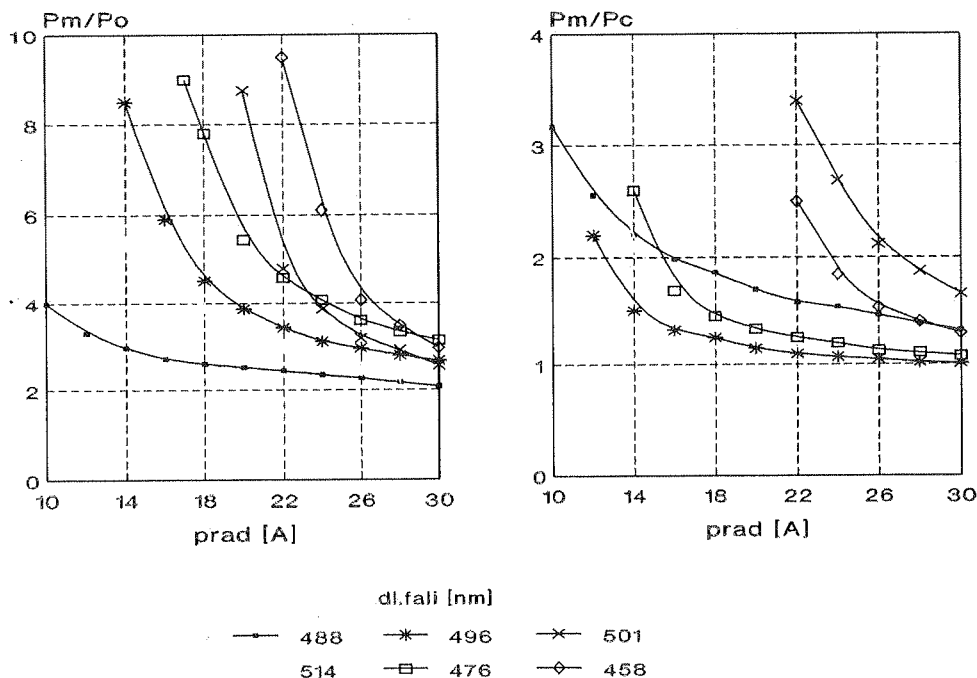


Rys. 4. Zależność mocy wyjściowej lasera P/P_o i sprawności energetycznej S/S_o zredukowanych do ich wartości bez pola magnetycznego od natężenia prądu wyładowania dla pola magnetycznego cewki i zespołu magnesów trwałych

Z przedstawionego rysunku wynika, że w całym zakresie zmienności prądu wyładowania moc lasera z magnesami znacznie przewyższa moc lasera z cewką. Podobny charakter ma przebieg sprawności energetycznej, mimo obserwowanego poprzednio niekorzystnego wzrostu napięcia, który jest z nadmiarem kompensowany przez wzrost mocy.

odowania
ciśnieniu

Podobne zależności uzyskano przy pracy lasera na poszczególnych liniach widmowych, rozdzielonych przy pomocy umieszczonego w rezonatorze lasera pryzmatu Brewstera (rys. 5).



Rys. 5. Zależność mocy wyjściowej lasera z zespołem magnesów trwałych P_m , zredukowana do wartości bez pola magnetycznego P_o oraz do wartości z cewką P_c od natężenia prądu wyładowania dla poszczególnych linii widmowych. Linie widmowe umieszczono w kolejności malejącego współczynnika wzmocnienia

Przedstawione na rys. 3–5 wyniki pomiarów wskazują, że wartość przewyższenia mocy lasera w polu magnesów trwałych maleje ze wzrostem prądu oraz, że jest ona większa dla linii widmowych o mniejszym współczynniku wzmocnienia. Obserwowany również wyraźny wzrost mocy lasera z magnesami w porównaniu z jego mocą w jednorodnym polu magnetycznym cewki sugeruje, że przyczyną tego zjawiska jest prawdopodobnie fakt istnienia w polu magnesów obszarów poprzecznego do osi wyładowania pola magnetycznego o symetrii radialnej (rys. 2). Jest to praktycznie jedyna cecha, która różni pole magnesów od jednorodnego pola cewki.

4. ANALIZA WPŁYWU RADIALNEGO POLA MAGNETYCZNEGO NA WZROST MOCY LASERA

Stosowane w jonowych laserach argonowych wysoko prądowe, nisko ciśnieniowe wyładowanie łukowe poddane jest działaniu pola magnetycznego. Jest rzeczą oczywistą, że poprzeczne pole magnetyczne o symetrii radialnej B_r , nie ma wpływu na ruch ładunków elektrycznych w radialnym polu elektrycznym E_r , a zatem na współczynnik dyfuzji elektronów D_e . W związku z tym zachowana zostaje wartość koncentracji elektronów n_e oraz jej radialna symetria rozkładu. Z drugiej jednak strony, radialne pole magnetyczne oddziałuje na elektrony poruszające się w osio-

wym polu elektrycznym E_z , prostopadłym do pola B_r . Spiralny ruch elektronów wzdłuż osi z powoduje wzrost ilości ich zderżeń z atomami na jednostkę długości z , a zatem zmniejszenie ruchliwości elektronów b_e wzdłuż osi z .

Wpływ pola magnetycznego na zmniejszenie ruchliwości elektronów w kierunku prostopadłym do jego kierunku dany jest prostym wyrażeniem [2]:

$$b_{eB} = \frac{b_{eo}}{1 + b_{eo}^2 \cdot B^2}, \quad (5)$$

gdzie: b_{eo} — ruchliwość elektronów bez pola magnetycznego,

b_{eB} — ruchliwość elektronów w obecności pola magnetycznego.

Zmodyfikowane wyrażenie Langevina określające wielkość ruchliwości elektronów b_{eo} , uwzględniające wpływ natężenia pola elektrycznego E_z na wartość ruchliwości dla sprężystych zderzeń elektronów z atomami ma postać [3]:

$$b_{eo} = \left[\frac{3}{2 \cdot \sqrt{\pi}} \right]^{0.25} \cdot \sqrt{\frac{e \cdot \lambda_e \cdot \sqrt{2 \cdot m_e / m_j}}{m_e \cdot E_z}}, \quad (6)$$

gdzie: m_e , m_j — masa elektronu i jonu

λ_e — średnia droga swobodna elektronu

Do obliczeń ruchliwości wykorzystano zależność doświadczalną natężenia pola elektrycznego E_z od prądu wyładowania I i ciśnienia napelnienia p :

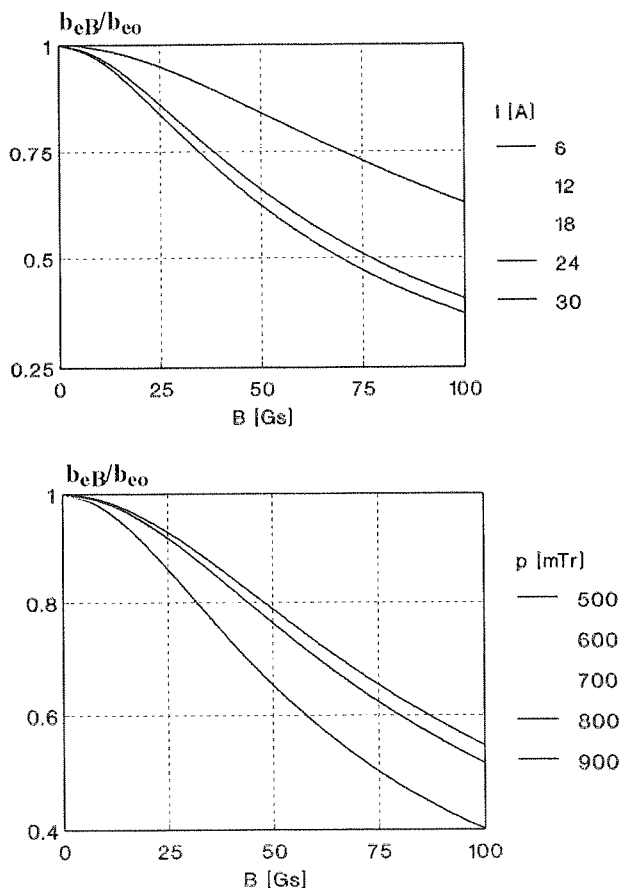
$$E_z(I, p) = (-0.076 \cdot p^3 + 0.15 \cdot p^2 - 0.062 \cdot p + 0.0055) \cdot I + 0.73 \cdot p + 2 \quad (7)$$

[V/cm, A, Tr]

Wpływ poprzecznego do osi wyładowania pola magnetycznego o symetrii radialnej na ruchliwość elektronów w wyładowaniu łukowym w argonie, obliczony z powyższych wyrażeń przedstawia rys. 6.

Z przeprowadzonych obliczeń wynika, że ruchliwość elektronów silnie maleje wraz ze wzrostem indukcji pola magnetycznego. Spadek ten jest większy dla małych ciśnień i dużych prądów wyładowania. Głównymi nośnikami prądu płynącego przez rurę wyładowczą są elektrony. Udział prądu jonowego jest niewielki i nie przekracza kilku procent wartości prądu elektronowego [3]. Można więc z dużą dokładnością przyjąć, że gęstość prądu płynącego przez rurę wyładowczą $j = e \cdot n_e \cdot b_e \cdot E_z$. Występujące poprzeczne pole magnetyczne zmniejsza znacznie wartość ruchliwości elektronów b_e . Dla zachowania tej samej, co bez pola magnetycznego wartości gęstości prądu, przy założonej stałej koncentracji elektronów n_e konieczny jest wzrost natężenia osiowego pola elektrycznego E_z , powodujący jednoczesny wzrost temperatury elektronowej T_e .

Dla określenia temperatury elektronowej niskociśnieniowego wyładowania łukowego w argonie przyjęto teorię swobodnego spadku jonu na ściankę Tonksa-Langmuira [4]. Wartość temperatury elektronowej T_e wyznaczono z warunku równowagi między ilością cząstek naładowanych wytwarzanych w procesie jonizacji zderzeniowej i traconych w wyniku procesów zubożniania na ściance rury wyładowczej.



Rys. 6. Zależność ruchliwości elektronów w kierunku osi wyładowania zredukowana do jej wartości bez pola magnetycznego b_{eB}/b_{e0} od indukcji pola magnetycznego B dla różnych wartości prądu wyładowania I . Obliczenia przeprowadzono dla ciśnienia argonu $p = 500$ mTr, prądu wyładowania $I = 20$ A i promienia kapilary wyładowczej $R_k = 1.2$ mm

$$(p_{ef} \cdot R_k)_{TL} \cdot \frac{Z}{p_{ef}} \cdot \left(\frac{m_j}{k \cdot T_e} \right)^{1/2} = 1.092, \quad (8)$$

gdzie: Z — szybkość jonizacji określona z równania produkcji jonów [5]:

$$Z = \frac{2 \cdot a \cdot m_j \cdot p_{ef}}{e \cdot \sqrt{\pi}} \cdot \left(\frac{2 \cdot k \cdot T_e}{m_j} \right)^{3/2} \cdot \exp\left(\frac{e \cdot V_i}{k \cdot T_e}\right) \cdot \left(1 + \frac{1}{2} \cdot \frac{e \cdot V_i}{k \cdot T_e}\right), \quad (9)$$

gdzie: a — stała (dla $Ar = 0.71$ [1/Vcm]) V_i — potencjał jonizacji (dla $Ar = 15.7$ V)

p_{ef} — efektywne ciśnienie Ar [Tr]

k — stała Boltzmanna.

W obliczeniach przyjęto [6] parametry plazmy typowe dla wyładowania łukowego lasera argonowego: temperaturę atomową T_a i efektywne ciśnienie argonu w kapilarze wyładowczej p_{ef} :

$$T_a = 300 \cdot (1 + 0.9 \cdot I \cdot R_k^{1.5}) \quad [\text{K}, \text{A}] \quad p_{ef} = p \cdot \frac{300}{T_a} \quad (10)$$

gdzie: p — ciśnienie napełnienia argonu [Tr].

Radialny rozkład potencjału w kapilarze wyładowczej wynosi wg Hernqvista [7]:

$$V(r/R_k) = -1.155 \cdot \frac{k \cdot T_e}{e} \cdot \left(1 - \sqrt{1 - (r/R_k)^2}\right) \quad (11)$$

i wynikający z niego ściśle wg prawa Boltzmanna rozkład koncentracji elektronów:

$$n_e(r/R_k) = n_{eo} \cdot \exp\left(\frac{e \cdot V(r/R_k)}{k \cdot T_e}\right). \quad (12)$$

Wpływ radialnego pola magnetycznego na decydujące o mocy lasera natężenie osiowego pola elektrycznego E_z wyznaczono z warunku stałości prądu wyładowania bez pola magnetycznego I_{eo} i z polem magnetycznym I_{eB} :

$$I_{eo} = 2 \cdot \pi \cdot e \cdot n_{eo} \cdot b_{eo} \cdot E_z \cdot \int_0^{R_k} n(r/R_k) \cdot r \cdot dr \quad (13)$$

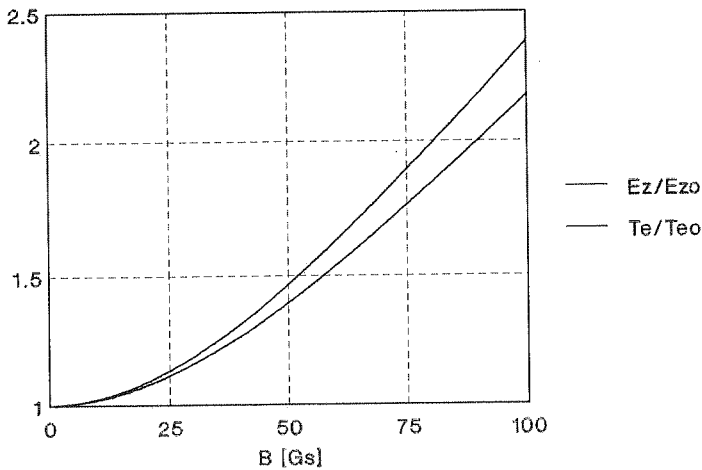
$$I_{eB} = 2 \cdot \pi \cdot e \cdot n_{eo} \cdot E_{zB} \cdot \int_0^{R_k} n(r/R_k) \cdot b_{eB}(r/R_k) \cdot r \cdot dr. \quad (14)$$

W wyrażeniach powyższych uwzględniono radialny rozkład koncentracji elektronów $n_e(\mathbf{r}/\mathbf{R}_k)$ oraz wynikający z radialnej, wg wzoru (4) zmienności pola magnetycznego rozkład ruchliwości elektronów $b_{eB}(\mathbf{r}/\mathbf{R}_k)$. W umieszczonych dalej wynikach obliczeń zmian parametrów plazmy od natężenia radialnego pola magnetycznego podano jego wartość na ścianie rury wyładowczej $\mathbf{r} = \mathbf{R}_k$. Natężenie osiowego pola elektrycznego E_z , związane bezpośrednio z napięciem na rurze laserowej U jest jednoznacznie określone przez temperaturę elektronową T_e [5]:

$$E_z(T_e) = 2 \cdot \sqrt{\frac{1.03 \cdot k \cdot T_e \cdot \left[V_j + \left(2.7 + \ln \sqrt{\frac{m_j}{2 \cdot m_e}} \right) \cdot \frac{k \cdot T_e}{e} \right]}{e \cdot \lambda_e \cdot R_k \cdot \sqrt{\pi}}}. \quad (15)$$

Numeryczne rozwiązanie układu równań (5–15) pozwoliło na określenie względnych zmian E_z i T_e od indukcji radialnego pola magnetycznego (rys. 7).

Wykonane obliczenia wskazują na silny wpływ radialnego pola magnetycznego na wzrost temperatury elektronowej i natężenia osiowego pola elektrycznego. Pole to, przy stosowanych wielkościach od 50 do 70 Gs powoduje przyrost temperatury



Rys. 7. Względna zmiana natężenia osiowego pola elektrycznego E_z/E_0 i temperatury elektronowej T_e/T_0 od indukcji radialnego pola magnetycznego B . Obliczenia wykonano dla prądu wyładowania $I = 15$ A, ciśnienia argonu $p = 500$ mTr i promienia kapilary wyładowczej $R_k = 1.2$ mm

elektronowej odpowiednio od 40 do 80%. Wyniki obliczeń zależności zredukowanych wartości natężenia pola elektrycznego od prądu wyładowania i ciśnienia badanej rury laserowej oraz porównanie ich z wartościami eksperymentalnymi przedstawiono na rys. 8.

Dobra zgodność jakościowa i ilościowa doświadczalnych i teoretycznych zależności zmian natężenia osiowego pola elektrycznego od prądu wyładowania i ciśnienia wskazują, że przyjęty model wpływu radialnego pola magnetycznego na ruchliwość elektronów ma decydujący wpływ na parametry plazmy wyładowania łukowego w laserze. Analizę wpływu wzrostu temperatury elektronowej na moc wyjściową lasera przedstawiono poniżej.

Głównym procesem prowadzącym do obsadzenia górnego i dolnego poziomu laserowego jest proces dwustopniowego pobudzenia elektronowego. W wyniku tego procesu moc lasera jest w przybliżeniu proporcjonalna do iloczynu gęstości elektronów n_e i jonów n_j oraz do szybkości pobudzenia górnych poziomów laserowych $\langle \sigma \cdot v_e \rangle$ [8]:

$$P_L \propto n_e \cdot n_j \cdot \langle \sigma \cdot v_e \rangle, \quad (16)$$

gdzie: σ — przekrój czynny na wzbudzenie poziomu laserowego,
 v_e — prędkość elektronów.

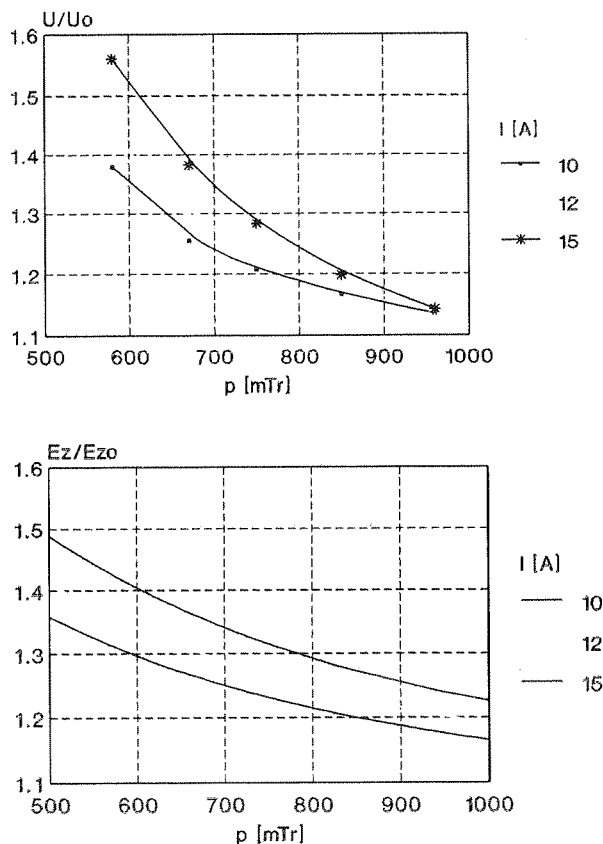
Przy określaniu względnych (w polu magnesów do bez pola) zmian mocy lasera gęstość ładunków nie jest istotna (zgodnie z przeprowadzonym poprzednio rozumowaniem pole radialne nie zmienia koncentracji elektronów). Po przyjęciu kolejnego założenia upraszczającego, że szybkość pobudzenia jest proporcjonalna do ilości elektronów o energii większej od potencjału pobudzenia poziomu energetycznego

Rys. 8
(dane
różne)

jonu
pobu

gdzie

R
fali
Otr
(4p²



Rys. 8. Zależność zredukowanych do wartości bez pola magnetycznego, napięcia na rurze laserowej U/U_0 (dane doświadczalne) i natężenia osiowego pola elektrycznego E_z/E_{z0} (wartości obliczone) od ciśnienia p dla różnych wartości prądu wyładowania I . Badania i obliczenia wykonano dla kapilary wyładowczej o $R_k = 0.9$ mm

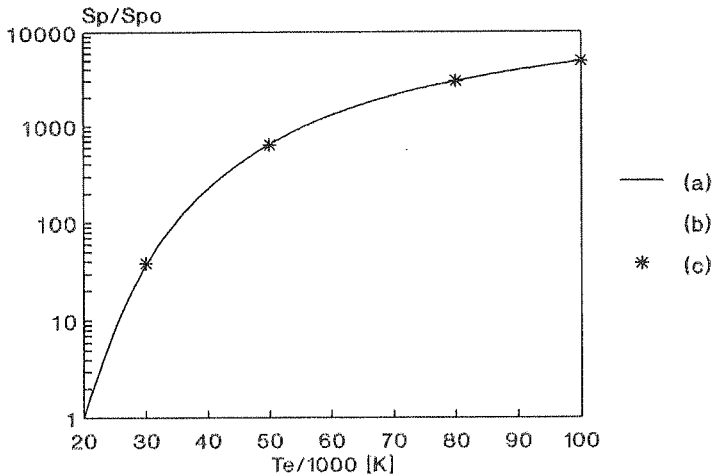
jonu wg wzoru Maxwella (9), zmiana mocy lasera zależy od zmian szybkości pobudzania poziomów laserowych.

$$\sigma \cdot v_e(T_e) > \propto \frac{4}{\sqrt{\pi}} \cdot \int_Y^{\infty} x^2 \cdot \exp(-x^2) \cdot \sigma(x) \cdot dx, \quad (17)$$

$$\text{gdzie: } Y = \sqrt{\frac{V_{wz} \cdot e}{k \cdot T_e}},$$

V_{wz} — potencjał wzbudzenia poziomu energetycznego

Rozważania przeprowadzono dla najsilniejszej linii lasera argonowego o długości fali 488 nm mającej największy udział procentowy w mocy wyjściowej lasera.. Otrzymano przedstawioną na rys. 9 zależność szybkości pobudzania górnego ($4p^2 D_{5/2}^0$) i dolnego ($4s^2 P_{3/2}$) poziomu laserowego z poziomu podstawowego jonu



Rys. 9. Zależność szybkości pobudzenia S_p/S_{po} górnych poziomów laserowych (a) i dolnych poziomów laserowych (b) od temperatury elektronowej T_e . Normalizacji dokonano dla szybkości pobudzenia przy $T_e = 20000$ K. Wartości obliczone wg (17) i podane przez Kitajewą (c)

($3p^5 P_{3/2}^0$) od temperatury elektronowej. Wartości przekrojów czynnych na pobudzenie przyjęto z pracy [10].

Jak wynika z przedstawionego wykresu, zgodność obliczeń z danymi literaturowymi jest dobra. Zależność inwersji obsadzeń od temperatury elektronowej można więc zapisać:

$$\Delta N(T_e) \propto \langle \sigma_a \cdot v_e(T_e) \rangle_a \cdot \tau_a - \langle \sigma_b \cdot v_e(T_e) \rangle_b \cdot \tau_b \cdot \frac{g_a}{g_b}, \quad (18)$$

gdzie: τ_a, τ_b — czasy życia górnego i dolnego poziomu laserowego,
 g_a, g_b — wagi statystyczne górnego i dolnego poziomu laserowego.

Zmiana współczynnika wzmocnienia k_B w obecności radialnego pola magnetycznego w porównaniu do współczynnika wzmocnienia k_o bez pola magnetycznego, ze względu na ich ścisłą proporcjonalność do wartości inwersji wynosi:

$$k_B = k_o \cdot \frac{\Delta N(T_{eB})}{\Delta N(T_e)}. \quad (19)$$

Wielkość i zmianę w funkcji prądu wyładowania współczynnika nienasyconego wzmocnienia k_o dla lasera pracującego bez pola magnetycznego na linii 488 nm oszacowano eksperymentalnie [1]:

$$k_o(I) \cdot L = 2.8 \cdot I - 25.1 \quad [\%, A].$$

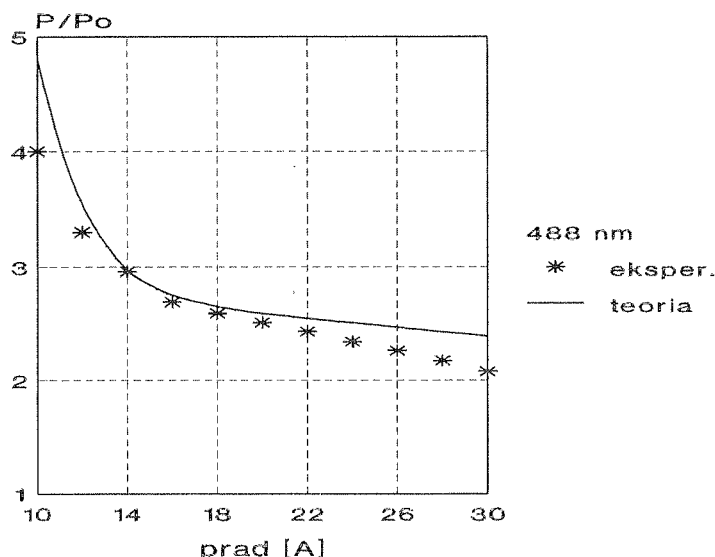
Wzrost współczynnika wzmocnienia w radialnym polu magnetycznym ma oczywisty wpływ na zwiększenie mocy lasera. W celu oszacowania wielkości tego wpływu

Rys. 10
wyjście

wykon
wielon
nienia
widmo

gdzie:

Poszcz
 Δv_D Δ
 Planck
 górneg
 R — st



Rys. 10. Zależność mocy wyjściowej lasera z zespołem magnesów trwałych P/P_0 , znormalizowana do mocy wyjściowej lasera pracującego bez pola magnetycznego w funkcji prądu wyładowania I obliczona wg 21. Przedstawiono porównanie z wynikami eksperymentów

wykorzystano elementarną teorię Rigroda [11], określającą moc lasera pracującego wielomodowo, przy zbliżonym do niejednorodnego mechanizmie nasycenia wzmocnienia. Parametr s określa stopień jednorodności poszerzonej niejednorodnie linii widmowej.

$$P_L = \frac{8 \cdot \pi \cdot h \cdot c}{\lambda^3} \cdot \frac{\Delta \nu_D}{\sqrt{\ln 2}} \cdot \frac{\Gamma^{-1}}{A_{ob}} \cdot \Phi(s \cdot X) \cdot T, \quad (20)$$

gdzie:

$$\Phi(s \cdot X) = \frac{\sqrt{\pi}}{2} \cdot s \cdot X \cdot \operatorname{erf}(\sqrt{\ln(s \cdot X)}), \quad s \approx 1 + 2 \cdot \sqrt{\frac{\ln 2}{\pi}} \cdot \frac{\Delta \nu_L}{\Delta \nu_D},$$

$$\Gamma = \frac{\left[\frac{g_a}{g_b} \cdot (\gamma_a - \gamma_{ab}) + \gamma_a \right]}{\gamma_a \cdot \gamma_b}, \quad \Delta \nu_D = \frac{2}{\lambda} \cdot \sqrt{\frac{2 \cdot R \cdot T_j \cdot \ln 2}{m_j}}.$$

Poszczególne symbole oznaczają:

$\Delta \nu_D$, $\Delta \nu_L$ — połówkowa szerokość dopplerowska i lorentzowska linii, h — stała Plancka, c — prędkość światła, X — parametr pobudzenia, γ_a , γ_b — stałe zaniku górnego i dolnego poziomu laserowego, T — transmisja zwierciadła wyjściowego, R — stała gazowa

W konsekwencji, względne zmiany mocy lasera P_B/P_o wywołane radialnym polem magnetycznym (rys. 10) przedstawia prosta zależność:

$$\frac{P_B}{P_o} = \frac{\Phi(s \cdot X_B)}{\Phi(s \cdot X_o)} \quad \text{gdzie: } X_B = \frac{k_B}{a}, \quad a \text{ — straty rezonatora.} \quad (21)$$

Wykonane obliczenia, mimo przyjętych założeń upraszczających wykazały dobrą zgodność z doświadczeniem. Tej dobrej zgodności sprzyjał fakt, że stanowiły one obliczenia porównawcze. Obliczenie wartości bezwzględnych mocy wyjściowej stanowi oczywiście dużo większy problem. Przeprowadzone obliczenia wykazały jednak wystarczającą zgodność w określeniu ilościowego wpływu pola magnesów na zwiększenie mocy lasera oraz wskazały wyraźnie na spadek przewyższenia mocy w miarę wzrostu prądu wyładowania.

Reasumując, dobra zgodność jakościowa i ilościowa obliczonych i doświadczalnych parametrów lasera (względne zmiany mocy i napięcia w funkcji prądu wyładowania i ciśnienia argonu) wskazuje na decydującą rolę zmian parametrów plazmy w radialnym polu magnetycznym. Można również z dużym prawdopodobieństwem stwierdzić, że przyjęty model wpływu tego pola na zmniejszenie ruchliwości elektronów i zwiększenie temperatury elektronowej jest prawidłowy i dominujący.

WNIOSKI

Obserwowany w przedstawionych wynikach eksperymentów silny wzrost mocy wyjściowej lasera i jego sprawności energetycznej w polu magnetycznym magnesów trwałych ma niewątpliwie duże znaczenie praktyczne. Jest on szczególnie istotny dla laserów argonowych chłodzonych powietrzem, które dotychczas pracują bez pola magnetycznego. Segmentowa konstrukcja proponowanego zespołu magnesów umożliwia umieszczenie go w radiatorze chłodzenia powietrznego [1], nie zmniejszając w zasadniczy sposób jego efektywności. W ten sposób można znacznie, co najmniej dwukrotnie zwiększyć moc wyjściową seryjnie produkowanych laserów argonowych chłodzonych powietrzem.

BIBLIOGRAFIA

1. P. Warda: *Wpływ pola magnetycznego pierścieniowych magnesów trwałych na parametry jonowego lasera argonowego*. Rozprawa doktorska, Politechnika Warszawska, Warszawa 1997
2. M.F. H o y a x: *The Low-Pressure Arc Positive Column in a Longitudinal Magnetic Field*. Am. J. Phys., 1968 36, No. 8, 726
3. W.I. G a p o n o w: *Elektronika*. PWN Warszawa 1965
4. I. L a g m u i r, L. T o n k s: *A General Theory of the Plasma of an Arc*. Phys. Rev., 1929, 34, 876
5. A. v o n E n g e l: *Ionized Gases*. Oxford University Press, London 1965
6. W.B. B r i d g e s, A.N. C h e s t e r, A.S. H a l s t e d, J.V. P a r k e r: *Ion Laser Plasmas*, Proc. of the IEEE, 1971, 59, No. 5, 724
7. K.G. H e r n q v i s t: *Effects of Transverse Magnetic Field for He-Cd Laser*. J. Appl. Phys., 1969, 40

8. K. Banse, H.R. Lüthi, W.H. Seelig: *Der Einfluss Axialer Magnetfelder auf den ArII-Hochleistungslaser*. Appl. Phys. 1974, 4, 141
9. A.E. Fotiadi, S.A. Fridrichow: *Wlijanije prodolnogo magnitnogo pola na izluczenie nieprerывnogo argonovogo lasera s oknami Briustera*. Zur. Techn. Fiz., 1969, 39, No. 9
10. P.W. Fielcan, M.M. Powcz: *Wosbuzhdenie jonow argona pri elektronno-atomnyh stolknovieniah*. Optika i Spektroskopia, 1970, 28, No. 2, 217
11. P.W. R i g r o d: *Gain Saturation and Output Power of Optical Masers*. J. Appl. Phys., 1963, 34, No. 9, 2602

J. KĘSIK, P. WARDA, W. WOLIŃSKI

POWER ENHANCEMENT OF ARGON ION LASER BY RING PERMANENT MAGNETS FIELD

S u m m a r y

A positive effect of application of permanent magnets field in argon ion laser has been observed. Some ring magnets producing quasi-homogeneous axial magnetic field increase several times the output power in particular for a small discharge current. This influence is stronger with respect to common magnet coil application. The reason of this phenomenon is caused by electron temperature increasing in radial magnetic field areas. Application of permanent magnets field should substantially enhance the output power of air cooled argon lasers.

Key words: argon laser, permanent magnets.

T
sowe
anal
meto
wej),
czny
omó

J
prote
uzysł
dwul
mian

Projektowanie optymalnych dolnoprzepustowych filtrów NOI o różnych stopniach licznika i mianownika transmitancji

FELICJA WYSOCKA-SCHILLAK

Zakład Elektroniki i Informatyki, Wyższa Szkoła Pedagogiczna, Bydgoszcz

Otrzymano 1998.11.23

Autoryzowano 1999.03.20

W pracy przedstawiono metodę projektowania dolnoprzepustowych filtrów NOI o różnych stopniach licznika i mianownika transmitancji i czebyszewowskim przebiegu charakterystyki amplitudowej. W metodzie tej zadanie zaprojektowania optymalnego filtra jest przekształcane w odpowiednie zadania programowania nieliniowego z ograniczeniami, które są następnie rozwiązywane metodą przesuwanej funkcji kary. Praca zawiera ponadto opis opracowanego programu FNOI oraz dwa przykłady obliczeniowe ilustrujące działanie tego programu.

Słowa kluczowe: filtry cyfrowe, filtry NOI, filtry optymalne, optymalizacja

1. WSTĘP

Tradycyjne metody projektowania filtrów o nieskończonej odpowiedzi impulsowej (w skrócie NOI) sprowadzają się do przekształcenia prototypowego filtra analogowego w odpowiedni filtr cyfrowy. Najczęściej stosowanymi tradycyjnymi metodami projektowania filtrów NOI są: metoda transformacji dwuliniowej (biliniowej), metoda niezmienności odpowiedzi impulsowej oraz metody oparte na numerycznym rozwiązywaniu równań różniczkowych. Wymienione metody są szczegółowo omówione m.in. w pracach [1–3].

Jeżeli chcemy uzyskać filtr cyfrowy NOI o dużej selektywności, odpowiednim prototypem analogowym jest filtr eliptyczny. Cyfrowy filtr eliptyczny NOI można uzyskać z prototypowego analogowego filtra eliptycznego metodą transformacji dwuliniowej. Transmitancja filtra eliptycznego NOI jest ilorzem dwóch wielomianów stopnia N zmiennej z^{-1} , gdzie N jest rzędem filtra.

Charakterystyka amplitudowa filtra eliptycznego ma przebieg równomiernie falisty o liczbie ekstremów równej $N + 1$ zarówno w pasmie przepustowym jak i w zaporowym. Charakterystyka ta stanowi najlepszą aproksymację w sensie Czebyszewa (lub minimaksu) idealnej charakterystyki amplitudowej filtra. Zastosowanie aproksymacji w sensie Czebyszewa umożliwia uzyskanie charakterystyki amplitudowej filtra o najbardziej stromym zboczu. Oznacza to, że dla filtra danego rzędu N , zadanej częstotliwości krańcowej pasma przepustowego oraz zadanych maksymalnych odchyleniach charakterystyki amplitudowej od charakterystyki idealnej w pasmach przepustowym i zaporowym, szerokość pasma przejściowego jest najmniejsza. Filtry zaprojektowane z zastosowaniem aproksymacji w sensie czebyszewowskim nazywane są w literaturze filtrami optymalnymi.

W praktyce zastosowanie filtrów eliptycznych nie zawsze jest jednak najlepszym rozwiązaniem. W przypadku tych filtrów zdarza się bowiem, że filtr dokładnie spełnia wymagania dotyczące przebiegu charakterystyki amplitudowej w pasmie przepustowym, natomiast odpowiednie wymagania w pasmie zaporowym spełnione są z bardzo dużym zapasem, lub też odwrotnie. W tego rodzaju sytuacjach bardziej efektywnym rozwiązaniem jest zastosowanie filtra o transmitancji posiadającej licznik i mianownik o różnych stopniach.

Przyjęcie różnych stopni licznika i mianownika transmitancji jest też korzystnym rozwiązaniem przy projektowaniu filtrów wąskopasmowych i szerokopasmowych. W przypadku filtrów wąskopasmowych wygodnie jest przyjąć wyższy stopień mianownika i zmniejszyć stopień licznika. W przypadku filtrów szerokopasmowych sytuacja jest dokładnie odwrotna, tzn. należy zmniejszyć stopień mianownika, a zwiększyć stopień licznika [4, 5]. Biorąc z kolei pod uwagę przebieg charakterystyki opóźnienia grupowego oraz skutki skończonej długości słowa maszynowego, należy dążyć do tego, aby liczba biegunów była możliwie mała [4].

Filtrów optymalnych o różnych stopniach licznika i mianownika transmitancji nie da się jednak zaprojektować metodą transformacji dwuliniowej. W literaturze opisane są różne metody numeryczne projektowania tego rodzaju filtrów [4–15]. Część z tych metod oparta jest na zmodyfikowanym algorytmie Remeza, w większości pozostałych wykorzystuje się natomiast programowanie liniowe lub nieliniowe. Metody prezentowane w literaturze są z reguły stosunkowo skomplikowane, a ich zastosowanie jest niekiedy dość ograniczone. Metody te też nie zawsze są bieżne [4, 6, 10]. W przypadku filtrów NOI nie ma jednej, ogólnie uznanej i stosowanej metody projektowania, takiej jak metoda Parksá-McClellana w przypadku filtrów o skończonej odpowiedzi impulsowej (w skrócie SOI).

W niniejszej pracy zaproponowano metodę projektowania optymalnych filtrów cyfrowych NOI o różnych stopniach licznika i mianownika transmitancji. Zadanie zaprojektowania takiego filtra jest przekształcane w odpowiednie zadanie programowania nieliniowego z ograniczeniami, które jest następnie rozwiązywane metodą przesuwanej funkcji kary. Praca zawiera również opis opracowanego programu FNOI oraz przykłady obliczeniowe.

2. SFORMUŁOWANIE PROBLEMU

W ogólnym przypadku transmitancja $H(z)$ filtru NOI może być przedstawiona w postaci [2]:

$$H(z) = \frac{W_1(z)}{W_2(z)} = \frac{a_0 + a_1 z^{-1} + a_2 z^{-2} + \dots + a_N z^{-N}}{1 + b_1 z^{-1} + b_2 z^{-2} + \dots + b_M z^{-M}}, \quad (1)$$

gdzie: a_k, b_j — współczynniki rzeczywiste.

Transmitancję $H(z)$ można również wyrazić następująco [5]:

$$H(z) = C z^{M-N} \frac{\prod_{i=1}^N (z - c_i)}{\prod_{i=1}^M (z - d_i)}. \quad (2)$$

Aby współczynniki a_k i b_j występujące we wzorze (1) były rzeczywiste, współczynnik C musi być liczbą rzeczywistą. Ponadto zera c_i oraz bieguny d_i muszą być rzeczywiste, lub też muszą występować w postaci par sprzężonych liczb zespolonych. Jeżeli wszystkie bieguny d_i znajdują się wewnątrz okręgu jednostkowego, czyli w obszarze $|z| < 1$, to filtr o transmitancji $H(z)$ jest filtrem stabilnym.

Bieguny i zera występujące w punkcie $z = 0$ nie mają wpływu na charakterystykę amplitudową, wnoszą natomiast liniową składową do charakterystyki fazowej [3].

Transmitancja określona na okręgu jednostkowym (tzn. dla $|z| = 1$), czyli funkcja

$$H(e^{j\omega}) = H(z)|_{z=e^{j\omega}}, \quad (3)$$

nazywana jest charakterystyką częstotliwościową lub amplitudowo-fazową układu. Funkcja $H(e^{j\omega})$ jest funkcją zespoloną. Charakterystyka amplitudowa $A(\omega)$ oraz fazowa $\varphi(\omega)$ układu są wyrażone następująco:

$$A(\omega) = |H(e^{j\omega})|, \quad (4)$$

$$\varphi(\omega) = \arg H(e^{j\omega}). \quad (5)$$

W przypadku filtru dolnoprzepustowego aproksymowana charakterystyka amplitudowa $A_d(\omega)$ filtru idealnego ma postać

$$A_d(\omega) = \begin{cases} 1, & 0 \leq \omega \leq \omega_p \\ 0, & \omega_s \leq \omega \leq \pi \end{cases}, \quad (6)$$

gdzie ω_p i ω_s są odpowiednio pulsacjami krańcowymi pasma przepustowego $P = [0, \omega_p]$ i zaporowego $S = [\omega_s, \pi]$.

Zadanie zaprojektowania optymalnego filtru cyfrowego NOI sprowadza się do takiego doboru współczynników a_k i b_j występujących we wzorze (1), lub co jest równoważne, takiego doboru zer c_i oraz biegunów d_i występujących we wzorze (2), aby otrzymana charakterystyka amplitudowa filtru była możliwie najlepszą aprok-

symacją w sensie Czebyszewa charakterystyki $A_d(\omega)$. W dalszej części pracy będziemy zajmować się poszukiwaniem zer c_i oraz biegunów d_i transmitancji $H(z)$ filtru.

Z znacznej większości metod projektowania opisywanych w literaturze zakłada się, że zera transmitancji leżą na okręgu jednostkowym, gdyż przy takim położeniu zer uzyskuje się maksymalne tłumienie w pasmie zaporowym. Takim rozkładem zer charakteryzują się również transmitancje cyfrowych filtrów eliptycznych, Butterwortha i Czebyszewa. W przypadku, gdy wszystkie zera transmitancji leżą na okręgu jednostkowym, współczynniki licznika transmitancji są symetryczne, w związku z czym implementacja filtru jest łatwiejsza [8].

Konsekwencją przyjęcia ograniczenia, że wszystkie zera transmitancji leżą na okręgu jednostkowym, jest jednak fakt, że dla danych wartości M i N oraz dla danych częstotliwości krańcowych pasma przepustowego i zaporowego istnieje pewna minimalna wartość amplitudy zafalowania w pasmie przepustowym $\delta_{p\min}$, którą jeszcze można uzyskać. Wartości $\delta_{p\min}$ odpowiada z kolei określona maksymalna wartość amplitudy zafalowania w pasmie zaporowym $\delta_{s\max}$ [5, 8].

W niektórych metodach projektowania dopuszcza się, oprócz zer transmitancji leżących na okręgu jednostkowym, również istnienie jednego lub kilku rzeczywistych zer transmitancji w przedziale $(0, 1)$ na płaszczyźnie z . Występowanie jednego zera tego rodzaju powoduje zmniejszenie wartości ekstremum charakterystyki amplitudowej występującego przy częstotliwości $f=0$ [4]. W wyniku wprowadzenia kilku zer tego rodzaju uzyskuje się natomiast zwiększenie liczby zafalowań charakterystyki amplitudowej w pasmie przepustowym oraz zmniejszenie wartości amplitud tych zafalowań poniżej $\delta_{p\min}$. W dalszych rozważaniach przyjmujemy, że oprócz zer transmitancji leżących na okręgu jednostkowym dopuszcza się istnienie jednego lub kilku rzeczywistych zer transmitancji w przedziale $(0, 1)$ na płaszczyźnie z .

Przyjęcie założenia, że zera transmitancji leżą na okręgu jednostkowym, czyli że $|c_i| = 1$, powoduje, że część rzeczywista i część urojona zespolonego zera związane są następującą zależnością

$$[\operatorname{Re}(c_i)]^2 + [\operatorname{Im}(c_i)]^2 = 1. \quad (7)$$

Przy takim założeniu, w celu wyznaczenia wartości zera wystarczy jedynie określenie jego części rzeczywistej, lub bez jego części urojonej. W dalszych rozważaniach będziemy wyznaczać część rzeczywistą zer transmitancji leżących na okręgu jednostkowym.

Oznaczmy przez $\mathbf{Y} = [y_1, y_2, \dots, y_R]^T$ — wektor o współrzędnych y_k zdefiniowanych następująco:

$$\begin{aligned} y_i &= \operatorname{Re}(c_i), & i &= 1, 2, \dots, N_1, \\ y_{N_1+i} &= c_i, & i &= 1, 2, \dots, N_2, \\ y_{N_1+N_2+i} &= \operatorname{Re}(d_i), & i &= 1, 2, \dots, M_1, \\ y_{N_1+N_2+M_1+i} &= \operatorname{Im}(d_i), & i &= 1, 2, \dots, M_1, \\ y_{N_1+N_2+2M_1+i} &= d_i, & i &= 1, 2, \dots, M_2, \\ y_{N_1+N_2+2M_1+M_2+1} &= C. \end{aligned}$$

gdzie N_1 jest liczbą par sprzężonych zer zespolonych, N_2 jest liczbą zer rzeczywistych, M_1 jest liczbą par sprzężonych biegunów zespolonych, M_2 jest liczbą biegunów rzeczywistych. Wektor $\mathbf{Y}$ może być również przedstawiony w następującej postaci dogodnej do dalszych rozważań

$$\mathbf{Y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_R \end{bmatrix} = \begin{bmatrix} \mathbf{Y}_1 \\ \mathbf{Y}_2 \end{bmatrix} = [\mathbf{Y}_1, \mathbf{Y}_2]^T, \quad (9)$$

gdzie: $\mathbf{Y}_1 = [y_1, y_2, \dots, y_{N_1+N_2}]^T$, $\mathbf{Y}_2 = [y_{N_1+N_2+1}, y_{N_1+N_2+2}, \dots, y_R]^T$,
 $R = N_1 + N_2 + 2M_1 + M_2 + 1$.

Należy zauważyć, że współrzędne tak określonego wektora $\mathbf{Y}_1$ są częściami rzeczywistymi zespolonych zer transmitancji oraz zerami rzeczywistymi. Współrzędne wektora $\mathbf{Y}_2$ są natomiast odpowiednio częściami rzeczywistymi i urojonymi zespolonych biegunów transmitancji oraz biegunami rzeczywistymi, a ponadto jedną ze współrzędnych wektora $\mathbf{Y}_2$ stanowi stała C .

Oznaczmy następnie przez $A(\omega, \mathbf{Y})$ charakterystykę amplitudową otrzymaną ze wzoru (4) przy przyjęciu biegunów i zer transmitancji o wartościach określonych przez wektor $\mathbf{Y}$. Funkcja błędu aproksymacji $E(\omega, \mathbf{Y})$ może być przedstawiona w postaci:

$$E(\omega, \mathbf{Y}) = [A(\omega, \mathbf{Y}) - A_d(\omega)] \quad (10)$$

Zadanie aproksymacji sformułujemy w sposób następujący¹⁾.

Dla zadanych wartości $N = 2N_1 + N_2$, $M = 2M_1 + M_2$, ω_p , ω_s oraz zadanej amplitudy zafalowań δ_1 charakterystyki $A(\omega, \mathbf{Y})$ w pasmie przepustowym, znaleźć wektor $\mathbf{Y} = [y_1, y_2, \dots, y_R]^T$, który minimalizuje wartość

$$\varepsilon = \max_{\omega \in P \cup S} |E(\omega, \mathbf{Y})|, \quad (11)$$

przy jednoczesnym spełnieniu warunków stabilności

$$\forall_{i=1, 2, \dots, M} |d_i| < 1. \quad (12)$$

3. PRZEKSZTAŁCENIE PROBLEMU W ZADANIE OPTYMALIZACJI

Rozpatrywane zadanie projektowania filtra można przekształcić w odpowiednie zadanie optymalizacji z ograniczeniami. W tym celu, przy założonych odpowiednich wartościach początkowych współrzędnych wektora $\mathbf{Y}$, podzielmy pasmo przepus-

¹⁾ Nie jest to jedyne możliwe sformułowanie zadania aproksymacji. Zamiast amplitudy zafalowań δ_1 można określić amplitudę zafalowań charakterystyki $A(\omega, \mathbf{Y})$ w pasmie zaporowym. Można też podać wymagane wartości zafalowań charakterystyki amplitudowej w pasmie przepustowym i zaporowym, a nie precyzować pulsacji krańcowej pasma przepustowego lub zaporowego. Przykłady tak sformułowanych zadań aproksymacji można znaleźć np. w pracach [8, 16].

towe $P = [0, \omega_p]$ i zaporowe $S = [\omega_s, \pi]$ na podprzedziały θ_j , $j = 1, 2, \dots, J + K$, zdefiniowane w taki sposób, że w każdym z tych podprzedziałów funkcja błędu ma dokładnie jedno ekstremum lokalne.

Filtry optymalne posiadają równomiernie falistą charakterystykę amplitudową o liczbie przerzutów funkcji błędu co najmniej równej $M + 1$ w pasmie przepustowym i $N + 1$ w pasmie zaporowym [4]. W związku z powyższym, w większości przypadków charakterystyka amplitudowa ma również $M + 1$ ekstremów w pasmie przepustowym i $N + 1$ ekstremów w pasmie zaporowym. Taka liczba ekstremów występuje zawsze w przypadku, gdy zarówno M jak i N są liczbami nieparzystymi.

W przypadku, gdy M jest liczbą nieparzystą, N jest liczbą parzystą i $M > N$, może pojawić się dodatkowe minimum przy częstotliwości $f = 0,5$, które nie jest związane z przerzutem funkcji błędu. Jeżeli natomiast M jest liczbą parzystą, N jest liczbą nieparzystą i $M < N$, może z kolei wystąpić dodatkowe maksimum przy częstotliwości $f = 0$. W bardzo rzadkich przypadkach, wartość tego maksimum jest taka sama, jak wartość pozostałych maksimów występujących w pasmie przepustowym. Mamy wówczas do czynienia z przypadkiem $M + 2$ przerzutów funkcji błędu w pasmie przepustowym. Tego typu rozwiązanie jest nazywane w literaturze rozwiązaniem *ponadrownomiernie falistym* [4].

Jeżeli M i N są liczbami parzystymi i $M > N$, to może wystąpić dodatkowe minimum przy częstotliwości $f = 0,5$, które nie jest związane z przerzutem funkcji błędu. W przypadku, gdy $M < N$ może natomiast pojawić się dodatkowe maksimum przy częstotliwości $f = 0$, a ponadto istnieje możliwość uzyskania rozwiązania *ponadrownomiernie falistego* [4].

Niech przedziały θ_j , $j = 1, 2, \dots, J + K$, będą określone w sposób następujący:

— w pasmie przepustowym

$$\begin{aligned}\theta_1 &= [0, \omega_1), \\ \theta_j &= [\omega_{j-1}, \omega_j), \quad j = 2, 3, \dots, J - 1, \\ \theta_J &= [\omega_{J-1}, \omega_p),\end{aligned}\tag{13}$$

— w pasmie zaporowym

$$\theta_{J+1} = [\omega_{J+2i-2}, \omega_{J+2i-1}], \quad i = 1, 2, \dots, K,\tag{14}$$

gdzie: J i K są liczbami przedziałów θ_j występujących odpowiednio w pasmie przepustowym i zaporowym, a $\omega_1 < \omega_2 < \dots < \omega_{J+2K-1}$ są unormowanymi pulsacjami, dla których spełnione są warunki:

— w pasmie przepustowym

$$(-1)^{m+j} \left. \frac{dA(\omega, Y)}{d\omega} \right|_{\omega=\omega_j} < 0, \quad j = 1, 2, 3, \dots, J-1\tag{15}$$

gdzie:

$$m = \begin{cases} 0 & \text{— gdy dla pulsacji } \omega = 0 \text{ istnieje maksimum lokalne,} \\ 1 & \text{— gdy dla pulsacji } \omega = 0 \text{ istnieje minimum lokalne;} \end{cases}$$

— w pasmie zaporowym

$$\left. \frac{dA(\omega, \mathbf{Y})}{d\omega} \right|_{\omega = \omega_{J+2i-2}} > 0, \quad i = 1, 2, \dots, K, \quad (16)$$

$$\left. \frac{dA(\omega, \mathbf{Y})}{d\omega} \right|_{\omega = \omega_{J+2i-1}} < 0, \quad i = 1, 2, \dots, K. \quad (17)$$

W celu określenia wartości ekstremów funkcji błędu $\Delta \tilde{E}_i(\mathbf{Y})$ wyznaczmy następnie wartości największych odległości $\Delta A_i(\mathbf{Y})$, $i = 1, 2, \dots, J + K + 4$ pomiędzy charakterystykami $A_d(\omega)$ i $A(\omega, \mathbf{Y})$ w poszczególnych przedziałach θ_j , $j = 1, 2, \dots, J + K$, oraz wartości odległości pomiędzy tymi charakterystykami dla pulsacji krańcowych ($\omega = 0$, $\omega = \omega_p$, $\omega = \omega_s$, $\omega = \pi$) w sposób następujący:

— w pasmie przepustowym — w przypadku, gdy dla pulsacji $\omega = 0$ występuje minimum lokalne

$$\Delta A_1(\mathbf{Y}) = A_d(0) - A(0, \mathbf{Y}) \quad (18)$$

$$\Delta A_i(\mathbf{Y}) = \max_{\omega \in \theta_{i-1}} (A(\omega, \mathbf{Y}) - A_d(\omega)), \quad i = 2, 4, 6, \dots, \quad (i \leq J + 1); \quad (19)$$

$$\Delta a_i(\mathbf{Y}) = \max_{\omega \in \theta_{i-1}} (A_d(\omega) - A(\omega, \mathbf{Y})), \quad i = 3, 5, 7, \dots, \quad (i \leq J + 1); \quad (20)$$

$$\Delta A_{J+2}(\mathbf{Y}) = A_d(\omega_p) - A(\omega_p, \mathbf{Y});$$

— w pasmie przepustowym — w przypadku, gdy dla pulsacji $\omega = 0$ występuje maksimum lokalne

$$\Delta A_1(\mathbf{Y}) = A(0, \mathbf{Y}) - A_d(0) \quad (21)$$

$$\Delta A_i(\mathbf{Y}) = \max_{\omega \in \theta_{i-1}} (A_d(\omega) - A(\omega, \mathbf{Y})), \quad i = 2, 4, 6, \dots, \quad (i \leq J + 1); \quad (22)$$

$$\Delta a_i(\mathbf{Y}) = \max_{\omega \in \theta_{i-1}} (A(\omega, \mathbf{Y}) - A_d(\omega)), \quad i = 3, 5, 7, \dots, \quad (i \leq J + 1); \quad (23)$$

$$\Delta A_{J+2}(\mathbf{Y}) = A_d(\omega_p) - A(\omega_p, \mathbf{Y}); \quad (24)$$

— w pasmie zaporowym

$$\Delta A_{J+3}(\mathbf{Y}) = A(\omega_s, \mathbf{Y}) - A_d(\omega_s), \quad (25)$$

$$\Delta A_i(\mathbf{Y}) = \max_{\omega \in \theta_{i-3}} (A(\omega, \mathbf{Y}) - A(\omega)), \quad i = J + 4, J + 5, \dots, J + K + 3. \quad (26)$$

Ponadto, jeżeli stopień licznika N jest liczbą parzystą i występuje maksimum przy pulsacji $\omega = \pi$, należy jeszcze dodatkowo uwzględnić

$$\Delta A_{J+K+4}(\mathbf{Y}) = A(\pi, \mathbf{Y}) - A_d(e^{j\pi}). \quad (27)$$

Wyznaczone wartości największych odległości $\Delta A_i(\mathbf{Y})$ są odpowiednio równe wartościom ekstremów funkcji błędu $\Delta \tilde{E}_i(\mathbf{Y})$, a więc $\Delta \tilde{E}_i(\mathbf{Y}) = \Delta A_i(\mathbf{Y})$, $i = 1, 2, \dots, J+K+4$.

Zadanie zaprojektowania filtra o równomiernie falistej charakterystyce amplitudowej, czyli o równomiernie falistej funkcji błędu zdefiniowanej wzorem (10), sprowadza się do znalezienia takiego wektora $\mathbf{Y} = [y_1, y_2, \dots, y_R]^T$, dla którego spełnione są warunki:

— w pasmie przepustowym

$$\Delta \tilde{E}_k(\mathbf{Y}) = \Delta \tilde{E}_l(\mathbf{Y}), \quad (28)$$

gdzie: $k, l = 1, 2, \dots, J+2$ — gdy wyrównywane są ekstrema $\Delta \tilde{E}_1, \Delta \tilde{E}_2, \dots, \Delta \tilde{E}_{J+2}$, lub $k, l = 2, 3, \dots, J+2$ — gdy wyrównywane są ekstrema $\Delta \tilde{E}_2, \Delta \tilde{E}_3, \dots, \Delta \tilde{E}_{J+2}$, co zależy od wartości N oraz M ;

— w pasmie zaporowym

$$\Delta \tilde{E}_k(\mathbf{Y}) = \Delta \tilde{E}_l(\mathbf{Y}) \quad (29)$$

gdzie: $k, l = J+3, J+4, \dots, J+K+3$ — gdy wyrównywane są ekstrema $\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+K+3}$, lub $k, l = J+3, J+4, \dots, J+K+4$ — gdy wyrównywane są ekstrema $\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+K+4}$, co również zależy od wartości N oraz M .

W przypadku filtra optymalnego liczba przerzutów funkcji błędu jest równa co najmniej $M+1$ w pasmie przepustowym i $N+1$ w pasmie zaporowym [4, 8]. Z jednoznaczności rozwiązania zadania aproksymacji w sensie Czebyszewa danej funkcji funkcją wymierną wynika, że jeżeli uda się wyznaczyć taki wektor $\mathbf{Y}$, dla którego funkcja błędu ma przebieg równomiernie falisty o liczbie przerzutów równej co najmniej $M+1$ w pasmie przepustowym i $N+1$ w pasmie zaporowym, to filtr o transmitancji określonej przez ten wektor jest filtrem optymalnym.

W celu rozwiązania sformułowanego powyżej zadania wprowadźmy teraz następujące dwie funkcje: $\tilde{X}_1(\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2})$ oraz $\tilde{X}_2(\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2})$, które przy spełnieniu odpowiednio (28) i (29) osiągają minimum. Przyjmijmy, że funkcje te mają następujące postacie:

$$\tilde{X}_1(\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2}) = \sum_{i=L_1}^{J+2} (\Delta \tilde{E}_i - \tilde{S})^2, \quad (30)$$

przy czym

$$\tilde{S} = \frac{1}{J+3-L_1} \sum_{i=L_1}^{J+2} \Delta \tilde{E}_i \quad (31)$$

jest średnią arytmetyczną ekstremów funkcji błędu $\Delta \tilde{E}_i$, $i = L_1, L_1 + 1, \dots, J + 2$, gdzie

$$L = \begin{cases} 1 & \text{gdy wyrównywane są ekstrema } \Delta \tilde{E}_1, \Delta \tilde{E}_2, \dots, \Delta \tilde{E}_{J+2}, \\ 2 & \text{gdy wyrównywane są ekstrema } \Delta \tilde{E}_2, \Delta \tilde{E}_3, \dots, \Delta \tilde{E}_{J+2}, \end{cases}$$

oraz

$$\tilde{X}_2(\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2}) = \sum_{i=J+3}^{J+L_2} (\Delta \tilde{E}_i - \tilde{S})^2, \quad (32)$$

przy czym

$$\tilde{S} = \frac{1}{L_2 - 2} \sum_{i=J+3}^{J+L_2} \Delta \tilde{E}_i \quad (33)$$

jest średnią arytmetyczną ekstremów funkcji błędu $\Delta \tilde{E}_i$, $i = J + 3, J + 4, \dots, J + L_2$, gdzie

$$L_2 = \begin{cases} K + 3 & \text{gdy wyrównywane są ekstrema } \Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+K+3}, \\ K + 4 & \text{gdy wyrównywane są ekstrema } \Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+K+4}, \end{cases}$$

Tak określone funkcje $\tilde{X}_1$ i $\tilde{X}_2$:

- są nieujemnymi funkcjami ciągłymi odpowiednio zmiennych $\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2}$ w przypadku funkcji $\tilde{X}_1$ oraz zmiennych $\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2}$ w przypadku funkcji $\tilde{X}_2$,
- przyjmują wartość zero wtedy i tylko wtedy gdy $\Delta \tilde{E}_{L_1} = \Delta \tilde{E}_{L_1+1} = \dots = \Delta \tilde{E}_{J+2}$ w przypadku funkcji $\tilde{X}_1$ lub gdy $\Delta \tilde{E}_{J+3} = \Delta \tilde{E}_{J+4} = \dots = \Delta \tilde{E}_{J+L_2}$ w przypadku funkcji $\tilde{X}_2$,
- nie posiadają lokalnych minimów.

Należy zauważyć, że wartości ekstremów $\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2}$ oraz $\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2}$ są funkcjami wektora $\mathbf{Y} = [\mathbf{Y}_1, \mathbf{Y}_2]^T$. Wartości ekstremów $\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2}$ znajdujących się w pasmie przepustowym są zależne przede wszystkim od położenia biegunów transmitancji, czyli od współrzędnych wektora $\mathbf{Y}_2$. Położenie zer transmitancji ma mały wpływ na wartości tych ekstremów. W przypadku ekstremów $\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2}$ znajdujących się w pasmie zaporowym sytuacja jest dokładnie odwrotna. Wartości tych ekstremów zależą przede wszystkim od położenia zer transmitancji, czyli od współrzędnych wektora $\mathbf{Y}_1$, a współrzędne wektora $\mathbf{Y}_2$ mają na te wartości mały wpływ. Problem wyrównania ekstremów można więc rozpatrywać oddzielnie w pasmie przepustowym i zaporowym.

Warunek opisany zależnością (28) jest równoważny warunkowi $\tilde{X}_1(\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2}) = 0$, a warunek opisany zależnością (29) jest równoważny warunkowi $\tilde{X}_2(\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2}) = 0$. W związku z powyższym, problem wyrównania ekstremów $\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2}$ można przekształcić w następujące równoważne zadanie optymalizacji z ograniczeniami:

Znaleźć wektor $\mathbf{Y}_2$, który dla ustalonego wektora $\mathbf{Y}_1$ minimalizuje funkcję celu $\tilde{X}_1(\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2})$ przy następujących warunkach ograniczających

$$\Delta \tilde{E}_k(\mathbf{Y}) > 0, \quad k = L_1, L_1 + 1, \dots, J + 2, \quad (34)$$

$$\forall_{i=1, 2, \dots, M_1} \sqrt{(y_{N_1+N_2+i})^2 + (y_{N_1+N_2+M_1+i})^2} < 1. \quad (35)$$

$$\forall_{i=1, 2, \dots, M_2} y_{N_1+N_2+2M_1+i} < 1 \quad (36)$$

$$\forall_{\omega \in (\omega_p, \omega_s)} A(\omega, \mathbf{Y}) - 1 < 0. \quad (37)$$

Podobnie problem wyrównania ekstremów $\Delta \tilde{E}_{J+3}$, $\Delta \tilde{E}_{J+4}, \dots$, $\Delta \tilde{E}_{J+L_2}$ można przekształcić w następujące równoważne zadanie optymalizacji z ograniczeniami:

Znaleźć wektor $\mathbf{Y}_1$, który dla ustalonego wektora $\mathbf{Y}_2$ minimalizuje funkcję celu $\tilde{X}_2(\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2})$ przy następujących warunkach ograniczających²⁾

$$\Delta \tilde{E}_k(\mathbf{Y}) > 0, \quad k = J + 3, J + 4, \dots, J + L_2. \quad (38)$$

Należy zauważyć, że ograniczenia (35) oraz (36) wynikają z warunków stabilności filtru, a ich spełnienie zapewnia uzyskanie filtru stabilnego. Spełnienie ograniczenia (37) zapewnia natomiast kontrolę przebiegu funkcji aproksymującej $A(\omega, \mathbf{Y})$ w przedziale (ω_p, ω_s) , tak aby nie przyjmowała ona wartości większych od 1.

Ze względu na fakt, że w pierwszym z wymienionych zadań optymalizacji wektor $\mathbf{Y}_2$ jest traktowany jako parametr, a wektor $\mathbf{Y}_1$ jako zmienna, natomiast w drugim — wektor $\mathbf{Y}_1$ jest traktowany jako parametr, a wektor $\mathbf{Y}_2$ jako zmienna, przed przystąpieniem do rozwiązywania tych zadań należy zdecydować, które z nich należy rozwiązywać jako pierwsze³⁾.

W celu określenia prawidłowej kolejności rozwiązywania tych zadań zbadano wpływ zmian wartości współrzędnych wektora $\mathbf{Y}_2$ na przebieg charakterystyki amplitudowej w pasmie zaporowym oraz wpływ zmian wartości współrzędnych wektora $\mathbf{Y}_1$ na przebieg charakterystyki amplitudowej w pasmie przepustowym dla kilkunastu różnych transmitacji. Okazało się, że zmiany wartości współrzędnych wektora $\mathbf{Y}_1$ mają stosunkowo mniejszy wpływ na przebieg charakterystyki amplitudowej w pasmie przepustowym, niż zmiany wartości współrzędnych wektora $\mathbf{Y}_2$ na przebieg charakterystyki amplitudowej w pasmie zaporowym. Z tego właśnie powodu postanowiono jako pierwsze rozwiązywać zadanie poszukiwania wektora $\mathbf{Y}_2$ minimalizującego funkcję $\tilde{X}_1(\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2})$ przy założonym wektorze początkowym $\mathbf{Y}_1$, a następnie znając rozwiązanie tego zadania $\mathbf{Y}_2^*$, rozwiązywać zadanie poszukiwania wektora $\mathbf{Y}_1$ minimalizującego funkcję $\tilde{X}_2(\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2})$ przy $\mathbf{Y}_2 = \mathbf{Y}_2^*$. Ponieważ zmiany wartości współrzędnych wektora $\mathbf{Y}_1$

²⁾ Należy zauważyć, że jeżeli w pasmie zaporowym $A_d(\omega) = 0$, to warunki te są zawsze spełnione.

³⁾ Przeprowadzone próby jednoczesnej minimalizacji funkcji $\tilde{X}_1(\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2})$ i $\tilde{X}_2(\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2})$ poprzez przyjęcie syntetycznego wskaźnika jakości typu $F(\tilde{X}_1, \tilde{X}_2) = k_1 \tilde{X}_1(\Delta \tilde{E}_{L_1}, \Delta \tilde{E}_{L_1+1}, \dots, \Delta \tilde{E}_{J+2}) + k_2 \tilde{X}_2(\Delta \tilde{E}_{J+3}, \Delta \tilde{E}_{J+4}, \dots, \Delta \tilde{E}_{J+L_2})$, gdzie: k_1, k_2 — współczynniki, nie dały zadowalających rezultatów.

powodują jednak pewną zmianę przebiegu charakterystyki amplitudowej w pasmie przepustowym, należy jeszcze powtórnie rozwiązać zadanie poszukiwania wektora $\mathbf{Y}_2$ minimalizującego funkcję $\tilde{X}_1(\Delta\tilde{E}_{L_1}, \Delta\tilde{E}_{L_1+1}, \dots, \Delta\tilde{E}_{J+2})$ przy $\mathbf{Y}_1 = \mathbf{Y}_1^*$, gdzie $\mathbf{Y}_1^*$ stanowi rozwiązanie poprzedniego zadania. Niekiedy może również zaistnieć potrzeba powtórnego rozwiązania zadania poszukiwania wektora $\mathbf{Y}_1$ minimalizującego funkcję $\tilde{X}_2(\Delta\tilde{E}_{J+3}, \Delta\tilde{E}_{J+4}, \dots, \Delta\tilde{E}_{J+L_2})$.

Ponieważ w sformułowanym uprzednio zadaniu minimalizacji funkcji $\tilde{X}_1(\Delta\tilde{E}_{L_1}, \Delta\tilde{E}_{L_1+1}, \dots, \Delta\tilde{E}_{J+2})$ liczba warunków jest o jeden mniejsza od liczby zmiennych, można wprowadzić dodatkowy warunek, że wszystkie $\Delta\tilde{E}_{L_1}, \Delta\tilde{E}_{L_1+1}, \dots, \Delta\tilde{E}_{J+2}$ ekstrema funkcji błędu są równe pewnej zadanej liczbie δ_1 określającej amplitudę zafalowań charakterystyki $A(\omega, \mathbf{Y})$ w pasmie przepustowym.

Obydwa rozpatrywane zadania optymalizacji można rozwiązać jedną z metod poszukiwania minimum z ograniczeniami. Z podanych w literaturze [17, 18] 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 [17, 18].

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 filtru NOI, dla przyjętego punktu startowego charakterystyka $A(\omega, \mathbf{Y})$ powinna mieć występujące na przemian po sobie minima i maksima lokalne, co jest niezbędne do określenia przedziałów θ_j , $j = 1, 2, \dots, J+K$. Wyznaczenie takiego punktu startowego stanowi więc pierwszy, istotny krok w procedurze projektowania.

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

$$E_2(\omega, \mathbf{Y}) = \sum_{j=1}^{J_1} [A(\omega_j, \mathbf{Y}) - A_d(\omega)]^2, \quad (39)$$

gdzie ω_j jest jedną z J_1 równoodległych wartości pulsacji z przedziału $[0, \pi]$. 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 w przypadku tego zadania dobór punktu startowego nie jest szczególnie istotny, a ponadto procedura jego rozwiązywania jest szybko zbieżna.

Należy zauważyć, że dla określonego przyjętego zbioru danych wejściowych i przy uwzględnieniu odpowiednich ograniczeń, zadanie czebyszewowskiej aproksymacji charakterystyki amplitudowej filtru NOI może nie mieć rozwiązania. W ta-

kim przypadku należy zmienić przyjęty stopień licznika i/lub mianownika transmitancji i powtórnie przeprowadzić obliczenia.

4. OPIS PROGRAMU FNOI

Na podstawie zaproponowanej metody projektowania filtrów NOI o czebyszewskich charakterystykach amplitudowych został opracowany program FNOI. Program ten został napisany w języku Fortran 77, przy czym do kompilacji wykorzystano kompilator PowerStation firmy Microsoft (wersja 4.0).

Danymi wejściowymi dla programu są:

- liczby N_1 oraz N_2 (przyjęto, że jeżeli $N_2 > 1$, to istnieje $N_2 - 1$ rzeczywistych zer transmitancji w przedziale $(0, 1)$ na płaszczyźnie z);
- liczby M_1 oraz M_2 ;
- wartości unormowanych częstotliwości $f_p = \omega_p/2\pi$ oraz $f_s = \omega_s/2\pi$;
- wartości parametru IW informującego, czy chcemy otrzymać wydruki kolejnych współrzędnych wyznaczonej charakterystyki amplitudowej $A(\omega, \mathbf{Y})$ ($IW = 1$), czy też nie ($IW = 0$);
- wartość parametru EPS, zwanego kryterium zakończenia procedury minimalizacji bez ograniczeń [18] (procedura ta kończy swoje działanie, gdy norma gradientu funkcji celu jest mniejsza od EPS).

W programie FNOI wykorzystano odpowiednio zmodyfikowaną procedurę GSPRIE opracowaną w Instytucie Automatyki Politechniki Warszawskiej [18]. Procedura ta umożliwia rozwiązywanie zadań programowania nieliniowego bez ograniczeń metodą gradientu sprzężonego Polaka-Ribiery, przy czym gradient funkcji celu jest estymowany numerycznie.

Do określenia wartości funkcji celu $\tilde{X}_1(\mathbf{Y})$ oraz $\tilde{X}_2(\mathbf{Y})$ konieczna jest znajomość wartości największych odległości $\Delta A_i(\mathbf{Y})$, $i = 1, 2, \dots, J+K+4$ pomiędzy charakterystykami $A_d(\omega)$ i $A(\omega, \mathbf{Y})$ w poszczególnych przedziałach θ_j , $j = 1, 2, \dots, J+K$. Wartości ekstremów funkcji $A(\omega, \mathbf{Y})$ w przedziałach θ_j wyznaczono posługując się metodą złotego podziału [17, 18]. 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.

Po przeprowadzeniu obliczeń przy użyciu programu FNOI jako wyniki końcowe otrzymujemy wartości współrzędnych wektora $\mathbf{Y}$, określające wartości zer i biegunów transmitancji oraz stałą C .

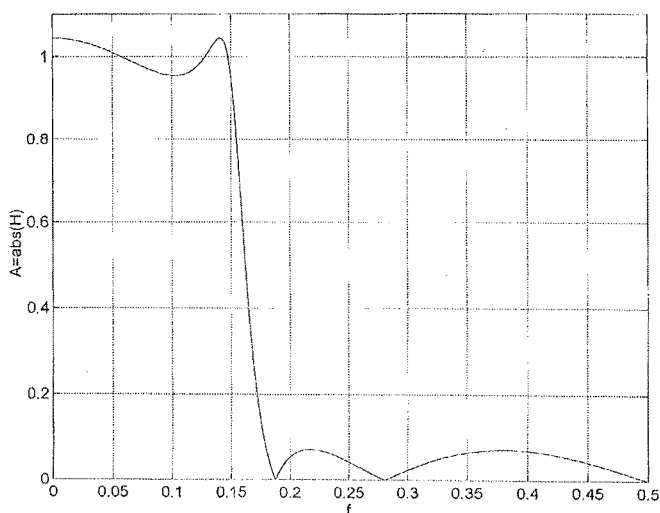
Istnieje ponadto możliwość drukowania:

- kolejnych wartości współrzędnych wyznaczonej charakterystyki amplitudowej $A(\omega, \mathbf{Y})$;
- wartości współrzędnych ekstremów funkcji błędu $\Delta \tilde{E}_i(\mathbf{Y})$, $i = 1, 2, \dots, J+K+4$.

Przy użyciu programu FNOI przeprowadzono projektowanie szeregu filtrów NOI o czebyszewskich charakterystykach amplitudowych. Obliczenia wykonywano przy użyciu mikrokomputera IBM PC Pentium z zegarem 233 MHz. Czas

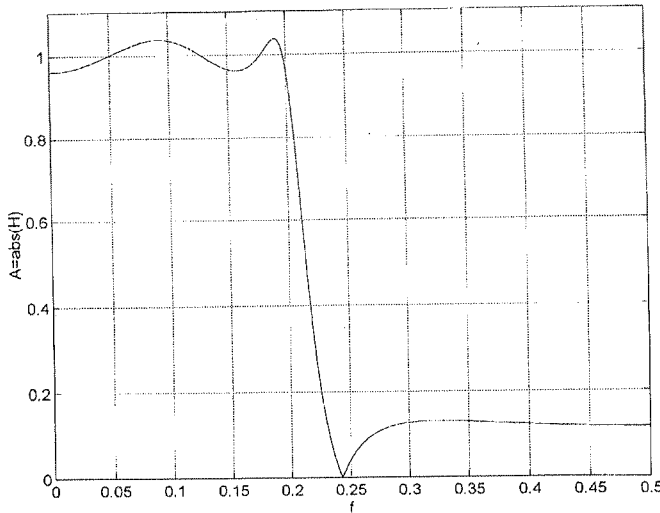
wykonywania obliczeń w przypadku filtrów o stopniach M oraz N mniejszych od 10 wahał się od kilkunastu do kilkudziesięciu sekund, przy czym czas ten wyraźnie zwiększał się w miarę wzrostu liczb M i N oraz w miarę zmniejszania się wartości parametru EPS.

W celu zademonstrowania przykładowych wyników otrzymanych przy użyciu programu FNOI, zaprojektowano filtr dolnoprzepustowy dla następujących danych $N = 5$ ($N_1 = 4$, $N_2 = 1$), $M = 3$ ($M_1 = 2$, $M_2 = 1$), $f_p = 0,15$, $f_s = 0,18$ i $\delta_1 = 0,045$. Do obliczeń przyjęto wartość parametru $\text{EPS} = 10^{-6}$. Przebieg otrzymanej charakterystyki amplitudowej $A(\omega, \mathbf{Y})$ dla powyższych danych jest pokazany na rysunku 1. Jak widać na tym rysunku, przebieg charakterystyki $A(\omega, \mathbf{Y})$ jest równomiernie falisty w pasmie przepustowym i zaporowym. Funkcja błędu ma w tym przypadku 4 przerzuty w pasmie przepustowym oraz 6 przerzutów w pasmie zaporowym. Charakterystyka $A(\omega, \mathbf{Y})$ jest więc charakterystyką filtru optymalnego.



Rys. 1. Przebieg otrzymanej charakterystyki amplitudowej $A(\omega, \mathbf{Y})$ dla danych: $N = 5$, $M = 3$, $f_p = 0,15$, $f_s = 0,18$ i $\delta_1 = 0,045$

Aby przeprowadzić porównanie wyników otrzymanych przy użyciu programu FNOI z wynikami prezentowanymi w literaturze, zaprojektowano filtr dla następujących danych: $N = N_1 = 2$, $M = M_1 = 4$, $f_p = 0,20$, $f_s = 0,23$, $\delta_1 = 0,0365$. Jest to przykład obliczeniowy zaczerpnięty z pracy Martineza i Parksa [8], w której opisana jest metoda projektowania optymalnych filtrów NOI oparta na zmodyfikowanym algorytmie Remeza. Podobnie jak w poprzednim przykładzie, do obliczeń przyjęto wartość parametru $\text{EPS} = 10^{-6}$. Przebieg otrzymanej charakterystyki amplitudowej $A(\omega, \mathbf{Y})$ dla powyższych danych jest pokazany na rysunku 2. Charakterystyka $A(\omega, \mathbf{Y})$ jest charakterystyką filtru optymalnego. Należy zauważyć, że w tym przypadku M i N są liczbami parzystymi oraz $M > N$, w związku z czym charak-



Rys. 2. Przebieg otrzymanej charakterystyki amplitudowej $A(\omega, Y)$ dla danych: $N = 2$, $M = 4$, $f_p = 0,20$, $f_s = 0,23$, $\delta_1 = 0,0365$

terystyka $A(\omega, Y)$ posiada dodatkowe minimum przy częstotliwości $f = 0,5$, które nie jest związane z przerzutem funkcji błędu.

W pracy Martineza i Parksa podane są wartości minimów i maksimów charakterystyki amplitudowej podniesionej do kwadratu. Wartości te wynoszą odpowiednio $A_{\min}(kw) = 0,926$ oraz $A_{\max}(kw) = 1,074$ w pasmie przepustowym oraz $A_{\max}(kw) = 0,017$ w pasmie zaporowym, co odpowiada amplitudom zafalowań charakterystyki amplitudowej $\delta_1 = 0,03534$ w pasmie przepustowym oraz $\delta_2 = 0,1304$ w pasmie zaporowym. W przypadku charakterystyki amplitudowej $A(\omega, Y)$ otrzymanej w wyniku projektowania filtru przy użyciu programu FNOI wartości amplitud zafalowań wynoszą: $\delta_1 = 0,0365$ i $\delta_2 = 0,1318$. Porównując otrzymane wartości δ_1 i δ_2 można stwierdzić, że wyniki projektowania filtrów metodą zaproponowaną w niniejszej pracy są praktycznie takie same, jak wyniki uzyskane przez Martineza i Parksa.

ZAKOŃCZENIE

W przedstawionej pracy zaproponowano metodę projektowania filtrów cyfrowych NOI o różnych stopniach licznika i mianownika transmitancji i czebyszewowskim przebiegu charakterystyki amplitudowej. W przeciwieństwie do wielu metod opisanych w literaturze, w metodzie tej są bezpośrednio wyznaczane zera i bieguny transmitancji. W metodach opisanych m.in. w pracach [4, 5, 8, 10, 14, 15] wyznaczane są natomiast współczynniki charakterystyki amplitudowej podniesionej do kwadratu. Następnie przeprowadza się faktoryzację wielomianów licznika i mianownika wyznaczonej funkcji wymiernej i odpowiednio wybiera się zera i bieguny,

przy czym, w celu otrzymania filtru stabilnego, wybiera się tylko bieguny leżące wewnątrz okręgu jednostkowego. Zaletą zaproponowanej metody jest więc możliwość pominięcia w procesie projektowania filtru dwóch etapów, a mianowicie faktoryzacji oraz wyboru zer i biegunów.

Przedstawiona metoda, po wprowadzeniu niewielkich modyfikacji, może zostać również wykorzystana do projektowania filtrów NOI, których charakterystyka amplitudowa $A(\omega, Y)$ aproksymuje pewną zadaną charakterystykę, inną niż przedziałami stałą. Ponadto, przy użyciu metod podobnych do opisanej w niniejszej pracy, można również projektować korektory fazy [19, 20] oraz filtry SOI [21, 22].

BIBLIOGRAFIA

1. S. K. Mitra, J. F. Kaiser (ed.): *Handbook for Digital Signal Processing*. John Wiley & Sons, Inc., New York 1993
2. A. Wojtkiewicz: *Elementy syntezy filtrów cyfrowych*. WNT, Warszawa 1984
3. A. V. Oppenheim, R. W. Schaffer: *Cyfrowe przetwarzanie sygnałów*. WKŁ, Warszawa 1979
4. L. B. Jackson: *An improved Martinez/Parks algorithm for IIR design with unequal numbers of poles and zeros*. IEEE Trans. Signal Processing, vol. 42, No. 5, May 1994, s. 1234–1238
5. T. Saramäki: *Design of optimum recursive digital filters with zeros on the unit circle*. IEEE Trans. Acoustics, Speech, and Signal Processing, vol. ASSP-31, No. 2, April 1983, s. 450–458
6. X. Zhang, H. Iwakura: *Design of IIR digital filters based on eigenvalue problem*. IEEE Trans. Signal Processing, vol. 44, No. 6, June 1996, s. 1325–1332
7. A. Aikhaïry: *Design of optimal IIR filters with arbitrary magnitude*. IEEE Trans. Circuits and Syst.-II: Analog and Digital Signal Processing, vol. 42, No. 9, Sept. 1995, s. 618–620
8. H. G. Martinez, T. W. Parks: *Design of recursive digital filters with optimum magnitude and attenuation poles on the unit circle*. IEEE Trans. Acoustics, Speech, and Signal Proc., vol. ASSP-26, No. 2, April 1978, s. 150–156
9. F. Argenti, E. Del Re: *Design of IIR eigenfilters in the frequency domain*. IEEE Trans. Signal Processing, vol. 46, No. 6, June 1998, s. 1694–1698
10. F. Brophy, A. C. Salazar: *Synthesis of spectrum shaping digital filters of recursive design*. IEEE Trans. Circuits Syst., vol. CAS-22, No. 3, March 1975, s. 197–204
11. C. Charalambous: *Minimax optimization of recursive digital filters using recent minimax results*. IEEE Trans. Acoust., Speech, Signal Processing, vol. ASSP-23, No. 4, Aug. 1975, s. 333–345
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, s. 98–112
13. E. Devleeschouwer, F. Grenet: *An efficient procedure for the design of a large class of analog and digital filters*. IEEE Trans. Circuits and Syst.-II: Analog and Digital Signal Processing, vol. 39, No. 1, Jan. 1992, s. 65–69
14. D. Dudgeon: *Recursive filter design using differential correction*. IEEE Trans. Acoust., Speech, Signal Processing, vol. ASSP-22, No. 6, Dec. 1974, s. 443–447
15. L. R. Rabiner, N. Y. Graham, H. D. Helm: *Linear programming design of IIR digital filters with arbitrary magnitude function*. IEEE Trans. Acoust., Speech, Signal Processing, vol. ASSP-22, No. 2, April 1974, s. 117–123
16. A. H. Gray, J. D. Markel: *A computer program for designing digital elliptic filters*. IEEE Trans. Acoustics, Speech, and Signal Proc., vol. ASSP-24, No. 6, Dec. 1976, s. 529–538

17. W. Findeisen, J. Szymanowski, A. Wierzbicki: *Teoria i metody obliczeniowe optymalizacji*. PWN, Warszawa 1980
18. T. Kręglewski, T. Rogowski, A. Ruszczyński, J. Szymanowski: *Metody optymalizacji w języku FORTRAN*. PWN, Warszawa 1984
19. F. Wysocka-Schillak: *Design of delay equalisers using equiripple error criterion*. Proceedings of First International Symposium on Communication Systems & Digital Signal Processing. Sheffield Hallam University (W. Brytania), 6–8 April 1998, Vol. 1, s. 250–253
20. F. Wysocka-Schillak: *Projektowanie cyfrowego korektora fazy NOI przy użyciu metod optymalizacji nieliniowej*. Kwartalnik Elektroniki i Telekomunikacji, 1998, t. 44, z. 2, s. 93–105
21. F. Wysocka-Schillak: *Chebyszewska aproksymacja charakterystyki amplitudowej dolnoprzepustowego filtru cyfrowego SOI o liniowej fazie*. Kwartalnik Elektroniki i Telekomunikacji, 1992, t. 38, z. 1, s. 211–223
22. F. Wysocka-Schillak, T. Wysocki: *On a design of nonlinear phase FIR digital filters with equiripple error function*. W: „Digital Signal Processing for Communication Systems”. Kluwer Academic Publishers. Boston 1997, s. 225–230

F. WYSOCKA-SCHILLAK

DESIGN OF OPTIMUM IIR LOW-PASS FILTERS
WITH UNEQUAL NUMERATOR AND DENOMINATOR ORDERS
OF THE TRANSFER FUNCTION

S u m m a r y

A method for designing IIR low-pass filters with optimum magnitude in the Chebyshev sense and with the transfer function having different numerator and denominator orders is presented. In the method, the considered filter design problem is transformed into equivalent constrained minimization problems, which are solved using a penalty function shifting method. A description of the computer program FNOI is included. Finally, two design examples illustrating the application of the method are given.

Key words: digital filters, IIR filters, optimum filters, optimization.

SPIS TREŚCI

J. Halawa: Wpływ zmian parametrów obiektu opisanego transmitancją	5
J. Halawa: Komputerowo wspomagany dobór nastaw regulatorów dla obiektów bez wyrównania	9
A. Mazur: Decymacja jako efektywna metoda konstrukcji zbiorów m-sekwencji oraz innych grup kodów pseudolosowych	17
J. Werewka, S. Żaba: Szeregowanie wiadomości w rozproszonych systemach czasu rzeczywistego wykorzystujących magistrale miejscowe	25
A. Drwał, J. Werewka: Projektowanie struktur sieci magistral miejscowych dla rozproszonych systemów sterowania	51
D. Kania: Pojemność kodowa programowalnych transkoderów opartych na strukturze typu PAL z różną liczbą termów dołączonych do poszczególnych wyjść	73
F. Wysocka-Schillak: Projektowanie optymalnych dolnoprzepustowych filtrów NOI o różnych stopniach licznika i mianownika transmitancji	85

CONTENTS

J. Halawa: The influence of change of the parameter of the object described by transfer function	5
J. Halawa: Computer aided method to determine the parameters of controllers for objects without steady-state	9
A. Mazur: Decimation as an effective method of finding sets of m-sequences and other groups of pseudorandom codes	17
J. Werewka, S. Żaba: Message scheduling in distributed real time systems based on fieldbus ...	25
A. Drwał, J. Werewka: Design of a fieldbus network structure for distributed control systems ..	51
D. Kania: Coding capacity of PAL-based programmable transcoder with uneven number terms per output	73
F. Wysocka-Schillak: Design of optimum IIR low-pass filters with unequal numerator and denominator orders of the transfer function	85

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 pisany 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 milimetre wave satelite, 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, pp. 141 — 148

Materiał ilustracyjny: Rysunki powinny być wykonane wyraźnie, na papierze gładkim, lub milimetrowym w formacie nie mniejszym niż 9×12 cm. Mogą być także w postaci wydruku komputerowego. Fotografie lub diapozytywy przyjmowane są raczej czarno-białe, lub barwne w formacie nie przekraczającym 10×15 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 Communittee for Electronics and Telecommunications of the Polish Academy of Sciences and is a direct continuation for 45 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 Englis, but title, the estensive summary and figure caption must be bilingual. Bilingual summary must shown 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 mrgin or uswing 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 Affilltions and relevant addresses.

All text pages should be numbered. Regular papers should have a lenght 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 Standarisation) 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 interlined.
- 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 footnotesw should be written between round brackets and upper.
- Mathematical details, anciliary 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 retrival 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 1999 is \$ 90 including postage for institutions. Subscriptions should be sent to the publisher, Polish Scientific Publishers PWN Ltd, Journal Division, Miodowa 10, 00-250 Warsaw, POLAND, fax (48) (22) 826 09 50, (48) (22) 826 71 63, with a copy by fast to Electronics and Telecommunications Quarterly 00-665 Warsaw, ul. Nowowiejska 15/19, r. 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.