

Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie
Wydział Elektrotechniki, Automatyki, Informatyki i Elektroniki

Rozprawa doktorska
Janusz Gozdecki

**Wpływ kształtowania ruchu na
opóźnienia w sieci IP w architekturze
Usług Zróżnicowanych dla reguł
przesyłania pakietów *EF PHB***

Promotor:
Prof. dr hab. inż. Zdzisław Papir

Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie
Wydział Elektrotechniki, Automatyki, Informatyki i Elektroniki
Katedra Telekomunikacji

Al. Mickiewicza 30, 30-059 Kraków, Polska
Tel. +48 12 6345582
Fax +48 12 6342372

www.agh.edu.pl
www.eaiie.agh.edu.pl
www.kt.agh.edu.pl



Copyright © Janusz Gozdecki, 2010
Wszystkie prawa zastrzeżone

Podziękowania

Szczególne podziękowania chciałbym złożyć promotorowi mojej rozprawy doktorskiej prof. Zdzisławowi Papirowi za merytoryczną opiekę nad rozprawą. Cenne uwagi prof. Zdzisława Papira nie narzucały rozwiązań, ale w istotny sposób przyczyniły się do nadania ostatecznego kształtu rozprawy. Dziękuję także za cierpliwość podczas sprawdzania i omawiania kolejnych, często bardzo różniących się między sobą, wersji rozprawy.

Pomimo, że poniższa rozprawa doktorska została napisana przeze mnie samodzielnie, to powstawała ona w konkretnym środowisku, środowisku życzliwych mi osób, które to środowisko także wywarło wpływ na ostateczny kształt tejże rozprawy. Chciałbym podziękować wszystkim współpracownikom w Katedrze Telekomunikacji za miłą i pełną motywacji do pracy atmosferę i wszelką pomoc jaka spotkała mnie z ich strony.

Chciałem podziękować także mojej rodzinie za wsparcie i wyrozumiałość.

Streszczenie

Rozprawa poświęcona jest analizie korzystnego wpływu kształtowania ruchu za pomocą algorytmu ciekącego wiadra na opóźnienia maksymalne, jakich mogą doznawać pakiety przy transmisji przez sieć IP o architekturze Usług Zróżnicowanych (DiffServ) obsługiwane przez reguły przesyłania pakietów EF PHB. EF PHB został zdefiniowany do budowania usług sieciowych o najwyższych parametrach jakościowych, wymagających bezstratnej transmisji pakietów i gwarancji opóźnień maksymalnych.

W rozprawie autor przedstawia wprowadzenie do zagadnień jakości usług w sieciach TCP/IP z naciskiem położonym na architektury sieci wspierające QoS (*Quality of Service*) (w szczególności architekturę Usług Zróżnicowanych) oraz podstawowe problemy, jakie napotyka się w tych architekturach. Autor przedstawia stan dotychczasowych badań rozważanego zagadnienia obejmujący: opis architektur sieci IP implementujących mechanizmy QoS oraz klas usług zdefiniowanych w tych architekturach, przegląd istniejących teorii dotyczących szacowania maksymalnych opóźnień w sieciach DiffServ oraz rozwiązania umożliwiające zmniejszenie tych opóźnień.

W centralnej części rozprawy autor definiuje matematyczny model opisujący opóźnienia maksymalne w sieci DiffServ z zastosowaniem algorytmu ciekącego wiadra do kształtowania ruchu w każdym łączu sieci DiffServ dla strumieni obsługiwanych przez EF PHB. Dla stworzonego modelu sieci autor sformułował i udowodnił centralne dla rozprawy twierdzenie opisujące opóźnienie maksymalne pomiędzy brzegami sieci. Analiza wzorów twierdzenia pozwoliła na wyodrębnienie szczególnych przypadków, dla których można było określić parametry sieci DiffServ i parametry algorytmu ciekącego wiadra zapewniające zmniejszenie opóźnień maksymalnych pomiędzy brzegami sieci.

W rozprawie autor wykazał, że w sieciach DiffServ, w których szybkość obsługi agregatów EF PHB jest mniejsza od przepustowości łącz, wprowadzenie kształtowania ruchu algorytmem ciekącego wiadra dla agregatów EF PHB w każdym węźle sieci

umożliwia zmniejszenie opóźnień maksymalnych. Wyprowadzone przez autora twierdzenia mogą zostać zastosowane do oceny opóźnień maksymalnych przy projektowaniu sieci operatorskich zgodnych z architekturą DiffServ dla usług sieciowych zbudowanych za pomocą reguł przesyłania pakietów EF PHB.

W odróżnieniu od sposobów zmniejszenia opóźnień maksymalnych znanych z literatury zastosowanie kształtowania ruchu wyłącznie dla agregatu EF nie wpłynie znacząco na dodatkowe obciążenie obliczeniowe ruterów w sieci DiffServ oraz nie wymaga ingerencji w format nagłówka IP.

Abstract

Most of the investigations in the area of Next Generation (NG) networks consider Internet Protocol (IP) as the ultimate mean for integrating access networks of different technologies with core networks and as a convergence protocol for versatile applications. Packet-based transport is the preferred solution due to the expected benefits in terms of deployment, management and maintenance costs. However, the departure from a circuit-switched operation towards a packet-based one brings advantages, but also poses substantial difficulties.

The Quality of Service (QoS) architecture in IP networks is required to provide the resource reservation guarantees that allow the differentiation and prioritization of flows, and deployment of advanced audio and video multimedia services. So far two QoS architectures have been standardised within an Internet Engineering Task Force (IETF) organization: Integrated Services (IntServ) with a per micro-flow Quality of Service (QoS) guarantees, and Differentiated Services (DiffServ) with aggregate-based QoS mechanisms.

The most crucial problem of QoS provisioning is the efficient support for delay-sensitive applications, with voice communication prevailing, which is a dominant factor for the overall success of the packet-based services. Managing a pure IP system based on DiffServ involves a problem of how to support per-domain services for transport of flow aggregates with a given QoS, especially for a Guaranteed Rate service. Per-domain services support data exchange by mixing traffic of different users and applications in traffic aggregates. Therefore, different aggregates are required to support delay-sensitive traffic and delay-tolerant traffic. Here, it can be necessary to support strict edge-to-edge QoS guarantees using aggregate scheduling, like in a Guaranteed Rate service.

The research on supporting multimedia services in DiffServ networks led to define an Expedited Forwarding Per Hop Behaviour (EF PHB) mechanism for servicing real-time applications. Strict deterministic control of the end-to-end delay for delay-sensitive

applications can be based on mathematical formulation of delay bounds for aggregate-based scheduling proposed by Charny, Le Boudec, and Jiang. The referred delay bounds are the base of EF PHB definition in a DiffServ architecture.

Several research proposals aim to obtain better delay bounds for DiffServ networks than the referred above. It is however done at the expense of more elaborate schedulers while preserving aggregate scheduling. They are mainly based on information carried in a packet and used by a scheduling algorithm in core routers. Examples include a concept of "dampers" and a Static Earliest Time First (SETF) scheduling algorithm.

Author proposes to use hop-by-hop aggregate-based shaping to decrease delay bounds for network services based on EF PHB. Author's findings are fundamental to understand how aggregate-based shaping can influence performance of DiffServ networks. In the dissertation author presents a theorem which is valid for such scheduling algorithms like Weighted Fair Queuing (WFQ), Self-Clocked Fair Queuing (SCFQ), and Virtual Clock (VC) used in networks with FIFO aggregate-base aggregation. In the dissertation delay bounds within the boundaries of a DiffServ network with hop-by-hop EF PHB aggregate leaky bucket shaping are considered. First a model of network architecture is derived. Based on this model author formulates a theorem which describes the end-to-end delay bound for EF PHB based services in a DiffServ network where hop-by-hop EF PHB aggregate leaky bucket shaping is applied. Author proves that the application of aggregate-based shaping in a general network architecture with FIFO Guaranteed Rate aggregate scheduling decreases the deterministic end-to-end delay bound for EF PHB based services in DiffServ networks. The statement is true when a service rate of EF PHB aggregate is lower than throughput of a link. Applying a leaky bucket shaping for flows aggregates in each node of the DiffServ domain avoids excessive jitter accumulation in the network.

Numerical comparison of the delay bounds for the example network shows that aggregate leaky bucket shaping can decrease the delay bound more than 80%. The decreasing of delay bound is in the expense of leaky bucket shaping implementation for an EF PHB aggregate in each node, what has no influence on scalability of network, but only little increases computation load in network nodes.

Spis treści

PODZIĘKOWANIA	III
STRESZCZENIE.....	V
ABSTRACT	VII
SPIS TREŚCI	IX
WYKAZ SKRÓTÓW	XI
WYKAZ SYMBOLI MATEMATYCZNYCH.....	XV
1 WPROWADZENIE – JAKOŚĆ USŁUG W SIECIACH IP	1
1.1 ZAKRES, MOTYWACJA I TEZA ROZPRAWY	6
1.2 ORYGINALNE OSIĄGNIĘCIA AUTORA.....	7
1.3 STRUKTURA ROZPRAWY	7
1.4 PUBLIKACJE I WSPÓŁPRACA NAUKOWA	8
2 OGRANICZANIE OPÓŹNIEŃ MAKSYMALNYCH W	
ARCHITEKTURZE USŁUG ZRÓŻNICOWANYCH	11
2.1 ARCHITEKTURA USŁUG ZINTEGROWANYCH	11
2.1.1 <i>Usługa gwarantowana GS</i>	14
2.1.2 <i>Usługa kontrolowanego obciążenia CL</i>	15
2.2 ARCHITEKTURA USŁUG ZRÓŻNICOWANYCH – MODEL DIFFERENTIATED	
SERVICES	16
2.2.1 <i>Klasy usług w architekturze DiffServ</i>	20
2.2.2 <i>Domeny i regiony sieci DiffServ</i>	22

2.3	OPÓŹNIENIA MAKSYMALNE W DOMENIE DIFFSERV DLA AGREGATÓW EF	23
2.3.1	<i>Szacowanie deterministycznych opóźnień maksymalnych w sieci DiffServ ...</i>	24
2.3.2	<i>Metody zmniejszania opóźnień maksymalnych w sieci DiffServ.....</i>	29
2.4	ZAKRES ROZPRAWY	33
3	WYZNACZENIE OPÓŹNIEŃ MAKSYMALNYCH W SIECI USŁUG ZRÓŻNICOWANYCH Z KSZTAŁTOWANIEM AGREGATÓW EF PHB ALGORYTMEM CIEKNĄCEGO WIADRA	35
3.1	MODEL SIECI.....	36
3.2	OPÓŹNIENIA MAKSYMALNE W SIECI DIFFSERV Z KSZTAŁTOWANIEM CHARAKTERYSTYKI RUCHU ZA POMOCĄ ALGORYTMU CIEKNĄCEGO WIADRA DLA AGREGATÓW EF PHB	38
3.2.1	<i>Dowody twierdzeń 3-1, 3-2, 3-3, 3-4 i 3-5</i>	46
3.3	ANALIZA UZYSKANYCH WYNIKÓW	61
3.3.1	<i>Porównanie opóźnień w klasycznej sieci DiffServ z rozwiązaniem zaproponowanym przez autora w twierdzeniu 3-3 dla EF PHB</i>	63
3.3.2	<i>Porównanie opóźnień w klasycznej sieci DiffServ z rozwiązaniem zaproponowanym przez autora w twierdzeniu 3-4 dla EF PHB</i>	76
3.4	PODSUMOWANIE ROZDZIAŁU.....	86
4	PODSUMOWANIE ROZPRAWY	89
5	SPIS LITERATURY	91
ANEKS A	ZASTOSOWANIE TEORII „NETWORK CALCULUS” DO OBLICZANIA MAKSYMALNYCH OPÓŹNIEŃ W SIECI – WYBRANE TWIERDZENIA ALGEBRY „MIN-PLUS”	95
ANEKS B	PORÓWNANIE RÓWNAŃ (2-7) I (2-13).....	99
ANEKS C	PORÓWNANIE RÓWNAŃ (2-13), (2-17) I (2-16)	103

Wykaz skrótów

AF	<i>Assured Forwarding</i>
ATM	<i>Asynchronous Transfer Mode</i>
BA	<i>Behaviour Aggregate</i>
BE	<i>Best-Effort</i>
BGP	<i>Border Gateway Protocol</i>
BR	<i>Border Router</i>
CL	<i>Controlled Load Service</i>
CR	<i>Core Router</i>
DHCP	<i>Dynamic Host Configuration Protocol</i>
DiffServ	<i>Differentiated Services</i>
DNS	<i>Domain Name Service</i>
DSCP	<i>DiffServ Code Point</i>
EDF	<i>Earliest Deadline First</i>
EF	<i>Expedited Forwarding</i>
ER	<i>Edge Router</i>
FIFO	<i>First In First Out</i>

FTP	<i>File Transfer Protocol</i>
GR	<i>Guaranteed Rate</i>
GS	<i>Guaranteed Service</i>
HTTP	<i>Hyper Text Transfer Protocol</i>
ICMP	<i>Internet Control Message Protocol</i>
IETF	<i>Internet Engineering Task Force</i>
IMAP	<i>Internet Message Access Protocol</i>
IntServ	<i>Integrated Services</i>
IP	<i>Internet Protocol</i>
MPLS	<i>Multi Protocol Label Switching</i>
MTU	<i>Maximum Transfer Unit</i>
NRT	<i>No real time</i>
NTP	<i>Network Time Protocol</i>
LB	<i>Leaky Bucket</i>
OSPF	<i>Open Shortest Path First</i>
PDB	<i>Per Domain Behaviour</i>
PHB	<i>Per Hop Behaviour</i>
POP	<i>Post Office Protocol</i>
PSRG	<i>Packet Scale Rate Guarantee</i>
QoS	<i>Quality of Service</i>
RADIUS	<i>Remote Authentication Dial In User Service</i>
RSVP	<i>Resource Reservation Protocol</i>
RT	<i>Real Time</i>
RTCP	<i>Real Time Control Protocol</i>
RTP	<i>Real-time Transport Protocol</i>
SCFQ	<i>Self Clocked Fair Queuing</i>
SETF	<i>Static Earliest Time First</i>

S-HTTP	<i>Secure-HTTP</i>
SIP	<i>Session Initiation Protocol</i>
SLA	<i>Service Level Agreement</i>
SMTP	<i>Simple Mail Transport Protocol</i>
TCP	<i>Transmission Control Protocol</i>
TFTP	<i>Trivial File Transfer Protocol</i>
TOS	<i>Type of Service</i>
TR	<i>Transit Router</i>
VC	<i>Virtual Clock</i>
VPN	<i>Virtual Private Network</i>
WFQ	<i>Weighted Fair Queuing</i>
WWW	<i>World-Wide Web</i>

Wykaz symboli matematycznych

α	współczynnik wykorzystania pasma
β	ograniczenie sumy rozmiarów wiader strumieni przepływających przez łącze zdefiniowana wzorem (2-6)
Δ	stała inteligentnego bufora
ε	współczynnik rozmiaru wiadra w algorytmie ciekącego wiadra
η	współczynnik szybkości obsługi w algorytmie ciekącego wiadra
ρ_l	szybkość obsługi pakietów w algorytmie ciekącego wiadra należących do agregatu EF PHB w łączu l
σ_l	rozmiar wiadra w algorytmie ciekącego wiadra dla agregatu EF PHB w łączu l
A	krzywa ograniczająca ruch wejściowy dla agregatu strumieni
a	krzywa ograniczająca ruch wejściowy dla pojedynczego strumienia
B	rozmiar bufora
b	rozmiar wiadra w algorytmie ciekącego wiadra
D	opóźnienie maksymalne pomiędzy brzegami sieci DiffServ
D_i	opóźnienie maksymalne dla pakietów strumienia i podlegającego rezerwacji na całej ścieżce w modelu IntServ

D^{path}	opóźnienie maksymalne dla pakietów strumienia danych transmitowanego pomiędzy brzegami domeny DiffServ na ściśle określonej ścieżce
C_l	przepustowość łącza l
d_l	maksymalne opóźnienie pakietu w łączu l
ded	różnica pomiędzy czasem życia pakietu (<i>deadline</i>) a jego aktualnym czasem emisji
E_l	błąd obsługi agregatu zależny od implementacji mechanizmów obsługi w łączu l
e	opóźnienie w węźle nieobciążonym
F	błąd aproksymacji modelu przepływowego (strumieniowego)
f_n	wielkość zdefiniowana w równaniu (2-5)
G	błąd aproksymacji modelu przepływowego (strumieniowego)
GR	zmienna zdefiniowana wzorem (2-12)
H	maksymalna liczba węzłów, przez które przechodzi pakiet pomiędzy dwoma dowolnymi ruterami brzegowymi domeny DiffServ
i	indeks
k	liczba łączy wchodzących do węzła n
L	zbiór wszystkich łączy l w rozważanej domenie DiffServ
l	indeks łącza
M_l	maksymalny rozmiar pakietu w łączu l
M_n	suma maksymalnych rozmiarów pakietów od wszystkich łączy wchodzących do węzła n
M	maksymalny rozmiar pakietu w sieci
N	zbiór wszystkich węzłów w sieci
n	indeks węzła
O_l	zbiór wszystkich strumieni przechodzących przez łącze l i należących do agregatu o najwyższym priorytecie

P_n	maksymalna szybkość ruchu wchodzącego do węzła n
p	maksymalna szybkość strumienia danych
$PATH$	zbiór łączy, przez które przechodzą pakiety strumienia danych pomiędzy brzegami domeny DiffServ
Q_{del}	opóźnienie <i>end-to-end</i> w modelu IntServ
R_l	szybkość obsługi pakietów należących do strumienia lub agregatu danych w łączy l
r	szybkość obsługi w algorytmie ciekącego wiadra
S_l	krzywa obsługi dla agregatu strumieni przechodzących przez łączy l i należących do agregatu o najwyższym priorytecie
s	krzywa obsługi dla strumienia
T	stała czasowa
T_l	stała czasowa dla łączy l
t	czas
t^a	czas przybycia pakietu do węzła
t_j^{we}	czas przybycia ostatniego bitu pakietu j do węzła sieciowego
t_j^{wy}	czas opuszczenia węzła sieciowego przez ostatni bit pakietu j
u	zmienna zdefiniowana wzorem (2-10)
v	zmienna zdefiniowana wzorem (3-18)
Z	stała charakterystyczna dla danego łączy i zależna od mechanizmów szeregowania i konfiguracji urządzenia (równanie (2-3))

1 Wprowadzenie – jakość usług w sieciach IP

Obecnie sieć Internet składa się z wielu heterogenicznych sieci zbudowanych z wykorzystaniem różnych technik łącza danych. Do komunikacji pomiędzy aplikacjami w takiej sieci stosuje się protokół IP (*Internet Protocol*). Analiza protokołu IP wskazuje na brak mechanizmów odpowiadających za jakość realizowanej usługi sieciowej, które można by w sposób efektywny zastosować dla obsługi aplikacji [1]. Protokół IP został skonstruowany w taki sposób, aby osiągnąć maksymalną prostotę i skalowalność nawet kosztem obniżenia wydajności sieci i jakości obsługi. Protokół IP nie korzysta z właściwości protokołów niższych warstw i oferuje zawodną, bezpołączeniową usługę sieciową, która jest narażona na utratę pakietów, zmianę kolejności pakietów i ich duplikację. Ponadto opóźnienia w transmisji danych, związane z kolejkowaniem w buforach urządzeń sieciowych, zwiększają się wraz ze wzrostem obciążenia sieci.

Z uwagi na niedostatek solidnych gwarancji jakościowych tradycyjny model usługi dostarczany przez protokół IP jest określany jako *Best-Effort* (BE) i do prawidłowego działania musi on być uzupełniany o dodatkowe protokoły warstw wyższych jak TCP (*Transmission Control Protocol*), które są niezbędne dla uzyskania niezawodności w połączeniach *end-to-end*.

Klasyfikacja usług opartych o protokół IP nie została jeszcze zatwierdzona przez żadne forum standaryzacyjne. Przedstawiony w tabeli 1-1 podział aplikacji jest wzorowany na opracowaniu [2], a został przeprowadzony na podstawie analizy standardów takich jak: [3], [4], [5], [6]. Aplikacje zostały podzielone na cztery klasy:

- 1) usługi klient-serwer – np. tradycyjne usługi Internetowe jak poczta elektroniczna, FTP, WWW,
- 2) usługi strumieniowania jednokierunkowego – transmisja dźwięku i wideo na przykład szeroko wykorzystywana przez stacje radiowe transmitujące przez Internet,

- 3) konwersacyjne usługi multimedialne – np. systemy konferencyjne wykorzystujące informację dźwiękową, wideo i dane oparte o protokoły T.120, H.323 i SIP,
- 4) usługi związane z utrzymaniem sieci – np. DNS, protokoły wyboru tras, DHCP.

Tabela 1-1: Klasyfikacja usług internetowych

Kategoria usług	Główne atrybuty – wymagania	Przykładowe aplikacje	Przykładowe protokoły
1. Klient-server	Zależności czasowe: NRT Symetria: Usługi asymetryczne. Architektura: punkt-punkt. Charakterystyka ruchu: VBR, ABR. Sposób obsługi: <i>“store and forward”</i> . Użytkownik: zarówno człowiek jak i maszyna.	Poczta elektroniczna Transmisja plików Telnet WWW	TCP, SMTP, POP, IMAP, FTP, TFTP, TELNET, SSH, HTTP, HTTPS
2. Strumieniowanie jednokierunkowe	Zależności czasowe: RT, brak ścisłych wymagań na małe opóźnienia, ograniczenia na fluktuacje opóźnienia – <i>jitter</i> . Symetria: Usługi wysoce asymetryczne, Architektura: punkt-punkt, punkt – wiele punktów. Charakterystyka ruchu: CBR, VBR. Sposób obsługi: <i>“store and forward”</i> . Użytkownik: klient.	Strumieniowanie audio Strumieniowanie video Strumieniowanie rozsiewcze	UDP, RTP, RTSP, RTCP
3. Usługi konwersacyjne	Zależności czasowe: RT, małe. opóźnienia transmisyjne, małe fluktuacje opóźnienia. Symetria: Usługi symetryczne. Architektura: punkt-punkt, punkt – wiele punktów. Charakterystyka ruchu: CBR, VBR, ABR. Sposób obsługi: konwersacyjny. Użytkownik: klient.	Telefonia IP, Telekonferencje, Wideo-konferencje, Gry sieciowe.	UDP, RTP, RTCP, RTSP, SIP, T.120, H.323, H.261, H.263
4. Usługi związane z utrzymaniem sieci	Zależności czasowe: NRT. Symetria: Usługi symetryczne. Architektura: punkt-punkt, punkt – wiele punktów. Charakterystyka ruchu: VBR, ABR. Sposób obsługi: sygnalizacja. Użytkownik: ISP, <i>Content Provider</i> .	DNS – system nazw domen, NTP- <i>network time protocol</i> , <i>Dynamic Host Control Protocol</i> , RADIUS, Dobór trasy	TCP, UDP, ICMP, IGMP, DNS, DHCP, NTP, OSPF, BGP

RT (*real time*) – usługi czasu rzeczywistego, NRT (*no RT*) – usługi niewymagające obsługi w czasie rzeczywistym

Przedstawiona klasyfikacja dotyczy usług korzystających z protokołu IP. Do klasy 1 zaliczamy klasyczne usługi Internetowe, które muszą być wspierane przez klasę 4 odpowiadającą za prawidłowe funkcjonowanie Internetu. Klasę 2 i 3 tworzą specyficzne usługi, które w Internecie funkcjonują wyłącznie w wydzielonych sieciach, ponieważ obowiązujące zasady obsługi ruchu BE nie spełniają podstawowych warunków dla implementacji klasy 2 i 3. Obecnie mogą one działać jedynie w tych obszarach sieci, gdzie gwarancje jakości transmisji uzyskuje się przez odpowiednie planowanie sieci wiążące się z jej znacznym przewymiarowaniem.

Protokół TCP dla zapewnienia bezstratnej transmisji wykorzystuje retransmisję pakietów, co dodatkowo wpływa na zwiększenie opóźnień transmisyjnych. Dla tradycyjnych usług internetowych, jak np. *File Transfer Protocol* (FTP) służący do transmisji plików, *Hyper Text Transfer Protocol* (HTTP) do przeglądania stron WWW (*World Wide Web*) czy *Simple Mail Transport Protocol* (SMTP) przeznaczony do transportu poczty elektronicznej, model transmisyjny typu BE jest do przyjęcia. Jednak wiele nowych aplikacji wymaga od systemu transmisyjnego ograniczonych czasów opóźnień oraz ograniczenia zmienności opóźnień (fluktuacji opóźnień) w czasie. Przykładowo, dla systemów wykorzystujących dobrej jakości dźwięk wymaga się, aby całkowite opóźnienie w połączeniu *end-to-end* nie przekraczało 150 ms. Pakiety z głosem przychodzące z opóźnieniem większym od 150 ms powinny być odrzucane. W tym przypadku retransmisja zgubionych pakietów realizowana przez protokół TCP tylko pogarsza warunki transmisyjne w łączu, powodując dodatkowy, zbędny ruch w sieci. Dla aplikacji wrażliwych na opóźnienia ten model obsługi jest nieodpowiedni nawet przy niskim wykorzystaniu łącz transmisyjnych.

Problem zmienności opóźnień pakietów – wynikających z kolejkowania w buforach urządzeń sieciowych i strat pakietów w sieci IP – może być eliminowany przez aplikacje, które w miarę możliwości adaptują się do obciążenia panującego w sieci – zwykle jest to realizowane przez buforowanie danych u odbiorcy. Umożliwia to kompensację zmienności opóźnień, ale też tylko do pewnego krytycznego obciążenia sieci. Nadal jednak istnieje potrzeba wdrażania nowych aplikacji, wymagających gwarantowanego pasma transmisyjnego i stałych opóźnień (np. dla usług konwersacyjnych z tabeli 1-1). Poszerzające się obszary zastosowań sieci IP powodują, że nieodzownym staje się zapewnienie gwarancji jakości usług świadczonych przez protokół IP, szczególnie dla interaktywnych aplikacji multimedialnych. Wprowadzenie takich aplikacji wymaga rozszerzenia obecnego modelu sieci IP o dodatkowe klasy usług, zapewniające gwarancje QoS jakości usług dla warstw wyższych.

Obecne wymagania jakości transmisji realizuje się przez odpowiednie planowanie sieci, eliminujące natłoki w sieci i przez stosowanie coraz to bardziej wydajnych technik transmisyjnych w warstwach niższych jak ATM, MPLS czy optyczne systemy transmisyjne. Brak współpracy w zakresie QoS pomiędzy protokołem IP a niższymi warstwami może jednak powodować, że planowanie i nowe techniki transmisyjne nie

zapewnią żądanej jakości transmisji. Uzyskanie gwarancji jakości usług w sieci IP wymaga rozszerzenia istniejącego modelu obsługi o nowe klasy usług umożliwiające gwarancje jakości usługi sieciowej.

Na obecnym etapie badań nad QoS w sieciach IP, w ramach działalności organizacji IETF, zdefiniowano dwie architektury sieci umożliwiające podział obsługiwanego ruchu na klasy i obsługę tych klas ze zróżnicowaną jakością. Są to architektury:

- Usług Zintegrowanych (Integrated Services – IntServ) [7],
- Usług Zróżnicowanych (Differentiated Services – DiffServ) [8].

Architektury te rozszerzają model BE sieci IP o nowe klasy obsługi umożliwiające realizację usługi sieciowej z żadaną jakością.

Podstawowym zadaniem opracowywanych architektur sieci IP z gwarancją jakości QoS usługi sieciowej jest umożliwienie realizacji usług z kategorii 2 i 3 (wymienionych w tabeli 1-1) w globalnej sieci Internet. Optymalna architektura będzie obsługiwać wszystkie rodzaje aplikacji w jednej sieci, niezależnie od jej rozmiaru i liczby węzłów składowych.

Klasy usług implementowane w modelu Usług Zintegrowanych umożliwiają obsługę każdego rodzaju aplikacji, z uwzględnieniem usług czasowo krytycznych (czasu rzeczywistego), dla których przeznaczona jest gwarantowana usługa sieciowa (*Guaranteed Service*) [9]. Usługa ta umożliwia aplikacji sterowanie parametrami transmisji, takimi jak na przykład opóźnienie całkowite pakietów w sieci. Jednak przyjęty sposób obsługi pakietów w węzłach sieci powoduje, że architektura Usług Zintegrowanych posiada problemy skalowalności w sieciach szkieletowych. Problemy te wiążą się z przyjętym modelem obsługi zgłoszeń – dla każdego **strumienia danych** generowanego przez aplikację, który identyfikuje się przez adres, port i protokół transportowy nadawcy i odbiorcy, należy w każdym routerze na drodze pomiędzy nadawcą i odbiorcą przechowywać informacje opisujące rezerwację zasobów.

Rezerwacja zasobów dla strumienia danych w węźle sieciowym polega na zagwarantowaniu pasma (gwarancja braku odrzucania pakietów lub prawdopodobieństwo odrzucania pakietów na ustalonym poziomie) i maksymalnego opóźnienia pakietów w węźle. Rezerwacja zasobów dla modelu Usług Zintegrowanych jest opisana przez parametry określające pasmo gwarantowane w węźle, klasę jakości obsługi oraz współczynniki określające stałe i zmienne opóźnienie w węźle. Ponadto, dla usługi sieciowej przeznaczonej dla aplikacji o najwyższych wymaganiach jakościowych, każdy strumień danych powinien być obsługiwany przez niezależną kolejkę (bufor pakietów). Powyższe wymagania powodują, że powstaje stosunkowo duży narzut sygnalizacyjny do przenoszenia informacji o parametrach rezerwacji, routery w sieci muszą przechowywać dużo informacji o pojedynczych połączeniach oraz potrzebują dedykowanych zasobów procesora do przetwarzania i obsługi pakietów należących do pojedynczych strumieni

danych podlegających rezerwacji. Wymagania te powodują, że rutery dotychczas stosowane w sieciach nie są w stanie obsługiwać dużej ilości równoległych rezerwacji jednocześnie. Opisany wyżej model nie sprawdza się szczególnie w sieciach rdzeniowych, gdzie operatorzy specjalizują się w przenoszeniu ruchu pomiędzy innymi operatorami sieci i zwykle nie mają kontraktów z pojedynczymi abonentami sieci, a musieliby zapewnić jakość obsługi dla poszczególnych strumienia od pojedynczego abonenta.

Podsumowując, architektura IntServ dostarcza aplikacjom możliwość sterowania jakością połączenia QoS w sieci w relacji *end-to-end*, ale niezależna obsługa wszystkich połączeń powoduje, że architektura ta wymaga w węzłach sieci dużych zasobów przeznaczonych na przetwarzanie, przechowywanie informacji o rezerwacjach i obsługę pakietów zgodnie z parametrami rezerwacji. Stwarza to poważny problem skalowalności modelu IntServ, szczególnie w obszarze szkieletowym sieci.

Problemy skalowalności sieci IntServ zostały wyeliminowane w architekturze Usług Zróżnicowanych, gdzie dedykowana obsługa dla pojedynczych strumieni danych możliwa jest tylko na brzegu sieci, natomiast wewnątrz sieci wszystkie strumienie danych są przydzielone do ograniczonej liczby agregatów DiffServ. W praktyce do oznaczenia agregatów przeznaczono sześciobitowe pole DSCP (*DiffServ Code Point*) w nagłówku protokołu IP; w protokole IPv4 są to bity w polu TOS (*Type of Service*), w IPv6 DSCP jest reprezentowane przez informacje w polu *Service Class* [10]. Strumienie danych należące do tego samego agregatu DiffServ są obsługiwane z tą samą jakością, zwykle przez zastosowanie jednej kolejki do obsługi wszystkich strumieni danych należących do tego agregatu DiffServ, bez jakiegokolwiek możliwości zróżnicowania jakości obsługi dla pojedynczych strumieni danych. W sieci strumień danych aplikacji jest podzielony na pakiety. Strumień pakietów transmitowanych w jednokierunkowym łączy sieci DiffServ i oznaczonych tym samym kodem DSCP nazywamy **agregatem DiffServ** (*DiffServ Behaviour Aggregate*). Dla uproszczenia, agregat DiffServ w dalszej części rozprawy będzie nazywany **agregatem**.

Zasada obsługi pakietów polegająca na grupowaniu pakietów pochodzących od różnych aplikacji w agregaty powoduje, że w architekturze Usług Zróżnicowanych pojawiły się problemy dotyczące gwarancji parametrów obsługi typu *end-to-end* w sieci dla aplikacji wymagających deterministycznych ograniczeń na opóźnienia maksymalne mierzone pomiędzy brzegami domeny DiffServ, czyli na przykład dla aplikacji kategorii konwersacyjnej z tabeli 1-1 i dla wszystkich aplikacji czasu rzeczywistego. Sieć wyposażona w mechanizmy DiffServ, w której zaimplementowano ten sam zestaw reguł przesyłania pakietów we wszystkich łączach oraz te same zasady obsługi ruchu na brzegu tej sieci nazywamy **domeną DiffServ**. Do obsługi strumieni danych aplikacji czasu rzeczywistego w architekturze Usług Zróżnicowanych przeznaczona jest, specjalnie do tego celu zaprojektowana, reguła przesyłania pakietów EF PHB (*Expedited Forwarding Per Hop Behaviour*). **Reguła przesyłania pakietów** definiuje mechanizmy obsługi pakietów zaimplementowane w pojedynczym węźle sieci DiffServ dla danego agregatu. W rozdziale

2.2 autor omawia także inne reguły przesyłania pakietów PHB, przeznaczone dla innych typów aplikacji. Z obecnie zdefiniowanych reguł przesyłania pakietów EF PHB dostarcza mechanizmy obsługi o najwyższych parametrach jakościowych. Implementacja w węzłach sieci kilku reguł przesyłania pakietów wymaga zastosowania algorytmu szeregowania, który będzie w sposób niezależny obsługiwać każdą z nich. W rozprawie autor rozważa klasę algorytmów szeregowania typu Gwarantowanej Szybkości (GR), których dwoma podstawowymi parametrami są: szybkość obsługi danej klasy ruchu oraz stała charakteryzująca opóźnienie obsługi.

Badania obsługi agregatów według reguły przesyłania pakietów EF PHB dowodzą [11][12][13], że w węzłach sieci opóźnienia maksymalne rosną wykładniczo ze wzrostem rozmiarów sieci, a maksymalne wykorzystanie pasma w sieci jest uzależnione od rozmiarów i architektury sieci. Dla zapewnienia małych opóźnień pakietów w sieci DiffServ należy znacznie przewymiarować pasmo wymagane przez aplikacje czasu rzeczywistego.

1.1 Zakres, motywacja i teza rozprawy

Rozprawa poświęcona jest badaniom możliwości zmniejszenia opóźnień maksymalnych pomiędzy brzegami sieci w architekturze Usług Zróżnicowanych. Celem rozprawy jest analiza wpływu kształtowania ruchu EF PHB za pomocą algorytmu ciekącego wiadra (LB) na opóźnienia maksymalne, jakich doznają pakiety przy transmisji przez sieć IP o architekturze Usług Zróżnicowanych.

Wprowadzenie kształtowania ruchu dla agregatów strumieni danych w każdym węźle sieci zapobiega nadmiernej kumulacji pakietów w węzłach sieci, powodującej możliwość wystąpienia w sieci niekontrolowanego wzrostu opóźnienia.

Następująca teza została postawiona i zostanie udowodniona w rozprawie:

Zastosowanie algorytmu ciekącego wiadra do kształtowania agregatów EF PHB w każdym węźle sieci IP w architekturze Usług Zróżnicowanych umożliwia zmniejszenie opóźnień maksymalnych pomiędzy brzegami domeny dla usług sieciowych zbudowanych za pomocą reguł przesyłania pakietów EF PHB pod warunkiem, że szybkość obsługi agregatów EF PHB jest mniejsza od przepustowości łączy.

Możliwość zapewnienia gwarancji pasma dla strumieni danych oraz kontrolowania opóźnień maksymalnych transmisji pakietu przez domenę DiffServ dla usług zbudowanych z zastosowaniem reguły przesyłania pakietów EF PHB umożliwia implementację usługi *Guaranteed Service* [9] pomiędzy brzegami sieci DiffServ bez konieczności implementacji dodatkowych mechanizmów w węzłach sieci.

Obecna tendencja do oferowania coraz szerszego zestawu usług w sieciach IP,

włączając w to telefonię internetową, telewizję internetową, wideo na żądanie powoduje, że operatorzy w coraz większym stopniu powinni być zainteresowani mechanizmami umożliwiającymi jednoczesną kontrolę opóźnień w swoich sieciach oraz zwiększanie ruchu w łączach bez potrzeby zwiększania ich przepustowości. Obecne sieci z zaimplementowaną jedną klasą obsługi ruchu typu *Best-Effort* możemy traktować jako specjalny przypadek sieci DiffServ.

1.2 Oryginalne osiągnięcia autora

Do najważniejszych oryginalnych osiągnięć rozprawy autor zalicza:

- zbudowanie modelu sieci DiffServ z kształtowaniem ruchu algorytmem ciekącego wiadra dla agregatów obsługiwanych przez regułę przesyłania pakietów EF PHB,
- sformułowanie i dowód ogólnego twierdzenia opisującego wpływ kształtowania ruchu algorytmem ciekącego wiadra na opóźnienie maksymalne w domenie DiffServ dla usług sieciowych zbudowanych z wykorzystaniem EF PHB,
- wyznaczenie rozwiązań ogólnego twierdzenia dla szczególnych przypadków opisujących, dla jakich warunków opóźnienia maksymalne w sieci DiffServ z kształtowaniem ruchu są mniejsze lub większe niż opóźnienia maksymalne w odpowiadającej jej sieci bez kształtowania ruchu,
- ilościowe oszacowanie wpływu kształtowania ruchu algorytmem ciekącego wiadra na opóźnienia maksymalne dla wybranych przypadków,
- możliwość praktycznego zastosowania uzyskanych wyników do projektowania sieci operatorskich.

1.3 Struktura rozprawy

Rozprawa jest podzielona na 4 rozdziały. Rozdział pierwszy „Jakość usług w sieciach IP” poświęcony jest wprowadzeniu do zagadnień jakości usług w sieciach TCP/IP z naciskiem położonym na architektury sieci wspierające QoS (*Quality of Service*) i podstawowe problemy, jakie napotyka się w tych architekturach. W tym rozdziale autor także prezentuje bardziej szczegółowo rozwiązywane zadanie badawcze oraz tezę rozprawy i najważniejsze osiągnięcia rozprawy.

Rozdział 2 „Ograniczenia opóźnień maksymalnych w architekturze Usług Różnicowanych” przedstawia stan dotychczasowych badań zagadnienia obejmujący:

- architektury sieci IP implementujących mechanizmy QoS i opisu klasy usług implementowanych przez te architektury,
- opis istniejących teorii dotyczących maksymalnych opóźnień, jakich może doznać pakiet przy przejściu przez domenę DiffServ,
- rozwiązania umożliwiające ograniczenie maksymalnych opóźnień w sieci DiffServ dla EF PHB.

Rozdział 3 „Wyznaczenie opóźnień maksymalnych w sieci Usług Zróżnicowanych z kształtowaniem agregatów algorytmem ciekącego wiadra dla EF PH” opisuje maksymalne opóźnienia w sieci z zastosowaniem algorytmu ciekącego wiadra do kształtowania ruchu w każdym łączy sieci DiffServ dla agregatów EF PHB. Rozdział 3.1 przedstawia model sieci DiffServ wykorzystywany w następnych rozdziałach do obliczania opóźnień maksymalnych. W rozdziale 3.2 autor opisuje wpływ kształtowania agregatów strumieni EF PHB algorytmem ciekącego wiadra na maksymalne opóźnienie. Autor zdefiniował jedno ogólne twierdzenie opisujące opóźnienia maksymalne w domenie DiffServ oraz cztery twierdzenia, które opisują maksymalne opóźnienia dla różnych zakresów wartości parametrów algorytmu ciekącego wiadra. W rozdziale 3.3 zostało przedstawione porównanie opóźnienia maksymalnego w sieci DiffServ, w której stosuje się algorytm ciekącego wiadra zaproponowany przez autora z opóźnieniami maksymalnymi w klasycznej sieci DiffServ.

Rozdział 4 stanowi podsumowanie osiągnięć uzyskanych w rozprawie. Rozprawę uzupełniają aneksy. W Aneksie A autor przedstawia najważniejsze twierdzenia aparatu matematycznego „*Network calculus*” wykorzystywane w rozprawie. Aneks B i C uzupełniają przedstawione w rozdziale 2.3 porównanie rozwiązań umożliwiających ograniczenie opóźnień maksymalnych w sieci DiffServ dla EF PHB.

1.4 Publikacje i współpraca naukowa

Wybrane wyniki rozprawy zostały opublikowane przez autora w dwóch artykułach naukowych:

- 1) J. Gozdecki, A. Jajszczyk, and R. Stankiewicz, "Quality of Service Terminology in IP Networks", IEEE Communications Magazine, March 2003, str. 153-159.
- 2) J. Gozdecki, "Decreasing delay bounds for a DiffServ network using leaky bucket shaping for EF PHB aggregates", HET-NETs 2010, Zakopane, Polska, 14-16 January 2010.

Pierwszy z nich, który ukazał się w IEEE Communications Magazine, został poświęcony terminologii związanej z jakością usług w sieciach IP. Artykuł obejmuje

omówienie definicji jakości usług w sieciach IP różnych organizacji standaryzacyjnych, opisanie najważniejszych właściwości architektur sieci Usług Zintegrowanych i Usług Zróżnicowanych oraz związanych z nimi pojęć.

W drugim artykule autor przedstawia najważniejsze osiągnięcie rozprawy, które udowadnia tezę stawianą przez autora w rozprawie.

Doświadczenia zdobywane podczas pracy nad doktoratem autor wykorzystywał w projektach międzynarodowych, gdzie zajmował się zagadnieniami związanymi z jakością usług w sieciach IP. Dwa najważniejsze projekty, które autor realizował podczas pracy nad doktoratem to:

- “Designing Advanced Interfaces for the Delivery and Administration of Location independent Optimised personal Services II (DAIDALOS II)”, project IST-2006-020943, w ramach VI Programu Ramowego Unii Europejskiej, 2006-2008
- “CARrier grade MEsh Networks (CARMEN)”, project FP7 214994, w ramach VII Programu Ramowego Unii Europejskiej, 2008-2010

W powyższych projektach autor zajmował się głównie integracją rozwiązań IP QoS w sieciach bezprzewodowych. Najważniejsze wyniki tych prac można znaleźć w następujących artykułach konferencyjnych:

- 1) J. Gozdecki, S. Sargento, D. Wagner, J. Rocha, F. Mitrano, J. Mödeker, “Multicast Mobility in Heterogeneous Technologies: Experimental Validation”, IEEE GLOBECOM 2009, Honolulu, Hawaii, USA, 30 November – 4 December 2009.
- 2) J. Gozdecki, S. Sargento, D. Wagner, J. Rocha, J. Mödeker, “Mobility with QoS in Broadcast Unidirectional Technologies: Experimental Validation”, IEEE ICC 2010 - Wireless and Mobile Networking Symposium ('ICC'10 WNS'), Cape Town, South Afrika, 23 – 27 May 2010, artykuł przyjęty do publikacji w materiałach konferencyjnych i IEEE Exploire.

Artykuły dotyczą integracji mechanizmów QoS w sieci o mieszanej architekturze Usług Zintegrowanych i Usług Zróżnicowanych, do realizacji jakości usług multicastowych i rozgłoszeniowych dla użytkownika ruchomego.

2 Ograniczanie opóźnień maksymalnych w architekturze Usług Zróżnicowanych

Rozdział ten przedstawia:

- przegląd architektur sieci IP implementujących mechanizmy QoS oraz związanych pojęć,
- metody obliczania maksymalnych opóźnień, jakich może doznać pakiet przy przejściu przez domenę DiffServ dla EF PHB,
- rozwiązania umożliwiające ograniczenie maksymalnych opóźnień w sieci DiffServ dla EF PHB.

2.1 *Architektura Usług Zintegrowanych*

Architektura Usług Zintegrowanych (*Integrated Services*) była pierwszym modelem sieci IP z gwarancją jakości QoS opracowanym przez organizację IETF zajmującą się standaryzacją Internetu. Podstawowymi założeniami tego modelu są [14]:

- możliwość kontrolowania zasobów wpływających na jakość transmisji (gwarantowane pasmo, maksymalne opóźnienia) bezpośrednio przez aplikację,
- niezależność obsługi strumieni danych w sieci – transmisja strumieni nie zakłóca wzajemnych parametrów transmisyjnych.

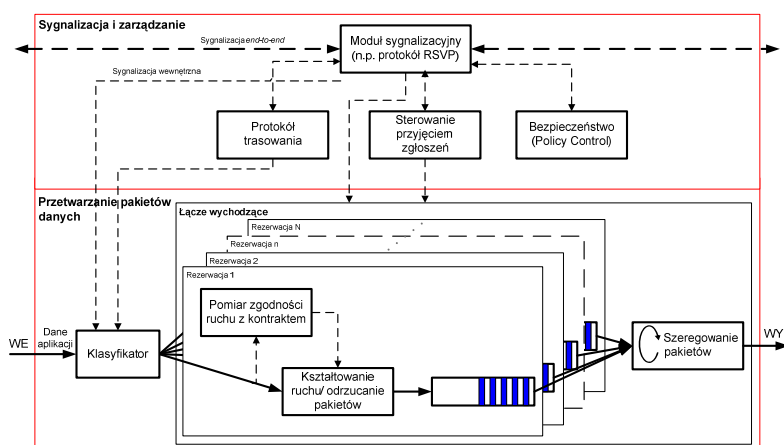
Podstawą działania modelu IntServ są zdefiniowane w tym modelu klasy usług. Obecnie są to dwie klasy usług, ale ich liczba nie została w żaden sposób ograniczona:

- **usługa gwarantowana GS** (*Guaranteed Service*) [9],

- usługa kontrolowanego obciążenia CL (*Controlled Load Service*) [15].

Powyższe zdefiniowane klasy usług wraz z protokołem sygnalizacyjnym **RSVP** (*Resource ReserVation Protocol*) [16] są podstawowymi blokami funkcjonalnymi składającymi się na model IntServ.

Klasy usług GS i CL definiują sposób, w jaki strumień danych mają być obsługiwane w pojedynczym węźle, a w powiązaniu z protokołem sygnalizacyjnym dają aplikacji możliwość sterowania jakością połączenia QoS w relacji *end-to-end*. Sterowanie jakością transmisji odbywa się dla strumienia lub strumieni danych objętych rezerwacją zasobów identyfikowaną przez sesję i transmitowanych w jednym kierunku – od nadawcy do odbiorcy. Sesja jest określona przez adres IP, protokół transportowy i port identyfikujący odbiorcę danych. Pakiety należące do tej samej klasy usług we wszystkich węzłach sieci są obsługiwane w pewien szczególny sposób, zależny właśnie od tej klasy. Dla strumieni danych GS lub CL nie zestawia się specjalnych połączeń – są one transmitowane w trybie bezpołączeniowym tak jak w klasycznym protokole IP. Dla każdej sesji należy przydzielić odpowiednie zasoby i obsługiwać ją według parametrów QoS dla niej przypisanych, tak aby inne strumienie danych nie wpływały na jej obsługę. Przez zasoby QoS rozumiemy czas procesora i pamięć w buforze kolejki (lub kolejek), jakie zostały przydzielone do obsługi danego strumienia danych transportowanego przez sieć.



Rysunek 2-1: Obsługa pakietów w routerze IntServ

Zasoby należy przydzielić w taki sposób, aby pakiety należące do obsługiwanego strumienia były obsługiwane w wymaganym czasie (nie czekały zbyt długo w kolejce na obsługę, co ma wpływ na opóźnienie pakietów w węźle) oraz były obsługiwane z odpowiednią szybkością (pasmo przydzielone do obsługi strumienia), tak aby ilość pakietów obsłużonych w jednostce czasu i prawdopodobieństwo odrzucenia pakietów były

zgodne z umową (zwaną dalej kontraktem) pomiędzy operatorem sieci a klientem korzystającym z usług sieci.

Na rysunku 2-1 przedstawiono logiczny model routera sieci IntServ. Architektura możemy podzielić na dwie warstwy. Warstwa sygnalizacji i zarządzania odpowiedzialna jest za konfigurację węzła zgodnie z żądaniami przesyłanymi przez protokół sygnalizacyjny, w tym za rezerwację zasobów dla strumieni aplikacji. Warstwa przetwarzania pakietów danych odpowiedzialna jest za obsługę pakietów zgodną z skonfigurowanymi parametrami rezerwacji dla danego strumienia danych określonego przez parametry sesji sygnalizacyjnej. Pakiety na wejściu do routera są klasyfikowane na podstawie reguł trasowania i parametrów strumienia do odpowiednich buforów w łączach wyjściowych (dla uproszczenia na rysunku 2-1 przedstawiono tylko jedno łącze wyjściowe). Dla strumienia danych należącego do jednej sesji sygnalizacyjnej przydzielony jest dedykowany bufor, tak aby obsługa niezależnych strumieni nie wpływała na siebie nawzajem. Na wyjściu z łącza pracuje algorytm szeregowania pakietów, który pobiera pakiety z buforów i wysyła je do następnego routera w sieci; także w tym elemencie routera ważne jest, aby obsługa poszczególnych strumieni nie miała wpływu na siebie, a jednocześnie parametry obsługi były zgodne z parametrami rezerwacji dla każdego ze strumieni.

Żądania przydzielenia odpowiedniej klasy usługi sieciowej wraz z parametrami rezerwacji dla strumieni danych aplikacji mogą być przekazywane do routerów przez protokoły zarządzające siecią lub przy użyciu dedykowanych protokołów. Najbardziej znanym protokołem do rezerwacji zasobów w sieciach IntServ jest RSVP [17], który został opracowany specjalnie do sygnalizacji QoS w architekturze IntServ i jest z nią nieodłącznie kojarzony. Żądania rezerwacji określają odpowiedni poziom zasobów sieciowych, które muszą być przydzielone dla obsługi danego strumienia danych.

Każdy router w sieci zgodnej z IntServ dla prawidłowego działania całego modelu musi implementować mechanizm sterowania przyjęciem zgłoszeń (*admission control*), który pozwoli obsługiwać nowe połączenia tylko w wypadku wolnych lokalnych zasobów routera. Przy kontroli przyjęcia nowego zgłoszenia należy rozważyć informacje dostarczane z systemów zdalnych w odniesieniu do charakterystyki ruchowej, jakiej będzie podlegał nowy strumień danych. Jednym z parametrów charakteryzujących ruch jest maksymalny rozmiar paczki danych przesyłany w strumieniu, który powinien być mniejszy od wielkości pola danych w protokole łącza danych MTU (*Maximum Transmission Unit*). Model sieci Internet z integracją usług opracowany przez IETF zakłada, że tylko pakiety, które nie są fragmentowane w niższych warstwach, mogą podlegać rezerwacji parametrów QoS.

Mechanizmy QoS dla modelu *IntServ* są implementowane w każdym węźle sieci zgodnej z modelem IntServ. Każdy węzeł implementujący model *IntServ* ma taką samą architekturę logiczną, natomiast może się różnić szczegółami implementacyjnymi.

Podział ruchu na klasy umożliwia zróżnicowaną obsługę strumieni danych

pochodzących od różnych aplikacji. W modelu IntServ różnicowanie obsługi pojedynczych strumieni odbywa się dwuetapowo. Najpierw strumień przydzielony jest do jednej z trzech zdefiniowanych klas: GS, CL lub BE. Następnie pakiety w zależności od adresu odbiorcy i nadawcy oraz opcjonalnie portów są obsługiwane zgodnie z parametrami rezerwacji i właściwościami danej klasy. Strumienie danych należące do GS i CL powinny być obsługiwane w ten sposób, aby ruch generowany przez jeden strumień nie wpływał na jakość obsługi innych strumieni. Dla klasy BE wszystkie pakiety traktowane są jednakowo i dlatego ta klasa nie zapewnia żadnych gwarancji jakości obsługi. Poniżej zostaną przedstawione właściwości klasy GS i CL, które realizują gwarancje jakości obsługi.

2.1.1 Usługa gwarantowana GS

Klasa ta jest przeznaczona dla aplikacji, które wymagają transmisji w czasie rzeczywistym, np. dla programów transmitujących dane wideo i sterujące nim oraz są wrażliwe na opóźnienia w przesyłaniu pakietów z danymi.

Klasa GS zapewnia:

- gwarantowaną szerokość pasma transmisyjnego,
- ograniczenia dotyczące maksymalnego opóźnienia w relacji nadawca – odbiorca,
- brak strat pakietów spowodowanych kolejkowaniem w buforach urządzeń sieciowych.

Każdy ruter dla strumienia, który ma podlegać rezerwacji według klasy GS, przydziela pasmo R i bufor o rozmiarze B . Operacja ta jest wykonywana przez aproksymację modelu usługi typu *fluid* [17], zatem strumień efektywnie ma do dyspozycji dedykowane łącze o szerokości pasma transmisyjnego R pomiędzy nadawcą i odbiorcą. W idealnym modelu typu *fluid*, ruch generowany przez strumień danych odpowiadający przepływności¹ paczki danych o szybkości r (*token bucket rate*) i wielkości paczki danych b (*bucket depth*), ma opóźnienie graniczne b/R , pod warunkiem, że $R \geq r$. Ponieważ routery realizują tylko pewną aproksymację modelu idealnego, dla jego uzupełnienia wprowadzono dwie wielkości F i G , opisujące błędy powstałe wskutek aproksymacji. W wyniku tego graniczne opóźnienie opisywane jest przez wyrażenie $b/R + F/R + G$. Jednak w klasie usługi GS maksymalne opóźnienie jest obliczane z użyciem parametru p (*peak rate*), określającego maksymalną szybkość strumienia danych, co w efekcie powoduje redukcję opóźnienia granicznego. W dodatku, efekt podziału strumienia na pakiety wymaga, aby do obliczeń była brana pod uwagę maksymalna wielkość pakietu M . Te dodatkowe czynniki powodują, że wyrażenie ograniczające opóźnienia kolejkowania dla

¹ Przez przepływność autor rozumie natężenie strumienia danych (informacji) mierzone w bitach na sekundę.

strumienia pakietów i podlegającego rezerwacji na całej ścieżce przybiera postać [9] [18][19]:

$$D_i = \frac{(b-M)(p-R)}{R(p-r)} + \frac{(M+F_{\text{tot}})}{R} + G_{\text{tot}} \quad (2-1)$$

przy założeniu, że $p \geq R \geq r$ oraz:

$$D_i = \frac{(M+F_{\text{tot}})}{R} + G_{\text{tot}} \quad (2-2)$$

dla przypadku, gdy $R \geq p \geq r$. Zmienne F_{tot} i G_{tot} reprezentują sumę błędów F i G dla wszystkich ruterów wzdłuż ścieżki strumienia danych. Wzory (2-1) i (2-2) umożliwiają wyznaczenie maksymalnych opóźnień powstających w sieci pomiędzy terminalami wymieniającymi dane.

Podczas procesu przyporządkowania klasy usługi GS dla konkretnego strumienia danych, ruter musi być poinformowany o charakterystyce ruchu T_{spec} (*traffic characteristics*) dla tego strumienia i jednocześnie o charakterystyce rezerwacji R_{spec} (*reservation characteristics*). T_{spec} i R_{spec} są to struktury danych. Dla obliczenia zasobów lokalnych w routerze znajdującym się wzdłuż ścieżki strumienia danych, wystarczających do zagwarantowania bezstratnej usługi (bez utraty pakietów), wymagana jest znajomość parametrów F_{tot} i G_{tot} , reprezentujących sumę współczynników F i G , odpowiednio dla każdego rutera wzdłuż ścieżki pakietów od ostatniego punktu zmieniającego charakterystykę ruchu.

Usługa typu GS musi być kontrolowana w każdym punkcie sieci dla zapewnienia zgodności generowanego ruchu z T_{spec} . Zwykle procedura kontrolująca działa na zasadzie obsługi pakietów niezgodnych z T_{spec} przez klasę BE. Jeśli dostępne są procedury oznaczające pakiety, należy zaznaczyć je jako niepodlegające rezerwacji GS, tak aby inne routery także traktowały je jako ruch typu BE, bez analizy ruchu jaki one generują.

2.1.2 Usługa kontrolowanego obciążenia CL

W odróżnieniu od usługi GS, usługa kontrolowanego obciążenia CL (*Controlled Load Service*) [15] nie dostarcza tak mocnych, ilościowych gwarancji dotyczących jakości usługi transmisyjnej. Jeśli dany strumień należy do klasy CL, ruter przydziela dla niego takie zasoby, aby strumień był obsługiwany z jakością dostępną w sieci nieobciążonej nadmiernie. Wraz ze wzrostem obciążenia sieci jakość usługi CL nie powinna się pogarszać w przeciwieństwie do usługi BE. Dla porównania, w usłudze tej jakość transmisji znacznie obniża się wraz ze wzrostem obciążenia sieci. Ruter implementujący usługę CL musi być wyposażony w mechanizmy kontrolne, sprawdzające czy pakiety należące do strumienia obsługiwanego w klasie CL spełniają parametry zadeklarowane w kontrakcie ruchowym. Wszystkie pakiety niespełniające warunków rezerwacji należy

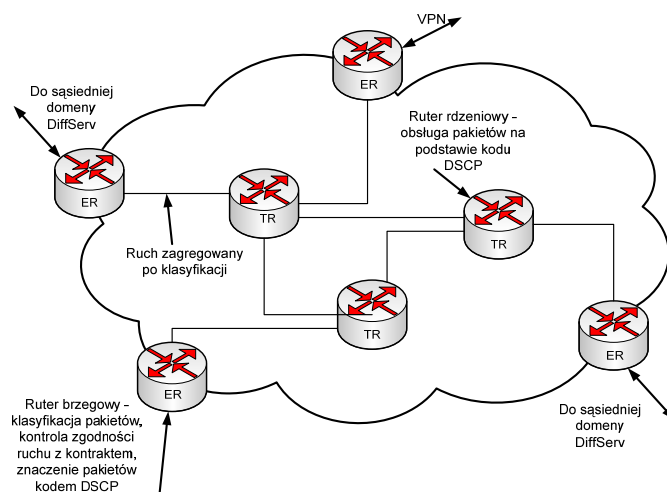
obsługiwać w klasie BE. Usługa CL jest przeznaczona dla aplikacji tolerujących pewien poziom strat pakietów i opóźnień w sieci.

2.2 Architektura usług zróżnicowanych – model *Differentiated Services*

Model *Differentiated Services* (DiffServ) [8] jest alternatywnym modelem wspierającym QoS w sieci IP opracowanym przez IETF. Główną przyczyną podjęcia prac nad tą architekturą były problemy ze skalowalnością modelu IntServ.

W modelu DiffServ głównym założeniem jest skalowalność architektury, która gwarantuje niezależność zasobów wymaganych do świadczenia usług transmisyjnych z wymaganą jakością od liczby niezależnych strumieni obsługiwanych w tej sieci i rodzaju obsługiwanego ruchu oraz rozmiaru sieci implementującej architekturę DiffServ. Założenia te zostały spełnione przez następujące rozwiązania:

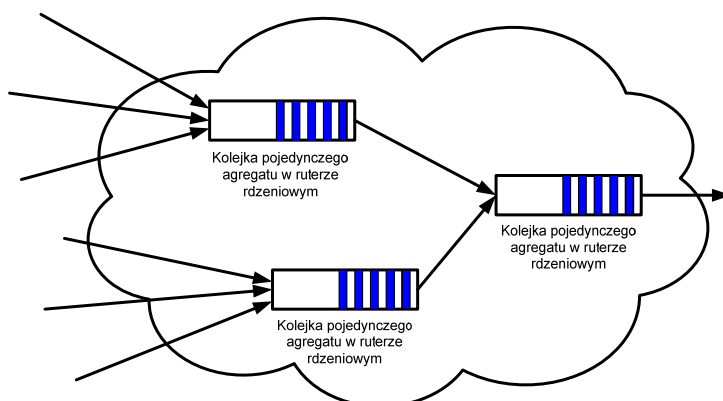
- podziału sieci implementującej model DiffServ na część brzegową i rdzeniową, z maksymalnym uproszczeniem funkcjonalności węzłów rdzeniowych,
- ograniczenie liczby klas usług maksymalnie do 64,
- rezygnacja z protokołu sygnalizacyjnego.



Rysunek 2-2: Architektura modelu DiffServ
(TR (*Transit Router*) – ruter rdzeniowy, ER (*Edge Router*) – ruter brzegowy, VPN (*Virtual Private Network*) – wirtualna sieć prywatna)

Zgodnie ze specyfikacją zawartą w dokumencie [8] architektura *Differentiated Services* bazuje na prostym modelu, w którym ruch wchodzący do sieci jest klasyfikowany i ewentualnie kształtowany zgodnie z kontraktem ruchowym oraz oznaczany jako należący do agregatu DiffServ (*DiffServ Behaviour Aggregate* – DiffServ BA) tylko na brzegu sieci w routerze brzegowym ER (rysunek 2-2).

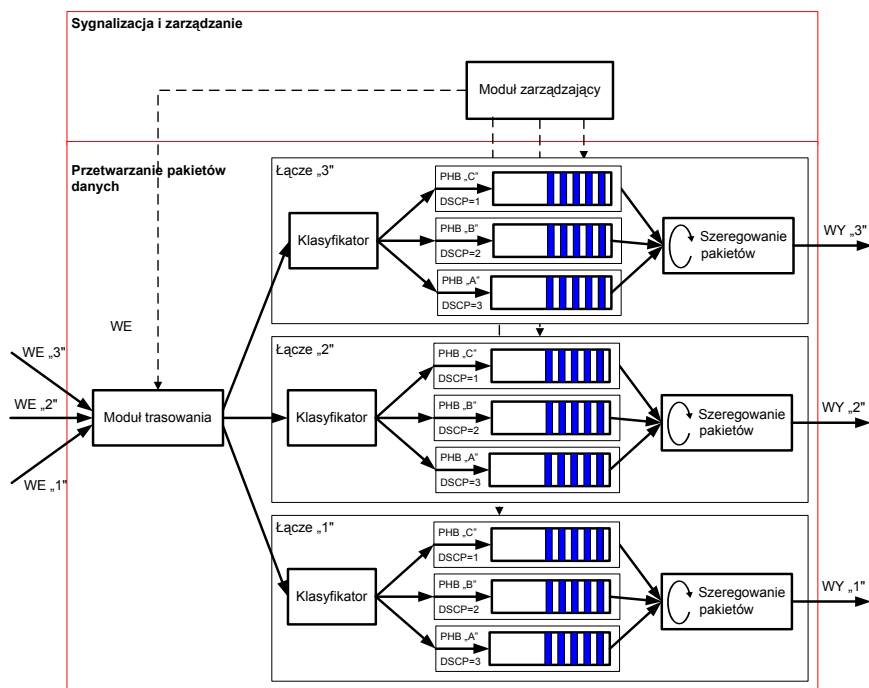
Agregat DiffServ, lub w kontekście tejże rozprawy krótko agregat, tworzą wszystkie pakiety w jednokierunkowym łączu mające tą samą wartość kodu DSCP (*DiffServ Code Point*) w nagłówku protokołu IP. Oznacza to, że gwarancja jakości dotyczy jednego kierunku transmisji, podobnie jak w modelu IntServ. Węzły rdzeniowe TR (*transit routers* nazywane w literaturze także *core routers*) klasyfikują i przekazują pakiety tylko na podstawie kodu DSCP. Oznacza to, że pakiety o tej samej wartości kodu DSCP wewnątrz sieci będą obsługiwane z tym samym poziomem jakości przez te same kolejki, tak jak to zostało schematycznie przedstawione na rysunku 2-3; każdy agregat reprezentowany przez unikalną wartość kodu DSCP jest obsługiwany przez jedną kolejkę w każdym z routerów. Algorytm szeregowania obsługujący poszczególne kolejki musi zagwarantować odpowiedni poziom obsługi dla każdej z nich, odpowiadający parametrom obsługi poszczególnych agregatów. Tego typu szeregowanie określa się jako *aggregate class-based scheduling*.



Rysunek 2-3: Obsługa pakietów typu *aggregate class-based scheduling*. Wszystkie obsługiwane strumienie należą do jednego agregatu oznaczonego jednym kodem DSCP i obsługiwanych przez jedną kolejkę w każdym z routerów

Na rysunku 2-4 przedstawiono architekturę logiczną routera rdzeniowego sieci DiffServ. Na wejściu do routera pakiety przychodzące z różnych łączy są przesyłane do odpowiednich łączy wychodzących na podstawie reguł trasowania. W łączu wychodzącym pakiety są klasyfikowane i przekazywane do odpowiednich buforów implementujących odpowiedni PHB na podstawie kodu DSCP. Na wyjściu łącza pracuje moduł zajmujący

się szeregowaniem pakietów i przesyłaniem do następnego rutera. Moduł szeregujący po pobraniu pakietu z kolejki może także buforować pakiety i umieszczać je w łączy według algorytmu FIFO lub w innej kolejności. Algorytm szeregowania musi zapewnić niezależną obsługę każdego z buforów.

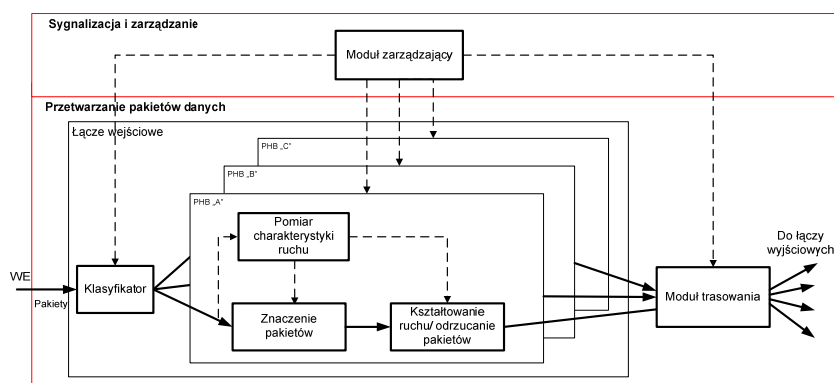


Rysunek 2-4: Architektura logiczna rutera rdzeniowego DiffServ
(PHB – Per Hop Behaviour, DSCP – DiffServ Code Point)

Na rysunku 2-5 przedstawiono moduły wymagane do przetwarzania pakietów na wejściu do rutera brzegowego DiffServ. Pakiety na wejściu do rutera są poddawane klasyfikacji i pomiarowi ruchu generowanego przez każdy strumień pakietów. Strumień pakietów na wejściu do sieci jest definiowany na podstawie umowy pomiędzy operatorem sieci DiffServ a jego klientem i może być utworzony na podstawie dowolnego pola lub zestawu pól w nagłówkach protokołów pakietów IP. Np. strumień może zostać zdefiniowany jako zbiór pakietów spełniających następujące reguły: {adres źródłowy = 149.156.114.3, protokół transportowy = TCP, port aplikacji = 22}. Następnie na podstawie umowy oraz pomiarów ruchu generowanego przez strumień pakiety są przypisywane do odpowiednich agregatów DiffServ i znaczone odpowiednim kodem DSCP. Jeśli charakterystyka ruchu generowanego przez strumień klienta wykracza poza

parametry umowy, strumień należy ukształtować w ten sposób, aby był zgodny z umową – także w umowie powinno być określone jak należy traktować pakiety wykraczające poza uzgodnioną charakterystykę ruchu, czy pakiety powinny być odrzucane, czy obsługiwane np. jak ruch BE. Dopuszcza się także możliwość, że dany strumień będą tworzyły pakiety wchodzące do węzła brzegowego różnymi łączami.

Następnie oznaczone pakiety są wysyłane do odpowiednich łączy wyjściowych na podstawie reguł trasowania. W łączach wyjściowych pakiety są obsługiwane tylko na podstawie kodów DSCP, tak jak w routerze rdzeniowym, bez uwzględnienia informacji o przynależności pakietu do danego strumienia wejściowego. Podsumowując, architektura logiczna routera brzegowego składa się z modułów przedstawionych na rysunku 2-5 w każdym z łączy wejściowych do routera oraz logiki routera rdzeniowego przedstawionego na rysunku 2-4.



Rysunek 2-5: Przetwarzanie pakietów na wejściu do routera brzegowego DiffServ

Podstawowe różnice w przetwarzaniu pakietów pomiędzy routerem w architekturze DiffServ i routerem w architekturze IntServ są następujące:

- w architekturze IntServ pakiety są klasyfikowane i obsługiwane na podstawie parametrów pojedynczych rezerwacji zestawianej przez aplikację w każdym węźle sieci, a w architekturze DiffServ klasyfikacja pakietów dokonywana jest tylko na brzegu sieci,
- w architekturze IntServ brak podziału na węzły brzegowe i rdzeniowe, co oznacza, że wszystkie węzły mają taką samą architekturę,
- w architekturze DiffServ budowa węzła rdzeniowego jest znacznie prostsza, ponieważ liczba kolejek (buforów w węźle) zależy od liczby obsługiwanych klas a nie od liczby rezerwacji jak w modelu IntServ.

2.2.1 Klasy usług w architekturze DiffServ

W architekturze *Differentiated Services* niezależne strumienie obsługiwane w tej samej klasie są agregowane i doświadczają tego samego poziomu jakości obsługi QoS. Dla każdego agregatu DiffServ w węźle sieciowym definiuje się regułę przesyłania pakietów PHB (*Per Hop Behaviour*). Reguły PHB są podstawowymi blokami, z których można tworzyć usługi sieciowe. Dotychczas zdefiniowano dwa podstawowe rodzaje PHB:

- *Expedited Forwarding* PHB [20][21][22],
- *Assured Forwarding* PHB Group [23].

Reguła *Expedited Forwarding* PHB (EF PHB) została zaprojektowana dla usług o niewielkich opóźnieniach i niewielkiej ich zmienności oraz gwarantowanym pasmie transmisyjnym w obrębie sieci DiffServ. Dla tej reguły jest zarezerwowany jeden kod DSCP. Za pomocą reguły EF PHB można realizować m.in. usługę gwarantowaną GS z modelu IntServ lub usługę postrzeganą na brzegach sieci jako dedykowane połączenie typu punkt-punkt.

Dokument [20] nie opisywał matematycznie zachowania reguły EF PHB. Ponieważ EF PHB jest przeznaczona do obsługi ruchu o najwyższej jakości, definicja w [20] nie była wystarczająca. Po wielu badaniach prowadzonych nad opóźnieniami w sieci DiffServ zostały przedstawione dwa dokumenty RFC opisujące matematycznie czasy obsługi pakietów obsługiwanych przez EF PHB. Pierwszy z nich [21] posiada status informacyjny, natomiast status rekomendacji uzyskał dokument [22].

Dokument [21] precyzuje zachowanie agregatu PHB definiując, że niezależnie od warunków panujących w sieci maksymalna różnica czasu obsługi pakietu przy przejściu przez węzeł musi być ograniczona według następującego równania:

$$d_j - e_j \leq Z \frac{MTU}{R}, \quad (2-3)$$

gdzie d_j i e_j oznaczają czasy obsługi tego samego pakietu j odpowiednio, gdy węzeł jest obciążony ruchem i w stanie nieobciążonym. Powyższe równanie nie dotyczy zmian opóźnień, jakich doznają pakiety podczas przejścia przez węzeł przy ustalonym obciążeniu węzła – nie jest to fluktuacja opóźnienia. MTU jest maksymalnym rozmiarem pakietu, jaki może być transmitowany w łączy bez fragmentacji, a R jest szybkością obsługi agregatu EF PHB w danym łączy. Współczynnik Z jest stałą charakterystyczną dla danego łączy i zależy od mechanizmów szeregowania i konfiguracji rutera. Równanie (2-3) musi być spełnione dla pakietów niewykraczających poza kontrakt ruchowy, ponieważ pakiety niespełniające warunków kontraktu powinny być odrzucane lub obsługiwane przez klasę BE.

Aktualnie obowiązująca definicja EF PHB w RFC3246 [22] wprowadza ściśle matematyczną definicję czasów obsługi pakietów w węźle implementującym ten typ

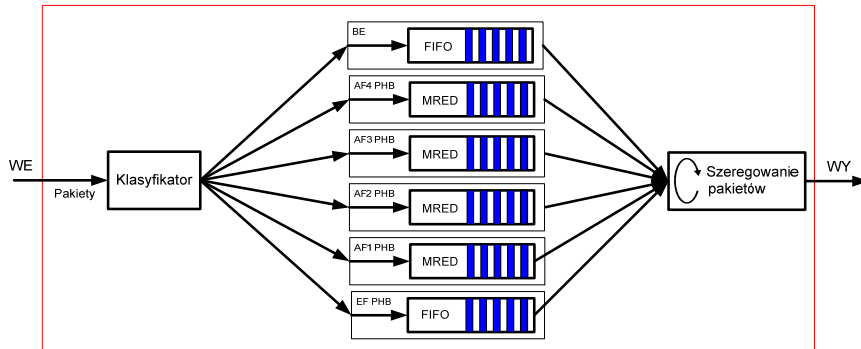
agregowanej obsługi pakietów w modelu DiffServ. Nowa formalna definicja EF PHB [22] wprowadza pojęcie gwarancji parametrów obsługi na skali pakietów – „*Packet Scale Rate Guarantee*” (PSRG). Węzeł sieciowy oferuje dla agregatu EF obsługę typu PSRG z opóźnieniem E , jeśli dla każdego pakietu j należącego do agregatu EF spełniona jest następująca nierówność:

$$t_j^{wy} \leq f_j + E \quad \text{dla wszystkich } j > 0 \quad (2-4)$$

gdzie f_j definiuje się iteracyjnie jako:

$$\begin{aligned} f_0 &= 0 & t_0^{wy} &= 0 \\ f_j &= \max\{t_j^{we}, \min(t_{j-1}^{wy}, f_{j-1})\} + \frac{m_j}{R} \end{aligned} \quad (2-5)$$

t_j^{wy} oznacza czas opuszczenia węzła sieciowego przez ostatni bit pakietu j , t_j^{we} jest czasem przybycia ostatniego bitu pakietu j do węzła sieciowego, m_j oznacza rozmiar pakietu j , natomiast R jest konfigurowaną szybkością obsługi pakietów należących do agregatu EF. E jest błędem obsługi agregatu EF i oznacza najgorszy przypadek różnicy czasu obsługi pomiędzy rzeczywistym czasem t_j^{wy} a estymowanym idealnym czasem f_j opuszczenia węzła przez pakiet j dla wszystkich pakietów należącym do agregatu EF.



Rysunek 2-6: Przykładowa konfiguracja buforów dla łącza wychodzącego w routerze rdzeniowym sieci DiffServ obsługującej EF PHB, grupę AF PHB i ruch BE (MRED – *Multi Random Early Detection*, EF – *Expedited Forwarding*, AF – *Assured Forwarding*, BE – *Best Effort*, FIFO – *First In First Out*)

Jeśli węzeł sieciowy implementuje regułę EF PHB dla agregatu strumieni, musi dla tego agregatu zapewniać obsługę typu PSRG z szybkością R i z opóźnieniem E , gdzie R jest konfigurowaną szybkością obsługi agregatu EF, a opóźnienie E jest błędem obsługi zależnym od konfiguracji sprzętowej i programowej danego węzła. Znajac wartość błędu E w każdym węźle sieci można określić maksymalne opóźnienie pakietów przy przejściu

przez sieć na podstawie wzoru (2-4).

Powyższa definicja EF PHB umożliwia budowę sieci DiffServ oferującej usługę gwarantowaną GS, która jest jedną z trzech podstawowych usług modelu IntServ, przeznaczoną do obsługi aplikacji o najwyższych wymaganiach jakościowych. Matematyczny opis parametrów umożliwia także weryfikację parametrów świadczonych usług sieciowych na podstawie pomiarów w sieci, co jest bardzo istotne z punktu widzenia obsługi i zarządzania siecią oraz parametrów technicznych zawieranych umów SLS (*Service Level Specification*) [24] pomiędzy operatorem sieci a użytkownikiem.

Grupa *Assured Forwarding* PHB (AF PHB) umożliwia zagwarantowanie różnych poziomów obsługi pakietów w obrębie domeny DiffServ. Obecnie zdefiniowano cztery klasy obsługi z trzema poziomami prawdopodobieństwa odrzucenia pakietu w każdej klasie. Tak zdefiniowany zestaw PHB dla oznaczenia pakietów wymaga dwunastu wartości kodów DSCP. Poziomy prawdopodobieństwa odrzucenia pakietu służą do nadawania priorytetów pakietom wewnątrz jednej klasy obsługi – w razie przeciążenia z buforu pierwsze są usuwane pakiety o największym prawdopodobieństwie odrzucenia. Poziom jakości obsługi w każdej klasie zależy głównie od zasobów przydzielonych dla tej klasy w węźle i jest zależny tylko od konfiguracji sieci DiffServ.

Przykładową konfigurację implementującą opisaną powyżej PHB dla łącza wychodzącego w routerze rdzeniowym sieci DiffServ przedstawia rysunek 2-6.

2.2.2 Domeny i regiony sieci DiffServ

Reguły PHB w modelu DiffServ nie są określone ilościowo, ale wyłącznie jakościowo i mogą różnić się w zależności od domeny. **Domena DiffServ** jest zbiorem węzłów implementujących architekturę DiffServ, mających te same mechanizmy PHB dla obsługiwanych agregatów DiffServ i zarządzanych według tych samych zasad. Wszystkie węzły w domenie muszą obsługiwać zagregowane strumienie danych w ten sam sposób. Domena ma dobrze zdefiniowany brzeg składający się z węzłów brzegowych, które klasyfikują pakiety, sprawdzają zgodność charakterystyk ruchowych ze specyfikacją SLS (*Service Level Specification*) [24] i ewentualnie kształtują te charakterystyki, tak aby były zgodne z SLS oraz oznaczają pakiety odpowiednimi wartościami w polu DSCP. Klasyfikacja i obsługa pakietów w węzłach wewnątrz domeny odbywa się tylko na podstawie wartości pola DSCP.

Domeny w sieci są niezależne od siebie, ale mogą ze sobą współpracować. Obszar sieci składający się z pewnej liczby domen ściśle współpracujących ze sobą jest nazywany **regionem DiffServ**.

Zbiór tych samych reguł PHB i reguł klasyfikacji strumieni na brzegu sieci dla jednego agregatu w domenie DiffServ tworzy PDB (*Per Domain Behaviour*) [25]. Przez PDB rozumie się poziom QoS usługi transmisji danych przy przejściu pakietów należących do agregatu przez domenę DiffServ. Dotychczas zostało zaproponowanych

kilka specyfikacji PDB, jak na przykład:

- *Virtual Wire* PDB zbudowany w oparciu o EF PHB [26],
- *Assured Rate* PDB korzystający z AF PHB [27],

które jednak nie stały się oficjalnymi standardami organizacji IETF. Obecnie jedyną udaną próbą zdefiniowania PDB jest dokument RFC 3663 [28] określający zasady obsługi ruchu o priorytecie niższym niż ruch *Best-Effort*.

W podsumowaniu należy stwierdzić, że nie udało się zunifikować usług dla architektury DiffServ korzystających z EF i AF PHB, czyli dla ruchu, który z zasady jest przeznaczony do obsługi ruchu o wyższym priorytecie niż *Best-Effort*. Oznacza to, że definicja usług, czyli definicja PDB, w domenie operatora zależy od jego polityki wewnętrznej.

W modelu *DiffServ* nie ma zdefiniowanego protokołu sygnalizacyjnego. Zasady klasyfikacji strumieni danych i ich charakterystyki ruchowe są konfigurowane w węzłach domeny przez narzędzia do zarządzania siecią lub protokoły sygnalizacyjne nieobjęte pracami nad architekturą sieci DiffServ. Parametry obsługi danych są konfigurowane na podstawie umowy SLA (*Service Level Agreement*) [8] pomiędzy klientem a administracją sieci, a dokładnie rzecz biorąc na podstawie specyfikacji SLS, która jest częścią SLA.

2.3 Opóźnienia maksymalne w domenie DiffServ dla agregatów EF

Bieżący rozdział poświęcony jest krytycznej analizie dotychczasowych badań nad opóźnieniami maksymalnymi w domenie sieci DiffServ dla agregatów strumieni danych obsługiwanych przez EF PHB. EF PHB został zaprojektowany dla aplikacji o najwyższych wymaganiach jakościowych, reprezentowanych przez kategorię usług konwersacyjnych z tabeli 1-1. Najważniejszym parametrem sieci pakietowych dla takich usług jest opóźnienie transmisyjne i zmienność tego opóźnienia. Opóźnienie transmisyjne dzieli się na stałe opóźnienie propagacji w łączach fizycznych i zmienne opóźnienie powstające w buforach urządzeń sieciowych. Zmienność opóźnienia określana jest przez różnicę pomiędzy opóźnieniem maksymalnym i minimalnym. Zalecenia dotyczące definiowania i wyznaczania zmienności opóźnienia można znaleźć w rekomendacji ITU-T Y.1541 [29]. Oczywiście, bardzo ważnym parametrem jakości usług sieciowych jest poziom strat pakietów, ale dla EF PHB zakłada się, że straty są bardzo małe i nie mają wpływu na działanie aplikacji.

Ponieważ opóźnienia mają bardzo ważny wpływ na działanie aplikacji sieciowych, zalecenie G.114 [30] wprowadza rekomendacje dotyczące opóźnień maksymalnych dla aplikacji, których jakość postrzegana przez użytkownika zależy od właściwości interaktywnych aplikacji. Opóźnienia maksymalne powinny być uwzględniane na etapie

planowania sieci i mają następujące wartości:

- od 0 do 150 ms: akceptowalne dla większości aplikacji użytkownika,
- 150 do 400 ms: akceptowalne pod warunkiem, że operator (sieci lub usługi) jest świadomy wpływu czasu transmisji na jakość transmisji aplikacji użytkownika,
- powyżej 400 ms: nieakceptowane dla celów planowania sieci.

Najprostszym przykładem aplikacji, dla której należy uwzględnić powyższe ograniczenia opóźnień jest telefonia IP.

W następnych podrozdziałach zostaną przedstawione wzory opisujące maksymalne opóźnienie w domenie DiffServ dla EF PHB wraz z ich omówieniem oraz sposoby ograniczenia tych opóźnień przez wprowadzanie nowych mechanizmów obsługi pakietów w węzłach.

2.3.1 Szacowanie deterministycznych opóźnień maksymalnych w sieci DiffServ

Pierwsze wyniki prac związanych z szacowaniem deterministycznych opóźnień maksymalnych w domenie DiffServ zostały przedstawione przez A. Charny [11] i J.-Y. Le Boudeca [31]. W ich wspólnej pracy [31] autorzy wyznaczyli deterministyczne ograniczenie na maksymalne opóźnienie pakietów pomiędzy dowolnymi ruterami brzegowymi domeny DiffServ obsługiwanych przez EF PHB. Prace A. Charny i J.-Y. Le Boudeca [12] znalazły również swoje odbicie w matematycznym opisie EF PHB [22].

Poniżej zostanie przedstawiony model sieci i wzór na obliczanie opóźnień maksymalnych dla domeny DiffServ przedstawiony w pracy [31].

Zakładamy, że dana jest sieć, w której wszystkie węzły są urządzeniami z buforami na wyjściu i implementującymi dyscyplinę szeregowania pakietów typu *Strict Priority* dla klas strumieni. Dyscyplina szeregowania *Strict Priority* obsługuje pakiety należące do klasy o niższym priorytecie tylko wtedy, gdy nie ma pakietów o wyższym priorytecie, a pakiety należące do jednej klasy są obsługiwane w kolejności przybycia do kolejki. Poziom jakości obsługi zależy od przynależności do klasy usług. Pakiety należące do tej samej klasy są buforowane w pojedynczej kolejce i są obsługiwane niezależnie od pakietów należących do innych klas. Pakiety do domeny sieci DiffServ wchodzi w ruterach brzegowych typu *ingress* – IGR, natomiast opuszczają sieć w ruterach brzegowych nazwanych *egress* – EGR. Pakiety wchodzące w tym samym węźle sieci IGR oraz opuszczające sieć przez ten sam węzeł EGR i należące do tej samej klasy tworzą strumień danych [31]. Tak zdefiniowany strumień danych może składać się z niezależnych strumieni danych od różnych aplikacji. Zakłada się, że w sieci została zaimplementowana obsługa przynajmniej dwóch agregatów. Przedstawione wzory dotyczą parametrów obsługi agregatu o najwyższym priorytecie, który może reprezentować parametry obsługi dla usług zbudowanych z użyciem EF PHB.

Zakładamy, że każdy strumień i przechodzący przez domenę DiffServ na wejściu do domeny jest zgodny z parametrami ciekącego wiadra (r_i, b_i) , gdzie r_i oznacza szybkość bitową strumienia, natomiast b_i jest rozmiarem wiadra algorytmu ciekącego wiadra dla strumienia i . Załóżmy, że O_l jest zbiorem wszystkich strumieni przechodzących przez łącze l oraz L jest zbiorem wszystkich łączy l w rozważanej domenie DiffServ. Zakładamy, że ruch wchodzący, obsługiwany z najwyższym priorytetem w każdym z łączy l , nie przekracza założonego współczynnika wykorzystania pasma α , co oznacza, że $\sum_{i \in O_l} r_i \leq \alpha R_l$, gdzie R_l jest szybkością transmisji gwarantowaną dla klasy o najwyższym priorytecie w łączy l , zwaną także szybkością obsługi. Dla algorytmu szeregowania *Strict Priority* R_l jest równe przepustowości łącza² C_l . Wymaga się także, aby istniało ograniczenie na wartość sumy rozmiarów wiader dla wszystkich strumieni przechodzących przez dane łącze l , a więc dla wszystkich łączy $l \in L$ istnieje β spełniające następującą nierówność:

$$\beta \geq \frac{\sum_{i \in O_l} b_i}{R_l}. \quad (2-6)$$

β interpretujemy jako maksymalny czas potrzebny na wysłanie wszystkich pakietów, które znajdują się w kolejce EF w dowolnej chwili pracy sieci, w dowolnym węźle domeny, ale pod warunkiem, że każdy z obsługiwanych strumieni jest zgodny z kontraktem.

Oznaczmy przez H maksymalną liczbę węzłów, przez które przechodzi pakiet pomiędzy dwoma dowolnymi ruterami brzegowymi. W tak zdefiniowanym systemie ograniczenie na maksymalne opóźnienie pakietu przy przejściu pomiędzy brzegami sieci wynosi

$$D = \frac{H}{1 - (H - 1)u\alpha} (E + u\beta) \quad (2-7)$$

pod warunkiem, że

$$\alpha < \min_l \frac{P_n}{(P_n - R_l)(H - 1) + R_l}. \quad (2-8)$$

P_n oznacza maksymalną przepływność ruchu wchodzącego do węzła n dla całej klasy ruchu o najwyższym priorytecie. E jest wyrażone za pomocą następującego wzoru:

² Przez przepustowość łącza C_l autor rozumie maksymalną ilość informacji mierzonej w bitach, jaka może być przesyłana przez łącze w jednostce czasu.

$$E = \max_l \frac{M_l}{R_l}. \quad (2-9)$$

M_l jest maksymalnym rozmiarem pakietu w łączu l spełniającym zależność $M_l \leq MTU$, gdzie MTU (*Maximum Transfer Unit*) jest maksymalnym rozmiarem pakietu, który może zostać przeniesiony przez łącze bez fragmentacji. Zmienna u określona jest przez następujące równanie:

$$u = \max_l \{u_l\}, \quad (2-10)$$

gdzie $u_l = \frac{P_n - R_l}{P_n - \alpha R_l}$.

Y. Jiang w swojej pracy [13] rozwinął teorię przedstawioną w [31] na wszystkie algorytmy kolejkowania należące do klasy Gwarantowanej Szybkości GR (*Guaranteed Rate*) – równanie (2-11) – oraz uwzględnił efekt wpływu podziału strumienia danych na pakiety przy definicji krzywej opisującej ruch wejściowy. W pracy [31] ruch wejściowy był traktowany jako strumień ciągły – model typu „fluid”.

Algorytm szeregowania należy do klasy Gwarantowanej Szybkości GR [32] z szybkością obsługi R i opóźnieniem E dla danego strumienia danych, jeśli gwarantuje, że pakiet j strumienia jest obsługiwany w czasie nie dłuższym niż

$$t_j = GR(j) + E, \quad (2-11)$$

gdzie E jest stałą i zależy od algorytmu szeregowania, a GR jest zdefiniowane rekurencyjnie jako:

$$GR(0) = 0$$

$$GR(j) = \max\{t_j^{wc}, GR(j-1)\} + \frac{m_j}{R}, \quad (2-12)$$

Przy czym m_j jest rozmiarem pakietu j , natomiast t_j^{wc} jest czasem przybycia pakietu j do węzła. Definicję GR spełnia wiele algorytmów szeregowania – najbardziej znane to VC (*Virtual Clock*) [33], WFQ (*Weighted Fair Queuing*) czy SCFQ (*Self Clocked Fair Queuing*) [34]. Tabele opisujące szybkość obsługi R i opóźnienie E dla wybranych algorytmów szeregowania zostały zamieszczona w [13] i [32]. W pracy [31] zostało udowodnione, że wyniki uzyskane przy zastosowaniu algorytmów szeregowania spełniających definicję klasy GR mają także zastosowanie do definicji EF PHB w [22] zdefiniowanej wzorem (2-4).

Y. Jiang w pracy [13] posługując się takim samym modelem sieci jak dla wzoru (2-7), udowodniono, że maksymalne opóźnienie w domenie DiffServ dla pakietów obsługiwanych przez EF PHB wynosi:

$$D = \frac{H}{1 - (H-1)u\alpha} (E' + u\beta), \quad (2-13)$$

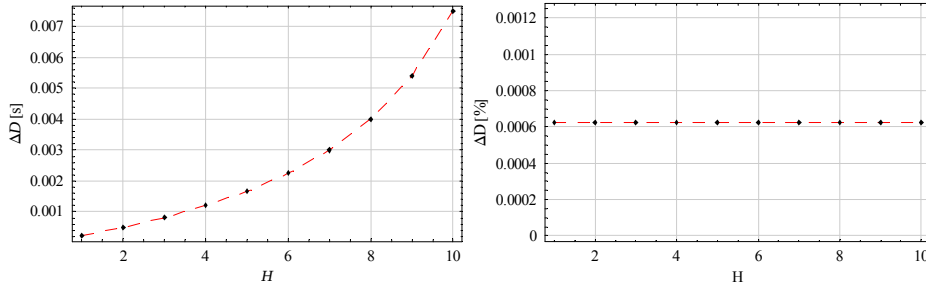
gdzie

$$E' = \max_i \left\{ \frac{(1-u_i)M_n}{R_i} + \frac{M_l}{R_i} + E_l \right\}. \quad (2-14)$$

E_l jest błędem zależnym od wyboru algorytmu kolejkowania, przykładowo dla *Strict Priority* $E_l = \frac{M_l}{C_l}$, gdzie C_l jest przepustowością łącza l . $M_n = \sum_{i=1}^k M_i$, gdzie k jest liczbą łączy wchodzących do węzła, z którego wychodzi łącze l . Pozostałe zmienne są zdefiniowane tak samo jak w równaniu (2-7). Wzór (2-13) jest prawdziwy dla

$$\alpha < \min_i \frac{P_n}{(P_n - R_i)(H-1) + R_i}. \quad (2-15)$$

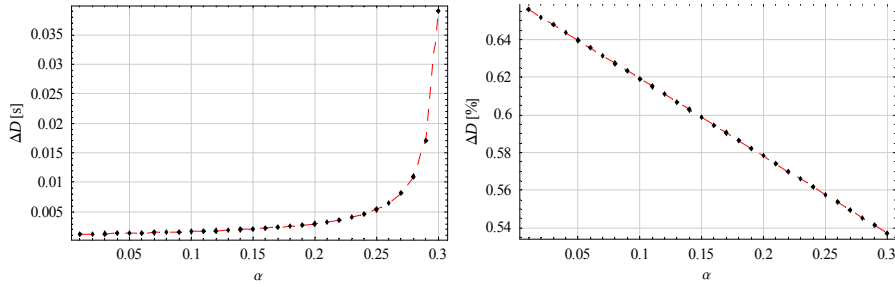
Wzór na opóźnienie maksymalne (2-13) wyprowadzony przez Y. Jianga daje dokładniejsze oszacowanie opóźnień maksymalnych niż wzór (2-7) zdefiniowany przez J.-Y Le Boudeca, ponieważ uwzględnia pakietowy charakter ruchu wejściowego. Wzór (2-13) posiada także szersze zastosowanie niż wzór (2-7), ponieważ można go stosować dla wszystkich algorytmów należących do klasy GR.



Rysunek 2-7: Porównanie opóźnień maksymalnych dla równań (2-7) i (2-13) w zależności od rozmiaru sieci H (lewy wykres przedstawia różnicę bezwzględną w sekundach, prawy wykres przedstawia różnicę względną) (wartości parametrów sieci przedstawia tabela 2-1)

Na rysunkach 2-7 i 2-8 autor przedstawia porównanie równań (2-7) i (2-13) dla parametrów przedstawionych w tabeli 2-1. Maksymalne opóźnienie jest większe dla wzoru (2-13). Różnica pomiędzy wartościami maksymalnego opóźnienia zwiększa się w funkcji H dla wartości bezwzględnych, natomiast względna wartość różnicy między opóźnieniami obliczona ze wzoru $\Delta D(H) = \frac{D(H)_{(2-13)} - D(H)_{(2-7)}}{D(H)_{(2-13)}} 100\%$ jest stała i nie

zależy od H – prawy wykres na rysunku 2-7.



Rysunek 2-8: Porównanie opóźnień maksymalnych dla równań (2-7) i (2-13) w zależności od współczynnika wykorzystania pasma α (lewy wykres przedstawia różnicę bezwzględną w sekundach, prawy wykres przedstawia różnicę względną) (wartości parametrów sieci zostały przedstawione w tabeli 2-1)

Różnica opóźnień pomiędzy rozważanymi wzorami w zależności od α przedstawiona jest na rysunku 2-8. Wraz ze wzrostem współczynnika wykorzystania pasma α bezwzględna różnica opóźnień rośnie, natomiast względna różnica zmniejsza swoją wartość. Względna różnica pomiędzy opóźnieńmi w funkcji α została wyznaczona ze wzoru $\Delta D(\alpha) = \frac{D(\alpha)_{(2-13)} - D(\alpha)_{(2-7)}}{D(\alpha)_{(2-13)}} 100\%$.

Tabela 2-1 Parametry wykresów dla rysunków 2-7 i 2-8

Parametr / rysunek	R [b/s]	C_l [b/s]	M_l [B]	β [s]	M_n [B]	P_n [b/s]	H	α	E_l [s]
Rysunek 2-7	1×10^8	1×10^8	1500	0.048	6000	4×10^8	-	0.1	1.2×10^{-4}
Rysunek 2-8	1×10^8	1×10^8	1500	0.048	6000	4×10^8	5	-	1.2×10^{-4}

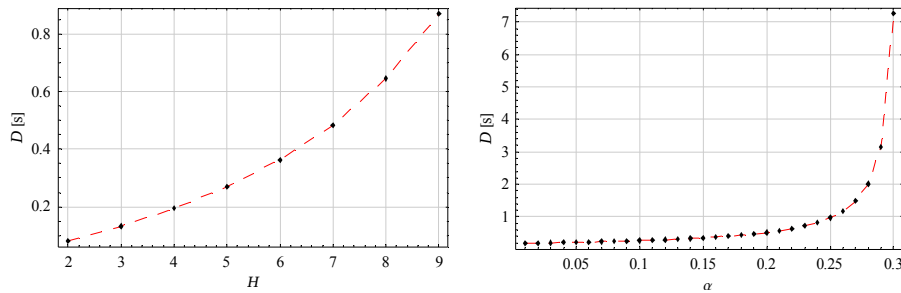
Znajomość maksymalnych opóźnień przejścia pakietu przez domenę DiffServ dla EF PDB umożliwia implementację usługi *Guaranteed Service* [9] w domenie DiffServ bez konieczności implementacji dodatkowych mechanizmów w węzłach sieci.

W pracy [31] zostało udowodnione, że dla wartości współczynnika wykorzystania pasma α większej od wartości dopuszczalnej (wyznaczonej wzorem (2-8)) opóźnienie maksymalne może być nieskończenie duże. Przyczyną tego stanu jest możliwość „sklejania” się pakietów i powstawania dużych paczek pakietów w sieci, jeśli stosujemy szeregowanie typu FIFO dla agregatów strumieni. W sieci DiffServ agregaty z węzła na

węzeł mogą mieć różną strukturę (składać się z różnych strumieni) i nawet w przypadku, gdy suma szybkości strumieni wchodzących w skład agregatu nie przekracza szybkości obsługi tego agregatu w każdym łączu sieci, opóźnienie maksymalne przy przejściu pakietów przez sieć może być nieskończenie duże. Dla szeregowania typu FIFO obsługującego agregaty strumieni danych na opóźnienie pakietów w danym węźle mają wpływ także te pakiety, które nie dzielą wspólnie łącza, a nawet pakiety, które w danej chwili opuściły już sieć.

2.3.2 Metody zmniejszania opóźnień maksymalnych w sieci DiffServ

W poprzednim podrozdziale zostały przedstawione matematyczne metody szacowania opóźnień maksymalnych w sieci o architekturze DiffServ dla EF PHB oraz warunki, dla których opóźnienia maksymalne istnieją. Z wzorów (2-8) i (2-15) wynika, że tylko część pasma przydzielonego dla EF PHB może zostać wykorzystana do transmisji danych, ponieważ jak to zostało udowodnione w [31], jeśli współczynnik wykorzystania pasma α nie spełnia warunków określonych równaniem (2-8) lub (2-15), to opóźnienie maksymalne może być nieskończenie duże.

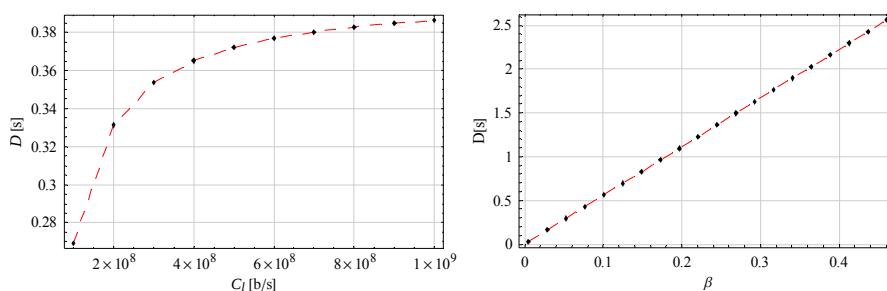


Rysunek 2-9: Opóźnienia maksymalne dla sieci opisanej parametrami z tabeli 2-1 wyznaczone przez równanie (2-13) w zależności od rozmiaru sieci H (lewy wykres – parametry jak dla rysunku 2-7) i współczynnika wykorzystania pasma α (prawy wykres – parametry jak dla rysunku 2-8)

Na rysunku 2-9 autor przedstawia zależności opóźnienia maksymalnego (wzór (2-13)) od rozmiaru sieci H i współczynnika wykorzystania pasma α dla parametrów sieci z tabeli 2-1. Opóźnienia maksymalne rosną wykładniczo dla obydwu rozpatrywanych parametrów oraz są stosunkowo duże.

Dla wzoru (2-13) opóźnienia maksymalne rosną także w przypadku zwiększania parametru β (lewy wykres na rysunku 2-10), parametru k (rysunek B-6), maksymalnego rozmiaru pakietu w łączu M_l (rysunek B-7). O ile zwiększenie opóźnienia maksymalnego nie budzi wątpliwości dla przypadków wymienionych powyżej, to wzrost opóźnienia maksymalnego w przypadku zwiększenia przepustowości łącza C_l (prawy wykres na

rysunku 2-10) może takie wątpliwości wzbudzić. Dzieje się tak, ponieważ wraz ze wzrostem przepustowości łącza do węzła może napłynąć większa liczba pakietów w jednostce czasu, co potencjalnie wpływa na wzrost wybuchowości ruchu na wejściu do węzła. Natomiast opóźnienie maksymalne można zmniejszyć zwiększając dostępną szybkość obsługi R dla agregatu obsługiwane przez EF PHB (rysunek B-5). Powyższe właściwości opóźnienia maksymalnego dla EF PHB w sieci DiffServ powodują, że ważne stają się rozwiązania mające na celu ograniczenie tych opóźnień, nawet za cenę wprowadzenia nowych mechanizmów, natomiast ważne jest, aby te dodatkowe mechanizmy, zgodnie z architekturą DiffServ, działały nie na pojedynczych strumieniach, ale na agregatach strumieni danych.



Rysunek 2-10: Zależność opóźnienia maksymalnego dla równania (2-13) w zależności od przepustowości łącza C_l (lewy rysunek) i parametru β (prawy rysunek) dla parametrów sieci przedstawionych w tabeli 2-1

Istnieje kilka propozycji zmniejszania maksymalnych opóźnień w sieciach zgodnych z architekturą DiffServ, w których obsługę ruchu wewnątrz sieci prowadzi się dla agregatów strumieni danych, a tylko na brzegu sieci prowadzi się sterowanie przyjęciem zgłoszeń i sprawdzenie zgodności parametrów strumieni z kontraktem ruchowym. Propozycje te dotyczą głównie budowy bardziej rozbudowanych algorytmów kolejkowania i szeregowania zgłoszeń.

Jednym z takich pomysłów jest koncepcja algorytmu szeregowania pakietów *Static Earliest Time First* (SETF) rozwijana przez Z.-L. Zhang [35][36]. Polega ona na wpisywaniu do pakietów wchodzących do sieci informacji o czasie przybycia pakietu do węzła brzegowego. Wewnątrz sieci, w ruterach rdzeniowych, pakiety należące do agregatu EF są obsługiwane w kolejności wyznaczonej przez informacje o czasie obsługi zapisane w pakiecie; pakiety, które przybyły do sieci wcześniej, są obsługiwane w pierwszej kolejności. Dla takiego systemu maksymalne opóźnienie w sieci jest ograniczone zgodnie ze wzorem

$$D = (E + \beta) \frac{1 - (1 - \alpha)^H}{\alpha(1 - \alpha)^{H-1}}. \quad (2-16)$$

Równanie (2-16) obowiązuje dla współczynnika wykorzystania pasma $\alpha < 1$, H jest maksymalną liczbą węzłów jakie musi pokonać pakiet przy przejściu przez sieć, E jest opóźnieniem algorytmu szeregowania klasy GR, natomiast parametr β jest zdefiniowany równaniem (2-6).

Równanie (2-16) jest prawdziwe, jeśli informacja o czasie przybycia pakietu do sieci jest zapisana w pakiecie z nieograniczoną dokładnością. W rzeczywistej sieci informacje o czasie przybycia można zapisywać w pakiecie tylko ze skończoną dokładnością i dla takiego przypadku maksymalne opóźnienie ma wartość leżącą w przedziale pomiędzy wartościami wyznaczonymi przez wzory (2-7) i (2-16)[35]. Przy implementacji algorytmu SETF należy także pamiętać o wymaganej zmianie formatu pakietu IP w celu dołożenia pola dla przenoszenia informacji o czasie oraz o synchronizacji czasu we wszystkich węzłach sieci.

Następnym rozwiązaniem umożliwiającym osiągnięcie mniejszych maksymalnych opóźnień w sieci DiffServ jest koncepcja „inteligentnych” buforów (*damper*) [37][38]. Jest to rozszerzenie wcześniej opisywanego rozwiązania. Polega ona na kompensacji zmiany opóźnienia powstającego w danym węźle w węźle następnym, przez który przechodzą pakiety. Dla tego rozwiązania zastosowanego dla agregatu obsługiwane przez EF PHB maksymalne opóźnienie pomiędzy brzegami sieci jest mniejsze niż w przypadku ogólnym przedstawionym wcześniej i jego wartość jest skończona dla wszystkich współczynników wykorzystania pasma mniejszych od 1.

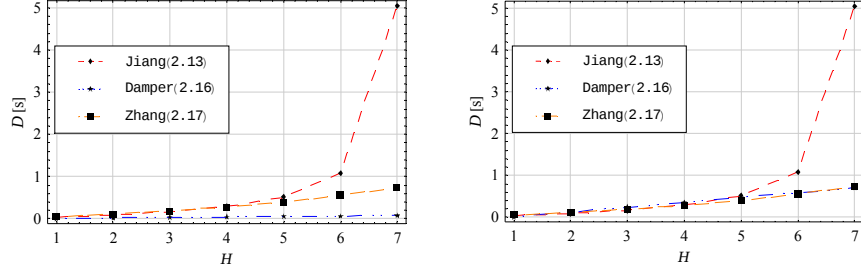
W skład systemu obsługi pakietów wchodzi inteligentny bufor i odpowiedni mechanizm szeregowania pakietów. Do poprawnego działania opisywanego algorytmu, na wyjściu z węzła sieciowego, wymagany jest mechanizm szeregowania typu EDF (*Earliest Deadline First*). EDF obsługuje pakiety na podstawie znaczników czasu skojarzonych z pakietami znajdującymi się w kolejce, nazwanych czasem życia pakietów (*deadline*). Sposób przydzielania tych znaczników może być różny i zależy od implementacji algorytmu EDF. W algorytmie szeregowania najpierw są obsługiwane pakiety z najmniejszą wartością znacznika czasu. Zakładamy, że każdy pakiet przenosi informacje o różnicy *ded* pomiędzy czasem życia pakietu a jego aktualnym czasem emisji (wyjścia z kolejki), jeśli ten czas jest większy od zera, a zero w przeciwnym wypadku. Inteligentny bufor, pracujący na wejściu do węzła sieciowego, określa czas, po którym pakiet będzie dostępny dla mechanizmu szeregowania pakietów z przedziału $[t^a + ded - \Delta, t^a + ded]$, gdzie Δ jest stałą inteligentnego bufora, a t^a określa czas przybycia pakietu do węzła. Pakiety są przetrzymywane w inteligentnym buforze zgodnie z wyznaczonym czasem dostępności a następnie przekazywane do mechanizmu szeregowania. Przy zastosowaniu mechanizmu inteligentnego bufora w każdym węźle sieci ograniczenie na maksymalne

opóźnienie w domenie DiffServ dane jest następującym wzorem:

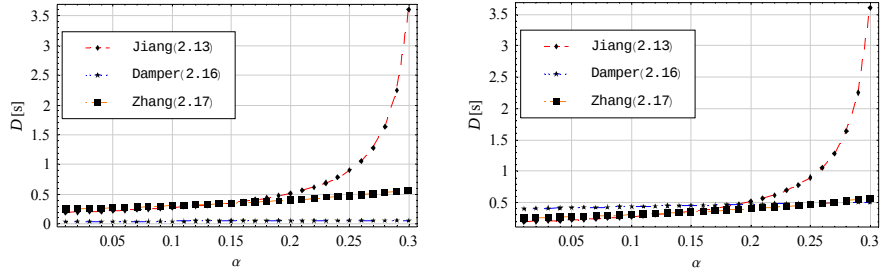
$$D = E + (H-1)\Delta(1+u\alpha) + u\beta\alpha. \quad (2-17)$$

Wzór obowiązuje dla każdego $0 < \alpha \leq 1$. E jest zdefiniowane tak jak we wzorze (2-9), u zgodnie ze wzorem (2-10), natomiast β dla wszystkich łączy l musi spełniać warunek zdefiniowany we wzorze (2-6).

Główną zaletą zastosowania inteligentnego bufora jest istnienie maksymalnego opóźnienia w całym zakresie współczynnika wykorzystania pasma $0 < \alpha \leq 1$. Dla małych wartości stałej Δ opóźnienia maksymalne są mniejsze w całym zakresie badanych parametrów od wartości wyznaczonych wzorem (2-13).



Rysunek 2-11: Porównanie opóźnień maksymalnych dla równań (2-13), (2-17) i (2-16) w zależności od rozmiaru sieci H dla $\Delta = 0.01$ (lewy wykres) i $\Delta = 0.1$ (prawy wykres) (pozostałe parametry zostały przedstawione w tabeli 2-2)



Rysunek 2-12: Porównanie opóźnień maksymalnych dla równań (2-13), (2-17) i (2-16) w zależności od współczynnika wykorzystania pasma α dla $\Delta = 0.01$ (lewy wykres) i $\Delta = 0.1$ (prawy wykres) (pozostałe parametry zostały przedstawione w tabeli 2-2)

Na rysunkach 2-11 i 2-12 autor przedstawia porównanie równań (2-13), (2-16) i (2-17). Dla uproszczenia przyjęto sieć zbudowaną z węzłów o takich samych parametrach z zastosowanym algorytmem szeregowania priorytetowego (*Strict Priority*) dla EF PHB, czyli przyjęto, że całe dostępne pasmo jest przydzielone dla EF PHB; inne agregaty są obsługiwane tylko wtedy, gdy kolejka EF jest pusta. Zakresy wartości na osiach odczytanych

są ograniczone przez warunek zapewniający istnienie maksymalnego opóźnienia dla wzoru (2-13).

Rysunek 2-11 przedstawia wartość opóźnienia maksymalnego w domenie DiffServ w zależności od maksymalnej liczby węzłów H , przez które przechodzi pakiet pomiędzy dwoma dowolnymi ruterami brzegowymi i dla dwóch różnych wartości współczynnika Δ . Na rysunku 2-12 przedstawiono zależności wartości opóźnienia maksymalnego w domenie DiffServ od współczynnika wykorzystania pasma α dla tych samych wartości współczynnika Δ , co poprzednio. Opóźnienie wyznaczone przy pomocy równania (2-13) jest mniejsze od opóźnień wyznaczonych ze wzoru (2-16) (dla odpowiednich wartości współczynnika Δ także dla wzoru (2-17)) dla małych wartości współczynnika wykorzystania pasma α , ale dla wyższych wartości tego współczynnika następuje bardzo szybki wzrost opóźnień, który nie jest obserwowany dla równań (2-16) i (2-17). Ponadto wzory (2-16) i (2-17) są prawdziwe dla $0 < \alpha \leq 1$, co nie zostało pokazane na rysunkach.

Tabela 2-2: Parametry wykresów dla rysunków 2-11 i 2-12

Parametr / rysunek	R_l [b/s]	C_l [b/s]	M_l [B]	P_n [b/s]	β [s]	M_n [B]	E_l [s]	H	α
Rysunek 2-11	1×10^7	1×10^7	1500	4×10^7	0.048	6000	1.2×10^{-3}	-	0.2
Rysunek 2-12	1×10^7	1×10^7	1500	4×10^7	0.048	6000	1.2×10^{-3}	5	-

2.4 Zakres rozprawy

W rozdziale 2.3.1 autor przedstawił podstawowe wzory dotyczące szacowania opóźnień maksymalnych dla reguły przesyłania EF PHB w sieci DiffServ. Obsługa pakietów należących do agregatu obsługiwane przez EF PHB w węzle sieci DiffServ została zdefiniowana przez rekomendacje IETF ścisłym wzorem matematycznym (2-4). Wzór matematyczny jest wymagany, ponieważ usługa sieciowa budowana w oparciu o EF PHB powinna umożliwić działanie aplikacji wymagających gwarancji na maksymalne opóźnienia transmisyjne i brak strat przesyłanych pakietów. Definicja gwarantuje, że w pojedynczym węzle implementującym EF PHB maksymalny czas obsługi pakietu będzie ograniczony i niezależny od obsługi innych pakietów należących do EF PHB oraz do innych PHB zaimplementowanych w danym węzle. J.-Y. Le Boudec i Y. Jiang w swoich pracach zdefiniowali maksymalne opóźnienie dla usługi zbudowanej z użyciem EF PHB dla domeny DiffServ, dla której ruch jest kontrolowany tylko na wejściu, natomiast wewnątrz domeny poziom obsługi zależy tylko od przynależności pakietów do danego

agregatu PHB. Wzory wyprowadzone przez wyżej wymienionych autorów wyznaczają maksymalne opóźnienie w sieci DiffServ dla EF PHB i jednocześnie wprowadzają warunek na maksymalne obciążenie łączy w domenie DiffServ, który musi być spełniony, aby usługę korzystającą z EF PHB można było zaimplementować w sieci DiffServ; wzory (2-7) i (2-13) definiujące maksymalne opóźnianie są prawdziwe tylko wtedy, gdy obciążenie łączy w całej domenie spełnia warunek dany wzorem (2-8).

W rozdziale 2.3.2 autor przedstawia dwa znane z literatury rozwiązania zmniejszające opóźnienia maksymalne dla EF PHB w sieci DiffServ, zachowujące szeregowanie pakietów dla klas ruchu. Jednak mają one poważne ograniczenia implementacyjne: wprowadzają one dodatkowe mechanizmy, które znacznie wpływają na złożoność całego systemu z powodu dodania nowego pola do pakietu IP oraz wymagają we wszystkich węzłach sieci synchronizacji czasu z dużą dokładnością.

Autor postawił przed sobą zadanie analizy wpływu algorytmu kształtowania ruchu na maksymalne opóźnienia w sieci DiffServ dla agregatów EF. Wprowadzenie kształtowania zgodnego z algorytmem ciekącego wiadra we wszystkich węzłach domeny DiffServ dla reguł przesyłania pakietów EF PHB, może wprowadzić dodatkowe opóźnienia, ale nie pozwoli na nadmierne „sklejanie” się pakietów w sieci według mechanizmu opisanego w [31], a tym samym spowoduje zmniejszenie maksymalnych opóźnień pomiędzy brzegami domeny sieci DiffServ. Proces „sklejania” się pakietów w sieci DiffServ powoduje znaczne zwiększenie wybuchowości obsługiwanego ruchu.

W odróżnieniu od sposobów ograniczenia opóźnień maksymalnych przedstawionych w rozdziale 2.3.2 zastosowanie kształtowania ruchu wyłącznie dla agregatu EF nie wpłynie znacząco na dodatkowe obciążenie obliczeniowe ruterów w sieci DiffServ oraz nie wymaga ingerencji w format nagłówka IP.

Autor nie znalazł w dostępnej literaturze badań analitycznych oraz symulacyjnych wpływu kształtowania ruchu na maksymalne opóźnienia dla agregatów EF PHB w sieci DiffServ.

3 Wyznaczenie opóźnień maksymalnych w sieci Usług Zróżnicowanych z kształtowaniem agregatów EF PHB algorytmem ciekącego wiadra

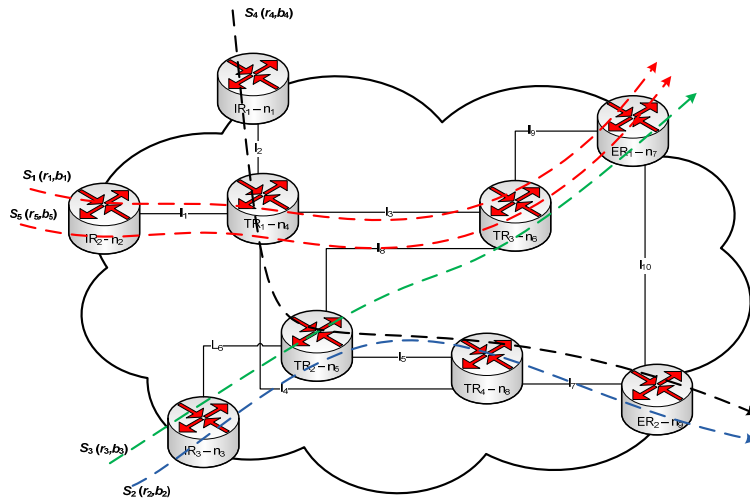
W poniższym rozdziale autor przedstawia opracowane przez siebie twierdzenia i wzory służące do szacowania opóźnień maksymalnych w sieci DiffServ, w której stosuje się kształtowanie ruchu dla agregatów EF PHB algorytmem ciekącego wiadra (LB). Przez **agregat EF PHB** autor rozumie agregat obsługiwany przez regułę przesyłania pakietów EF PHB w łączy sieci. W rozdziale 3.1 został zdefiniowany model matematyczny sieci DiffServ, służący do szacowania maksymalnych opóźnień w rozdziałach następnych. Model ten został tak opracowany, aby odpowiadał modelowi sieci, dla którego zostały wyprowadzone wzory w rozdziale 2.3, co umożliwia bezpośrednie porównania wyników pomiędzy rozwiązaniami opisanymi w rozdziale 2.3 i rozwiązaniami zaproponowanymi przez autora w rozdziale 3. W rozdziale 3.2 zostały wyprowadzone wzory na maksymalne opóźnienia dla EF PDB w rozważanej sieci przy zastosowaniu kształtowania ruchu algorytmem ciekącego wiadra dla agregatów EF PHB. Rozdział 3.3 przedstawia porównanie opóźnień maksymalnych dla agregatów EF PHB w sieci DiffServ z zastosowanym kształtowaniem zaproponowanym przez autora z opóźnieniami maksymalnymi w klasycznej sieci DiffServ bez zastosowanego kształtowania ruchu dla agregatów EF PHB. W podrozdziale 3.4 autor podsumowuje wyniki bieżącego rozdziału.

Zastosowana metodyka postępowania przewiduje wykorzystanie aparatu matematycznego stosowanego pierwotnie głównie do sieci o architekturze Usług Zintegrowanych i użycie go do modelu Usług Zróżnicowanych. Aparat ten znany pod nazwą „*Network calculus*” jest aparatem matematycznym stosowanym do analizy kresów górnych (a także dolnych) parametrów sieci pakietowych [38][40], jak opóźnienia maksymalne czy wymagane rozmiary buforów w ruterach. Aparat ten znalazł także

zastosowanie do wyznaczania wielkości maksymalnych parametrów dla pojedynczych węzłów lub całej sieci w architekturze Usług Zróżnicowanych [38]. W Aneksie A zostały zebrane najważniejsze twierdzenia algebry „*min-plus*” [38], które są częścią „*Network calculus*”, i do których autor będzie się odwoływał w bieżącym rozdziale. „*Network calculus*” jest wciąż rozwijanym aparatem matematycznym rozszerzanym o coraz to nowe możliwości, jak na przykład zastosowanie do modelowania procesów stochastycznych [39].

3.1 Model sieci

Autor zakłada dowolną topologię sieci z węzłami buforującymi pakiety na wyjściu, implementującymi algorytm kolejkowania typu *First-In-First-Out* (FIFO) dla rozwiązywanej klasy ruchu oraz algorytmem szeregowania typu *Guaranteed Rate* (GR) – wzór (2-11). Przez **klasę ruchu** autor rozumie zbiór strumieni danych w sieci, których pakiety są oznaczone takim samym kodem DSCP. Oznaczmy przez N zbiór wszystkich węzłów sieciowych n , przez L zbiór wszystkich łączy l w rozważanej sieci. Zakładamy, że maksymalny rozmiar pakietu w rozważanej sieci nie przekracza M . Maksymalny rozmiar pakietu w łączy l oznaczmy przez M_l , $M_l \leq M$.



Rysunek 3-1: Architektura sieci zdefiniowana w rozdziale 3.1 z przykładowymi strumieniami priorytetowymi. Oznaczenia: IR – ruter *ingress*, ER – ruter *egress*, TR – ruter rdzeniowy, $l_i \in L$ – łączy l_i , $n_i \in N$ – węzeł n_i , $S_i(r_i, b_i)$ – strumień i o parametrach (r_i, b_i)

Ruch przychodzący do sieci wchodzi węzłami typu „*ingress*”, natomiast opuszcza sieć w węzłach typu „*egress*” – rysunek 3-1. W sieci istnieje co najmniej jedna grupa strumieni danych przepływających pomiędzy brzegami sieci, która jest obsługiwana przez algorytm szeregowania typu GR. W tym miejscu należy zaznaczyć, że dla algorytmu szeregowania *Strict Priority* tylko klasa o najwyższym priorytecie jest obsługiwana zgodnie z definicją GR.

Algorytm szeregowania GR pierwotnie został zdefiniowany dla strumieni modelu Usług Zintegrowanych, ale może on być łatwo rozszerzony na szeregowanie pakietów należących do jednej klasy ruchu, jeśli będziemy traktować strumienie należące do jednej klasy ruchu w łączu jako pojedynczy strumień danych, co zostało wykazane w [13]. Zbiór strumieni danych należący do jednej klasy ruchu w łączu nazywamy agregatem DiffServ, zgodnie z definicją agregatu DiffServ wprowadzoną w rozdziale 1. W opisywanym modelu opóźnienia gwarantowane przez algorytm szeregowania GR dla agregatu są takie same jak zostały zdefiniowane dla pojedynczych strumieni danych.

W rozprawie autor zajmuje się opisem i parametrami obsługi tylko dla jednej, wybranej klasy ruchu obsługiwanej w każdym węźle sieci zgodnie z algorytmem GR, dalej nazywanej **klasą priorytetową**, natomiast pozostałe klasy ruchu są traktowane jako **ruch tła** (*background traffic*). Ruch tworzony przez pakiety należące do klasy priorytetowej nazywamy **ruchem priorytetowym**. Wszystkie obliczenia i wyniki dotyczą tylko klasy priorytetowej, ale uwzględniają wpływ ruchu tła na parametry obsługi tejże klasy. W danej sieci może istnieć wiele klas ruchu spełniających definicję klasy priorytetowej i w takim wypadku wybieramy jedną z nich, a pozostałe traktujemy jako ruch tła. W każdym łączu sieci pakiety należące do klasy priorytetowej (zwanej także **agregatem priorytetowym**), są buforowane w pojedynczej kolejce, niezależnie od kolejek obsługujących inne klasy – zgodnie z mechanizmem *aggregate class-based scheduling* przedstawionym na rysunku 2-3.

Strumieniem ruchu (w odróżnieniu od strumienia danych), lub krótko strumieniem, nazywamy zbiór pakietów, które należą do tej samej klasy ruchu, przechodzą przez sieć tą samą ścieżką, czyli przechodzą przez te same węzły sieci od węzła *ingress* do węzła *egress*. Wszystkie strumienie należące do klasy priorytetowej będą nazywane **strumieniami priorytetowymi**. Zakładamy, że i -ty strumień priorytetowy o parametrach (r_i, b_i) jest ograniczony na wejściu do sieci w taki sposób, aby był zgodny z algorytmem ciekącego wiadra o szybkości obsługi r_i i rozmiarze wiadra b_i . Strumień i może się składać z jednego lub większej liczby strumieni danych, ale nie stawiamy założeń, co do charakterystyki ruchowej poszczególnych strumieni wchodzących w skład strumienia i – ważne jest tylko, aby cały zbiór strumieni tworzący strumień i był ograniczony na wejściu do sieci przez parametry (r_i, b_i) . Dopuszcza się także, aby przez tę samą ścieżkę w sieci mógł przechodzić więcej niż jeden strumień (strumienie S_1 i S_5 na rysunku 3-1). Przez O_l oznaczmy zbiór wszystkich strumieni priorytetowych w łączu l .

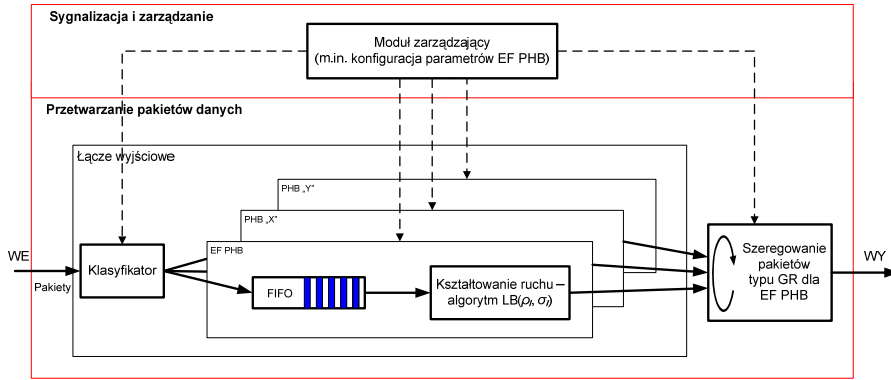
Przedstawiony powyżej model spełnia założenia reguł przekazywania pakietów EF

PHB dla modelu Usług Zróżnicowanych oraz jest zgodny z modelami przedstawionymi w [31] i [13]. W [12] zostało udowodnione, że definicja PSRG dla EF PHB – wzór (2-4) – jest mocniejsza od definicji GR – wzór (2-11) – co implikuje, że wszystkie zależności (w szczególności wartości maksymalne opóźnień) wyznaczone dla węzłów implementujących algorytmy szeregowania typu GR obowiązują także dla EF PHB zdefiniowanego w RFC 3246 [22].

3.2 Opóźnienia maksymalne w sieci DiffServ z kształtowaniem charakterystyki ruchu za pomocą algorytmu ciekącego wiadra dla agregatów EF PHB

Założmy, że dla wszystkich łączy $l \in L$ ilość ruchu priorytetowego generowanego przez strumień priorytetowy w jednym łączy spełnia następującą nierówność $\sum_{i \in O_l} r_i \leq \alpha R_l$, gdzie

R_l jest szybkością obsługi ruchu priorytetowego w łączy l , a $0 \leq \alpha \leq 1$ jest współczynnikiem wykorzystania pasma dla ruchu priorytetowego. Dodatkowo, wymagane jest, aby suma rozmiarów wiader strumieni przepływających przez każde z łączy była ograniczona przez $\sum_{i \in O_l} b_i \leq \beta R_l$.



Rysunek 3-2: Architektura łącza wyjściowego w węźle sieci DiffServ z zastosowanym kształtowaniem ruchu algorytmem ciekącego wiadra dla EF PHB

Przez H oznaczmy maksymalną liczbę węzłów, przez jakie musi przejść pakiet przesyłany pomiędzy dwoma dowolnymi węzłami na brzegach sieci. Przez R oznaczmy minimalną szybkość obsługi agregatu EF PHB (nazywanego w modelu sieci agregatem priorytetowym) w rozpatrywanej sieci: $R = \min\{R_l\}$. Jeśli w opisanej sieci zastosujemy

kształtowanie charakterystyki ruchu dla agregatu EF PHB w każdym węźle za pomocą algorytmu ciekącego wiadra o parametrach (ρ_l, σ_l) , tak jak to zostało przedstawione na rysunku 3-2, to maksymalne opóźnienie pakietu przy przejściu przez sieć jest opisane twierdzeniem 3-1³.

Twierdzenie 3-1: *Jeśli w opisanej sieci zastosujemy kształtowanie charakterystyki ruchu dla agregatu EF PHB w każdym łączu sieci $l \in L$ za pomocą algorytmu ciekącego wiadra o parametrach $(\rho_l = \eta C_l, \sigma_l = \varepsilon C_l)$, to maksymalne opóźnienie pakietu przy przejściu przez tę sieć jest ograniczone przez następującą nierówność:*

$$D \leq H \max_{l \in L} \{d_l\}, \quad (3-1)$$

gdzie d_l jest rozwiązaniem następującej nierówności:

$$d_l \leq h \left(\min \left\{ \alpha R_l (t + (H-1)d_l) + \beta R_l, P_n t + M_n, \rho_n t + \sigma_n \right\}, U(t_{T_l}) \min \{ \rho_l (t - T_l) + \sigma_l, R_l (t - T_l) \} \right) \quad (3-2)$$

i określa maksymalne opóźnienie w łączu $l \in L$.

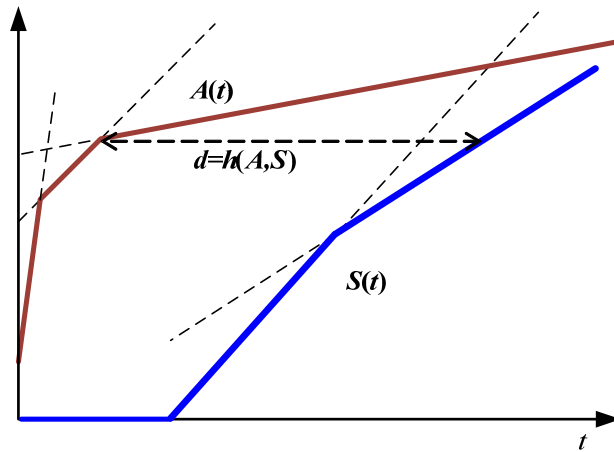
Funkcja $h(\cdot)$ zdefiniowana wzorem (A-6) wyznacza maksymalną odległość horyzontalną (odchylenie horyzontalne) pomiędzy krzywymi będącymi argumentami tej funkcji. Funkcje $\max(\cdot)$ i $\min(\cdot)$ wyznaczają odpowiednio maksimum i minimum ze swoich argumentów. Funkcja $U(\cdot)$ jest określona w następujący sposób

$$U(t_a) = \begin{cases} 0 & \text{dla } t \leq a \\ 1 & \text{dla } t > a \end{cases}. \quad M_n \text{ jest sumą maksymalnych rozmiarów pakietów z wszystkich}$$

przychodzących łączy do węzła n , $M_n = \sum_{i=1}^k M_l$, gdzie k jest liczbą łączy wchodzących do węzła, z którego wychodzi łącze l . Ponieważ w rozpatrywanym modelu sieci autor przyjął, że rozmiar pakietu dla wszystkich łączy w sieci jest ograniczony przez M , to w węźle n , $M_n = kM$. P_n oznacza maksymalną przepływność ruchu wchodzącego do węzła n dla całej klasy ruchu priorytetowego i wychodzącego przez łącze l . Dla rutera o dużej szybkości magistrali wewnętrznej P_n może być obliczone jako $P_n = \sum_{l=1}^k C_l$, gdzie C_l jest przepustowością łącza l przychodzącego do węzła n oraz k jest liczbą łączy przychodzących do węzła n . W przypadku, gdy maksymalna szybkość P_n nie może być oszacowana, można przyjąć, że $P_n = \infty$. P_n zawsze spełnia warunek $P_n > R_l$, przy czym szybkość obsługi agregatu R_l zależy od zastosowanego algorytmu szeregowania i na przykład dla algorytmu szeregowania *Strict Priority* dla klasy o najwyższym priorytecie

³ Dowody twierdzeń zostaną przedstawione w następnym podrozdziale.

$R_l = C_l$, dla algorytmu WFQ (*Weighted Fair Queuing*) R_l jest równy szybkości gwarantowanej dla danego agregatu. Zmienna $\rho_n = \sum_{i=1}^k \rho_i = \sum_{i=1}^k \eta_i C_i$ jest sumą szybkości obsługi w algorytmie ciekącego wiadra od wszystkich łączy przychodzących do węzła n . Analogicznie, $\sigma_n = \sum_{i=1}^k \sigma_i = \sum_{i=1}^k \varepsilon_i C_i$ jest sumą rozmiarów wiader w algorytmie ciekącego wiadra od wszystkich łączy przychodzących do węzła n . Jeśli w łączy l zachodziłby warunek $\sigma_l < M_l$, to wszystkie pakiety, których rozmiar jest większy od σ_l , byłyby blokowane przez algorytm ciekącego wiadra. Powyższa właściwość powoduje, że rozmiar wiadra w algorytmie ciekącego wiadra dla każdego łączy l musi spełniać warunek $\sigma_l \geq M_l$, co implikuje, że dla każdego węzła n musi zachodzić następująca nierówność $\sigma_n \geq M_n$. $T_l = \frac{M_l}{R_l} + E_l$, gdzie E_l jest stałą zależną od algorytmu szeregowania GR zastosowanego w łączy l .



Rysunek 3-3: Przykładowa krzywa obsługi $S(t)$ i krzywa ruchu przychodzącego $A(t)$ dla systemu z kształtowaniem ruchu algorytmem LB w węźle sieci DiffServ

Nie dla każdego zestawu parametrów sieci nierówność (3-2) posiada rozwiązanie. Interpretacja graficzna odchylenia horyzontalnego pomiędzy krzywymi z twierdzenia 3-1 została przedstawiona na rysunku 3-3. Krzywa $A(t) = \min\{P_n t + M_n, \rho_n t + \sigma_n, \alpha R_l(t + (H-1)d_l) + \beta R_l\}$ reprezentuje funkcję ograniczającą od góry ruch przychodzący, natomiast krzywa $S(t) = U(t_{T_l}) \min\{\rho_l(t - T_l) + \sigma_l, R_l(t - T_l)\}$ jest funkcją opisującą krzywą obsługi w łączy l , czyli funkcją określającą ile maksymalnie

bitów danych może zostać obsłużonych w łączu w zależności od czasu. Maksymalna odległość horyzontalna d pomiędzy tymi krzywymi określa maksymalne opóźnienie pakietów w łączu l .

W dalszej części tego rozdziału autor omawia przypadki szczególne twierdzenia 3-1, które pozwalają przedstawić wzór na opóźnienie maksymalne w zwartej postaci, tak jak to ma miejsce dla równania (2-13) i oszacować dla jakich parametrów sieci DiffServ opłacalne staje się zastosowanie mechanizmu kształtowania ruchu w każdym z węzłów. W tym przypadku przez słowo „opłacalne” autor rozumie przypadek, kiedy przez zastosowanie w sieci kształtowania ruchu zaproponowanego przez autora zostaną zmniejszone opóźnienia maksymalne lub/i będzie można zwiększyć wykorzystanie pasma zarezerwowanego dla EF PHB.

Jeśli w opisywanej sieci stosujemy kształtowanie charakterystyki ruchu dla agregatu EF PHB w każdym węźle za pomocą algorytmu ciekącego wiadra LB o parametrach $(\rho_l = \eta C_l, \sigma_l = \varepsilon C_l)$, to jedno z twierdzeń 3-2, 3-3, 3-4 albo 3-5 jest prawdziwe w zależności od wzajemnych relacji pomiędzy wartościami parametrów (ρ_l, σ_l) algorytmu ciekącego wiadra i wartościami parametrów opisujących sieć:

Twierdzenie 3-2: *Jeśli w opisywanej sieci dla każdego łącza $l \in L$ spełnione są następujące warunki:*

$$\rho_l \geq R_l \quad (3-3)$$

lub

$$\rho_l < R_l \quad (3-4)$$

i

$$\frac{\sigma_l R_l}{R_l - \rho_l} \geq \frac{P_n (\alpha R_l (H-1)d_l + \beta) - M_n \alpha R_l}{P_n - \alpha R_l}, \quad (3-5)$$

oraz spełniony jest wspólny warunek

$$\sigma_n \geq \frac{R_l ((P_n - \alpha) (H-1)d_l + \beta) + M_n (\rho_n - \alpha R_l)}{P_n - \alpha R_l} \quad (3-6)$$

gdzie d_l jest dane następującym wzorem:

$$d_l = \frac{H}{1 - (H-1)u_l \alpha} \left(\frac{(1-u_l)M_n}{R_l} + \frac{M_l}{R_l} + E_l + u_l \beta \right) \quad (3-7)$$

oraz współczynnik wykorzystania pasma w każdym węźle $l \in L$ spełnia następujący warunek

$$\alpha < \min \frac{P_n}{(P_n - R_l)(H-1) + R_l}, \quad (3-8)$$

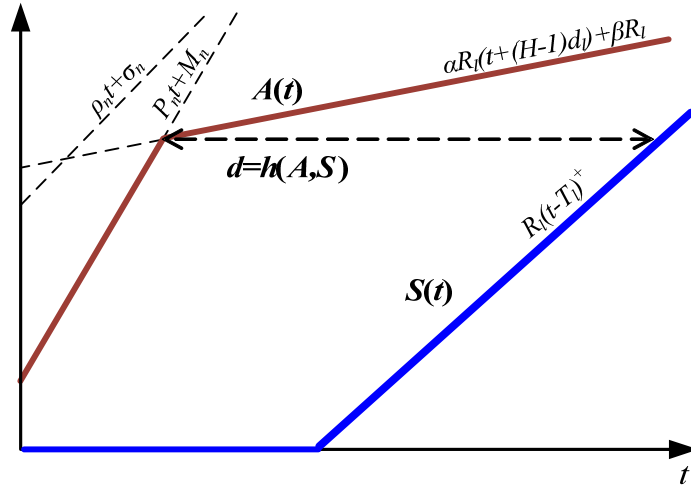
to zastosowanie kształtowania ruchu algorytmem ciekącego wiadra nie będzie miało wpływu

na maksymalne opóźnienie w sieci DiffServ dla reguł przesyłania pakietów EF PHB.

Dla twierdzenia 3-2 nierówność (3-2), którą należy rozwiązać aby obliczyć maksymalne opóźnienie w przypadku ogólnym, upraszcza się do następującej postaci:

$$d_i \leq h(\min\{\alpha R_i(t + (H-1)d_i) + \beta R_i, P_n t + M_n\}, R_i(t - T_i)^+). \quad (3-9)$$

Operator $(\cdot)^+$ jest zdefiniowany następująco: $(a)^+ = 0$ dla $a \leq 0$ oraz $(a)^+ = a$ dla $a > 0$. Rozwiązaniem powyższej nierówności, z uwzględnieniem, że $D \leq H \max\{d_i\}$, jest równanie (2-13) oraz warunek na współczynnik wykorzystania pasma dany równaniem (2-15). Oznacza to, że w takiej sieci maksymalne opóźnienie jest określone wzorem (2-13), gdy spełniony jest warunek na współczynnik wykorzystania pasma dany wzorem (2-15), czyli ma taką samą wartość jak w odpowiadającej jej sieci DiffServ bez kształtowania ruchu. Na rysunku 3-4 autor przedstawia krzywą opisującą ruch wejściowy $A(t)$ i krzywą obsługi $S(t)$ dla tego przypadku. Należy zauważyć, że w tym przypadku nierówność (3-9) nie zawiera właściwego równania krzywej opisującej ruch wejściowy i właściwego równania krzywej obsługi dla całościowego ujęcia twierdzenia 3-2, a jedynie krzywe pomiędzy którymi występuje maksymalna odległość horyzontalna.



Rysunek 3-4: Krzywa obsługi $S(t)$ i krzywa ruchu przychodzącego $A(t)$ dla twierdzenia 3-2

Twierdzenie 3-3: Jeśli w rozpatrywanej sieci dla każdego łącza $l \in L$ parametry sieci i algorytmu ciekącego wiadra spełniają następujące warunki:

$$\rho_l \geq R_l \quad (3-10)$$

lub

$$\rho_l < R_l \quad (3-11)$$

i

$$\frac{\sigma_l R_l}{R_l - \rho_l} \geq \frac{R_l ((H-1)\alpha T_l + \beta) \rho_n + \sigma_n (\alpha(H-1)(v_l + \rho_n - \rho_n v_l) - 1)}{(1 - \alpha v_l (H-1))(\rho_n - \alpha R_l)} \quad (3-12)$$

i

$$\rho_l \geq \alpha R_l \quad (3-13)$$

oraz spełnione są wspólne warunki:

$$\sigma_n < \frac{R_l ((H-1)T_l \alpha + \beta)(P_n - \rho_n) + M_n (\rho_n + \alpha((H-2)R_l - \rho_n(H-1)))}{(H-2)R_l \alpha + P_n (1 + \alpha - H\alpha)} \quad (3-14)$$

i współczynnik wykorzystania pasma

$$\alpha < \min_{l \in L} \left(\frac{\rho_n}{(\rho_n - R_l)(H-1) + R_l}, \frac{P_n}{(P_n - R_l)(H-1) + R_l} \right), \quad (3-15)$$

to maksymalne opóźnienie D pakietów należących do EF PHB pomiędzy brzegami sieci jest ograniczone następującą nierównością:

$$D \leq \frac{H}{1 - (H-1)v\alpha} (E'' + v\beta). \quad (3-16)$$

W powyższym wzorze

$$E'' = \max_l \left\{ \frac{(1 - v_l)\sigma_n}{R_l} + \frac{M_l}{R_l} + E_l \right\}. \quad (3-17)$$

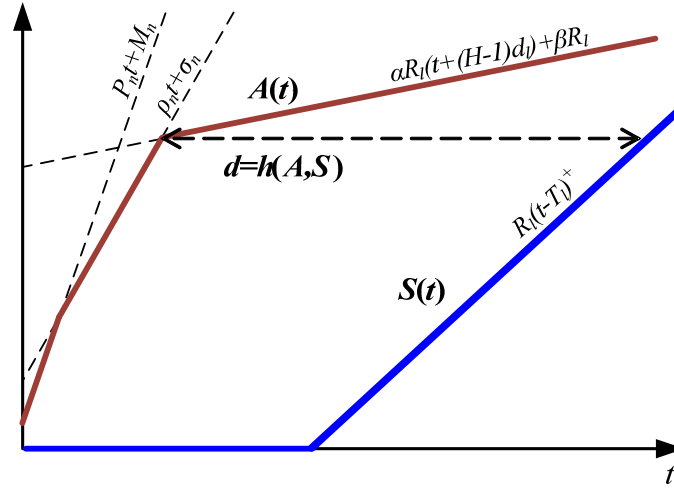
Zmienna v określona jest przez następujące równanie:

$$v = \max_l \{v_l\}, \text{ gdzie } v_l = \frac{\rho_n - R_l}{\rho_n - \alpha R_l}. \quad (3-18)$$

E_l jest błędem zależnym od wyboru algorytmu kolejowania i jest to zmienna zdefiniowana tak samo jak dla równania (2-15). Pozostałe zmienne są zdefiniowane w twierdzeniu 3-1.

Korzyści związane ze zmniejszeniem opóźnienia maksymalnego przez zastosowania kształtowania ruchu z parametrami odpowiadającymi twierdzeniu 3-3 można uzyskać jedynie wtedy, gdy $\forall l \in L \quad R_l < C_l$. Dla takiego przypadku można znaleźć taki zestaw parametrów $(\rho_l = \eta C_l, \sigma_l = \varepsilon C_l)$, że opóźnienie maksymalne jest mniejsze w porównaniu z analogicznym systemem bez kształtowania ruchu. W przypadku gdy $R_l = C_l$ to dla spełnionego warunku (3-13) opóźnienie maksymalne dla twierdzenia 3-3 będzie równe

lub wyższe niż w przypadku analogicznej sieci bez zastosowania kształtowania ruchu.



Rysunek 3-5: Krzywa obsługi $S(t)$ i krzywa ruchu przychodzącego $A(t)$ dla twierdzenia 3-3 dla warunku (3-13)

Na rysunku 3-5 autor przedstawia krzywą opisującą ruch wejściowy $A(t)$ i krzywą obsługi $S(t)$ dla twierdzenia 3-3. W tym przypadku maksymalna horyzontalna odległość występuje pomiędzy punktem przecięcia się krzywych $\rho_n t + \sigma_n$ i $\alpha R_l(t + (H-1)d_l) + \beta R_l$ oraz krzywą $R_l(t - T_l)^+$, co implikuje, że opóźnienie maksymalne w łączu l jest rozwiązaniem następującej nierówności:

$$d_l \leq h(\min\{\alpha R_l(t + (H-1)d_l) + \beta R_l, \rho_n t + \sigma_n\}, R_l(t - T_l)^+). \quad (3-19)$$

Twierdzenie 3-4: Jeśli w opisywanej sieci dla każdego łącza $l \in L$ spełnione są następujące warunki:

$$R_l \geq \rho_l \geq \alpha R_l, \quad (3-20)$$

i

$$\begin{aligned} \sigma_n \leq & \frac{M_n(\rho_n \rho_l + \alpha R_l(\rho_l(H-2) - \rho_n(H-1)))}{\rho_l(P_n + \alpha R_l(H-2)) - P_n \alpha R_l(H-1)} \\ & + \frac{R_l(P_n - \rho_n)(\alpha \rho_l T_l(H-1) + \beta \rho_l - \alpha \sigma_l(H-1))}{\rho_l(P_n + \alpha R_l(H-2)) - P_n \alpha R_l(H-1)} \end{aligned} \quad (3-21)$$

i

$$\frac{\sigma_l R_l}{R_l - \rho_l} < \frac{\rho_n R_l (\alpha R_l - \rho_n) ((H-1)T_l \alpha \rho_l + \beta \rho_l - \alpha \sigma_l (H-1))}{(\rho_n - \alpha R_l) (\alpha R_l ((H-1)(\rho_n - \rho_l) + \rho_l) - \rho_l \rho_n)} + \frac{\sigma_n ((H-2)R_l \alpha \rho_l + ((H-1)R_l \alpha (R_l \alpha - 1 - \rho_l) + \rho_l) \rho_n)}{(\rho_n - \alpha R_l) (\alpha R_l ((H-1)(\rho_n - \rho_l) + \rho_l) - \rho_l \rho_n)} \quad (3-22)$$

oraz współczynnik wykorzystania pasma w każdym łączu $l \in L$ spełnia następującą zależność:

$$\alpha \leq \min_l \left(\frac{\rho_l \rho_n}{((H-1)(\rho_n - \rho_l) + \rho_l) R_l}, \frac{\rho_l P_n}{((H-1)(P_n - \rho_l) + \rho_l) R_l} \right), \quad (3-23)$$

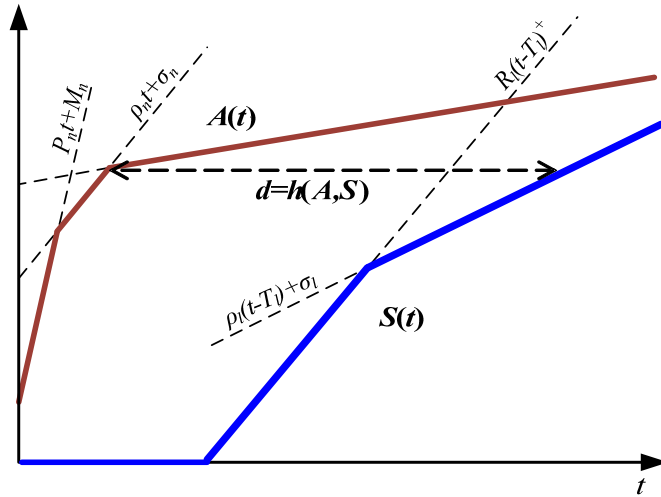
to maksymalne opóźnienie D pomiędzy brzegami sieci dla pakietów należących do strumieni priorytetowych jest ograniczone przez następującą nierówność:

$$D \leq H \max_l \{d_l\}, \quad (3-24)$$

gdzie

$$d_l \leq \frac{\sigma_n (\rho_l - \alpha R_l) + (\rho_l T_l - \sigma_l) (\rho_n - \alpha R_l) + \beta R_l (\rho_n - \rho_l)}{(\rho_n - \alpha R_l) \rho_l - \alpha R_l (H-1) (\rho_n - \rho_l)} \quad (3-25)$$

określa maksymalne opóźnienie w łączu l .



Rysunek 3-6: Krzywa obsługi i krzywa ruchu przychodzącego dla twierdzenia 3-4

W przypadku spełnienia warunków twierdzenia 3-4 opóźnienie maksymalne dla całej domeny DiffServ istnieje i pod warunkiem odpowiedniego doboru parametrów algorytmu ciekącego wiadra ($\rho_l = \eta C_l, \sigma_l = \varepsilon C_l$) i parametrów sieci może być mniejsze niż

w odpowiadającej jej sieci bez kształtowania ruchu, co zostanie wykazane w rozdziale 3.3.2.

Na rysunku 3-6 autor przedstawia krzywą opisującą ruch wejściowy $A(t)$ i krzywą obsługi $S(t)$ dla twierdzenia 3-4, w którym nierówność (3-2), którą należy rozwiązać aby obliczyć maksymalne opóźnienie, upraszcza się do następującej postaci:

$$d_l \leq h(\min\{\alpha R_l(t + (H-1)d_l) + \beta R_l, \rho_n t + \sigma_n\}, U(t_{T_l})(\rho_l(t - T_l) + \sigma_l)). \quad (3-26)$$

W tym przypadku maksymalna horyzontalna odległość występuje pomiędzy punktem przecięcia się krzywych $\rho_n t + \sigma_n$ i $\alpha R_l(t + (H-1)d_l) + \beta R_l$ oraz krzywą $\rho_l(t - T_l) + \sigma_l$.

Twierdzenie 3-5: *Dla parametru ρ_l algorytmu ciekącego wiadra spełniającego warunek*

$$\rho_l < \alpha R_l \quad (3-27)$$

wartość maksymalnego opóźnienia w domenie sieci DiffServ dla EF PHB nie jest ograniczona.

Jeśli szybkość obsługi algorytmu LB jest mniejsza od sumy przepływności strumieni wchodzących w skład agregatu priorytetowego, to nie można oszacować maksymalnego opóźnienia domenie DiffServ.

Należy jeszcze raz podkreślić, że założenia do twierdzeń 3-2, 3-3, 3-4 i 3-5 są komplementarne, co oznacza, że dobierając odpowiednio parametry algorytmu LB możemy wpływać na opóźnienie w sieci i decydować, które twierdzenie ma zastosowanie do obliczania opóźnienia maksymalnego. Twierdzenia 3-2, 3-3, 3-4 i 3-5 są szczególnymi przypadkami twierdzenia 3-1. Za pomocą twierdzeń 3-2, 3-3 i 3-4 można sprawdzić, jakie potencjalne korzyści można uzyskać przez zastosowanie kształtowania ruchu za pomocą algorytmu ciekącego wiadra w sieci DiffServ dla danych parametrów sieci. Zależności pomiędzy twierdzeniami 3-3, 3-4 i wzorem (2-15) zostaną przedstawione w rozdziale 3.3.

Dowody powyższych twierdzeń zostały przedstawione w następnym podrozdziale.

3.2.1 Dowody twierdzeń 3-1, 3-2, 3-3, 3-4 i 3-5

W pierwszym kroku autor wyprowadza wzór (3-2) do twierdzenia 3-1. Wzór ten oraz pozostałe wyprowadzone wzory pomocnicze są podstawą do wyprowadzenia pozostałych twierdzeń. W trakcie przeprowadzania dowodu twierdzeń autor korzysta z metody „time-stopping” [32] polegającej na pokazaniu niezależnie, że jeśli istnieje maksymalne opóźnienie, to dowodzone twierdzenie jest prawdziwe oraz że istnieje kres górny opóźnienia. Bardzo podobny sposób postępowania został także zastosowany w [31] dla wzoru (2-7) i w [13] dla dowodu wzoru (2-13).

Podczas wyprowadzania wzorów autor wykorzystał pakiet oprogramowania Mathematica™ do obliczeń symbolicznych w celu weryfikacji wyprowadzonych wzorów

i następnie do porównania uzyskanych wyników z wzorem (2-13), który określa opóźnienia maksymalne dla EF PHB w klasycznej sieci DiffServ.

3.2.1.1 Dowód twierdzenia 3-1

W pierwszym etapie autor wyprowadza wzory na krzywą ograniczającą ruch przychodzący do węzła (*arrival curve*), następnie zostanie wyznaczona krzywa opisująca proces obsługi ruchu w węźle (*service curve*) rozważanej sieci. Znając postać wyżej wymienionych krzywych i korzystając z twierdzeń teorii *Network Calculus* autor wyprowadza ogólne równanie, którego rozwiązanie, o ile istnieje, wyznacza wartość maksymalnego opóźnienia w sieci DiffServ z kształtowaniem ruchu za pomocą algorytmu ciekącego wiadra w każdym łączu sieci, dla EF PHB.

Założmy, że w każdym węźle w rozpatrywanej sieci istnieje opóźnienie maksymalne d' . Ponieważ autor zakłada, że ruch generowany przez każdy strumień i należący do agregatu priorytetowego na wejściu do sieci jest ograniczony przez algorytm ciekącego wiadra o parametrach (r_i, b_i) , to krzywa opisująca ruch wejściowy strumienia i ma postać $a_i(t) = r_i t + b_i$. Krzywa opisująca ruch wejściowy w łączu l dla agregatu priorytetowego w węźle brzegowym jest sumą krzywych $a_i(t)$ odpowiadających strumieniom przechodzącym przez łącze l oraz należących do agregatu priorytetowego i ma następującą postać:

$$A''(t) = \sum_{i \in O_l} a_i(t) = \sum_{i \in O_l} (r_i t + b_i). \quad (3-28)$$

Korzystając z wcześniejszych założeń w rozdziale 3.2 wynika, że dla każdego łącza $l \in L$ $\sum_{i \in S} (r_i t + b_i) \leq \alpha R_l t + \beta R_l$, co oznacza, że równanie (3-28) możemy zapisać jako

$$A''(t) \leq \alpha R_l t + \beta R_l. \quad (3-29)$$

Równanie krzywej ograniczającej ruch wchodzący do węzła n , w którym zachodzi najgorszy przypadek opóźnienia maksymalnego w sieci, przybierze następującą postać:

$$A'(t) \leq \min(P_n t + M_n, \rho_n t + \sigma_n, \alpha R_l(t + (H-1)d_l) + \beta R_l). \quad (3-30)$$

P_n jest maksymalną szybkością transmisji z wszystkich łączy wchodzących do węzła n , M_n jest sumą maksymalnych rozmiarów pakietów z wszystkich łączy przychodzących do węzła n . ρ_n jest sumą szybkości obsługi mechanizmu kształtowania ruchu ze wszystkich łączach przychodzących do węzła n , σ_n jest sumą rozmiarów wiader mechanizmów kształtowania ruchu ze wszystkich łączy przychodzących do węzła n . Składnik $(H-1)d_l$ w trzecim argumencie operatora $\min(\cdot)$ odpowiada za opóźnienie powstałe w $(H-1)$ węzłach poprzedzających węzeł, w którym jest obliczane opóźnienie maksymalne.

Krzywa obsługi ruchu w łączu l dla węzła implementującego algorytm szeregowania

typu GR zdefiniowanego tak, jak w równaniu (2-11) ma postać [40]

$$s_1(t) = R_l(t - T_l)^+ . \quad (3-31)$$

Gdzie $T_l = \frac{M_l}{R_l} + E_l$. E_l jest stałą zależną od algorytmu szeregowania zastosowanego w łączu l i tak na przykład dla algorytmu szeregowania *Strict Priority* dla kolejki o najwyższym priorytecie $E_l = \frac{M_l}{C_l}$, dla algorytmu szeregowania WFQ: $E_l = \frac{M_l}{R_{WFQ}}$, gdzie R_{WFQ} jest szybkością gwarantowaną obsługi danej klasy ruchu. E_l odpowiada stałej E z definicji algorytmów szeregowania klasy GR – wzór (2-11).

Dodatkowo w każdym łączu l stosujemy mechanizm kształtowania ruchu według algorytmu ciekącego wiadra z parametrami (ρ_l, σ_l) , dla którego krzywa obsługi ruchu ma postać [38][40]:

$$s_2(t) = \begin{cases} 0 & \text{dla } t \leq 0 \\ \rho_l t + \sigma_l & \text{dla } t > 0 \end{cases} \quad (3-32)$$

Na rysunku 3-7 przedstawiono schematycznie sposób obsługi ruchu w łączu wychodzącym z węzła dla reguły przesyłania pakietów EF PHB w rozpatrywanej sieci. Możliwe są dwa sposoby obliczenia opóźnień w łączu przedstawionym na rysunku 3-7. Pierwszy sposób polega na obliczeniu niezależnie opóźnień wprowadzanego przez moduł szeregowania pakietów według algorytmu GR i opóźnień wprowadzanego przez algorytm kształtowania ruchu typu LB, drugi na policzeniu wypadkowej krzywej obsługi dla obydwu mechanizmów obsługi istniejących w węźle i w następnym kroku na obliczeniu opóźnień w węźle. Jak zostało udowodnione w [38] lepsze rezultaty osiąga się korzystając z drugiej metody.



Rysunek 3-7: Obsługa ruchu z zastosowaniem mechanizmów kształtowania ruchu w łączu wychodzącym z węzła sieci

Wypadkowa krzywa obsługi dla systemu przedstawionego na rysunku 3-7, składającego się z modułu szeregowania pakietów typu GR i modułu kształtowania ruchu, ma następującą postać:

$$S(t) = (s_1 \otimes s_2)(t) . \quad (3-33)$$

Na podstawie definicji splotu danego równaniem (A-4) $S(t) = \min_{0 \leq s \leq t} \{s_2(t-s) + s_1(t)\}$

zostaną przeprowadzone obliczenia dla dwóch przypadków:

dla $0 \leq t \leq T_l$:

$$S(t) = \min_{0 \leq s \leq t} \{s_2(t-s) + s_1(s)\} = \min_{0 \leq s \leq t} \{s_2(t-s) + 0\} = s_2(0) + 0 = 0 \quad (3-34)$$

dla $T \leq t$:

$$\begin{aligned} S(t) &= \min_{0 \leq s \leq t} \{s_2(t-s) + s_1(s)\} = \min_{0 \leq s \leq T_l} \{\rho_l(t-s) + \sigma_l + 0\} \\ &\wedge \min_{T_l < s < t} \{\rho_l(t-s) + \sigma_l + R_l(s-T_l)\} \wedge \min_{s=t} \{s_2(t-s) + R_l(t-T_l)\} \\ &= \rho_l(t-T_l) + \sigma_l \wedge \min_{T_l \leq s < t} \{\rho_l t - \rho_l s + \sigma_l + R_l s - R_l T_l\} \wedge R_l(t-T_l) \\ &= \rho_l(t-T_l) + \sigma_l \wedge \{\rho_l t + \sigma_l - R_l T_l + \min_{T \leq s < t} \{R_l s - \rho_l s\}\} \wedge R_l(t-T_l) \\ &= \rho_l(t-T_l) + \sigma_l \wedge \rho_l t + \sigma_l - R_l T_l \wedge R_l(t-T_l) = \rho_l(t-T_l) + \sigma_l \wedge R_l(t-T_l) \end{aligned} \quad (3-35)$$

Podsumowując:

$$S(t) = \begin{cases} 0 & \text{dla } 0 \leq t \leq T_l \\ \min(\rho_l(t-T_l) + \sigma_l, R_l(t-T_l)) & \text{dla } t > T_l \end{cases} \quad (3-36)$$

Znając krzywą ograniczającą ruch wejściowy i krzywą obsługi w węźle dla najgorszego przypadku opóźnienie maksymalne uzyskuje się obliczając odchylenie horyzontalne pomiędzy tymi dwoma krzywymi według wzoru (A-6) [38]:

$$d_l \leq h(A'(t), S(t)) \quad (3-37)$$

Podstawiając wzory (3-30) i (3-36) do powyższej nierówności oraz stosując funkcję

zdefiniowaną następująco $U(t_a) = \begin{cases} 0 & \text{dla } t \leq a \\ 1 & \text{dla } t > a \end{cases}$, otrzymujemy następującą nierówność:

$$d_l \leq h \left(\min \{ \alpha R_l(t + (H-1)d_l) + \beta R_l, P_n t + M_n, \rho_n t + \sigma_n \}, U(t_{T_l}) \min \{ \rho_l(t-T_l) + \sigma_l, R_l(t-T_l) \} \right), \quad (3-38)$$

której rozwiązanie, o ile istnieje, określa maksymalne opóźnienie w łączy l . Znajdując w całej rozważanej sieci łączy w którym opóźnienie maksymalne jest największe i mnożąc je przez maksymalną liczbę hopów w tej sieci otrzymamy opóźnienie maksymalne D , jakie może doświadczyć pakiet przy przejściu przez tą domenę:

$$D \leq H \max_l \{d_l\}, \quad (3-39)$$

Wyprowadzone równania (3-38) i (3-39) są zgodne z równaniami (3-2) i (3-1), co kończy dowód twierdzenia 3-1.

Twierdzenia 3-2, 3-3 i 3-4 są szczególnymi rozwiązaniami twierdzenia 3-1. Integralnym składnikiem tych twierdzeń są warunki określające, dla jakich parametrów

sieci poszczególne twierdzenia są prawdziwe. W następnych podrozdziałach zostaną przeprowadzone dowody twierdzeń 3-2, 3-3 i 3-4 z wykorzystaniem wzorów z dowodu twierdzenia 3-1.

3.2.1.2 Dowód twierdzenia 3-2

Dowód twierdzenia 3-2 polega na wykazaniu, że jeśli warunki (3-6), (3-5), (3-7) i (3-8) są spełnione przez parametry sieci to kształtowanie ruchu algorytmem ciekącego wiadra nie będzie wpływało na opóźnienia maksymalne w sieci, a to oznacza że sieć będzie się zachowywać tak, jak gdyby mechanizm ten nie był zaimplementowany w węzłach sieci.

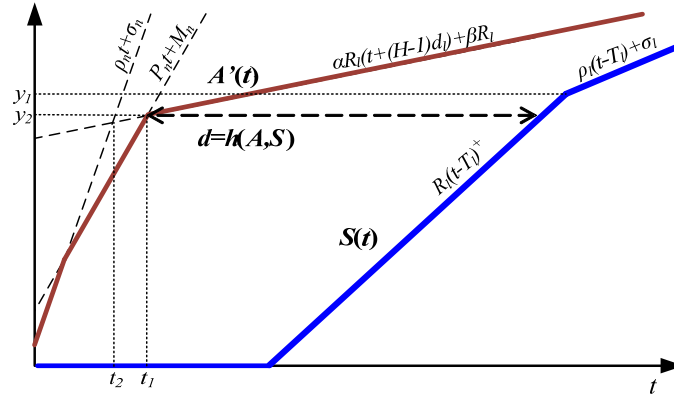
Zakładamy, że w systemie istnieje maksymalne opóźnienie d_l . Na rysunku 3-8 autor zamieszcza krzywą ograniczającą ruch wchodzący do łącza l oraz krzywą obsługi w łączu l dla twierdzenia 3-2, które mają ogólną postać daną równaniami (3-30) i (3-36), natomiast maksymalna odległość pomiędzy krzywą ograniczającą ruch wejściowy i krzywą obsługi jest mierzona pomiędzy krzywą

$$A'(t) \leq \min(P_n t + M_n, \alpha R_l(t + (H-1)d_l) + \beta R_l) \quad (3-40)$$

oraz krzywą obsługi daną następującym równaniem:

$$S(t) = R_l(t - T_l)^+, \quad (3-41)$$

czyli pomiędzy krzywymi jak w systemie bez zastosowanego kształtowania ruchu algorytmem ciekącego wiadra w każdym węźle sieci.



Rysunek 3-8: Wzajemne położenie krzywej ruchu przychodzącego $A'(t)$ i krzywej obsługi $S(t)$ w łączu dla twierdzenia 3-2

Krzywa ograniczająca ruch przychodzący do łącza ma następującą właściwość: kolejne odcinki i ostatnia półprosta tworzące krzywą ograniczającą ruch przychodzący muszą mieć

ze wzrostem t coraz mniejsze nachylenie względem osi t . Jest to spowodowane zastosowaniem funkcji minimum, która ma taką właściwość jeśli argumentami są proste typu $at + b$, gdzie współczynniki a i b są zawsze większe od zera.

Dla zachowania wzajemnych relacji krzywych $A'(t)$ i $S(t)$ zgodnie z rysunkiem 3-8 (wszystkie relacje rozpatrujemy tylko dla $t > 0$) muszą być spełnione trzy warunki:

- 1) Nachylenie krzywej $\alpha R_l(t + (H-1)d_l) + \beta R_l$ względem osi t będzie mniejsze niż nachylenie krzywych $\rho_n t + \sigma_n$ i $P_n t + M_n$. Ponieważ z założeń modelu $R_l \leq C_l$ co implikuje, że $P_n > \alpha R_l$, to wystarczającym warunkiem spełnienia tego założenia jest warunek (3-6) w twierdzeniu 3-2, dla którego zawsze zachodzi nierówność $\rho_n > \alpha R_l$.
- 2) Odcięta t_1 punktu przecięcia krzywej $P_n t + M_n$ z krzywą $\alpha R_l(t + (H-1)d_l) + \beta R_l$ będzie miała wartość większą bądź równą odciętej t_2 punktu przecięcia krzywej $\rho_n t + \sigma_n$ z krzywą $\alpha R_l(t + (H-1)d_l) + \beta R_l$. Warunek ten odpowiada nierówności (3-6) w twierdzeniu 3-2.
- 3) Jeśli $\rho_l < R_l$, to rzędna y_1 punktu przecięcia krzywych $P_n t + M_n$ i $\alpha R_l(t + (H-1)d_l) + \beta R_l$ musi mieć wartość mniejszą od rzędnej y_2 punktu przecięcia krzywych $\rho_l(t - T_l) + \sigma_l$ i $R_l(t - T_l)$. Warunek ten odpowiada nierówności (3-5) w twierdzeniu 3-2.

Dla warunku drugiego wartości odciętych t_1 i t_2 są następujące:

$$t_1 = \frac{\alpha R_l(H-1)d_l + \beta R_l - M_n}{P_n - \alpha R_l}, \quad t_2 = \frac{\alpha R_l(H-1)d_l + \beta R_l - \sigma_n}{\rho_n - \alpha R_l},$$
 co odpowiada następującej zależności:

$$t_1 - t_2 = \frac{\alpha R_l(H-1)d_l + \beta R_l - M_n}{P_n - \alpha R_l} - \frac{\alpha R_l(H-1)d_l + \beta R_l - \sigma_n}{\rho_n - \alpha R_l} \geq 0. \quad (3-42)$$

Rozwiązując nierówność 51 względem σ_n otrzymujemy warunek:

$$\sigma_n \geq \frac{R_l((P_n - \alpha R_l)(H-1)d_l + \beta R_l) + M_n(\rho_n - \alpha R_l)}{P_n - \alpha R_l}, \quad (3-43)$$

który jest zgodny z warunkiem (3-6) w twierdzeniu 3-2. Wzór na d_l autor wyprowadza bezpośrednio z równania (19) w [13] i ma on następującą postać:

$$d_l = \frac{H}{1 - (H-1)u_l \alpha} \left(\frac{(1-u_l)M_n}{R_l} + \frac{M_l}{R_l} + E_l + u_l \beta \right). \quad (3-44)$$

d_l można także w łatwy sposób uzyskać rozwiązując nierówność (3-9) przy

zastosowaniu wzorów (A-7), (A-8) i (A-9). Aby warunek dany wzorem (3-43) był prawdziwy to musi być także prawdziwa nierówność (2-15), bo tylko wtedy d_l jest opisane wzorem (3-44). W przeciwnym przypadku d_l jest nieograniczone, podobnie jak opóźnienie maksymalne dane wzorem (2-13), co zostało udowodnione w [31]. Podsumowując, warunki (3-7) i (3-8) wynikają z równań (2-15) i (2-13).

Dla warunku trzeciego autor wyprowadza wzór na nierówność (3-5). Wartość funkcji $S(t) = \min\{\rho_l(t - T_l) + \sigma_l, R_l(t - T_l)\}$ dla $t > T_l$ w punkcie załamania wynosi:

$$y_1 = S\left(\frac{\sigma_l}{R_l - \rho_l} + T_l\right) = \frac{\sigma_l R_l}{R_l - \rho_l}. \text{ Wartość funkcji (3-40) w punkcie załamania wynosi:}$$

$$y_2 = A'\left(\frac{\alpha R_l(H-1)d_l + \beta R_l - M_n}{P_n - \alpha R_l}\right) = \frac{P_n(\alpha R_l(H-1)d_l + \beta R_l) - M_n \alpha R_l}{P_n - \alpha R_l}. \text{ To oznacza, że dla}$$

twierdzenia 3-2 musi być spełniony warunek:

$$\frac{\sigma_l R_l}{R_l - \rho_l} \geq \frac{P_n(\alpha R_l(H-1)d_l + \beta) - M_n \alpha R_l}{P_n - \alpha R_l}, \quad (3-45)$$

który jest zgodny z warunkiem (3-5) w twierdzeniu 3-2. d_l dane jest wzorem (3-44).

Przy powyższych warunkach opóźnienie maksymalne w sieci z kształtowaniem ruchu algorytmem ciekącego wiadra będzie liczone pomiędzy krzywymi danymi wzorami (3-40) i (3-41), czyli będzie takie samo jak w sieci bez kształtowania ruchu algorytmem ciekącego wiadra. W takiej sieci muszą być spełnione warunki dotyczące istnienia opóźnienia maksymalnego, które są takie same jak w sieci bez kształtowania ruchu. Autor wykazał, że przy spełnionych warunkach danych wzorami (3-6), (3-5), (3-7), (3-8) i (3-9) opóźnienie maksymalne będzie takie samo jak w systemie bez kształtowania ruchu. Warunki (3-6), (3-5), (3-7) zapewniają, że to algorytm ciekącego wiadra będzie przekazywał pakiety niezwłocznie po ich otrzymaniu bez wprowadzania dodatkowego opóźnienia, a zatem będzie przezroczysty dla obsługiwanych pakietów.

3.2.1.3 Dowód twierdzenia 3-3

Dowód twierdzenia 3-3 składa się z dwóch części; w części pierwszej autor dowodzi, że jeśli istnieje skończony kres górny opóźnienia, to twierdzenie 3-3 jest prawdziwe (Część I), natomiast w drugiej autor wykazuje, że istnieje kres górny opóźnienia (Część II).

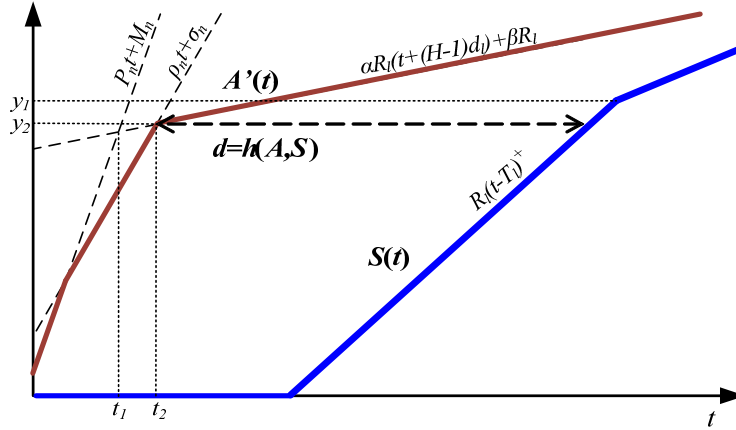
Część I. Wyprowadzenie wzoru na opóźnienie maksymalne dla twierdzenia 3-3.

Zakładamy, że najgorszy przypadek opóźnienia istnieje i jest ograniczony przez d_l . Na rysunku 3-9 autor przedstawia graficzną interpretację twierdzenia 3-3, które między innymi definiuje, jakie warunki musi spełniać sieć, aby opóźnienie maksymalne, czyli maksymalne horyzontalne odchylenie pomiędzy krzywą ograniczającą ruch przychodzący i krzywą obsługi było mierzone pomiędzy krzywymi:

$$A'(t) \leq \min(\rho_n t + \sigma_n, \alpha R_l(t + (H-1)d_l) + \beta R_l) \quad (3-46)$$

i

$$S(t) = R_l(t - T_l)^+ . \quad (3-47)$$



Rysunek 3-9: Wzajemne położenie krzywej ruchu przychodzącego $A'(t)$ i krzywej obsługi $S(t)$ w łączu dla twierdzenia 3-3

Dla zachowania wzajemnych relacji krzywych $A'(t)$ i $S(t)$ zgodnie z rysunkiem 3-9 (wszystkie relacje rozpatrujemy tylko dla $t > 0$) muszą być spełnione trzy warunki:

- 1) Nachylenie krzywej $\alpha R_l(t + (H-1)d_l) + \beta R_l$ względem osi t będzie mniejsze niż nachylenie krzywych $P_n t + M_n$ i $\rho_n t + \sigma_n$. Ponieważ w przyjętym modelu sieci $R_l \leq C_l$, to $P_n > \alpha R_l$, czyli nachylenie krzywej $P_n t + M_n$ jest zawsze większe od nachylenia krzywej $\alpha R_l(t + (H-1)d_l) + \beta R_l$. Natomiast wystarczającym warunkiem mniejszego nachylenia krzywej $\alpha R_l(t + (H-1)d_l) + \beta R_l$ względem krzywej $\rho_n t + \sigma_n$ jest warunek (3-13) w twierdzeniu 3-3, dla którego zawsze zachodzi nierówność $\rho_n > \alpha R_l$.
- 2) Odcięta t_1 punktu przecięcia krzywej $P_n t + M_n$ z krzywą $\alpha R_l(t + (H-1)d_l) + \beta R_l$ będzie miała wartość mniejszą bądź równą odciętej t_2 punktu przecięcia krzywej $\rho_n t + \sigma_n$ z krzywą $\alpha R_l(t + (H-1)d_l) + \beta R_l$. Warunek ten odpowiada nierówności (3-14) w twierdzeniu 3-3.

- 3) Rzędna y_1 punktu przecięcia krzywych $\rho_n t + \sigma_n$ i $\alpha R_l(t + (H-1)d_l) + \beta R_l$ musi mieć wartość mniejszą od rzędnej y_2 punktu przecięcia krzywych $\rho_l(t - T_l) + \sigma_l$ i $R_l(t - T_l)$. Warunek ten odpowiada nierówności (3-12) w twierdzeniu 3-3.

Warunek pierwszy dowodu twierdzenia 3-3 pozostaje taki sam jak w dowodzie twierdzenia 3-2, natomiast warunek drugi twierdzenia 3-3 jest przeciwny do warunku drugiego w dowodzie twierdzenia 3-2.

Wyznaczenie nierówności (3-14) na podstawie warunku drugiego do twierdzenia 3-3: Analogicznie jak w dowodzie twierdzenia 3-2 wartości odciętych t_1 i t_2 będą odpowiednio

$$\text{wynosiły: } t_1 = \frac{\alpha R_l(H-1)d_l + \beta R_l - M_n}{P_n - \alpha R_l}, \quad t_2 = \frac{\alpha R_l(H-1)d_l + \beta R_l - \sigma_n}{\rho_n - \alpha R_l}. \quad \text{Zgodnie z}$$

warunkiem drugim następująca nierówność musi być prawdziwa:

$$t_2 - t_1 = \frac{\alpha R_l(H-1)d_l + \beta R_l - \sigma_n}{\rho_n - \alpha R_l} - \frac{\alpha R_l(H-1)d_l + \beta R_l - M_n}{P_n - \alpha R_l} \geq 0. \quad (3-48)$$

Ponieważ prawdziwe jest założenie, że dla każdego łącza l $\rho_l \geq \alpha R_l$, to rozwiązując nierówność (3-48) względem σ_n otrzymujemy następujący warunek:

$$\sigma_n \leq \frac{R_l((P_n - \rho_n)\alpha(H-1)d_l + \beta) + M_n(\rho_n - \alpha R_l)}{P_n - \alpha R_l}, \quad (3-49)$$

gdzie d_l jest rozwiązaniem równania

$$d_l \leq h(\min\{\alpha R_l(t + (H-1)d_l) + \beta R_l, \rho_n t + \sigma_n\}, R_l(t - T_l)^+). \quad (3-50)$$

Ostateczna postać warunku (3-49) i rozwiązanie równania (3-50) zostanie przedstawiona w dalszej części dowodu.

Wyznaczenie nierówności (3-12) na podstawie warunku trzeciego do twierdzenia 3-3: Rzędne z warunku trzeciego do twierdzenia 3-3 odpowiadają punktom załamania krzywej $A'(t)$ i krzywej $S(t)$. Wartość funkcji $S(t) = \min\{\rho_l(t - T_l) + \sigma_l, R_l(t - T_l)\}$

$$\text{w punkcie załamania } y_1 = S\left(\frac{\sigma_l}{R_l - \rho_l} + T_l\right) = \frac{\sigma_l R_l}{R_l - \rho_l}. \quad \text{Wartość funkcji}$$

$A'(t) \leq \min(\rho_n \cdot t + \sigma_n, \alpha R_l(t + (H-1)d_l) + \beta R_l)$ w punkcie załamania wynosi:

$$y_2 = A'\left(\frac{\alpha R_l(H-1)d_l + \beta R_l - \sigma_n}{\rho_n - \alpha R_l}\right) = \frac{\rho_n(\alpha R_l(H-1)d_l + \beta R_l) - \sigma_n \alpha R_l}{\rho_n - \alpha R_l}. \quad \text{Zgodnie}$$

z warunkiem trzecim musi być spełniony następujący warunek

$$\frac{\sigma_l R_l}{R_l - \rho_l} \geq \frac{\rho_n(\alpha R_l(H-1)d_l + \beta R_l) - \sigma_n \alpha R_l}{\rho_n - \alpha R_l}, \quad (3-51)$$

gdzie d_l jest opóźnieniem maksymalnym w łączu l .

Opóźnienie maksymalne w dowolnym węźle sieci autor wyprowadza przez obliczenie horyzontalnego odchylenia pomiędzy krzywymi $A'(t)$ – wzór (3-46) i $S(t)$ – wzór (3-47), zgodnie ze wzorem (A-6):

$$d_l \leq d' = h(A', S). \quad (3-52)$$

$$d_l \leq d' = h\left(\min(\rho_n t + \sigma_n, \alpha R_l t + (\alpha R_l (H-1)d_l + \beta R_l)), R_l(t - T_l)^+\right). \quad (3-53)$$

Na podstawie wzorów (A-7), (A-8) i (A-9) otrzymujemy:

$$d_l \leq \frac{\sigma_n + \frac{(\alpha R_l (H-1)d_l + \beta R_l - \sigma_n)(\rho_n - R_l)}{\rho_n - \alpha R_l}}{R_l} + T_l \quad (3-54)$$

Przenosząc d_l na prawą stronę równania i przyjmując $v_l = \frac{\rho_n - R_l}{\rho_n - \alpha R_l}$ otrzymujemy:

$$d_l(1 - \alpha(H-1)v_l) \leq \frac{\sigma_n - v_l \sigma_n + \beta R_l v_l}{R_l} + T_l. \quad (3-55)$$

Wyrażenie w nawiasie po lewej stronie nierówności musi być większe od zera, aby powyższa nierówność miała sens ze względu na możliwość oszacowania wielkości d_l od góry. Rozwiązując nierówność $(1 - \alpha(H-1)v_l) > 0$ ze względu na współczynnik wykorzystania pasma α otrzymujemy, że musi on spełniać dla łącza l następującą nierówność:

$$\alpha \leq \frac{1}{(H-1)v_l}. \quad (3-56)$$

Po dalszych przekształceniach równania (3-55) otrzymujemy:

$$d_l \leq \frac{\sigma_n - v_l \sigma_n + \beta R_l v_l + R_l T_l}{(1 - \alpha(H-1)v_l)R_l}. \quad (3-57)$$

Opóźnienie maksymalne pomiędzy brzegami domeny DiffServ wynosi

$$D = H \max\{d_l\}. \quad (3-58)$$

Przyjmując następujące oznaczenia: $v = \max\{v_l\}$, $E'' = \max\left\{\frac{(1-v_l)\sigma_n}{R_l} + \frac{M_l}{R_l} + E_l\right\}$

oraz $T_l = \frac{M_l}{R_l} + E_l$ wzór (3-58) przybierze następującą postać:

$$D \leq \frac{H}{1 - (H-1)v\alpha} (E'' + v\beta), \quad (3-59)$$

co jest zgodne ze wzorami (3-16), (3-17) i (3-18) twierdzenia 3-3. Podstawiając do wzorów (3-49) i (3-51) wzór (3-57) otrzymujemy warunki jakie muszą spełnić wartości parametrów sieci, w tym parametry ciekącego wiadra (ρ_i, σ_i) , w każdym węźle sieci:

$$\sigma_n < \frac{R_i((H-1)T_i\alpha + \beta)(P_n - \rho_n) + M_n(\rho_n + \alpha((H-2)R_i - \rho_n(H-1)))}{(H-2)R_i\alpha + P_n(1 + \alpha - H\alpha)} \quad (3-60)$$

oraz

$$\frac{\sigma_i R_i}{R_i - \rho_i} \geq \frac{R_i((H-1)\alpha T_i + \beta)\rho_n + \sigma_n(\alpha(H-1)(v_i + \rho_n - \rho_n v_i) - 1)}{(1 - \alpha v_i(H-1))(\rho_n - \alpha R_i)}. \quad (3-61)$$

Wzory (3-60) i (3-61) są zgodne z warunkami opisanymi nierównościami (3-14) i (3-12) w twierdzeniu 3-3.

Nierówność (3-61) ze względu na współczynnik wykorzystania pasma jest funkcją homograficzną, która dla parametrów sieci zgodnej z przyjętym modelem posiada asymptotę pionową przecinającą oś odciętych w punkcie $\frac{P_n}{(P_n - R_i)(H-1) + R_i}$ i aby nierówność ta była spełniona współczynnik wykorzystania pasma musi mieć wartość mniejszą od wartości tego punktu. Uwzględniając warunek (3-56) wzór (3-59) jest prawdziwy tylko wtedy, gdy:

$$\alpha < \min_{i \in L} \left(\frac{\rho_n}{(\rho_n - R_i)(H-1) + R_i}, \frac{P_n}{(P_n - R_i)(H-1) + R_i} \right), \quad (3-62)$$

co jest zgodne ze wzorem (3-15) w twierdzeniu 3-3.

W trakcie dowodu autor pokazał, że przy założeniu, że w rozpatrywanej sieci maksymalne opóźnienie istnieje i są spełnione warunki (3-6), (3-62), (3-60) i (3-61), to opóźnienie maksymalne w sieci DiffServ dla agregatów EF PHB jest ograniczone nierównościami (3-59), co jest zgodne z twierdzeniem 3-3 i kończy pierwszą część dowodu tego twierdzenia.

Część II:

W tej części dowodu autor udowodni, że istnieje skończone opóźnienie maksymalne. Dowód korzysta z metody „time-stopping” [32]. Dla dowolnego czasu t rozważmy system wirtualny zbudowany według modelu rozważanej sieci. Załóżmy, że wszystkie źródła w naszym systemie wirtualnym zostały zatrzymane o czasie t . System wirtualny spełnia założenia z Części I dowodu, ponieważ ilość ruchu na wyjściu jest skończona. Załóżmy, że d_i jest najgorszym przypadkiem opóźnienia we wszystkich węzłach naszego systemu wirtualnego w czasie t . W Części I dowodu zostało udowodnione, że dla twierdzenia 3-3

dla każdego t opóźnienie maksymalne jest ograniczone przez d_l . Przyjmując $t \rightarrow \infty$ pokazujemy, że najgorszy przypadek opóźnienia maksymalnego nadal pozostaje ograniczony przez d' przy spełnionych warunkach do twierdzenia 3-3.

3.2.1.4 Dowód twierdzenia 3-4

Dowód twierdzenia 3-4, podobnie jak dla twierdzenia 3-3, składa się z dwóch części; w trakcie pierwszej części autor dowodzi, że jeśli istnieje skończony kres górny opóźnienia, to twierdzenie 3-4 jest prawdziwe (Część I), natomiast w drugiej części autor wykazuje, że istnieje kres górny opóźnienia (Część II). Dowód twierdzenia 3-4 przebiega bardzo podobnym tokiem myślowym do twierdzenia 3-3.

Część I. Wyprowadzenie wzoru na opóźnienie maksymalne dla twierdzenia 3-4.

Zakładamy, że najgorszy przypadek opóźnienia istnieje i jest ograniczony przez d_l . Na rysunku 3-10 autor przedstawia graficzną interpretację twierdzenia 3-4, które między innymi definiuje, jakie warunki musi spełniać sieć, aby opóźnienie maksymalne, czyli maksymalne horyzontalne odchylenie pomiędzy krzywą ograniczającą ruch przychodzący i krzywą obsługi było mierzone pomiędzy krzywymi:

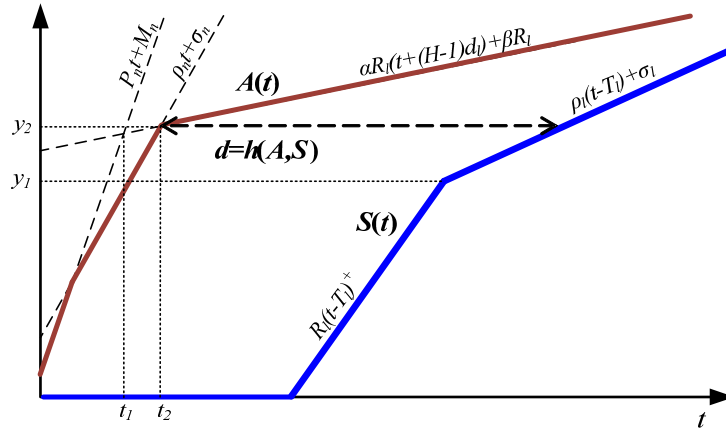
$$A'(t) \leq \min(\rho_n t + \sigma_n, \alpha R_l(t + (H-1)d_l) + \beta R_l) \quad (3-63)$$

i

$$S(t) = \rho_l(t - T_l) + \sigma_l. \quad (3-64)$$

Dla zachowania wzajemnych relacji krzywych $A'(t)$ i $S(t)$ zgodnie z rysunkiem 3-10 (wszystkie relacje rozpatrujemy tylko dla $t > 0$) muszą być spełnione trzy warunki:

- 1) Nachylenie krzywej $\alpha R_l(t + (H-1)d_l) + \beta R_l$ względem osi t będzie mniejsze niż nachylenie krzywych $P_n t + M_n$ i $\rho_n t + \sigma_n$. Ponieważ w przyjętym modelu sieci $R_l \leq C_l$, to $P_n > \alpha R_l$, czyli nachylenie krzywej $P_n t + M_n$ jest zawsze większe od nachylenia krzywej $\alpha R_l(t + (H-1)d_l) + \beta R_l$. Natomiast wystarczającym warunkiem mniejszego nachylenia krzywej $\alpha R_l(t + (H-1)d_l) + \beta R_l$ względem krzywej $\rho_n t + \sigma_n$ jest warunek (3-20) w twierdzeniu 3-4, dla którego zawsze zachodzi nierówność $\rho_n > \alpha R_l$.
- 2) Odcięta t_1 punktu przecięcia krzywej $P_n t + M_n$ z krzywą $\alpha R_l(t + (H-1)d_l) + \beta R_l$ będzie miała wartość mniejszą bądź równą odciętej t_2 punktu przecięcia krzywej $\rho_n t + \sigma_n$ z krzywą $\alpha R_l(t + (H-1)d_l) + \beta R_l$. Warunek ten odpowiada nierówności (3-21) w twierdzeniu 3-4.
- 3) Rzędna y_1 punktu przecięcia krzywych $\rho_n t + \sigma_n$ i $\alpha R_l(t + (H-1)d_l) + \beta R_l$ musi mieć wartość większą od rzędnej y_2 punktu przecięcia krzywych $\rho_l(t - T_l) + \sigma_l$ i $R_l(t - T_l)$. Warunek ten odpowiada nierówności (3-22) w twierdzeniu 3-4.



Rysunek 3-10: Wzajemne położenie krzywej ruchu przychodzącego $A(t)$ i krzywej obsługi $S(t)$ w łączu dla twierdzenia 3-4

Warunek pierwszy dowodu twierdzenia 3-4 pozostaje taki sam jak w dowodach twierdzeń 3-2 i 3-3, warunek drugi twierdzenia 3-4 jest analogiczny do warunku drugiego w dowodzie twierdzenia 3-3, natomiast warunek trzeci twierdzenia 3-4 jest przeciwny do warunku trzeciego w dowodzie twierdzenia 3-3.

Wyznaczenie nierówności (3-21) na podstawie warunku drugiego do twierdzenia 3-4.

Postępując tak jak w dowodzie twierdzenia 3-3 dla warunku drugiego, dla odpowiadającego mu warunku w dowodzie twierdzenia 3-4, otrzymamy, że dla każdego łącza musi być spełniony następujący warunek:

$$\sigma_n \leq \frac{R_l((P_n - \rho_n)\alpha(H-1)d_l + \beta) + M_n(\rho_n - \alpha R_l)}{P_n - \alpha R_l}, \quad (3-65)$$

gdzie d_l jest odległością horyzontalną pomiędzy krzywymi zdefiniowanymi wzorami (3-46) i (3-47). Opóźnienie d_l zostanie wyznaczone w dalszej części dowodu.

Wyznaczenie nierówności (3-22) na podstawie warunku trzeciego do twierdzenia 3-4 polega na zastosowaniu tego samego schematu postępowania jak w przypadku dowodu warunku trzeciego w twierdzeniu 3-3, z jedną różnicą, że w równaniu (3-51) musimy zastosować przeciwny znak, co prowadzi następującej nierówności:

$$\frac{\sigma_l R_l}{R_l - \rho_l} \leq \frac{\rho_n(\alpha R_l(H-1)d_l + \beta R_l) - \sigma_n \alpha R_l}{\rho_n - \alpha R_l}. \quad (3-66)$$

Następnie autor wyznaczy maksymalne opóźnienie dla twierdzenia 3-4. Zakładając, że warunki (3-20), (3-65) i (3-66) są spełnione opóźnienie maksymalne w węźle dla tego twierdzenia jest zdefiniowane przez następującą nierówność:

$$d_l \leq d' = h(U(t_{T_l})(\rho_l(t - T_l) + \sigma_l), \min(\rho_n t + \sigma_n, \alpha R_l t + (\alpha R_l(H-1)d_l + \beta R_l))) \quad (3-67)$$

Pierwsze wyrażenie w funkcji $h(\cdot)$ oznacza, że odległość będzie mierzona pomiędzy częścią funkcji $S(t)$, która pochodzi od mechanizmu kształtowania ruchu, natomiast drugi parametr funkcji h pozostaje taki sam jak w przypadku twierdzenia 3-3. Równanie (3-67) możemy przekształcić do następującej postaci:

$$d_l \leq d' = h\left(U(t_{T_l})\rho_l\left(t - \left(T_l - \frac{\sigma_l}{\rho_l}\right)\right), \min(\rho_n t + \sigma_n, \alpha R_l t + (\alpha R_l(H-1)d_l + \beta R_l))\right) \quad (3-68)$$

Następnie po zastosowaniu wzorów (A-7), (A-8) i (A-9) do powyższego równania i wykonaniu kilku przekształceń, wyrażenie na maksymalne opóźnienie d_l w łączu przyjmie postać:

$$d_l \leq \frac{\sigma_n(\rho_l - \alpha R_l) + (\rho_l T_l - \sigma_l)(\rho_n - \alpha R_l) + (\alpha R_l(H-1)d_l + \beta R_l)(\rho_n - \rho_l)}{\rho_l(\rho_n - \alpha R_l)} \quad (3-69)$$

Po przeniesieniu d_l na prawą stronę równania otrzymujemy:

$$d_l \left(1 - \frac{\alpha R_l(H-1)(\rho_n - \rho_l)}{\rho_l(\rho_n - \alpha R_l)}\right) \leq \frac{\sigma_n(\rho_l - \alpha R_l) + (\rho_l T_l - \sigma_l)(\rho_n - \alpha R_l) + \beta R_l(\rho_n - \rho_l)}{\rho_l(\rho_n - \alpha R_l)} \quad (3-70)$$

Wyrażenie w nawiasie po lewej stronie równania musi być większe od zera, aby powyższa nierówność miała sens, ze względu na możliwość ograniczenia od góry wartości d_l . Rozwiązując nierówność $1 - \frac{\alpha R_l(H-1)(\rho_n - \rho_l)}{\rho_l(\rho_n - \alpha R_l)} > 0$ ze względu na współczynnik wykorzystania pasma α otrzymujemy, że musi on spełniać dla każdego łącza l następującą nierówność:

$$\alpha \leq \frac{\rho_l \rho_n}{((H-1)(\rho_n - \rho_l) + \rho_l)R_l} \quad (3-71)$$

Po dalszych przekształceniach wzoru (3-70) otrzymujemy:

$$d_l \leq \frac{\sigma_n(\rho_l - \alpha R_l) + (\rho_l T_l - \sigma_l)(\rho_n - \alpha R_l) + \beta R_l(\rho_n - \rho_l)}{(\rho_n - \alpha R_l)\rho_l - \alpha R_l(H-1)(\rho_n - \rho_l)} \quad (3-72)$$

Powyższe równanie jest zgodne z równaniem (3-25). Warunki wyznaczone we wzorach (3-65) i (3-66) po podstawieniu obliczonego d_l z powyższego wzoru przyjmą

postać:

$$\sigma_n \leq \frac{M_n(\rho_n \rho_l + \alpha R_l(\rho_l(H-2) - \rho_n(H-1)))}{\rho_l(P_n + \alpha R_l(H-2)) - P_n \alpha R_l(H-1)} + \frac{R_l(P_n - \rho_n)(\alpha \rho_l T_l(H-1) + \beta \rho_l - \alpha \sigma_l(H-1))}{\rho_l(P_n + \alpha R_l(H-2)) - P_n \alpha R_l(H-1)} \quad (3-73)$$

i

$$\frac{\sigma_l R_l}{R_l - \rho_l} < \frac{\rho_n R_l(\alpha R_l - \rho_n)((H-1)T_l \alpha \rho_l + \beta \rho_l - \alpha \sigma_l(H-1))}{(\rho_n - \alpha R_l)(\alpha R_l((H-1)(\rho_n - \rho_l) + \rho_l) - \rho_l \rho_n)} + \frac{\sigma_n((H-2)R_l \alpha \rho_l + ((H-1)R_l \alpha(\rho_l - 1 - \rho_l) + \rho_l)\rho_n)}{(\rho_n - \alpha R_l)(\alpha R_l((H-1)(\rho_n - \rho_l) + \rho_l) - \rho_l \rho_n)} \quad (3-74)$$

Równania (3-73) i (3-74) są zgodne z równaniami (3-21) i (3-22) twierdzenia 3-4. Opóźnienie maksymalne pakietu przy transmisji przez domenę DiffServ jest wyznaczone przez następującą nierówność:

$$D \leq H \max_l \{d_l\}. \quad (3-75)$$

Nierówność (3-73) ze względu na współczynnik wykorzystania pasma jest funkcją homograficzną, która dla parametrów sieci zgodnej z przyjętym modelem posiada asymptotę pionową przecinającą oś odciętych w punkcie $\frac{\rho_l P_n}{(P_n(H-1) - \rho_l(H-2))R_l}$ i aby nierówność ta była spełniona współczynnik wykorzystania pasma musi mieć wartość mniejszą od wartości tego punktu. Uwzględniając warunek (3-71), wspólny warunek na ograniczenie współczynnika pasma przyjmie postać:

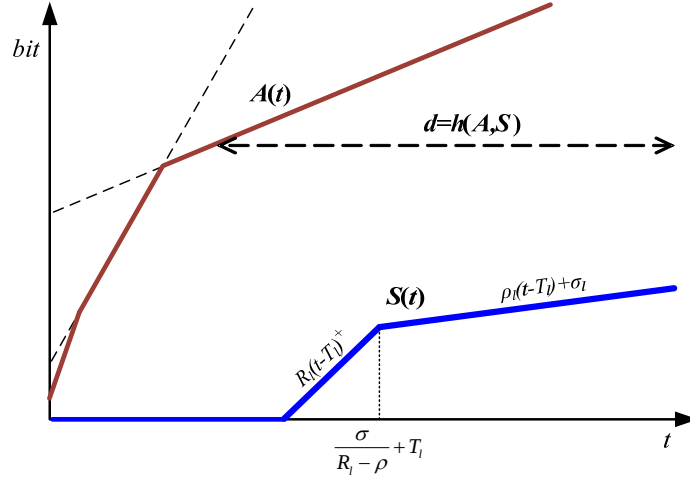
$$\alpha \leq \min_l \left(\frac{\rho_l \rho_n}{((H-1)(\rho_n - \rho_l) + \rho_l)R_l}, \frac{\rho_l P_n}{((H-1)(P_n - \rho_l) + \rho_l)R_l} \right). \quad (3-76)$$

Równania (3-72), (3-73), (3-74), (3-75) i (3-76) są zgodne z równaniami twierdzenia 3-4, co udowadnia pierwszą część dowodu dla tego twierdzenia, że jeśli opóźnienie maksymalne istnieje, to jego wartość jest zdefiniowana twierdzeniem 3-4, jeśli parametry sieci spełniają warunki (3-73), (3-74) i (3-76).

Część II:

W tej części dowodu autor udowodni, że skończone opóźnienie maksymalne istnieje. Dowód korzysta z metody „time-stopping” [32]. Dla dowolnego czasu t rozważmy system wirtualny zbudowany według modelu rozważanej sieci. Załóżmy, że wszystkie źródła w naszym systemie wirtualnym zostały zatrzymane o czasie t . System wirtualny spełnia założenia z Części I dowodu, ponieważ ilość ruchu na wyjściu jest skończona. Załóżmy, że d_l jest najgorszym przypadkiem opóźnienia we wszystkich węzłach naszego systemu

wirtualnego w czasie t . W części pierwszej dowodu zostało udowodnione, że dla twierdzenia 3-4 $d_l < d'$ dla każdego t . Przyjmując $t \rightarrow \infty$ pokazujemy, że najgorszy przypadek opóźnienia maksymalnego nadal pozostaje ograniczony przez d' .



Rysunek: 3-11 Dowód twierdzenia 3-5

3.2.1.5 Dowód twierdzenia 3-5

Dowód został przedstawiony na rysunku 3-11. Jak łatwo zauważyć, jeśli szybkość obsługi algorytmu ciekącego wiadra $\rho_l < \alpha R_l$, to nachylenie krzywej obsługi $S(t)$ dla $t > \frac{\sigma}{R_l - \rho} + T_l$ jest mniejsze od nachylenia krzywej ruchu przychodzącego $A(t)$, bo w tym wypadku $\rho_l < P_n$ i $\rho_l < \rho_n$. To oznacza, że dla $t \rightarrow \infty$ odległość horyzontalna pomiędzy tymi krzywymi też będzie dążyć do nieskończoności. Ponieważ maksymalna odległość horyzontalna pomiędzy krzywą obsługi $S(t)$ a krzywą ograniczającą ruch przychodzący $A(t)$ jest miarą maksymalnego opóźnienia, to maksymalne opóźnienie nie może być ograniczone w przypadku, gdy $\rho_l < \alpha R_l$, co dowodzi poprawności twierdzenia 3-5.

3.3 Analiza uzyskanych wyników

Celem analizy uzyskanych wyników jest porównanie opóźnień maksymalnych dla sieci DiffServ z zastosowanym kształtowaniem ruchu dla agregatów EF PHB z opóźnieniami maksymalnymi w klasycznej sieci DiffServ oraz pokazanie, że **zastosowanie algorytmu**

cieknącego wiadra do kształtowania agregatów EF PHB w każdym węźle sieci IP w architekturze Usług Zróżnicowanych umożliwia zmniejszenie opóźnień maksymalnych pomiędzy brzegami domeny dla usług sieciowych zbudowanych za pomocą reguł przesyłania pakietów EF PHB pod warunkiem, że szybkość obsługi agregatów EF PHB jest mniejsza od przepustowości łączy.

Twierdzenia wyprowadzone przez autora w rozdziale 3.2 obowiązują dla dowolnej architektury sieci, co powoduje, że omawiane twierdzenia znajdują zastosowanie dla EF PHB we wszystkich sieciach DiffServ. Należy zaznaczyć, że w szczególnych przypadkach sieci DiffServ opóźnienia maksymalne mogą być mniejsze od tych, które wyprowadził autor.

Twierdzenie 3-1 jest twierdzeniem ogólnym i jego analiza sprowadza się do interpretacji twierdzeń 3-2, 3-3, 3-4 i 3-5, które są rozwiązaniami twierdzenia dla szczególnych przypadków. To, które twierdzenie będzie miało zastosowanie dla sieci o danych parametrach zależy głównie od wartości parametrów algorytmu ciekącego wiadra, które zostaną zastosowane do kształtowania agregatów EF PHB, odniesionych do parametrów sieci. Porównanie obejmuje twierdzenia 3-3 i 3-4 ponieważ potwierdzają one tezę rozprawy. Twierdzenia 3-3 i 3-4 różnią się zarówno równaniami określającymi opóźnienie maksymalne i maksymalną wartość współczynnika pasma, jak i równaniami określającymi warunki, dla których te twierdzenia są spełnione. Przyczyna różnicy pomiędzy tymi dwoma twierdzeniami jest najbardziej widoczna przy interpretacji graficznej twierdzeń na rysunkach 3-5 i 3-6, gdzie autor wskazał pomiędzy jakimi częściami krzywej ruchu przychodzącego i krzywej obsługi występuje największa odległość horyzontalna, która określa opóźnienie maksymalne, odpowiednio dla twierdzenia 3-3 i 3-4.

Twierdzenia 3-2 i 3-5 pokazują dla jakich parametrów algorytmu ciekącego wiadra nie ma sensu stosować kształtowania ruchu dla agregatów EF PHB z punktu widzenia opóźnień maksymalnych i wykorzystania pasma w sieci DiffServ. Są one zaprzeczeniem tezy stawianej przez autora w rozprawie, a zostały zaprezentowane dla przedstawienia pełnego obrazu wpływu kształtowania ruchu agregatów EF PHB na opóźnienia maksymalne w sieciach DiffServ. Twierdzenie 3-2 nie jest porównywane, ponieważ w przypadku tego twierdzenia opóźnienie maksymalne w sieci z zastosowanym kształtowaniem ruchu dla agregatów EF PHB ma takie same właściwości jak w odpowiadającej jej klasycznej sieci DiffServ. Natomiast interpretacja twierdzenia 3-5 jest prosta i jednoznaczna: jeśli zastosujemy kształtowanie ruchu algorytmem ciekącego wiadra dla agregatu EF PHB, w którym szybkość obsługi jest mniejsza od sumy przepływności strumieni wchodzących w skład agregatu priorytetowego obsługiwanego przez EF PHB, to nie można oszacować maksymalnego opóźnienia domenie DiffServ dla tego agregatu.

3.3.1 Porównanie opóźnień w klasycznej sieci DiffServ z rozwiązaniem zaproponowanym przez autora w twierdzeniu 3-3 dla EF PHB

W poniższym rozdziale autor przedstawia porównanie wartości opóźnień maksymalnych w sieci z zastosowanym kształtowaniem ruchu dla agregatu EF PHB algorytmem ciekącego wiadra w każdym węźle sieci i spełniającej warunki twierdzenia 3-3 z wartościami opóźnień maksymalnych w klasycznej sieci DiffServ o tych samych parametrach. W klasycznej sieci DiffServ opóźnienia maksymalne i warunek na maksymalne obciążenie są opisane przez Y. Jianga równaniami (2-13) i (2-15), które zostały zamieszczone w rozdziale 2.3.1. Głównym zadaniem porównania jest pokazanie, dla jakich parametrów sieci z zastosowanym kształtowaniem ruchu spełniającym warunki twierdzenia 3-3 możliwe jest zmniejszenie opóźnień maksymalnych w sieci DiffServ.

W pierwszej części bieżącego rozdziału autor przeprowadzi analityczne porównanie, które potwierdza tezę stawianą przez autora rozprawy, pomiędzy równaniem (2-13) opisującym opóźnienie maksymalne pomiędzy brzegami klasycznej domeny DiffServ a równaniem (3-16) twierdzenia 3-3 opisującym opóźnienie maksymalne pomiędzy brzegami domeny DiffServ z zastosowaniem kształtowania pasma algorytmem ciekącego wiadra dla agregatów EF PHB. W porównaniu zostaną także użyte równania (2-14) i (2-10) uzupełniające równanie (2-13) oraz równania (3-17) i (3-18) twierdzenia 3-3. Dla ułatwienia przeprowadzenia analizy w tabeli 3-1 autor przypomina postać równań podlegających bezpośredniemu porównaniu.

Porównując równania (2-13) i (3-16) łatwo zauważyć, że forma tych równań jest podobna, natomiast różnią się one tylko dwoma zmiennymi: w równaniu (2-13) występują zmienne u i E' , którym odpowiadają zmienne v i E'' w równaniu (3-16). Zmienne E' (równanie (2-14)) i E'' (równanie (3-17)) są także podobne do siebie, a różnią się między sobą zmiennymi u i v oraz M_n i σ_n . Ponieważ zarówno równanie (2-13) jak i równanie (3-16) są funkcjami rosnącymi względem parametrów u i v , to przy założeniu że $\sigma_n = M_n$, o tym, które równanie będzie miało większą wartość zdecyduje, który z parametrów, u czy v , będzie miał wartość większą. Parametry u i v zostały zdefiniowane odpowiednio równaniami (2-10) i (3-18), które różnią się jednym parametrem: w równaniu (2-10) występuje zmienna P_n , która w równaniu (3-18) jest zastąpiona zmienną ρ_n .

Równania (2-10) i (3-18) są funkcjami rosnącymi względem parametrów odpowiednio P_n i ρ_n , co implikuje, że w sieci w której zastosowano kształtowanie ruchu spełniające warunki twierdzenia 3-3, aby wartość opóźnienia maksymalnego była mniejsza od wartości opóźnienia maksymalnego w odpowiadającej jej sieci bez kształtowania ruchu musi zachodzić warunek $\rho_n < P_n$. Na podstawie definicji zmiennych P_n i ρ_n powyższy warunek odpowiada warunkowi $\rho_i < C_i$. Opóźnienie maksymalne w sieci

z kształtowaniem ruchu będzie tym mniejsze od opóźnienia w odpowiadającej jej sieci bez kształtowania ruchu, im różnica pomiędzy C_l i ρ_l będzie większa. Biorąc pod uwagę warunek (3-10) twierdzenia 3-3, aby opóźnienie maksymalne w sieci z kształtowaniem ruchu mogło być mniejsze od opóźnienia maksymalnego w odpowiadającej jej sieci bez kształtowania ruchu, to musi być spełniony warunek $R_l < \rho_l < C_l$, co oznacza, że szybkość obsługi agregatu EF PHB musi być mniejsza od przepustowości łącza ale większa od szybkości obsługi agregatu EF PHB dla każdego łącza $l \in L$.

Tabela 3-1: Przypomnienie równań porównywanych w rozdziale 3.3.1

Porównywane równania Y. Jianga:	
$D = \frac{H}{1 - (H-1)u\alpha} (E' + u\beta), \text{ gdzie}$	(2-13)
$E' = \max_l \left\{ \frac{(1-u_l)M_n}{R_l} + \frac{M_l}{R_l} + E_l \right\} \text{ i}$	(2-14)
$u = \max_l \{u_l\}, \text{ gdzie } u_l = \frac{P_n - R_l}{P_n - \alpha R_l}.$	(2-10)
$\alpha < \min_l \frac{P_n}{(P_n - R_l)(H-1) + R_l}$	(2-15)
Porównywane równania J. Gozdeckiego:	
$D \leq \frac{H}{1 - (H-1)v\alpha} (E'' + v\beta) \text{ gdzie}$	(3-16)
$E'' = \max_l \left\{ \frac{(1-v_l)\sigma_n}{R_l} + \frac{M_l}{R_l} + E_l \right\} \text{ i}$	(3-17)
$v = \max_l \{v_l\}, \text{ gdzie } v_l = \frac{\rho_n - R_l}{\rho_n - \alpha R_l}.$	(3-18)
$\alpha < \min_l \left(\frac{\rho_n}{(\rho_n - R_l)(H-1) + R_l}, \frac{P_n}{(P_n - R_l)(H-1) + R_l} \right)$	(3-15)

Porównując równania (2-15) i (3-15) opisujące maksymalną wartość współczynnika wykorzystania pasma α odpowiednio w klasycznej sieci DiffServ i w sieci DiffServ z kształtowaniem pasma dla agregatów EF PHB spełniających warunki twierdzenia 3-3

łatwo zauważyć, że w sieci z kształtowaniem pasma wartość współczynnika wykorzystania pasma α nie może być większa od wartości współczynnika wykorzystania pasma dla klasycznej sieci DiffServ, ponieważ jednym z argumentów funkcji minimum w równaniu (3-15) jest równanie (2-15).

W tabeli 3-2 zostały przedstawione parametry sieci dla rysunków 3-12, 3-14 i 3-18. W porównaniu z siecią bez kształtowania ruchu sieć z kształtowaniem ruchu jest opisana dodatkowo przez dwa parametry: $\rho_l = \eta C_l$ i $\sigma_l = \varepsilon C_l$. W tabeli tej występują dwa rodzaje parametrów: zależnych od konfiguracji węzła lub łącza w tym węźle (takie jak: R_l , C_l , P_n , M_n , E_l , T_l) i parametrów, które mają taką samą wartość we wszystkich węzłach sieci (M , α , β , ε , η). W bieżącym rozdziale przez rozmiar sieci autor rozumie wartość zmiennej H , czyli maksymalną liczbę łączy, jaką pakiet musi przebyć, aby zostać przesłanym pomiędzy dowolnymi dwoma ruterami brzegowymi sieci. Dla rozważań w tym rozdziale możemy przyjąć, że tabela 3-2 przedstawia parametry dla sieci homogenicznej, w której wszystkie węzły mają taką samą konfigurację, lub dla węzła w sieci heterogenicznej, w którym występują największe opóźnienia.

Tabela 3-2: Parametry sieci dla rysunków 3-12, 3-13, 3-14, 3-18 i 3-19

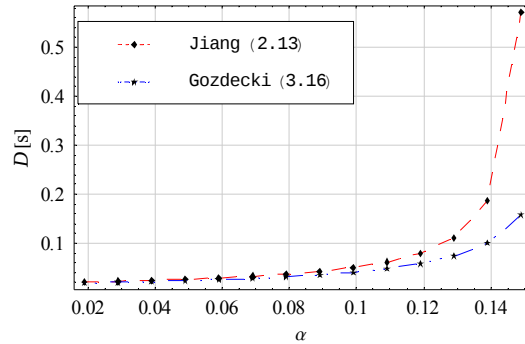
Parametr Rysunek	R_l [b/s]	C_l [b/s]	M [B]	P_n [b/s]	M_n [B]	H
Rysunek 3-12, 3-13	5×10^8	1×10^9	1500	6×10^9	9000	8
Rysunek 3-14	5×10^8	1×10^9	1500	6×10^9	9000	2+8
Rysunek 3-18, 3-19	5×10^8	1×10^9	1500	6×10^9	9000	2+8
Parametr / rysunek	E_l [s]	T_l [s]	ε [s]	η	α	β [s]
Rysunek 3-12, 3-13	0.000012	0.000036	0.000012	0.505	0.01+0.149	0.0024
Rysunek 3-14	0.000012	0.000036	0.000012	0.505	-	0.0024
Rysunek 3-18, 3-19	0.000012	0.000036	0.000012	0.505	0.149	0.0024

Tabela 3-2 jest wyjściowym zestawem parametrów do pokazania ogólnych mechanizmów odpowiedzialnych za opóźnienia w sieciach typu DiffServ z obsługą zagregowanych strumieni danych i te parametry mają zastosowanie także dla pozostałych rysunków w bieżącym rozdziale.

Wartość opóźnienia maksymalnego dla równania (2-13) i twierdzenia 3-3 silnie zależy od parametrów H , α , β , R_l i P_n oraz dodatkowo dla sieci z kształtowaniem ruchu we wszystkich węzłach sieci od parametrów R_l/C_l , ε i η . Parametry ε i η odpowiadają za

parametry algorytmu ciekącego wiadra oraz pośrednio odpowiadają za wartość stosunku szybkości obsługi agregatu priorytetowego R_i względem szybkości obsługi w algorytmie ciekącego wiadra, a więc za parametry R_i / ρ_i oraz R_i / ρ_n .

Równania (2-10) i (3-18) są funkcjami rosnącymi względem parametrów odpowiednio P_n i ρ_n , co implikuje, że w sieci w której zastosowano kształtowanie ruchu spełniające warunki twierdzenia 3-3, aby wartość opóźnienia maksymalnego była mniejsza od wartości opóźnienia maksymalnego w odpowiadającej jej sieci bez kształtowania ruchu musi zachodzić warunek $\rho_n < P_n$. Na podstawie definicji zmiennych P_n i ρ_n powyższy warunek odpowiada warunkowi $\rho_i < C_i$. Opóźnienie maksymalne w sieci z kształtowaniem ruchu będzie tym mniejsze od opóźnienia w odpowiadającej jej sieci bez kształtowania ruchu im różnica pomiędzy C_i i ρ_i będzie większa. Biorąc pod uwagę warunek (3-10) twierdzenia 3-3, aby opóźnienie maksymalne w sieci z kształtowaniem ruchu mogło być mniejsze od opóźnienia maksymalnego w odpowiadającej jej sieci bez kształtowania ruchu, to musi być spełniony warunek $R_i < \rho_i < C_i$, co oznacza, że szybkość obsługi agregatu EF PHB musi być mniejsza od przepustowości łącza, ale większa od szybkości obsługi agregatu EF PHB dla każdego łącza $l \in L$.

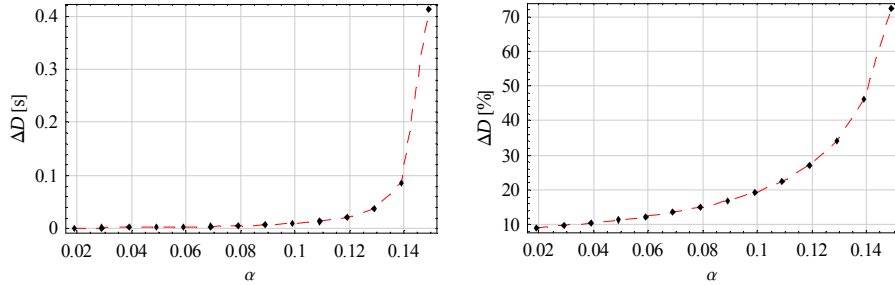


Rysunek 3-12: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od współczynnika wykorzystania pasma α

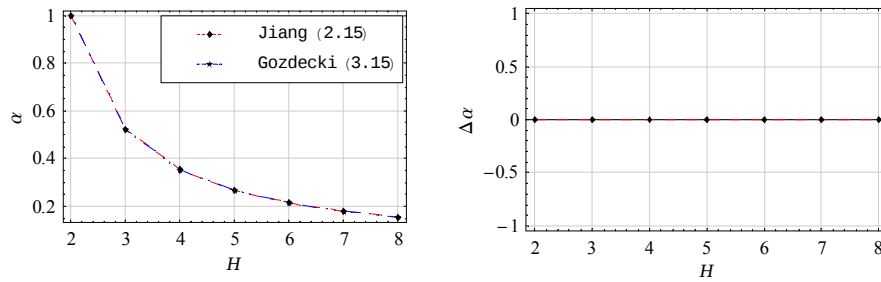
Rysunki 3-12 i 3-13 przedstawiają zależność opóźnienia maksymalnego od współczynnika wykorzystania pasma dla równań (2-13) i (3-16). W przypadku klasycznej sieci DiffServ opóźnienia maksymalne (równanie (2-13)) rosną homograficznie ze wzrostem współczynnika wykorzystania pasma α . Dla rozwiązania zaproponowanego przez autora w twierdzeniu 3-3 wzrost opóźnień ma ten sam charakter, ale opóźnienia maksymalne są znacznie mniejsze. Zysk w postaci obniżenia opóźnień maksymalnych jest tym większy im większe jest obciążenie sieci i sięga od około 6.5% dla $\alpha = 0.01$ do 90% dla $\alpha = 0.15$. Różnica bezwzględna została wyliczona według następującego wzoru

$$\Delta D(\alpha) = \frac{D(\alpha)_{(3-13)} - D(\alpha)_{(3-16)}}{D(\alpha)_{(3-13)}} 100\%.$$

Opisane powyżej właściwości rozwiązania zaproponowanego przez autora w twierdzeniu 3-3 powodują, że dla parametrów sieci podanych w tabeli 3-2 daje ono lepsze rezultaty od rozwiązania opisanego wzorem (2-13).



Rysunek 3-13: Różnica opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od współczynnika wykorzystania pasma α (lewy wykres różnica bezwzględna, prawy wykres różnica względna)



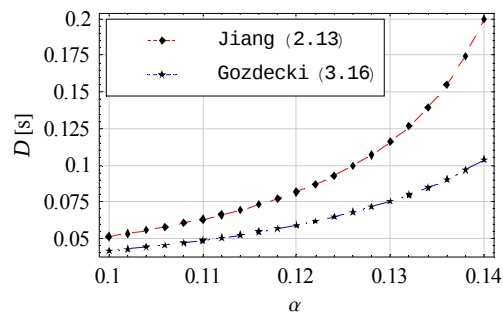
Rysunek 3-14: Porównanie współczynnika wykorzystania pasma dla równania (2-15) i twierdzenia 3-3 (równanie (3-15)) w zależności od rozmiaru sieci H

Na rysunku 3-14 autor przedstawia porównanie maksymalnej wartości współczynnika wykorzystania pasma α dla twierdzenia 3-3 – równanie (3-16) – i równania (2-15) w zależności od rozmiaru sieci H . Na prawym wykresie rysunku 3-14 autor przedstawia wykresy równań (2-15) i (3-16) w zależności od zmiennej H , natomiast lewy wykres tego rysunku przedstawia różnicę pomiędzy wartościami równań (2-15) i (3-16) wyrażoną w wartościach bezwzględnych $\Delta\alpha(H) = \alpha(H)_{(3-15)} - \alpha(H)_{(2-15)}$.

Z rysunku 3-14 wynika, że współczynnik wykorzystania pasma w obydwu systemach ma taką samą wartość. Ponieważ w systemie z kształtowaniem ruchu dla rozpatrywanego agregatu obsługa pakietów jest typu FIFO, tak jak w systemie bez kształtowania ruchu, to

wynik ten jest potwierdzeniem wyników pracy [31], w której zostało udowodnione, że dla wartości współczynnika wykorzystania pasma α większej od wartości maksymalnej wyznaczonej wzorem (2-8), opóźnienie maksymalne może być dowolnie duże.

Praca sieci na granicy maksymalnej wartości współczynnika wykorzystania pasma może być niekorzystna z powodu stosunkowo dużych zmian wartości opóźnienia maksymalnego przy niewielkiej zmianie współczynnika wykorzystania pasma. Ponadto, opóźnienia maksymalne dla współczynnika wykorzystania pasma osiągnącego swą wartość maksymalną są stosunkowo duże i dlatego jako kryterium doboru parametrów sieci może zostać wybrane opóźnienie maksymalne. Załóżmy, że dla sieci scharakteryzowanej w tabeli 3-2 maksymalne opóźnienie nie może przekroczyć 100 ms.

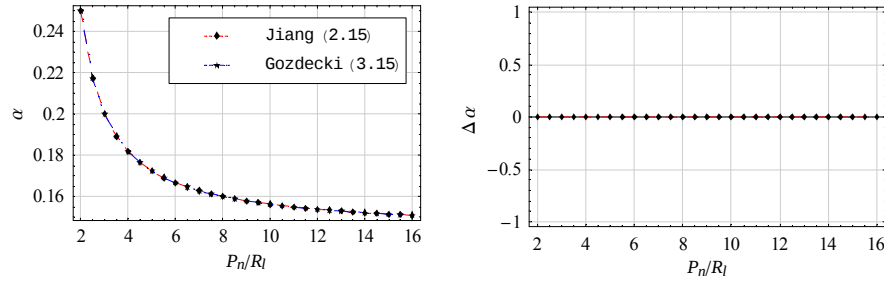


Rysunek 3-15: Opóźnienie maksymalne dla równania (2-13) i twierdzenia 3-3 w zależności od współczynnika wykorzystania pasma α

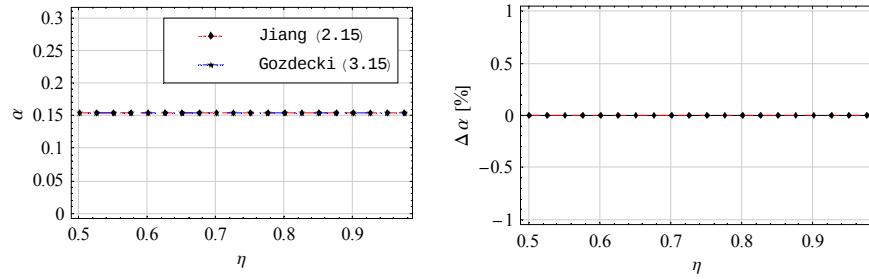
Z rysunku 3-15 dla takiego przypadku możemy odczytać, że maksymalna wartość współczynnika wykorzystania pasma dla sieci bez kształtowania ruchu wyniesie około 12.6%, natomiast w przypadku stosowania kształtowania ruchu jest to około 13.8%, a to oznacza, że maksymalne obciążenie sieci może być o około 9.5% większe dla systemu z kształtowaniem ruchu w stosunku do sieci bez kształtowania ruchu. Jeśli założymy, że opóźnienie maksymalne nie może przekroczyć 50 ms to zysk będzie wynosił około 12%. Z wyżej przedstawionych rezultatów wynika, że jeśli dla parametrów danej sieci bez kształtowania ruchu współczynnik wykorzystania pasma ma wartość mniejszą od wartości maksymalnej danej wzorem (2-15), ponieważ opóźnienia maksymalne przekraczają dla niego ustalony próg, to w odpowiadającej jej sieci z kształtowaniem ruchu współczynnik wykorzystania pasma będzie miał wartość większą dla tego samego progu opóźnienia maksymalnego. Warunek ten zachodzi tylko wtedy, gdy w łączach sieci szybkość obsługi agregatu EF PHB jest mniejsza od przepustowości łącz.

Rysunek 3-16 przedstawia wpływ stosunku parametrów P_n/R_l na maksymalną wartość współczynnika wykorzystania pasma α . Wraz ze wzrostem stosunku P_n/R_l maksymalna wartość współczynnika wykorzystania pasma maleje i jego minimum dla

$P_n \rightarrow \infty$ wynosi $\frac{1}{H-1}$.



Rysunek 3-16: Różnica maksymalnej wartości współczynnika wykorzystania pasma α dla twierdzenia 3-3 (równanie (3-15)) i równania (2-15) w zależności od stosunku P_n/R_l dla $H = 8$

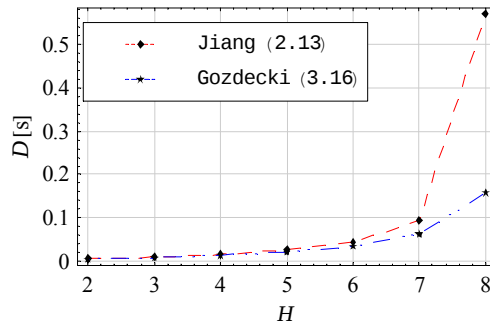


Rysunek 3-17: Porównanie maksymalnej wartości współczynnika wykorzystania pasma α dla równania (2-15) i twierdzenia 3-3 (równanie (3-15)) w zależności od współczynnika η – lewy wykres przedstawia wartości bezwzględne, prawy wykres przedstawia różnicę względną

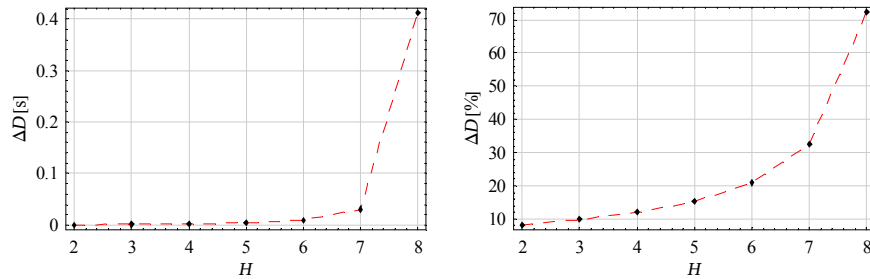
Rysunek 3-17 przedstawia zależność współczynnika wykorzystania pasma α dla równania (2-15) i twierdzenia 3-3 (równanie (3-15)) i w zależności od współczynnika η . Dla tego rysunku pozostałe parametry sieci zostały zdefiniowane w tabeli 3-2. Lewy wykres rysunku 3-17 przedstawia wartości bezwzględne dla równań (2-15) i (3-15), natomiast prawy wykres przedstawia różnicę względną pomiędzy tymi równaniami. Z wykresu wynika, że wartość parametru η nie wpływa na współczynnik wykorzystania pasma α dla parametrów sieci z tabeli 3-2, dla których twierdzenie 3-3 jest prawdziwe w zakresie $0.5 < \eta < 1$.

Na rysunkach 3-18 i 3-19 autor przedstawia porównanie opóźnień maksymalnych dla równań (2-13) i (3-16) w zależności od rozmiaru sieci H dla ustalonej wartości parametru $\alpha = 0.15$. Parametr α został tak dobrany, aby wykorzystanie pasma było jak

największe i równania (2-13) i (3-16) były jednocześnie prawdziwe w całym zakresie zmiennej H .



Rysunek 3-18: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od rozmiaru sieci H



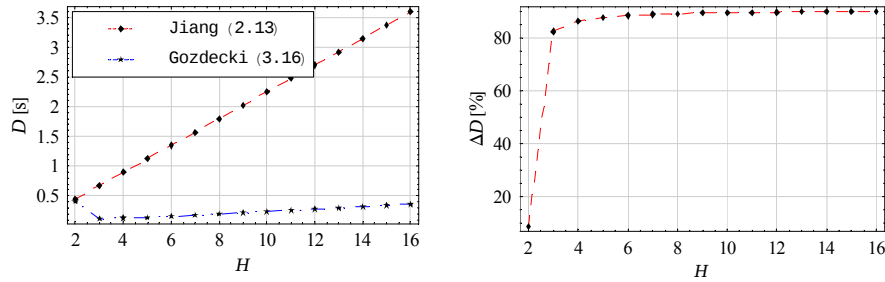
Rysunek 3-19: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od rozmiaru sieci H

Dla przypadku z rysunków 3-18 i 3-19 opóźnienie maksymalne w systemie z kształtowaniem ruchu jest mniejsze od opóźnienia dla systemu bez kształtowania ruchu. Opóźnienie maksymalne w granicach zmienności H w systemie zaproponowanym przez autora jest mniejsze od około 7% do 90% w porównaniu z siecią bez kształtowania ruchu. Największy zysk z zastosowania kształtowania ruchu jest dla $H=8$, dla którego to parametru wartość współczynnika wykorzystania pasma α zbliża się do maksymalnej dopuszczalnej wartości.

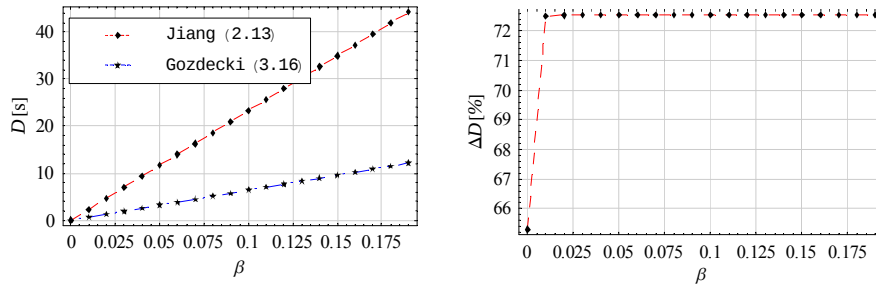
Na rysunku 3-20 autor przedstawia zależność opóźnienia maksymalnego dla równania (2-13) i twierdzenia 3-3 w zależności od rozmiaru sieci H dla współczynnika wykorzystania pasma α obliczanego dla każdego punktu według wzoru

$$\alpha(H) = 0.99 \frac{P_n}{(P_n - R_l)(H - 1) + R_l} . \quad (3-77)$$

Oznacza to, że dla każdej wartości H parametr α ma wartość 99% wartości dopuszczalnej dla równania (2-13). Pozostałe parametry sieci są takie same jak dla rysunków 3-12 i 3-18 (tabela 3-2). Z wykresu na rysunku 3-20 widzimy, że jeśli wartość współczynnika wykorzystania pasma α dla danego rozmiaru sieci H zbliża się do wartości maksymalnej, to zysk z zastosowania kształtowania ruchu jest stosunkowo duży dla całego zakresu zmiennej H i zwiększa się ze wzrostem tego parametru, ale charakter tego wzrostu jest zupełnie inny niż na rysunku 3-18.



Rysunek 3-20: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od rozmiaru sieci H dla współczynnika wykorzystania pasma α obliczanego dla każdego punktu według wzoru (3-77).



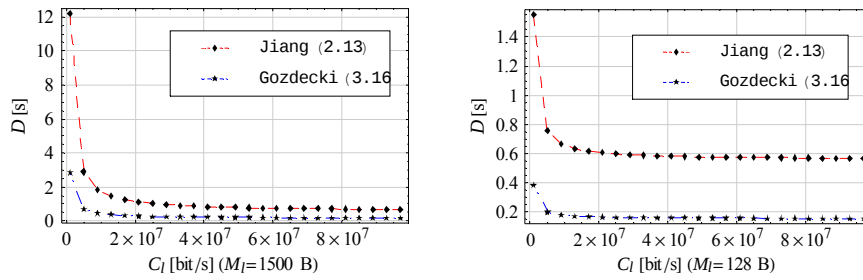
Rysunek 3-21: Porównanie wpływu wartości parametru β na opóźnienie maksymalne dla równania (2-13) i twierdzenia 3-3

Rysunek 3-21 przedstawia wpływ parametru β na opóźnienie maksymalne dla równania (2-13) i twierdzenia 3-3 – prawy wykres przedstawia wartości bezwzględne dla równań (2-13) i (3-16), natomiast lewy różnicę względną pomiędzy tymi równaniami liczoną według wzoru $\Delta D(\beta) = \frac{D(\beta)_{(2-13)} - D(\beta)_{(3-16)}}{D(\beta)_{(2-13)}} 100\%$. Pozostałe parametry sieci mają wartość zdefiniowaną w tabeli 3-2.

Wraz ze wzrostem parametru β opóźnienie maksymalne rośnie w przybliżeniu

liniowo zarówno dla równania (2-13) jak i (3-16), jednak w przypadku równania (3-16) wzrost opóźnienia jest znacznie mniejszy, co powoduje, że dla parametrów z tabeli 3-2 opóźnienia są o około 70% mniejsze dla systemu z kształtowaniem ruchu, niż w systemie, w którym nie zastosowano kształtowania ruchu zaproponowanego przez autora. Zysk jest obserwowalny w szerokim zakresie zmiennej β . Minimalna wartość zmiennej β wynosi 0.000012 co odpowiada rozmiarowi wiadra o wielkości jednego pakietu o rozmiarach MTU dla sieci Ethernet. Maksymalna wartość β jest znacznie większa niż zastosowany zakres tej zmiennej na rysunku.

Na rysunku 3-22 autor przedstawia porównanie maksymalnych opóźnień dla równania (2-13) i twierdzenia 3-3 w zależności od C_l i M_l przy założeniu że parametry H , α , β , ε i η mają stałą wartość określoną w tabeli 3-2, natomiast $R_l/C_l = 0.5$, $R_l/P_n = 0.0625$.

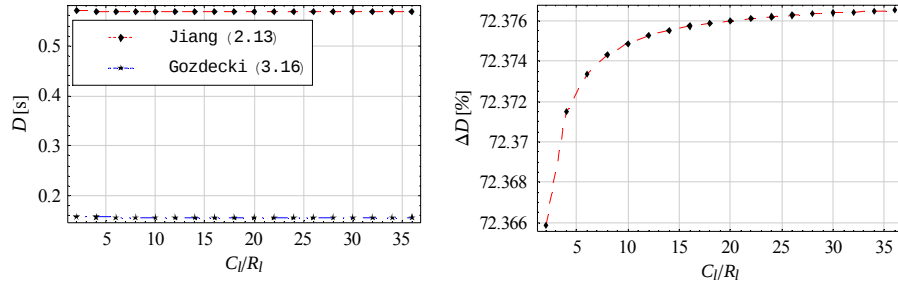


Rysunek 3-22: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od C_l . Lewy wykres dla $M_l = 1500$ B, prawy wykres dla $M_l = 128$ B

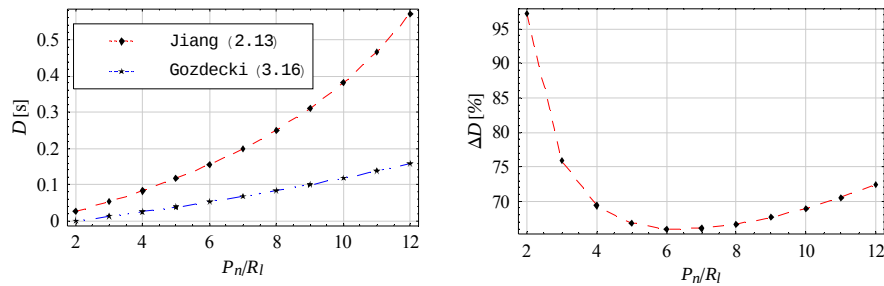
Dla prawego wykresu na rysunku 3-22 parametr $M_l = 1500$ B, dla lewego wykresu $M_l = 128$ B. Parametr M_l ma znaczący wpływ na opóźnienie maksymalne dla łącza o niskiej przepustowości (poniżej kilku Mbit/s dla sieci typu Ethernet, dla której maksymalny rozmiar pakietu wynosi 1500 B) i jest to związane z wartością stosunków M_l/C_l i M_l/R_l , które przy określonej wartości zmiennej C_l (dla rysunku 3-22 powyżej 6 Mbit/s) przestają mieć znaczący wpływ na opóźnienie maksymalne. Pomijając wpływ rozmiaru maksymalnego pakietu na rysunku 3-22, zwiększenie przepustowości łącza nie powoduje zmniejszenia opóźnień maksymalnych, jeśli parametry H , α , β , ε , η , R_l/C_l , R_l/P_n mają stałą wartość.

Rysunki 3-23 i 3-24 przedstawiają wpływ parametrów R_l/C_l i P_n/R_l na opóźnienia maksymalne dla równań (2-13) i (3-16). Dla rysunku 3-23 zmiana wartości parametru R_l/C_l odbywa się przez zmianę przepustowości łącza C_l , współczynnik $\eta = 1.01R_l/C_l$, natomiast pozostałe parametry mają wartości zdefiniowane w tabeli 3-2. Stała 1.01 została dobrana na podstawie rysunku 3-25, tak aby wyeksponować zalety stosowania

kształtowania ruchu. Jeśli szybkość obsługi agregatu priorytetowego jest równa przepustowość łącza, to wprowadzenie kształtowania ruchu nie zmniejszy opóźnień maksymalnych. Zmniejszenie opóźnień maksymalnych możemy uzyskać przez zastosowanie kształtowania ruchu jedynie w przypadku, gdy $C_l / R_l > 1$.



Rysunek 3-23: Wpływ stosunku C_l / R_l na opóźnienie maksymalne dla równania (2-13) i twierdzenia 3-3, przy założeniu, że $\eta = 1.01 R_l / C_l$ oraz wartości R_l i P_n są stałe

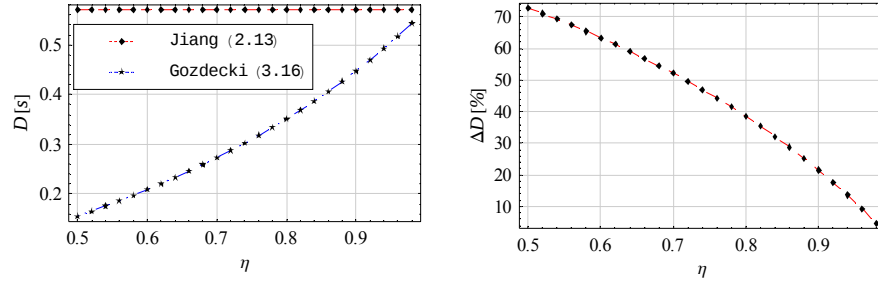


Rysunek 3-24: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od parametru P_n / R_l , przy założeniu, że wartości C_l i R_l są stałe

Na rysunku 3-24 autor przedstawia porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-3 w zależności od parametru P_n / R_l przy założeniu, że wartość C_l jest stała. Wzrost wartości stosunku P_n / R_l powoduje zwiększenie opóźnień maksymalnych – wzrost ten może mieć dwojakie źródło: zwiększenie liczby łączy wchodzących do węzła lub zwiększenie przepustowości łączy wchodzących do węzła – w obydwu tych przypadkach prawdopodobieństwo, że w krótkim odstępie czasu do łącza wychodzącego nadejdzie większa liczba pakietów rośnie, co powoduje, że w takim systemie potencjalnie opóźnienia maksymalne będą większe.

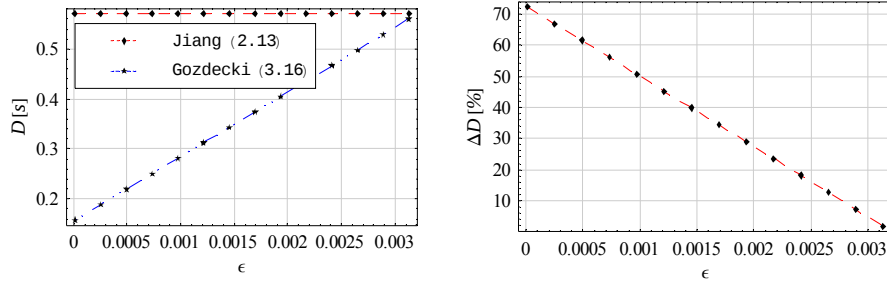
Mechanizm ciekącego wiadra zmniejsza możliwość sklejanie się pakietów należących do agregatów EF PHB w większe paczki danych, gdy szybkość obsługi agregatu EF PHB w łączach jest mniejsza od przepustowości łącz. Jeśli szybkość obsługi w algorytmie

cieknącego wiadra zmierza do wartości przepustowości łącza ($\eta \rightarrow 1$), wtedy opóźnienia maksymalne są takie same w systemie z i bez kształtowania ruchu, co zostało przedstawione na rysunku 3-25.



Rysunek 3-25: Wpływ parametru η na różnicę opóźnienia maksymalnego pomiędzy równaniem (2-13) i twierdzeniem 3-3

Na wykresach rysunku 3-25 zakres zmienności parametru η jest ograniczony nierównością $R_i < \eta C_i < C_i$, dla której, jeśli jest spełniony warunek opisany równaniem (3-15), to zawsze można dobrać parametry η i ε , tak by twierdzenie 3-3 było prawdziwe.



Rysunek 3-26: Wpływ parametru ε na opóźnienia maksymalne dla twierdzenia 3-3 i różnicę opóźnień pomiędzy równaniem (2-13) i twierdzeniem 3-3

Na rysunku 3-26 autor przedstawia porównanie wpływu parametru ε na opóźnienia maksymalne dla równania (2-13) i twierdzenia 3-3. Równanie (2-13) jest niezależne od zmiennej ε i zostało umieszczone na lewym wykresie rysunku 3-26 tylko dla porównania. Dla równania (3-16) najmniejsze opóźnienie jest uzyskiwane dla $\varepsilon = M / C_i$, co oznacza, że rozmiar bufora w algorytmie ciekącego wiadra $\sigma_i = M$. Jest to najmniejsza wartość jaką może przybrać rozmiar bufora w algorytmie ciekącego wiadra. Prawy wykres na rysunku 3-26 przedstawia różnicę względną pomiędzy opóźnieniami maksymalnym dla

równania (2-13) i (3-16) liczoną według wzoru $\Delta D(\varepsilon) = \frac{D(\varepsilon)_{(2-13)} - D(\varepsilon)_{(3-16)}}{D(\varepsilon)_{(2-13)}} 100\%$.

Podsumowując rozważania bieżącego podrozdziału, dla sieci DiffServ o dowolnej architekturze scharakteryzowanej przez zakresy wartości zmiennych z rysunków bieżącego rozdziału, najważniejsze zalety i wady rozwiązania zaproponowanego przez autora dla EF PHB w twierdzeniu 3-3 są następujące:

Zalety:

- Stosując kształtowanie ruchu algorytmem ciekącego wiadra dla agregatów EF PHB sieci DiffServ można zmniejszyć opóźnienie maksymalne w sieciach o różnych rozmiarach i konfiguracji pod warunkiem, że zastosowane algorytmy szeregowania w łączach obsługują EF PHB z szybkością mniejszą niż przepustowość łączy. Powyższe wymagania spełniają takie algorytmy jak WFQ, SCFQ oraz VC [13].
- Prosty algorytm ciekącego wiadra stosowany dla agregatu danych nie będzie wpływał na skalowalność sieci.
- Istnieje możliwość zwiększenia wykorzystania pasma dla EF PHB w sieci. Jeśli dla danej sieci bez kształtowania ruchu współczynnik wykorzystania pasma ma wartość mniejszą od wartości maksymalnej danej wzorem (2-15), ponieważ opóźnienia maksymalne przekraczają dla niego ustalony próg, to w odpowiadającej jej sieci z kształtowaniem ruchu współczynnik wykorzystania pasma będzie miał wartość większą dla tego samego progu opóźnienia maksymalnego. Warunek ten zachodzi tylko, gdy szybkość obsługi agregatu EF PHB jest mniejsza od szybkości łączy w węzłach sieci.

Wady:

- Kształtowanie ruchu za pomocą algorytmu ciekącego wiadra dla agregatów EF PHB w sieci DiffServ nie powoduje zmniejszenia opóźnień maksymalnych, jeśli w węzłach zostaną zastosowane algorytmy szeregowania, które obsługują EF PHB z szybkością przepustowości łączy. Powyższe warunki spełniają algorytmy FIFO oraz Strict Priority dla klasy o najwyższym priorytecie.
- Konieczna implementacja dodatkowych algorytmów kształtowania ruchu w węzłach dla agregatów w sieci DiffServ i uwzględnienie nowych parametrów w systemie do zarządzania siecią.

Zysk z zastosowania kształtowania ruchu w sieci jest tym większy, im:

- większe jest wykorzystanie pasma α ,

- większy jest rozmiar sieci mierzony parametrem H ,
- większa jest wartość parametru β charakteryzującego wybuchowość ruchu przychodzącego,
- wartość stosunku R_l / C_l jest mniejsza,
- wartość parametru P_n jest większa.

3.3.2 Porównanie opóźnień w klasycznej sieci DiffServ z rozwiązaniem zaproponowanym przez autora w twierdzeniu 3-4 dla EF PHB

W bieżącym rozdziale autor przedstawia porównanie wartości opóźnień maksymalnych w sieci DiffServ z zastosowanym kształtowaniem ruchu dla agregatu EF PHB algorytmem ciekącego wiadra w każdym węźle sieci i spełniającej warunki twierdzenia 3-4 z opóźnieniami maksymalnymi w klasycznej sieci DiffServ. Podobnie jak w rozdziale poprzednim celem porównania jest pokazanie, dla jakich parametrów sieci spełniających warunki twierdzenia 3-4 można zmniejszyć opóźnienie maksymalne w sieci DiffServ przez zastosowanie kształtowania ruchu dla agregatów EF PHB w każdym węźle sieci.

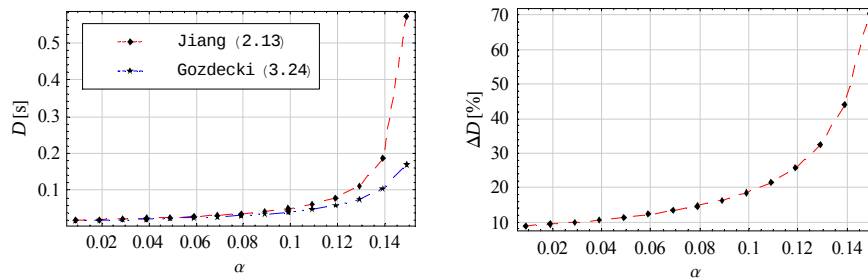
W tabeli 3-3 zostały przedstawione parametry sieci dla rysunków z bieżącego rozdziału. Dla uzyskania możliwości porównania wyników z wynikami z poprzedniego rozdziału parametry tabeli 3-3 różnią się od parametrów z tabeli 3-2 tylko parametrami odpowiedzialnymi za kształtowanie ruchu, tak aby zostały spełnione warunki twierdzenia 3-4. Z porównania parametrów z wyżej wymienionych tabel wynika, że możliwe jest zastosowanie różnych wariantów kształtowania ruchu w tej samej sieci DiffServ, zmieniając tylko wartości parametrów odpowiedzialnych za kształtowanie ruchu. Uwagi z rozdziału 3.3.1 odnośnie parametrów sieci obowiązują także dla tego rozdziału.

W tabeli 3-3 występują dwa rodzaje parametrów: zależnych od konfiguracji węzła lub łącza w tym węźle, takich jak R_l , C_l , P_n , M_n , E_l , T_l i parametrów, które mają zastosowanie we wszystkich węzłach sieci, takie jak M , α , β , ε i η . Dla rozważań w tym rozdziale możemy przyjąć, że tabela 3-3 przedstawia parametry dla sieci homogenicznej, w której wszystkie węzły mają taką samą konfigurację, lub dla węzła w sieci heterogenicznej, w którym występują największe opóźnienia. Autor pomija w ten sposób wpływ architektury sieci na opóźnienia, a wyniki mają zastosowanie do każdej sieci zgodnej z modelem z rozdziału 3.1. Tabela 3-3 jest wyjściowym zestawem parametrów służącym do pokazania ogólnych mechanizmów odpowiedzialnych za opóźnienia w sieciach typu DiffServ z kształtowaniem ruchu dla zagregowanych strumieni danych spełniających warunki do twierdzenia 3-4. Wartość opóźnienia maksymalnego dla

twierdzenia 3-4 zależy od szeregu parametrów. Podobnie jak dla twierdzenia 3-3, największy wpływ na opóźnienie maksymalne mają wartości następujących parametrów: α , H , β , R_l/C_l , R_l/P_n , ε i η .

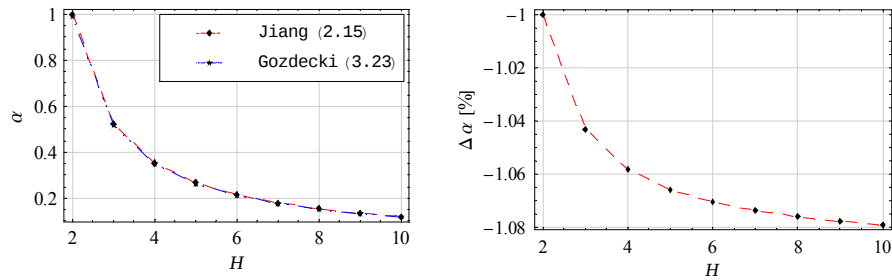
Tabela 3-3: Parametry sieci dla rysunków 3-27, 3-28, 3-32

Parametr Rysunek	R_l [b/s]	C_l [b/s]	M [B]	P_n [b/s]	M_n [B]	H
Rysunek 3-27	5×10^8	1×10^9	1500	6×10^9	9000	8
Rysunek 3-28	5×10^8	1×10^9	1500	6×10^9	9000	2÷8
Rysunek 3-32	5×10^8	1×10^9	1500	6×10^9	9000	2÷8
Parametr Rysunek	E_l [s]	T_l [s]	ε [s]	η	α	β [s]
Rysunek 3-27	0.000012	0.000036	0.000012	0.495	0.01÷0.149	0.0024
Rysunek 3-28	0.000012	0.000036	0.000012	0.495	-	0.0024
Rysunek 3-32	0.000012	0.000036	0.000012	0.495	0.149	0.0024

Rysunek 3-27: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od współczynnika wykorzystania pasma α

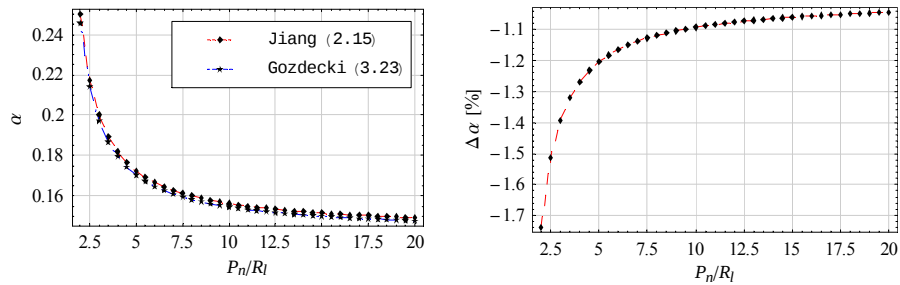
Rysunek 3-27 przedstawia porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od wartości współczynnika wykorzystania pasma α dla parametrów sieci przedstawionych w tabeli 3-3. W badanym zakresie zmienności parametru α opóźnienie maksymalne dla sieci z kształtowaniem ruchu jest mniejsze niż w klasycznej sieci DiffServ. Zysk z zastosowania kształtowania ruchu zmienia się od kilku procent dla $\alpha = 0.01$ do 70% dla α zbliżającego się do wartości maksymalnej. W obydwu przypadkach wzrost opóźnienia wraz ze wzrostem współczynnika wykorzystania pasma ma

charakter funkcji homograficznej.



Rysunek 3-28: Porównanie maksymalnej wartości współczynnika wykorzystania pasma α dla równania (2-13) i twierdzenia 3-4 w zależności od rozmiaru sieci H

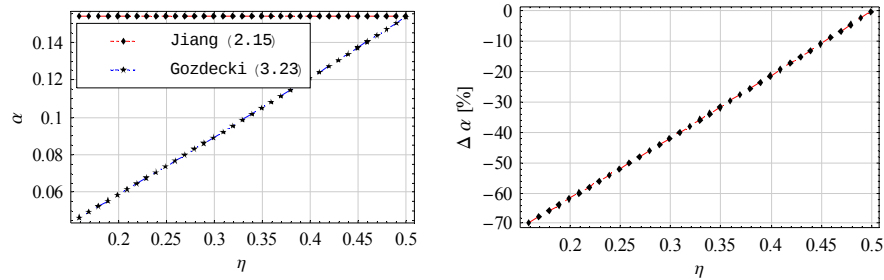
Na rysunku 3-28 autor przedstawia porównanie maksymalnej wartości współczynnika wykorzystania pasma α dla równania (2-13) i twierdzenia 3-4 w zależności od rozmiaru sieci H . Dla sieci o parametrach z tabeli 3-3 z kształtowaniem ruchu i spełniającej warunki twierdzenia 3-4 współczynnik wykorzystania pasma jest minimalnie mniejszy w porównaniu z siecią bez kształtowania ruchu a różnica względna wynosi około 1%.



Rysunek 3-29: Różnica maksymalnej wartości współczynnika wykorzystania pasma α dla twierdzenia 3-4 (równanie (3-23)) i równania (2-15) w zależności od stosunku P_n/R_l dla $H=8$

Na rysunku 3-29 autor przedstawia wpływ stosunku parametrów P_n/R_l na maksymalną wartość współczynnika wykorzystania pasma α . Wraz ze wzrostem stosunku P_n/R_l maksymalna wartość współczynnika wykorzystania pasma α maleje i osiąga minimum dla $P_n \rightarrow \infty$, dla której to wartości wynosi $\alpha = \frac{1}{H-1}$. Dla całego zakresu zmiennej P_n/R_l maksymalna wartość współczynnika wykorzystania pasma α dla sieci z kształtowaniem ruchu jest nieznacznie mniejsza niż dla sieci bez kształtowania ruchu.

Różnica ta maleje ze wzrostem parametru P_n / R_l .

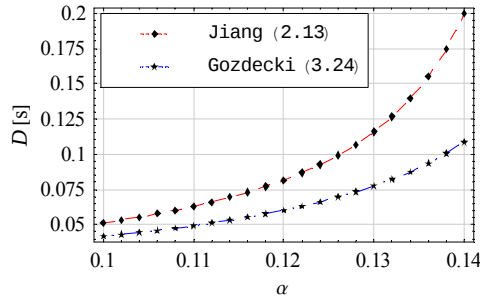


Rysunek 3-30: Porównanie maksymalnej wartości współczynnika wykorzystania pasma α dla równania (2-15) i twierdzenia 3-4 (równanie (3-23)) w zależności od współczynnika η
– lewy wykres przedstawia wartości bezwzględne, prawy wykres przedstawia różnicę względną

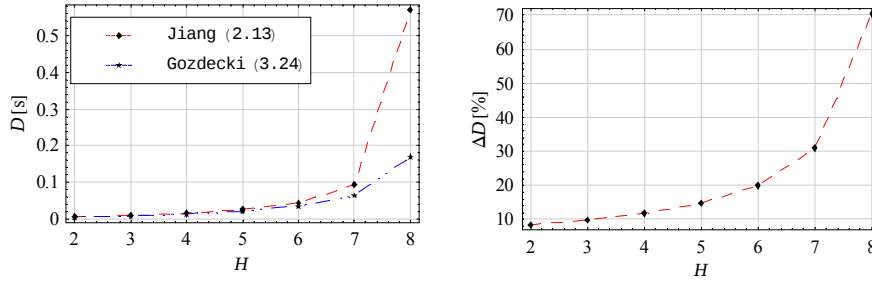
Na rysunku 3-30 autor przedstawia porównanie maksymalnej wartości współczynnika wykorzystania pasma α dla równania (2-15) i twierdzenia 3-4 (równanie (3-23)) w zależności od współczynnika η . Lewy wykres przedstawia wartości bezwzględne.

Wykres równania (2-15) jest niezależny od współczynnika η i został umieszczony jako wartość odniesienia. Prawy wykres przedstawia różnicę względną w procentach. W całym zakresie zmienności parametru η maksymalna wartość współczynnika wykorzystania pasma α jest mniejsza w systemie z kształtowaniem ruchu algorytmem ciekącego wiadra i wartość ta rośnie ze wzrostem parametru η . Wartość współczynnika wykorzystania pasma α w sieci z kształtowaniem ruchu zbliża się do α w sieci bez kształtowania ruchu dla $\eta \rightarrow R_l / C_l$, ale jest zawsze od niej mniejsza.

Dla twierdzenia 3-4 zachodzi analogiczna właściwość, jak dla twierdzenia 3-3: jeśli dla parametrów danej sieci bez kształtowania ruchu współczynnik wykorzystania pasma ma wartość mniejszą od wartości maksymalnej danej wzorem (2-15), ponieważ opóźnienia maksymalne przekraczają dla niego ustalony próg, to w odpowiadającej jej sieci z kształtowaniem ruchu współczynnik wykorzystania pasma będzie miał wartość większą. Warunek ten zachodzi tylko wtedy, gdy szybkość obsługi agregatu EF PHB w łączu jest mniejsza od przepustowości łącza w węzłach sieci. Dla zachowania analogii z rozdziałem poprzednim autor zamieszcza rysunek 3-31, na podstawie którego można wyznaczyć o ile współczynnik wykorzystania pasma w sieci z kształtowaniem ruchu będzie większy niż w klasycznej sieci DiffServ, jeśli dla sieci opisanej w tabeli 3-3 maksymalne opóźnienie nie może przekroczyć odpowiednio 100 ms i 50 ms. Wyniki są bardzo zbliżone do wyników z rozdziału poprzedniego. Porównując rysunki 3-15 i 3-31 można zauważyć, że opóźnienia maksymalne dla twierdzenia 3-3 są nieznacznie mniejsze od opóźnień dla twierdzenia 3-4.



Rysunek 3-31: Opóźnienie maksymalne dla równania (2-13) i twierdzenia 3-3 w zależności od współczynnika wykorzystania pasma α



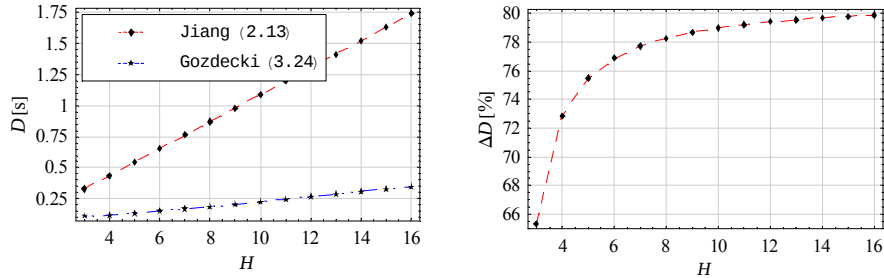
Rysunek 3-32: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od rozmiaru sieci H

Na rysunku 3-32 autor przedstawia porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od rozmiaru sieci H dla parametrów sieci zdefiniowanych w tabeli 3-3. Parametr α w tabeli 3-3 jest tak dobrany, aby wykorzystanie pasma było jak największe i równania (2-13) i (3-24) były prawdziwe w całym zakresie zmienności H . Opóźnienie maksymalne w granicach zmienności H w systemie zaproponowanym przez autora jest mniejsze od około 5% do 70% w porównaniu z siecią bez kształtowania ruchu. Podobnie jak dla twierdzenia największy zysk z zastosowania kształtowania ruchu jest dla $H = 8$, dla którego to parametru wartość współczynnika wykorzystania pasma α zbliża się do maksymalnej dopuszczalnej wartości.

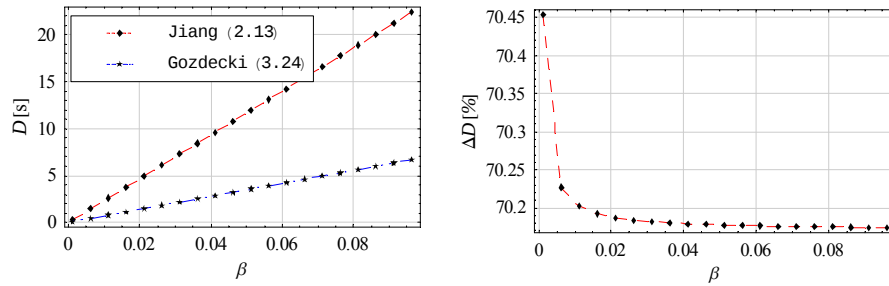
Na rysunku 3-33 autor przedstawia zależność opóźnienia maksymalnego dla równania (2-13) i twierdzenia 3-4 w zależności od rozmiaru sieci H dla współczynnika wykorzystania pasma α obliczanego dla każdego punktu według wzoru (3-23) pomnożonego przez stały współczynnik 0.99 według wzoru:

$$\alpha(H) = 0.99 \min \left(\frac{\rho_l \rho_n}{((H-1)(\rho_n - \rho_l) + \rho_l)R_l}, \frac{\rho_l P_n}{((H-1)(P_n - \rho_l) + \rho_l)R_l} \right). \quad (3-78)$$

Pozostałe parametry sieci są takie same jak dla rysunków 3-27 i 3-32 (tabela 3-3). Na wykresie z rysunku 3-20 widzimy, że jeśli wartość współczynnika wykorzystania pasma α dla danego rozmiaru sieci H zbliża się do wartości maksymalnej, to zysk z zastosowania kształtowania ruchu jest stosunkowo duży dla całego zakresu zmiennej H i zwiększa się ze wzrostem wartości tego parametru.



Rysunek 3-33: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od rozmiaru sieci H dla współczynnika wykorzystania pasma α obliczanego dla każdego punktu według wzoru (3-78)

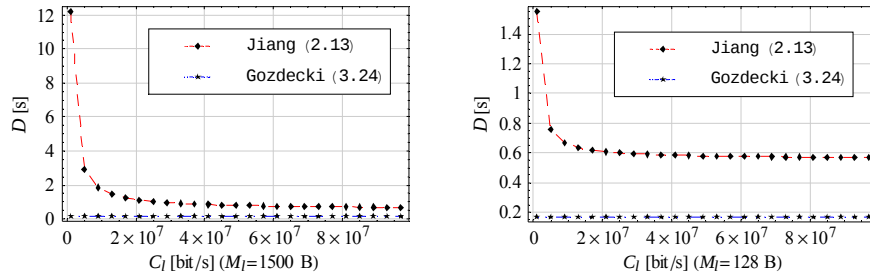


Rysunek 3-34: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od parametru β

Na rysunku 3-34 autor przedstawia zależność opóźnienia maksymalnego D od parametru β dla równania (2-13) i twierdzenia 3-4 – prawy wykres przedstawia wartości bezwzględne dla równań (2-13) i (3-24), natomiast lewy różnicę względną pomiędzy tymi równaniami liczoną według wzoru $\Delta D(\beta) = \frac{D(\beta)_{(3-13)} - D(\beta)_{(3-24)}}{D(\beta)_{(3-13)}} 100\%$. Pozostałe

parametry sieci mają wartość zdefiniowaną w tabeli 3-3. Wraz ze wzrostem parametru β opóźnienie maksymalne rośnie w przybliżeniu liniowo zarówno dla równania (2-13) jak i (3-24), jednak w przypadku równania (3-24) wzrost opóźnienia jest znacznie mniejszy, co powoduje, że dla parametrów z tabeli 3-2 opóźnienia są o około 80% mniejsze dla sieci

z kształtowaniem ruchu, niż w sieci, w której nie zastosowano kształtowania ruchu zaproponowanego przez autora. Zysk jest obserwowalny w całym zakresie zmiennej β .

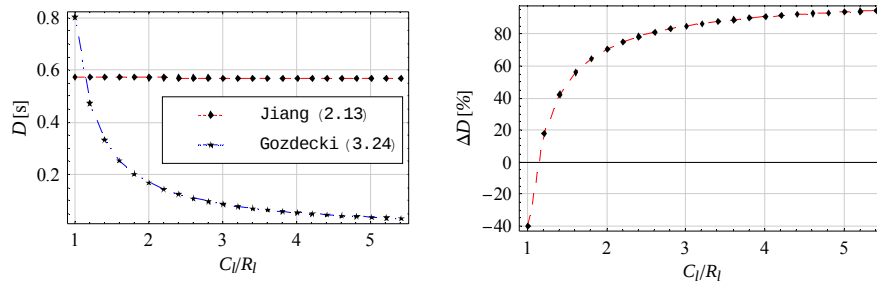


Rysunek 3-35: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od przepustowości łącza C_l dla sieci homogenicznej gdzie $P_n = k C_l$ i $R_l / C_l = 0.5$. Lewy wykres dla $M_l = 1500$ B, prawy wykres dla $M_l = 128$ B

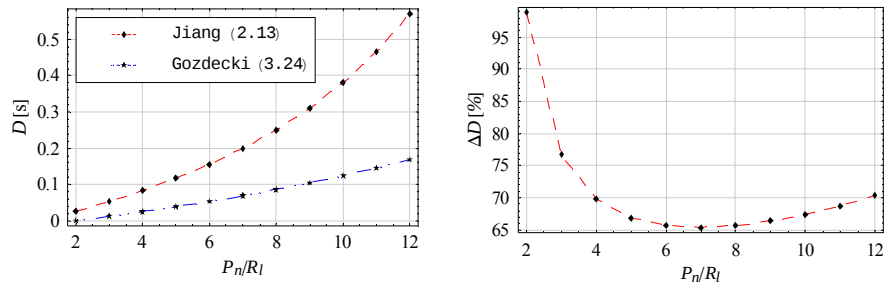
Na rysunku 3-35 autor przedstawia porównanie maksymalnych opóźnień dla równania (2-13) i twierdzenia 3-4 w zależności od C_l i M_l przy założeniu że parametry H , α , β , ε i η mają stałą wartość zdefiniowaną w tabeli 3-3, natomiast parametry $P_n = 8C_l$ i $R_l / C_l = 0.5$. Dla prawego wykresu na rysunku 3-35 parametr $M_l = 1500$ B, dla lewego wykresu $M_l = 128$ B. Parametr M_l ma znaczący wpływ na opóźnienie maksymalne dla łączy o niskiej przepustowości i jest to związane z wartością stosunków M_l / C_l i M_l / R_l , które powyżej określonej wartości zmiennej C_l (dla rysunku 3-35 powyżej 6 Mbit/s) przestają mieć znaczący wpływ na opóźnienie maksymalne. Dla łączy o odpowiednio niskich przepustowościach, wraz ze wzrostem maksymalnego rozmiaru pakietu opóźnienie maksymalne znacząco rośnie i jest mniejsze w sieci z kształtowaniem ruchu zaproponowanym przez autora. Pomijając wpływ maksymalnego rozmiaru pakietu na rysunku 3-35, zwiększenie przepustowości łączy w całej sieci nie powoduje zmniejszenia opóźnień maksymalnych, jeśli parametry H , α , β , ε , η , R_l / C_l mają wartość stałą, czyli jeśli ze wzrostem przepustowości łączy zwiększamy proporcjonalnie ruch przenoszony przez sieć.

Rysunek 3-36 przedstawia porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od stosunku C_l / R_l dla $\eta = 0.99 R_l / C_l$. R_l ma stałą wartość, która tak jak pozostałe parametry jest zdefiniowana w tabeli 3-3. Szybkość obsługi na rysunku zmienia się w takich granicach, że wpływ rozmiaru maksymalnego rozmiaru pakietu na opóźnienie maksymalne jest pomijalny i nie można go zaobserwować na rysunku 3-35. C_l jest parametrem charakteryzującym parametry obsługi w łączy l , który ma wpływ na wartość parametru P_n w węźle, do którego wchodzi łączy l . P_n jest parametrem, który charakteryzuje ruch wejściowy. Dla równania (2-13) wpływ zmian

parametru C_l jest pomijalny, natomiast jest on bardzo ważny dla twierdzenia 3-4, jeśli odniesiemy go do wartości zmiennej R_l . Dla poprzedniego rysunku 3-35 stosunek $C_l/R_l=2$ i w całym zakresie zmian C_l opóźnienie maksymalne dla sieci z kształtowaniem ruchu jest mniejsze. Jeśli $C_l/R_l=1$, czyli szybkość obsługi agregatu priorytetowego jest równa przepustowości łącza, to zastosowanie kształtowania ruchu zaproponowanego przez autora powoduje zwiększenie opóźnienia maksymalnego w sieci w porównaniu z siecią bez kształtowania ruchu dla agregatów DiffServ. Oznacza to, że zastosowanie kształtowania ruchu algorytmem ciekącego wiadra ma uzasadnienie tylko dla agregatów, których szybkość obsługi jest mniejsza od przepustowości łącza.



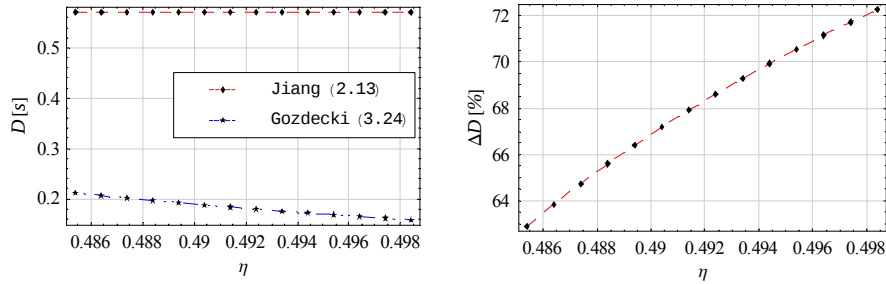
Rysunek 3-36: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od stosunku C_l/R_l dla $\eta=0,99 R_l/C_l$



Rysunek 3-37: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od parametru P_n/R_l , przy założeniu, że wartości C_l i R_l są stałe

Na rysunku 3-37 autor przedstawia porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od parametru P_n/R_l przy założeniu, że wartości C_l i R_l są stałe. Tak jak dla twierdzenia 3-3, tak i w przypadku twierdzenia 3-4 wzrost wartości stosunku P_n/R_l powoduje zwiększenie opóźnień maksymalnych. Wzrost ten może mieć dwojakie źródło: zwiększenie liczby łączy wchodzących do węzła lub

zwiększenie przepustowości łączy wchodzących do węzła. W obydwu tych przypadkach prawdopodobieństwo, że w krótkim odstępie czasu do łączy wychodzącego nadejdzie większa liczba pakietów, rośnie, co powoduje, że w takim systemie potencjalnie opóźnienia maksymalne będą większe. Mechanizm ciekącego wiadra zmniejsza możliwości sklejania się pakietów w większe paczki danych, gdy szybkość obsługi w łącach wychodzących jest mniejsza niż przepustowość łączy i to powoduje, że dla sieci z kształtowaniem ruchu zaproponowanym przez autora opóźnienia maksymalne są niższe niż w klasycznej sieci DiffServ.



Rysunek 3-38: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od parametru η

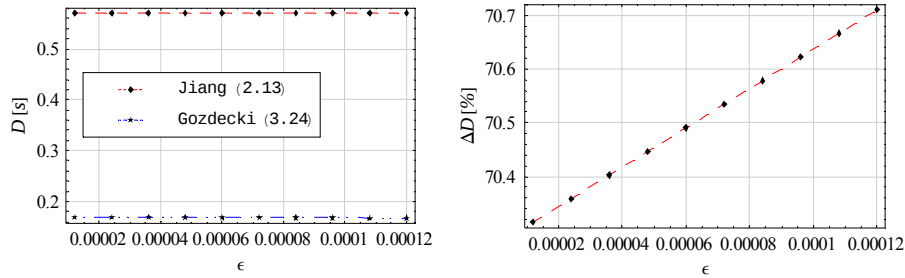
Rysunek 3-38 przedstawia zależność opóźnienia maksymalnego dla twierdzenia 3-4 od parametru η . Równanie (2-13) jest niezależne od parametru η i występuje na rysunku jako wartość odniesienia. Prawy wykres przedstawia wartości bezwzględne dla równań (2-13) i (3-24), natomiast lewy różnicę względną pomiędzy tymi równaniami liczoną według wzoru $\Delta D(\eta) = \frac{D(\eta)_{(2-13)} - D(\eta)_{(3-24)}}{D(\eta)_{(2-13)}} 100\%$. Pozostałe parametry mają wartość określoną

w tabeli 3-3. Parametr η jest zmieniany w wąskim zakresie od 0.485 do 0.999, bo tylko w takim zakresie twierdzenie 3-4 jest prawdziwe dla parametrów z tabeli 3-3. W badanym zakresie wraz ze wzrostem parametru η opóźnienia maksymalne dla twierdzenia 3-4 nieznacznie maleją. Zysk z zastosowania kształtowania ruchu zaproponowanego przez autora wynosi od około 65% do 72% w badanym zakresie zmiennej η .

Na rysunku 3-39 autor przedstawia zależność opóźnień maksymalnych dla twierdzenia 3-4 od parametru ε . Równanie (2-13) jest niezależne od parametru ε i występuje na rysunku jako wartość odniesienia. Prawy wykres przedstawia wartości bezwzględne dla równań (2-13) i (3-24), natomiast lewy różnicę względną pomiędzy tymi równaniami liczoną według wzoru $\Delta D(\varepsilon) = \frac{D(\varepsilon)_{(2-13)} - D(\varepsilon)_{(3-24)}}{D(\varepsilon)_{(2-13)}} 100\%$. Pozostałe

parametry mają wartość określoną w tabeli 3-3. Zakres zmienności parametru ε wynosi

od $\varepsilon = M / C_l$, co oznacza, że rozmiar bufora w algorytmie ciekącego wiadra $\sigma_l = M$, do wartości $\varepsilon = 10M / C_l$ i z dokładnością odnoszącą się do parametru M jest to cały zakres, w którym twierdzenie 3-4 jest prawdziwe dla parametrów z tabeli 3-3. Opóźnienie maksymalne dla twierdzenia 3-4 nieznacznie spada wraz ze wzrostem wartości ε i jest mniejsze od opóźnienia maksymalnego w sieci bez kształtowania ruchu zaproponowanego przez autora o około 70%.



Rysunek 3-39: Porównanie opóźnień maksymalnych dla równania (2-13) i twierdzenia 3-4 w zależności od parametru ε

Podsumowując porównanie opóźnień maksymalnych w klasycznej sieci DiffServ opisanych równaniem (2-13) z opóźnieniami maksymalnymi w sieci spełniającej warunki do twierdzenia 3-4 autor stwierdza, że dla sieci DiffServ o dowolnej architekturze scharakteryzowanej przez zakresy wartości zmiennych z analizy bieżącego rozdziału, najważniejsze zalety i wady rozwiązania zaproponowanego przez autora w twierdzeniu 3-4 są następujące:

Zalety:

- Stosując kształtowanie ruchu algorytmem ciekącego wiadra dla agregatów EF PHB sieci DiffServ można zmniejszyć opóźnienie maksymalne w sieciach o różnych rozmiarach i konfiguracji pod warunkiem, że zastosowane algorytmy szeregowania w łączach obsługują EF PHB z szybkością mniejszą niż przepustowość łączy.
- Prosty algorytm ciekącego wiadra stosowany dla agregatu danych nie będzie wpływał na skalowalność sieci.
- Istnieje możliwość zwiększenia wykorzystania pasma dla EF PHB w sieci, jeśli szybkość obsługi agregatu EF PHB w łączu jest mniejsza od przepustowości łącza w węzłach sieci. Jeśli dla danej sieci bez kształtowania ruchu współczynnik wykorzystania pasma ma wartość mniejszą od wartości maksymalnej danej wzorem (2-15), ponieważ opóźnienia maksymalne przekraczają dla niego ustalony

próg, to w odpowiadającej jej sieci z kształtowaniem ruchu współczynnik wykorzystania pasma będzie miał wartość większą dla tego samego progu opóźnienia maksymalnego.

Wady:

- Kształtowanie ruchu za pomocą algorytmu ciekącego wiadra dla agregatów EF PHB w sieci DiffServ nie powoduje zmniejszenia opóźnień maksymalnych, jeśli w węzłach zostaną zastosowane algorytmy szeregowania, które obsługują EF PHB z szybkością równą przepustowości łączy. Powyższe warunki spełnia algorytm *Strict Priority* dla klasy o najwyższym priorytecie.
- Konieczna implementacja dodatkowych algorytmów kształtowania ruchu w węzłach dla agregatów w sieci DiffServ i uwzględnienie nowych parametrów w systemie do zarządzania siecią.
- Mniejszy maksymalny współczynnik wykorzystania pasma niż w klasycznej sieci DiffServ.
- Opóźnienia maksymalne są większe niż dla rozwiązania zaproponowanego w rozdziale 3-3 dla analogicznej architektury sieci.
- Zakresy zmienności parametrów, dla których twierdzenie 3-4 jest prawdziwe są mniejsze niż dla twierdzenia 3-4.

Zysk z zastosowania kształtowania ruchu w sieci jest tym większy, im:

- większe jest wykorzystanie pasma α ,
- większy jest rozmiar sieci mierzony parametrem H ,
- większy jest parametr β charakteryzujący wybuchowość ruchu przychodzącego; zakres zmienności tego parametru jest znacznie ograniczony w porównaniu z twierdzeniem 3.3,
- stosunek R_l / C_l jest mniejszy,
- parametr P_n jest większy.

3.4 Podsumowanie rozdziału

W rozdziale 3 autor wyprowadził twierdzenie 3-1 opisujące opóźnienia maksymalne

w sieci DiffServ z kształtowaniem ruchu dla agregatów EF PHB za pomocą algorytmu ciekącego wiadra implementowanego w każdym węźle sieci. Model sieci, dla której ma zastosowanie zaproponowane twierdzenie, obejmuje dowolną architekturę zgodną z architekturą DiffServ z węzłami buforującymi pakiety na wyjściu oraz implementującymi algorytm szeregowania typu GR z obsługą typu FIFO.

W zależności od parametrów danej sieci i parametrów algorytmu ciekącego wiadra autor wyodrębnił cztery szczególne przypadki twierdzenia 3-1, które zostały opisane twierdzeniami od 3-2 do 3-5. Dla sieci spełniającej warunki twierdzenia 3-2 opóźnienia maksymalne są takie same jak w sieci bez zastosowanego kształtowania ruchu.

W przypadku sieci spełniających warunki twierdzeń 3-3 i 3-4 **możliwe jest znaczne zmniejszenie opóźnień maksymalnych w sieci DiffServ, ale pod warunkiem, że szybkość obsługi agregatów EF PHB w łączach będzie mniejsza od przepustowości łącz w węzłach sieci**. Powyższy warunek spełniają sieci, w których jako algorytm szeregowania zastosujemy przykładowo algorytmy VC, WFQ lub SCFQ. Natomiast powyższy warunek nie jest spełniony dla algorytmu *Strict Priority* dla kolejki o najwyższym priorytecie, która jest obsługiwana z szybkością przepustowości łącza oraz dla sieci, w których istnieje tylko jedna klasa ruchu obsługiwana z szybkością równą przepustowości łącza. W badanym zakresie zmiennych opóźnienia maksymalne dla twierdzenia 3.3 są mniejsze niż dla twierdzenia 3.4. Ponadto zakres stosowalności twierdzenia 3.4 wyznaczony przez zakresy zmiennych i warunki do tego twierdzenia jest węższy. Maksymalna wartość współczynnika wykorzystania pasma dla sieci spełniającej warunki twierdzenia 3.3 jest taka sama lub mniejsza jak w sieci bez kształtowania ruchu – dla przykładu sieci opisanej w tabeli 3.1 wartość ta jest taka sama. Natomiast dla sieci z kształtowaniem ruchu dla agregatów EF PHB spełniającej warunki twierdzenia 3.4 wartość maksymalna współczynnika wykorzystania pasma jest nieznacznie mniejsza niż w przypadku odpowiadającej jej sieci bez kształtowania ruchu. W szczególnym przypadku istnieje możliwość zwiększenia wykorzystania pasma w sieci dla EF PHB, jeśli zastosujemy kształtowanie ruchu dla agregatów EF PHB zgodne z twierdzeniami 3.3 lub 3.4. Jeśli dla danej sieci bez kształtowania ruchu współczynnik wykorzystania pasma ma wartość mniejszą od wartości maksymalnej z powodu opóźnień maksymalnych przekraczających zakładany próg, to w odpowiadającej jej sieci z kształtowaniem ruchu współczynnik wykorzystania pasma będzie miał wartość większą dla tego samego progu opóźnienia maksymalnego. Warunek ten zachodzi tylko wtedy, gdy szybkość obsługi agregatu EF PHB w łączu jest mniejsza od przepustowości łącza w węzłach sieci.

Twierdzenie 3-5 dowodzi, że jeśli zastosujemy kształtowanie ruchu algorytmem ciekącego wiadra dla EF PHB, w którym szybkość obsługi jest mniejsza od sumy przepływności strumieni wchodzących w skład agregatu priorytetowego obsługiwanego przez EF PHB, to nie można oszacować maksymalnego opóźnienia domenie DiffServ dla tego agregatu.

4 Podsumowanie rozprawy

Głównym celem rozprawy była analiza wpływu kształtowania ruchu za pomocą algorytmu ciekącego wiadra na opóźnienia maksymalne, jakich mogą doznawać pakiety przy transmisji przez sieć IP o architekturze Usług Zróżnicowanych z wykorzystaniem reguły przesyłania pakietów EF PHB. EF PHB został zdefiniowany do budowania usług sieciowych o najwyższych parametrach jakościowych, wymagających bezstratnej transmisji pakietów i gwarancji maksymalnych opóźnień przy przesyłaniu pakietów przez domenę sieci DiffServ. Ścisła definicja matematyczna EF PHB pozwala wyliczyć opóźnienie maksymalne dla usługi sieciowej zbudowanej z użyciem EF PHB pomiędzy brzegami sieci typu DiffServ, ale jednocześnie wprowadza ograniczenie na maksymalną wartość współczynnika wykorzystania pasma w sieci typu DiffServ oraz powoduje, że opóźnienia maksymalne w sieci rosną homograficznie wraz ze wzrostem obciążenia EF PHB.

Autor w postawionej tezie stwierdza, że: **„Zastosowanie algorytmu ciekącego wiadra do kształtowania agregatów EF PHB w każdym węźle sieci IP w architekturze Usług Zróżnicowanych umożliwia zmniejszenie opóźnień maksymalnych pomiędzy brzegami domeny dla usług sieciowych zbudowanych za pomocą reguł przesyłania pakietów EF PHB pod warunkiem, że szybkość obsługi agregatów EF PHB jest mniejsza od przepustowości łączy”**. Dowód tezy polegał na zbudowaniu modelu matematycznego opisującego sieć DiffServ dla EF PHB oraz wprowadzeniu do niego kształtowania ruchu dla agregatów EF PHB w każdym węźle sieci. Dla stworzonego modelu sieci autor sformułował i udowodnił centralne dla rozprawy twierdzenie opisujące opóźnienie maksymalne pomiędzy brzegami sieci. Analiza tego równania pozwoliła na wyodrębnienie szczególnych przypadków, które pozwoliły określić, dla jakich parametrów sieci DiffServ i parametrów algorytmu ciekącego wiadra można zmniejszyć opóźnienie maksymalne pomiędzy brzegami sieci. Na podstawie tej analizy **autor stwierdza, że postawiona w rozprawie teza jest prawdziwa**.

Do najważniejszych oryginalnych osiągnięć rozprawy autor zalicza:

- zbudowanie modelu sieci DiffServ z kształtowaniem ruchu dla agregatów EF we wszystkich węzłach sieci,
- sformułowanie i dowód ogólnego twierdzenia o opóźnieniach maksymalnych opisującego wpływ kształtowania agregatów ruchu algorytmem ciekącego wiadra na opóźnienie maksymalne w sieci DiffServ,
- wyznaczenie dla twierdzenia o opóźnieniach maksymalnych rozwiązań jego szczególnych przypadków opisujących, dla jakich parametrów opóźnienia maksymalne w sieci DiffServ z kształtowaniem ruchu są mniejsze lub większe niż w odpowiadającej jej sieci bez kształtowania ruchu.

Autor przedstawił także analizę ilościową dla przykładowej sieci (spełniającej założenia modelu), z której wynika, że zmniejszenie opóźnień w sieci DiffServ z kształtowaniem ruchu może być znaczne oraz przedstawił, jak poszczególne parametry sieci wpływają na ten spadek. W szczególnym przypadku, jeśli zastosujemy kształtowanie ruchu dla agregatów EF PHB, to istnieje także możliwość zwiększenia wartości współczynnika wykorzystania pasma dla usług zbudowanych w oparciu o EF PHB, pomimo że maksymalna wartość współczynnika wykorzystania pasma dla sieci z zastosowanym kształtowaniem ruchu algorytmem ciekącego wiadra nie może być większa niż w sieci bez kształtowania ruchu. Przypadek ten występuje wtedy, gdy dla danej sieci bez kształtowania ruchu współczynnik wykorzystania pasma ma wartość mniejszą od dopuszczalnej wartości z powodu przekroczenia założonego progu opóźnienia maksymalnego dla wartości maksymalnej współczynnika wykorzystania pasma – w odpowiadającej jej sieci z kształtowaniem ruchu współczynnik wykorzystania pasma będzie miał wartość większą dla tego samego progu opóźnienia maksymalnego.

Wyprowadzone przez autora twierdzenia mogą zostać zastosowane do oceny opóźnień maksymalnych przy projektowaniu sieci operatorskich zgodnych z architekturą DiffServ dla usług sieciowych zbudowanych za pomocą reguł przesyłania pakietów EF PHB. Uzyskane wyniki mają zastosowanie do obliczenia opóźnień maksymalnych dla dowolnej architektury sieci DiffServ, a w przypadku potrzeby wyliczenia opóźnienia maksymalnego dla konkretnej ścieżki w sieci DiffServ, można zastosować wzory cząstkowe przedstawione w dowodach wyprowadzonych twierdzeń i w ten sposób obliczyć opóźnienie maksymalne na wybranej ścieżce, które jest mniejsze lub równe od opóźnienia maksymalnego dla całej domeny DiffServ.

5 Spis literatury

- [1] Postel J. (ed.), "Internet Protocol - DARPA Internet Program Protocol Specification," RFC 791, USC/Information Sciences Institute, September 1981.
- [2] Marziale V., Vitaletti A., "A Framework for Internet QoS Requirements Definition and Evaluation: an experimental approach", IST MOBILE COMMUNICATIONS SUMMIT 2001, Barcelona 2001.
- [3] ITU-T Recomm. I.350 (03/93) "General Aspects of Quality of Service and Network Performance in Digital Networks, including ISDNs".
- [4] ETSI TR 003 "Network Aspects (NA), "General Aspects of Quality of Service (QoS) and Network Performance (NP)"
- [5] ETSI TS 123 107 v3.5.9 (2000-12) "Universal Mobile Telecommunications System (UMTS); QoS Concept and Architecture"
- [6] N° 11 Report from UMTS Forum "Enabling UMTS / Third Generation Services and Applications", UMTS Forum, October 2000
- [7] Braden R., Clark D., Shenker S., "Integrated Services in the Internet Architecture: an Overview", RFC 1633, June 1994
- [8] Blake S., Black D., Carlson M., Davies E., Wang Z., Weiss W., "An Architecture for Differentiated Services," RFC 2475, December 1998, www.ietf.org/rfc/rfc2475.txt
- [9] Shenker S., Partridge C., Guerin R., "Specification of Guaranteed Quality of service", RFC 2212, September 1997
- [10] Nichols K., Blake S., Baker F., Black D., "Definition of the Differentiated Services Field (DS Field) in the IPv4 and IPv6 Headers", RFC 2474, December 1998
- [11] Charny A., "Delay Bounds in a network with aggregate scheduling", Technical report, Cisco, 2000. [ftp://ftpeng.cisco.com/ftp/acharny/aggregate_delay_v4.ps](http://ftpeng.cisco.com/ftp/acharny/aggregate_delay_v4.ps).

- [12] Bennett J. C. R., Benson K., Charny A., Courtney W. F., Le Boudec J.-Y., "Delay Jitter Bounds and Packet Scale Rate Guarantee for Expedited Forwarding"
In ACM/IEEE Transactions on Networking, August 2002, Vol. 10, No. 4, pp: 529-540
- [13] Jiang Y., "Delay Bounds for a Network of Guaranteed Rate Servers with FIFO Aggregation", Proceedings of ICCCN 2002, New York, May 2002
- [14] Braden R., Clark D., Shenker S., "Integrated Services in the Internet Architecture: an Overview", RFC 1633, June 1994.
- [15] Wroclawski J., "Specification of the Controlled-Load Network Element Service", RFC 2211, September 1997
- [16] Braden R., Zhang L., Berson S., Herzog S., Jamin S.: RFC 2205 „Resource Reservation Protocol (RSVP) – Version 1 Functional Specification”, September 1997
- [17] Wroclawski J., "The Use of RSVP with IETF Integrated Services", RFC 2210, September 1997.
- [18] Parekh A., Gallagher R., "A Generalized Processor Sharing Approach to Flow Control – The Single Node Case", IEEE/ACM Trans. Networking, Vol. 1, No. 3, 1993
- [19] Parekh A., Gallagher R., "A Generalized Processor Sharing Approach to Flow Control – The Multiple Node Case", IEEE/ACM Trans. Networking, Vol. 2, No. 2, 1996
- [20] Jacobson V., Nichols K., Poduri K., "An Expedited Forwarding PHB", RFC 2598, June 1999.
- [21] Armitage G., Carpenter B., Casati A., Crowcroft J., Halpern J., Kumar B., Schnizlein J., "A Delay Bound alternative revision of RFC 2598", RFC 3248, March 2002.
- [22] Davie B., Charny A., Baker F., Bennett J.C.R., Benson K., Le Boudec J.Y., Chiu A., Courtney W., Davari S., Firoiu V., Kalmanek C., Ramakrishnan K. and Stiliadis D., "An Expedited Forwarding PHB (Per-Hop Behavior)", RFC 3246, March 2002.
- [23] Baker F., Heinanen J., Weiss W., Wroclawski J., "Assured Forwarding PHB Group", RFC 2597, <ftp://ftp.isi.edu/in-notes/rfc2597.txt>
- [24] Grossman D., "New Terminology and Clarifications for Diffserv", IETF RFC 3260, April 2002.
- [25] Nichols K., Carpenter B., „Definition of Differentiated Services Per Domain Behaviours and Rules for their Specification”, RFC 3086, April 2001

- [26] Jacobson V., Nichols K., Poduri K.: IETF Draft „The ‘Virtual Wire’ Per-Domain Behaviour”, July 2000
- [27] Seddigh N., Nandy B., Heinanen J.: IETF Draft „An Assured Rate Per-Domain Behaviour for Differentiated Services”, November 2000
- [28] Bless R., Nichols K., Wehrle K.: RFC 3662, "A Lower Effort Per-Domain Behavior (PDB) for Differentiated Services", December 2003.
- [29] ITU-T Recommendation Y.1541, "Network performance objectives for IP-based services", May 2002.
- [30] ITU-T Recommendation G.114, "One-way transmission time", May 2000.
- [31] Charny A., Le Boudec J.-Y., "Delay Bounds in a Network with Aggregate Scheduling", Proceedings of QOFIS, Berlin, October 2000
- [32] Goyal P., Vin H., "Generalized guaranteed rate scheduling algorithms: a framework," IEEE/ACM Trans. Networking, vol 5-4, pp. 561–572, August 1997.
- [33] Zhang L., "Virtual clock: A new traffic control algorithm for packet switching. ACM Transaction on Computer Systems, 9(2):101-124, maj 1991.
- [34] S.J. Golestani, "A self-clocked fair Queuing Scheme for High Speed Applications", Materiały konferencyjne INFOCOM'94, 1994
- [35] Zhang Z. L., Duan Z., Hou T. Y., "Fundamental trade-offs in aggregate packet scheduling", SPIE Vol. 4526, August 2001.
- [36] Verma D., Zhang H., Ferrari D., "Guaranteeing delay jitter bounds in packet switching networks," in Proceedings of Tricomm '91, Chapel Hill, April 1991, pp. 35–46.
- [37] Cruz R. L., "SCED+: efficient management of quality of service guarantees," in IEEE Infocom'98, San Francisco, March 1998.
- [38] Le Boudec J.-Y., Thiran P., „Network Calculus”, Springer Verlag LNCS 2050, June 2001.
- [39] Jiang Y., Liu Y., "Stochastic Network Calculus", ISBN: 978-1-84800-126-8, Springer, September 2008,
- [40] Le Boudec J.-Y., "Application of Network Calculus to Guaranteed Service Networks", In IEEE Transactions on Information Theory, May 1998, Vol. 44 No. 3, pp: 1087-1096
- [41] Firoiu V., Le Boudec J.-Y., Towsley D., Zhang Z.-L., "Theories and Models for Internet Quality of Service", Proceedings of the IEEE, special issue in Internet Technology, September 2002, Vol. 90 No. 9, 1565—1591.

- [42] Nichols K., Jacobson V., Zhang Z.-L., “A Two-bit Differentiated Services Architecture for the Internet”, RFC 2638, July 1999.

Aneks A Zastosowanie teorii „*Network calculus*” do obliczania maksymalnych opóźnień w sieci – wybrane twierdzenia algebry „*min-plus*”

Zadaniem tego rozdziału jest przedstawienie najważniejszych twierdzeń stosowanych do obliczania opóźnień maksymalnych w sieci DiffServ za pomocą aparatu matematycznego „*Network calculus*” i algebry „*min-plus*”.

Na początku autor przedstawi definicje krzywej ograniczającej strumień A (*arrival curve*) i definicje krzywej obsługi S (*service curve*) [38][40].

Rozważmy strumień danych wchodzący do systemu G i opisany przez funkcję $V(t)$, która jest równa liczbie transmitowanych bitów strumienia w przedziale czasu $[0, t)$. Przez $V^*(t)$ oznaczmy funkcję opisującą strumień wychodzący z systemu. Mając funkcje rosnące w szerokim sensie A , mówimy, że strumień jest ograniczony przez A wtedy i tylko wtedy, jeśli dla każdego $t' \leq t$ spełniony jest warunek

$$V(t) - V(t') \leq A(t - t'). \quad (A-1)$$

Mówimy, że system G oferuje dla strumienia krzywą obsługi S jeśli dla każdego $t \geq 0$ istnieje $t_0 \geq 0$ oraz $t_0 \leq t$ i spełniony jest warunek

$$V^*(t) - V(t_0) \leq S(t - t_0). \quad (A-2)$$

Krzywa obsługi dla dwóch stanowisk obsługi o krzywych obsługi odpowiednio $s_1(t)$ i $s_2(t)$ połączonych szeregowo dana jest wzorem:

$$S(t) = (s_1 \otimes s_2)(t), \quad (A-3)$$

gdzie operacja splotu $\otimes$ w algebrze „*min-plus*” zdefiniowana jest następująco [38]:

$$(f \otimes g)(t) = \min_{0 \leq s \leq t} \{f(t - s) + g(s)\}. \quad (A-4)$$

Poniżej zostaną przedstawione trzy podstawowe wzory algebry „*min-plus*” [31][38]:

1. Obliczanie maksymalnego rozmiaru kolejki – „*backlog bound*”

Dla danych funkcji $V(t)$ i $V^*(t)$, zdefiniowanych jak wyżej, kres górny rozmiaru

buforu w systemie, wymagany do bezstratnej obsługi, jest opisany przez odchylenie wertykalne wyżej wymienionych funkcji i jest zdefiniowany następująco:

$$V(t) - V^*(t) \leq \sup_{s \geq 0} \{A(s) - S(s)\}. \quad (\text{A-5})$$

2. Kres górny opóźnienia („*delay bound*”)

Kres górny opóźnienia [38][40] jest określony przez odchylenie horyzontalne pomiędzy krzywą opisującą ruch wejściowy do systemu $A(t)$, a krzywą obsługi tego systemu $S(t)$, zdefiniowane w [38][40] jako:

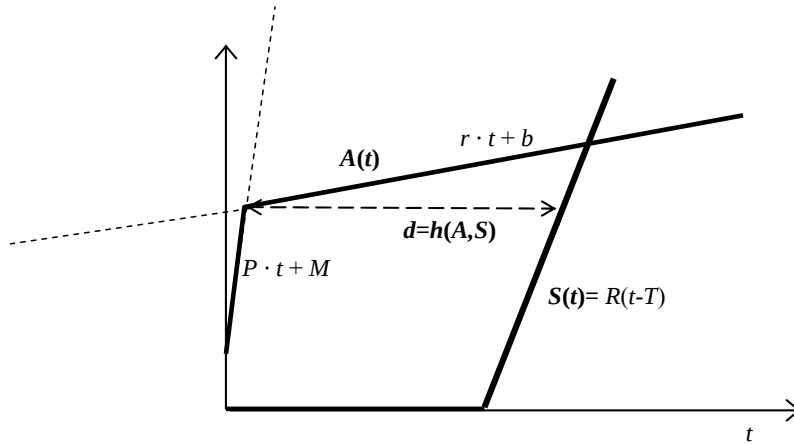
$$h(A, S) := \sup_{t \geq 0} [\inf \{T : T \geq 0 \text{ i } A(t) \leq S(t+T)\}]. \quad (\text{A-6})$$

System może być pojedynczym węzłem sieciowym lub składać się z wielu węzłów. Powyższe twierdzenie algebry *min-plus* [38] jest używana przy obliczaniu opóźnienia we wszystkich rozpatrywanych przypadkach. Do wyznaczenia odchylenia horyzontalnego autor skorzysta z gotowych wzorów przedstawionych w [41] (wzór (9) w [41]). Załóżmy, że mamy dane krzywe $A(t)$ i $S(t)$ zdefiniowane następująco

$$A(t) = \min(Pt + M, rt + b) \quad (\text{A-7})$$

i

$$S(t) = R(t - T)^+. \quad (\text{A-8})$$



Rysunek A-1: Definicja odchylenia horyzontalnego dla krzywych $A(t)$ i $S(t)$ zdefiniowanych wzorami (A-7) i (A-8)

Dla krzywych $A(t)$ i $S(t)$ odchylenie horyzontalne wynosi [41]:

$$h(A, S) = \frac{M + \frac{\beta - M}{P - r}(P - R)^+}{R} + T. \quad (\text{A-9})$$

Operator $(\cdot)^+$ jest zdefiniowany jako: $(a)^+ = 0$ dla $a \leq 0$ oraz $(a)^+ = a$ dla $a > 0$ lub $(a)^+ = \max(0, a)$.

$S(t)$ jest zgodny z algorytmem typu GR (wzór (2-11)) a $A(t)$ jest używana do opisu ruchu klasy *Guaranteed Rate* dla modelu sieci usług zintegrowanych. W modelu usług zróżnicowanych dla EF także wykorzystywane są algorytmy szeregowania pakietów zgodne z GR. Autor wykorzystuje wyniki otrzymane dla modelu usług zintegrowanych w modelu usług zróżnicowanych, podobnie jak inni autorzy [9][13] [31].

3. Krzywa ograniczająca strumień na wyjściu systemu („*output flow*”)[38]

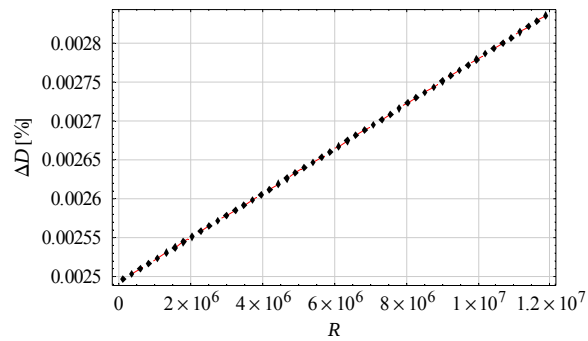
Jeśli strumień ograniczony przez krzywą $A(t)$ jest obsługiwany przez system, który oferuje krzywą obsługi $S(t)$, to na wyjściu systemu strumień jest ograniczony przez krzywą

$$A^* = A \circ S. \quad (\text{A-10})$$

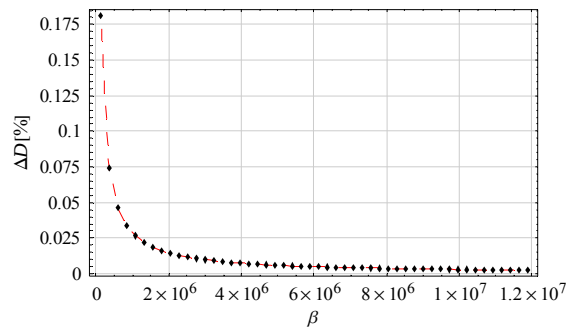
gdzie operacja $\circ$ jest zdefiniowana jako $(A \circ S)(t) = \sup_{s \geq 0} \{A(t+s) - S(s)\}$ i jest operacją odwrotną do operacji splotu w algebrze „min-plus”.

Aneks B Porównanie równań (2-7) i (2-13)

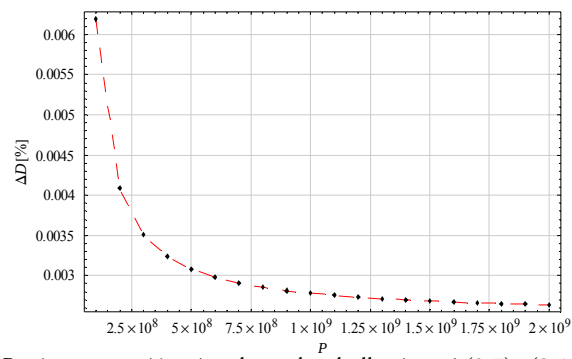
W dodatku przedstawiono rozszerzone porównanie równań (2-7) i (2-13) z rozdziału 2.3.1 dla parametrów sieci przedstawionych w tabeli 2-1 oraz wpływ parametrów R , M_l , k na opóźnienia maksymalne dla równania (2-13).



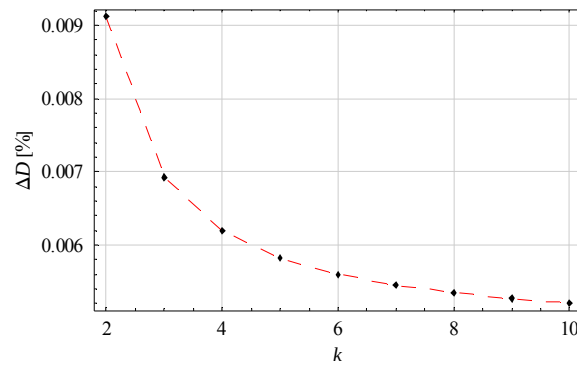
Rysunek B-1: Porównanie opóźnień maksymalnych dla równań (2-7) i (2-13) w zależności od szybkości obsługi R (parametry sieci zostały przedstawione w tabeli 2-1)



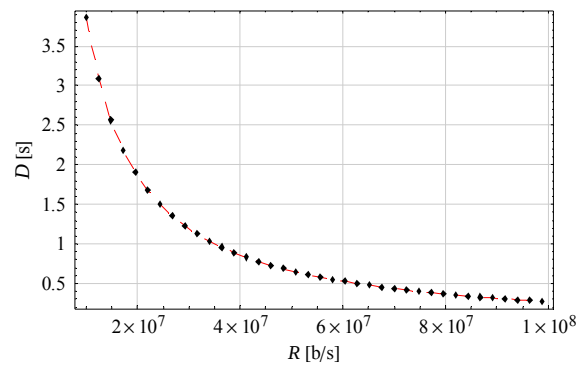
Rysunek B-2: Porównanie opóźnień maksymalnych dla równań (2-7) i (2-13) w zależności od parametru β (parametry sieci zostały przedstawione w tabeli 2-1)



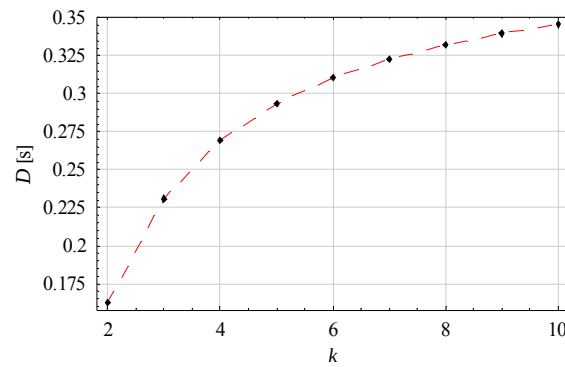
Rysunek B-3: Porównanie opóźnień maksymalnych dla równań (2-7) i (2-13) w zależności od przepustowości łącza P (pozostałe parametry zostały przedstawione w tabeli 2-1)



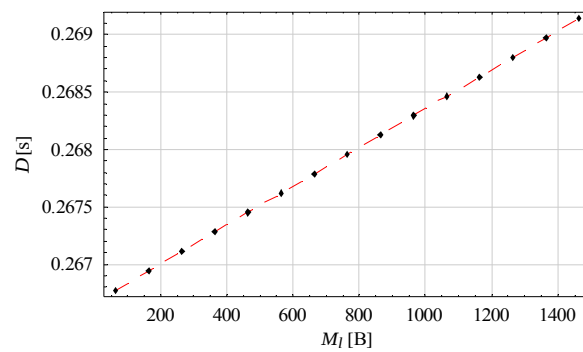
Rysunek B-4: Porównanie opóźnień maksymalnych dla równań (2-7) i (2-13) w zależności od współczynnika k (pozostałe parametry zostały przedstawione w tabeli 2-1)



Rysunek B-5: Zależność opóźnienia maksymalnego dla równania (2-13) w zależności od szybkości obsługi R (dla parametrów sieci przedstawionych w tabeli 2-1)



Rysunek B-6: Zależność opóźnienia maksymalnego dla równania (2-13) w zależności od współczynnika k (dla parametrów sieci przedstawionych w tabeli 2-1)



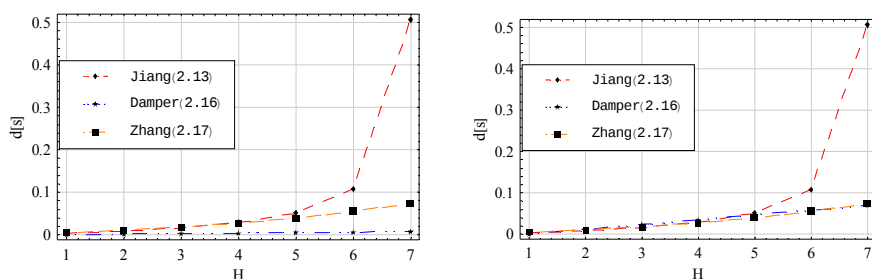
Rysunek B-7: Zależność opóźnienia maksymalnego dla równania (2-13) w zależności od maksymalnego rozmiaru pakietu M_l (dla parametrów sieci przedstawionych w tabeli 2-1)

Aneks C Porównanie równań (2-13), (2-17) i (2-16)

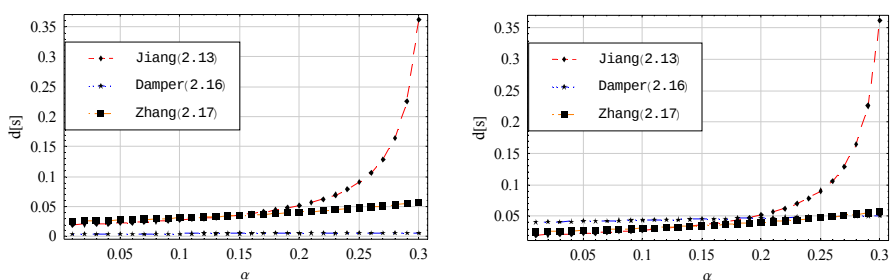
W dodatku przedstawiono porównanie równań (2-13), (2-17) i (2-16) dla parametrów sieci przedstawionych w tabeli C-1. Zwiększenie przepustowości łączy nie zmniejszy opóźnień maksymalnych, jeśli proporcjonalnie zwiększymy obciążenie sieci. Zakładamy, że gdy mamy dostępną większą przepustowość łączy, to obsługujemy więcej strumieni o zbliżonej charakterystyce ruchu, a zatem musimy zwiększyć proporcjonalnie rozmiar wiadra i w takim wypadku maksymalne opóźnienia pozostają na takim samym poziomie jak przy niższej wartości przepustowości łączy. Rozmiar wiadra ograniczający ruch wchodzący do sieci jest istotnym parametrem wpływającym na opóźnienia maksymalne w sieci.

Tabela C-1: Parametry wykresów dla rysunków C-1, C-2, C-3, C-4

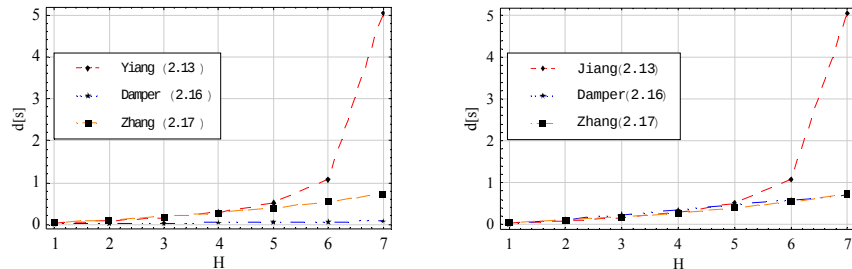
Parametr/ rysunek	R_l [b/s]	C_l [b/s]	M_l [B]	P_n [b/s]	β [B]	τ [b/s]	M_n [B]	E_l	H	α
Rysunek C-1	1×10^8	1×10^8	1500	4×10^8	60000	0.0048	6000	1.2×10^{-4}	-	0.05
Rysunek C-2	1×10^8	1×10^8	1500	4×10^8	60000	0.0048	6000	1.5×10^{-4}	5	-
Rysunek C-3	1×10^8	1×10^8	1500	4×10^8	600000	0.048	6000	1.2×10^{-4}	-	0.2
Rysunek C-4	1×10^8	1×10^8	1500	4×10^8	1×10^8	0.048	6000	1.2×10^{-4}	5	-



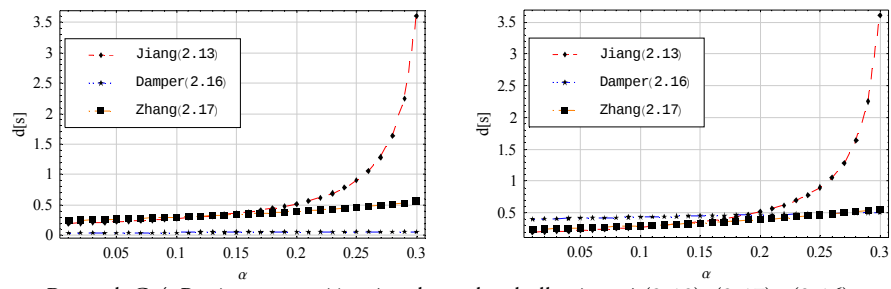
Rysunek C-1: Porównanie opóźnień maksymalnych dla równań (2-13), (2-17) i (2-16) w zależności od rozmiaru sieci dla $\Delta=0.001$ (lewy wykres) i $\Delta=0.01$ (prawy wykres) (pozostałe parametry zostały przedstawione w tabeli C-1)



Rysunek C-2: Porównanie opóźnień maksymalnych dla równań (2-13), (2-17) i (2-16) w zależności od α dla $\Delta=0.001$ (lewy wykres) i $\Delta=0.10$ (prawy wykres) (pozostałe parametry zostały przedstawione w tabeli C-1)



Rysunek C-3: Porównanie opóźnień maksymalnych dla równań (2-13), (2-17) i (2-16) w zależności od rozmiaru sieci dla $\Delta=0.01$ (lewy wykres) i $\Delta=0.1$ (prawy wykres) (pozostałe parametry zostały przedstawione w tabeli C-1)



Rysunek C-4: Porównanie opóźnień maksymalnych dla równań (2-13), (2-17) i (2-16) w zależności od α dla $\Delta=0.01$ (lewy wykres) i $\Delta=0.1$ (prawy wykres) (pozostałe parametry zostały przedstawione w tabeli C-1)