

Badanie możliwości i ograniczeń zarządzania ruchem sieciowym z wykorzystaniem protokołu MPLS we współczesnych sieciach teleinformatycznych

Dariusz Chaładyniak^{*}, Maciej Podsiadły[†]

Warszawska Wyższa Szkoła Informatyki

Abstrakt

W artykule przedstawiono technologię MPLS (ang. MultiProtocol Label Switching), która reprezentuje nowy poziom w rozwoju standardów wykorzystywanych w połączeniach technologii przełączania warstwy drugiej modelu odniesienia ISO/OSI z technologią routingu w warstwie trzeciej. Podstawowym celem procesu standaryzacyjnego MPLS jest stworzenie elastycznej struktury sieci, która będzie posiadała dużo większą efektywność niż dotychczasowe systemy sieci oraz będzie w dużo większym stopniu skalowalna. Analizując wyniki otrzymane w badaniach symulacyjnych, można stwierdzić, iż zastosowanie sieci MPLS jest w stanie zagwarantować optymalną jakość usług dla różnego rodzaju aplikacji.

Słowa kluczowe – protokół MPLS, inżynieria ruchu, jakość usług, opóźnienie, zmienność opóźnienia

1. Wprowadzenie

Technika przesyłania danych oparta na przełączaniu etykiet MPLS pełni dzisiaj zastosowanie jako narzędzie do sterowania ruchem oraz jakością parametrów QoS

^{*} E-mail: dchalad@wwsi.edu.pl

[†] E-mail: m_podsiadly@poczta.wwsi.edu.pl

w sieciach z protokołem IP. MPLS jako docelowy protokół w sieciach transmisji danych stanowi połączenie zalet protokołów IP i ATM:

- pozwala na rozszerzenie możliwości protokołu IP o mechanizmy kontroli jakości usług QoS;
- oferuje mechanizmy zabezpieczeń przed przeciążeniami i mechanizmy do zarządzania ruchem;
- realizuje koncepcje optymalnego doboru tras IP i nieskomplikowaną obsługę pakietów w węzłach przełączających, typową dla sieci ATM;
- może zagwarantować ułatwione kierowanie przepływem strumieni w sieci, oraz znacznie zwielokrotnić procesy komutacji w routerach;
- protokół MPLS może być stosowany w połączeniu z protokołami wyboru tras lub bez nich;
- MPLS umożliwia kierowanie pakietów przez dowolne węzły sieciowe, nie tylko dla ruchu pojedynczych połączeń, ale również dla ruchu zagregowanego VPN.

Rzeczywisty rozwój sieci teleinformatycznych wymusił zastosowanie pewnych modyfikacji protokołu IP na potrzeby różnych aplikacji. Biorąc pod uwagę rozmiary sieci globalnej i możliwości współpracy pomiędzy urządzeniami pochodzącymi od różnych producentów istotnie nie jest to łatwym zadaniem. Problemy napotymane są już podczas wdrażania usprawnionych i coraz bardziej wydajnych mechanizmów wykorzystywanych do sterowania i zarządzania ruchem. Należy zaznaczyć, że protokół IP nie posiada mechanizmów zabezpieczeń przed przeciążeniami. Spowodowane jest to metodą transportowania pakietów w sieciach wykorzystujących protokoły routingu, takich jak: OSPF czy RIP. Nie posiadają one mechanizmów do zarządzania ruchem, czego skutkiem jest stan, w którym pakiety kierowane w tym samym kierunku są przesyłane tą samą trasą, bez uwzględnienia warunków panujących w sieci. W tym przypadku niewykonalne staje się sterowanie ruchem w taki sposób, aby sieć była obciążana równomiernie i nie dochodziło do występowania przeciążeń na „obleganych” kierunkach. Niekorzystne są również mechanizmy retransmisji pakietów, co jednoznacznie dyskwalifikuje TCP/IP do przekazywania usług czasu rzeczywistego.

Protokół IP został zoptymalizowany pod kątem doboru jak najkrótszej trasy w sieci, a nie pod kątem wykorzystywania mechanizmów sterowania przepływem danych. Zatem rozwiązania wykorzystywane dotychczas dla sieci transportowych będą musiały zostać rozbudowane o mechanizmy wspomagające różnicowanie ruchu i zarządzanie nim. MPLS ma na celu zrewolucjonizować stan obecnie istniejących sieci i dostosować je do możliwości wprowadzania nowych usług. MPLS jako rozszerzenie IP w założeniach ma na celu zapewniać m.in.:

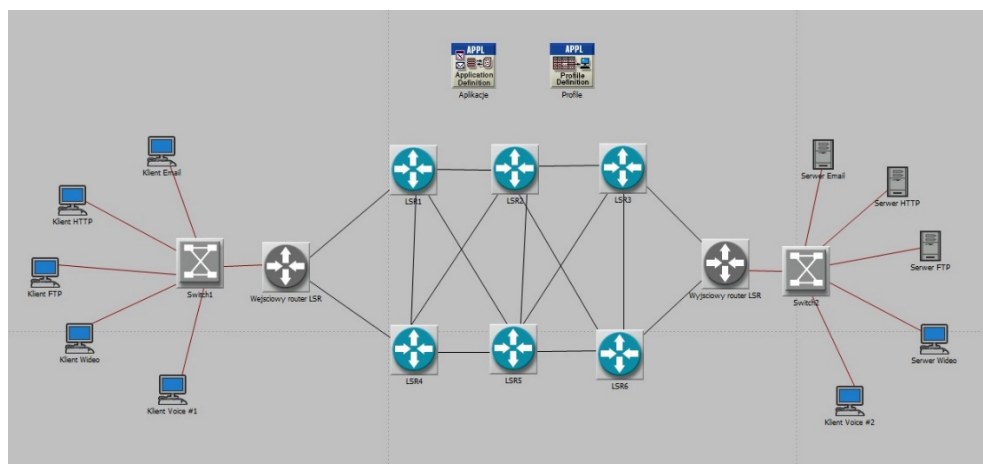
- etykietowanie przesyłanych pakietów, co daje możliwość zwiększenia efektywności oraz zarządzania ruchem w sieci;
- przesyłanie pakietów z taką samą przypisaną etykietą, za pomocą tych samych, ustalonych tunelach – ścieżkach wirtualnych (na wzór sieci ATM);
- realizację funkcji routingowych i funkcji przełączania w węzłach wykorzystujących technikę MPLS;
- usprawnić proces routingu oraz przesyłania pakietów przez sieć.

Zatem proponowany dla systemu IP protokół MPLS integruje zalety obu technik transmisji danych, czyli routingu w IP i zarządzania jakością QoS w ATM.

2. Analiza sieci IP i MPLS w środowisku symulacyjnym Riverbed

Założenia projektowe symulacji:

- utworzenie i konfiguracja sieci zawierającej pięć usług;
- konfiguracja MPLS na urządzeniach brzegowych sieci;
- reorganizacja sieci pod względem obsługi protokołu IP;
- utworzenie obsługi polityki QoS w sieci IP;
- ustawienie symulacji sieci;
- wybranie odpowiednich statystyk testowanych podczas symulacji;
- analiza wyników porównania działania sieci MPLS w stosunku do IP.

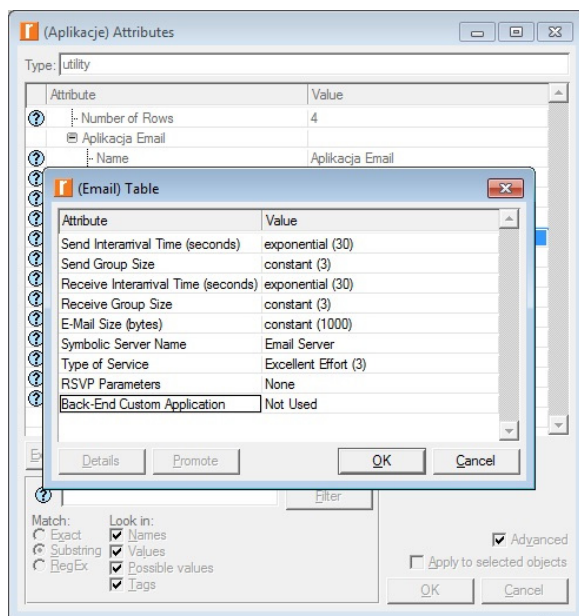


Rysunek 1. Analizowana topologia sieci MPLS w symulatorze Riverbed

Powyższa topologia sieci MPLS zawiera 8 routerów LSR, z czego jeden jest wejściowym routerem LSR, a drugi wyjściowym urządzeniem. Routery brzegowe symulują główny komponent domeny MPLS, czyli wejście / wyjście pakietów zawierających etykietę. Wewnątrz sieci urządzenia posiadają interfejsy dla WAN PPP, gdzie konstruują tablicę LFIB w celu rozgłoszenia i podmiany etykiet. Trasa LSP wewnątrz sieci tworzy się automatycznie. Aktywna konfiguracja LSP wymaga do sygnalizacji protokołu LDP, który wykorzystuje dynamiczne protokoły routingu do kalkulacji odpowiedniej trasy. Zewnętrzne routery LSR połączone są z przełącznikami podłączonymi do stacji roboczych i serwerów obsługujących różne usługi. Pięć stacji reprezentuje klientów usług: *email*, *http*, *ftp*, *wideo* oraz *voice*. Kolejne dwie stacje są klientem *voice* #2 oraz serwerem *wideo*. Trzy serwery obsługują usługi: *email*, *http* oraz *ftp*.

Dodatkowo schemat zawiera dwa magazyny danych: konfigurację aplikacji oraz odpowiadające im profile. Reorganizacja topologii pod względem obsługi protokołu IP polega na wyeliminowaniu wewnętrznych routerów LSR i zamianie urządzenia wejściowego i wyjściowego na podstawowe dwa routery obsługujące routing, z czego jeden z nich jest routerem QoS.

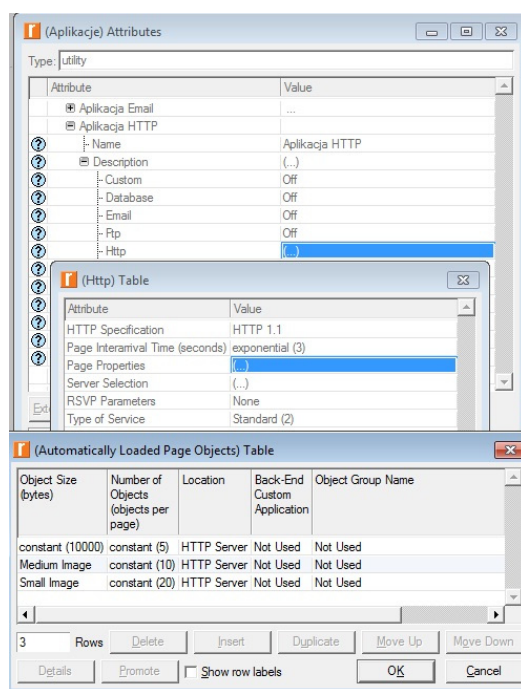
Pierwszym krokiem poprawnej konfiguracji usług jest utworzenie magazynu danych aplikacji i ustawienie atrybutów dla każdej z nich. Poniżej pokazano ustawienia aplikacji *email*.



Rysunek 2. Okno atrybutów usługi *email*

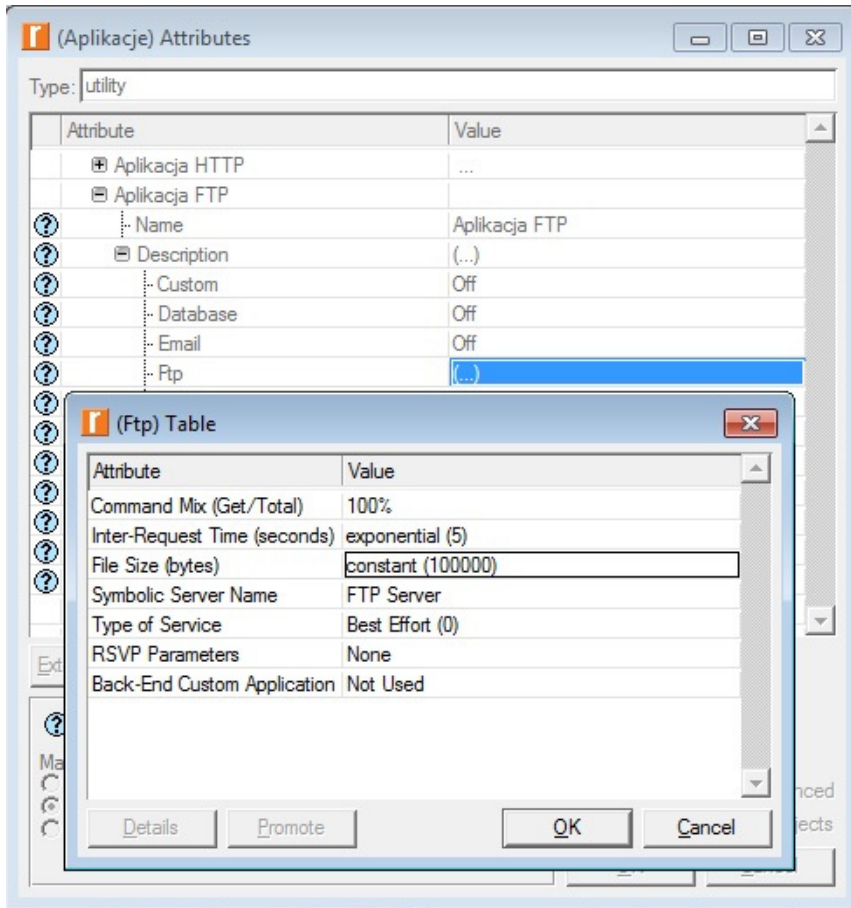
Wiadomości email są często wykorzystywane do komunikacji w organizacji, dlatego ustawiono parametr Type of Service (ToS) o wartości Excellent Effort, który jest najwyższym priorytetem zaraz za usługami czasu rzeczywistego. Ruch wykorzystuje protokół TCP, gdzie tworzy się sesja do generowania i wysyłania zapytań o wiadomość email o wartości 1000 bajtów.

Drugą usługą jest wykorzystanie protokołu *http*. Symuluje ona scenariusz otwierania stron webowych wewnątrz sieci. Wysyłane są zapytania do serwera o obiekty, w tym obrazy różnej wielkości poprzez protokół TCP. Wartość ToS jest ustawiona jako standard obsługi. Poniżej charakterystyka *http* wraz z wielkościami ładowanych obiektów na odpływających stronach.



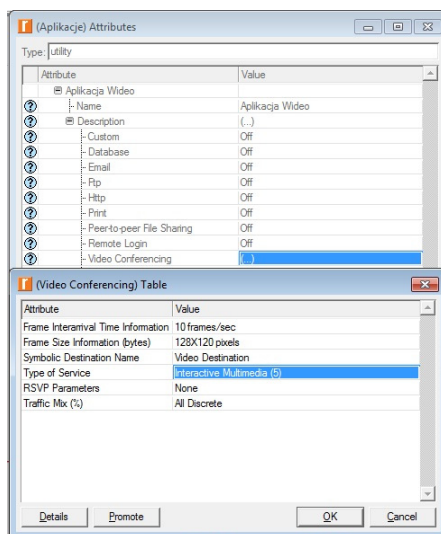
Rysunek 3. Okno atrybutów usługi *http*

Trzecią usługą jest generowanie ruchu *ftp* z serwera bezpośrednio do klienta, oparte na jego żądaniu pobierania danych. Stacja pobiera jeden plik per sesja TCP, natomiast serwer może zmieniać charakterystykę zależną od sesji. Usługa w symulowanej sieci nie jest oznaczona jako główna i posiada najmniejszy priorytet w polu ToS. Poniżej charakterystyka, zawierająca między innymi wielkość pliku około 97 KB.



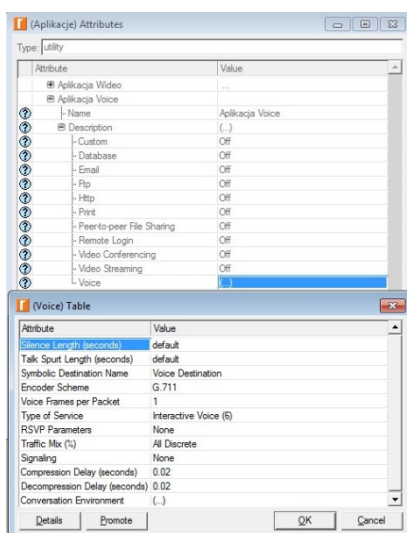
Rysunek 4. Okno atrybutów usługi *ftp*

Czwartą usługą jest wykorzystanie ruchu *video* poprzez protokół UDP. Generowany jest symulowany ruch około 1,2 Mbit/s. Wartość ToS posiada wysoki priorytet dla ruchu czasu rzeczywistego. Poniżej pokazana jest charakterystyka niskiej jakości wideokonferencji, zawierająca między innymi wielkość ramki oraz szybkość jej nadawania przez serwer.



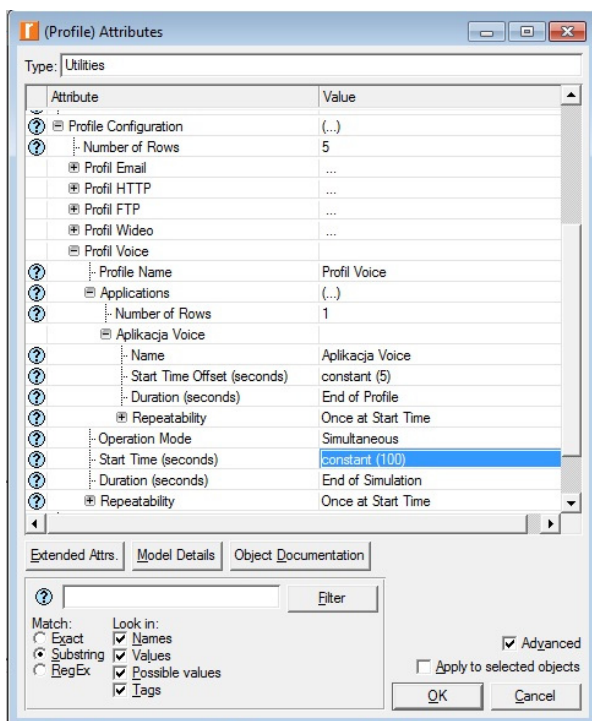
Rysunek 5. Okno atrybutów usługi *wideo*

Ostatnią usługą jest obsługa ruchu głosowego, zawierająca najwyższy priorytet pola ToS. Dzięki takiej wartości, istnieje możliwość dokładnej analizy wyników symulacji, jak opóźnienie czy zmienność opóźnienia. W przypadku tej usługi jak i wideo, ruch przepuszczany jest przez protokół UDP.



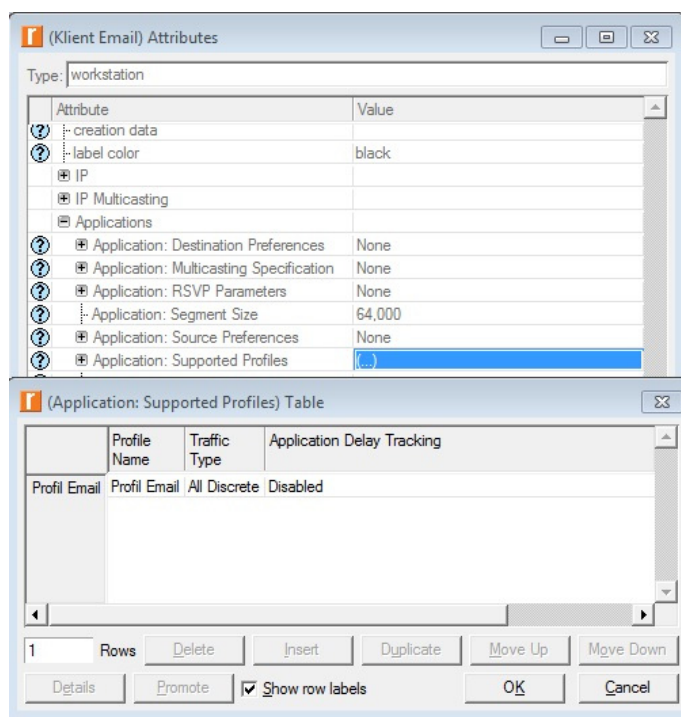
Rysunek 6. Okno atrybutów usługi *voice*

Kolejnym etapem badań jest utworzenie profili dla każdej usługi w magazynie danych profilowych. Każdy z nich zawiera takie same atrybuty dla opisanych aplikacji. Wszystkie usługi zostają uruchomione 100 sekund po starcie symulacji zdarzeń. Ważnym elementem jest tryb operacyjny, który pozwoli na jednoczesne uruchomienie wszystkich aplikacji. Poniżej zaprezentowano atrybuty dla wybranego profilu – ruchu głosowego.



Rysunek 7. Okno atrybutów profilu usługi *voice*

W kolejnym procesie przeprowadzonej symulacji jest przydzielenie odpowiedniego profilu lub serwisu do danej stacji roboczej i serwera. Na rysunku nr 8 widać przydzielenie profilu *email* do stacji roboczej. Ważnym aspektem jest możliwość ustawienia większej liczby profili do danego urządzenia, jednak w tym przypadku każda stacja robocza / serwer obsługuje jedną aplikację. Zaznaczono również opcję obsługi całego ruchu na bazie dyskretnych zdarzeń, inną charakterystyką jest przydzielenie danego ruchu dla całego profilu, czyli na przykład przy wykorzystaniu większej liczby urządzeń.



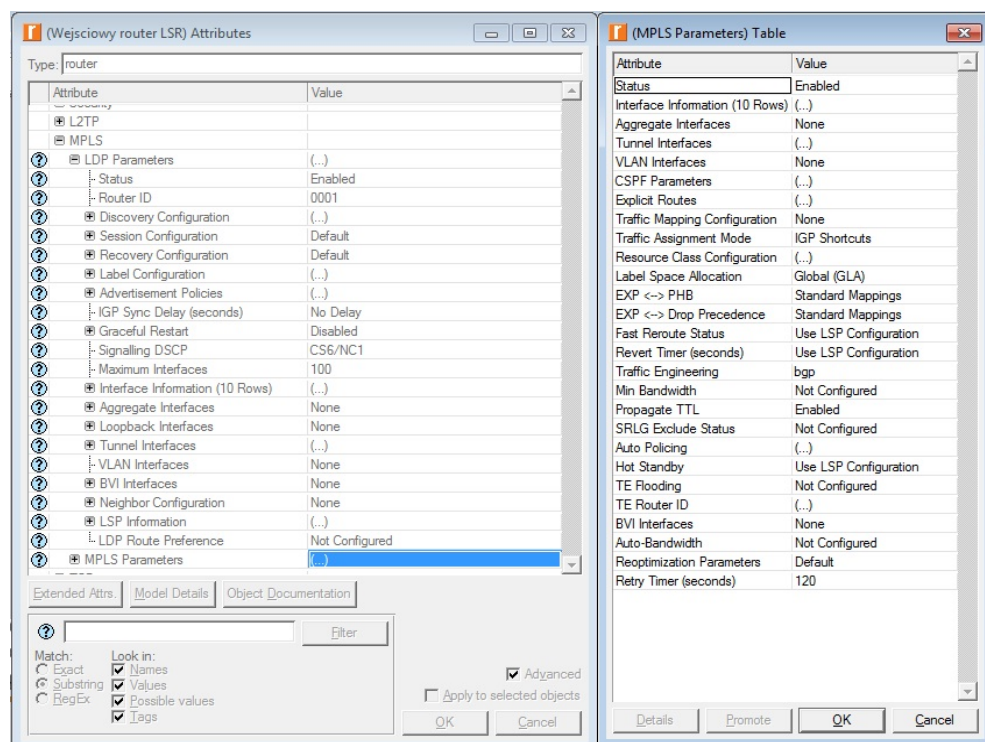
Rysunek 8. Okno atrybutów stacji obsługującej usługę *email*

Następny etap badań zawiera konfigurację atrybutów MPLS na routerach brzegowych symulacji. Poniżej najważniejsze ustawienia podzielone na funkcje protokołu LDP oraz MPLS:

- Status – włączona usługa.
- Discovery Configuration – ustawienie cyklicznych wiadomości w celu odkrycia sąsiadów bezpośrednio podłączonych do routera obsługujących etykietowanie pakietów.
- Session Configuration – domyślne atrybuty potrzebne do ustanowienia sesji LDP TCP w celu wymiany informacji o etykietach między przyległymi urządzeniami.
- Label Configuration – ustawienie atrybutów etykiety, w tym wykorzystanie wartości explicit NULL oraz akceptacji o rozgłoszenia danej etykiety.
- LSP Information – lista tunelowych interfejsów LSP skonfigurowanych na urządzeniu obsługującym LDP.
- Tunnel Interfaces – charakterystyka tunelu pod względem ruchu oraz przeznaczenia dla TE.

- Explicit Routes – ustawienie zapytań o administracyjnie skonfigurowane trasy.
- Fast ReRoute Status – ustawienie na podstawie domyślnej trasy LSP.
- Traffic Engineering – domyślne ustawienie wartości w tablicy BGP na urządzeniu wejściowym, zawierające tylko charakterystyki docelowe BGP.
- Propagate TTL – wartość TTL jest skopiowana z nagłówka pakietu IP.
- Reoptimization Parameters – optymalizacja właściwości tuneli.

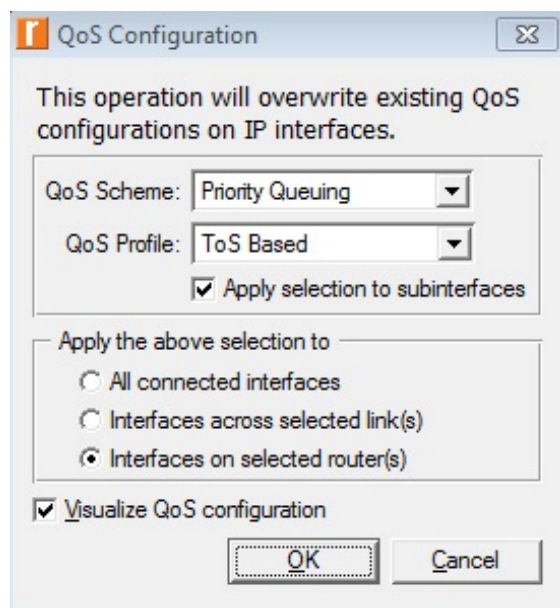
Na poniższym rysunku pokazano opisane atrybuty na przykładzie routera wejściowego LSR.



Rysunek 9. Okno atrybutów LDP / MPLS na routerze wejściowym LSR

Etap konfiguracji topologii sieci MPLS został zakończony. Rekonfiguracja modelu sieci IP wymaga usunięcia wewnętrznych routerów LSR. Brzegowe urządzenia są podłączone tak samo jak w zaprezentowanej topologii ze stacjami roboczymi i serwerami. Jeden z nich jest routerem QoS. Skonfigurowano na nim obsługę ruchu przy pomocy priorytetowej kolejki PQ. Mechanizm ten składa się z czterech kolejek

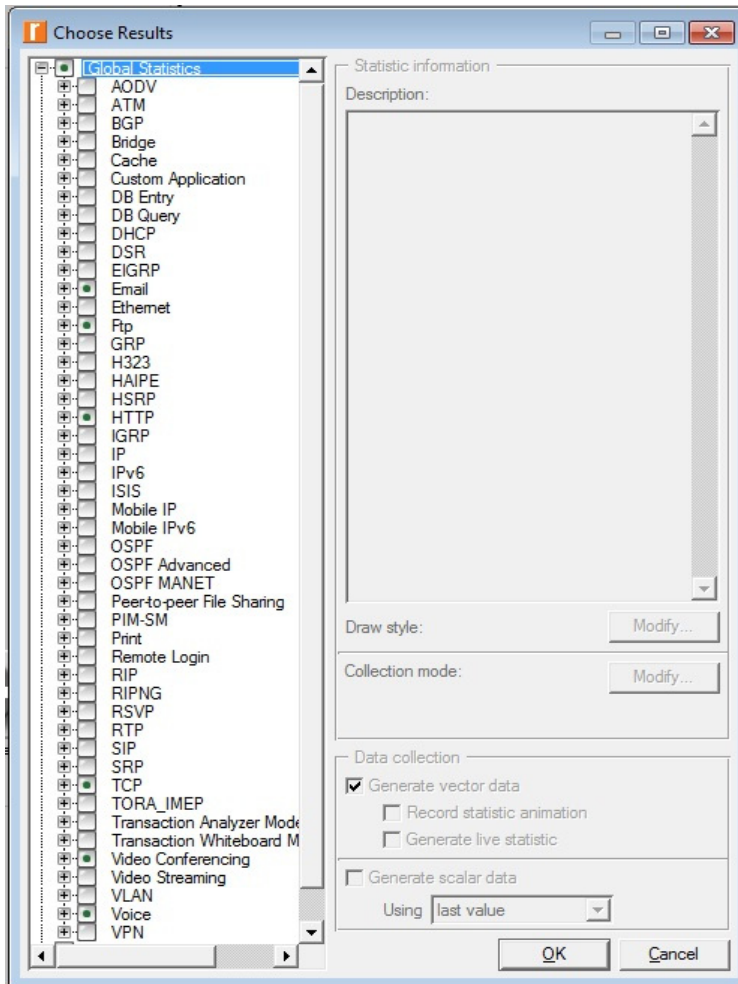
o różnym priorytecie. Główną charakterystyką decydującą o kolejności przesłania danego pakietu jest wartość pola ToS. Poniżej zaprezentowano konfigurację QoS na wszystkich interfejsach routera.



Rysunek 10. Konfiguracja QoS w sieci IP

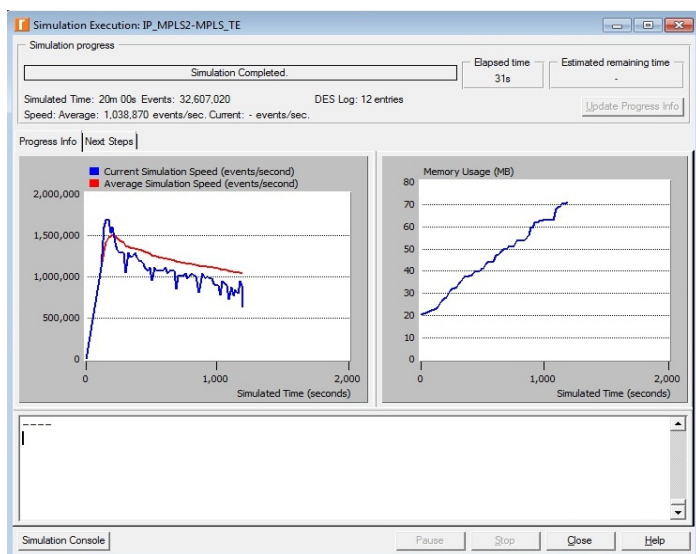
Obydwa modele sieci zostały poprawnie skonfigurowane. Następnym krokiem jest wybór indywidualnych statystyk symulacji. Dla każdej usługi wybrano odpowiednie schematy. Dla usługi *ftp* ustawiono średni czas reakcji na pobranie pliku z serwera, dla usługi *http* średni czas reakcji na odczyt obiektu na stronie. Z kolei dla usługi *email* wyznaczono średni czas reakcji na pobranie i wysłanie wiadomości.

Aplikacja obsługująca wideokonferencje zbiera statystyki dotyczące średniego czasu opóźnienia pakietów od źródła do miejsca docelowego wzdłuż sieci. Standardowo dla usług głosowych skonfigurowano statystyki dotyczące opóźnienia oraz zmienności opóźnienia. Ostatnimi parametrami symulacji jest określenie opóźnienia oraz liczby retransmisji TCP poprzez docelowy węzeł.



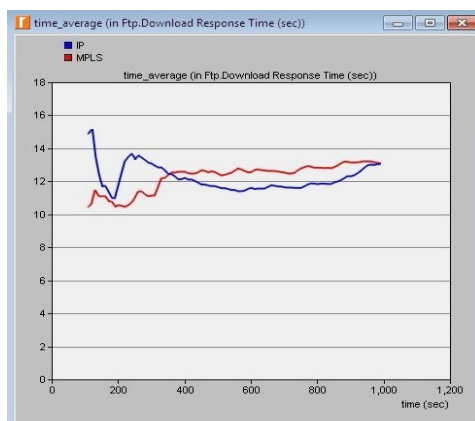
Rysunek 11. Wybór statystyk symulacji

Po wybraniu odpowiednich statystyk ustawiono zbieranie danych zaprezentowanych scenariuszy. Procesy symulacji trwają 1200 sekund. Podany czas nie ma odniesienia do rzeczywistości, gdzie realna symulacja trwała 31 sekund. Zarejestrowano ponad 32 miliony zdarzeń, o częstotliwości około 1 miliona zdarzeń na sekundę rzeczywistości. Takie wystąpienie stanowi dużą wartość, przybliżając dokładność obliczeń do rzeczywistości.



Rysunek 12. Okno wykonania symulacji

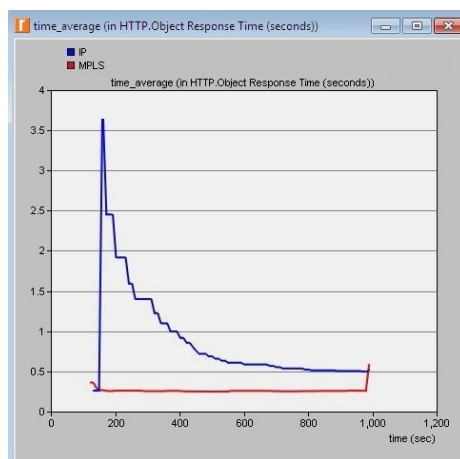
Po zakończonej symulacji należy przystąpić do analizy jej wyników. Rezultaty mogą być zaprezentowane na wiele sposobów. Poniżej przedstawiono porównanie zachowania sieci IP oraz MPLS w stosunku do danej usługi. Wykresy można dowolnie skonfigurować według predefiniowanych kryteriów. Istnieje możliwość ustawienia zakresu czasu, ilości zdarzeń lub przedstawienie wykresu pod różnymi postaciami: histogram, wykres punktowy, zakres liniowy, średnia, graf czy wykres warstwowy.



Rysunek 13. Średni czas reakcji pobierania z serwera *ftp*

Powyższa charakterystyka reprezentuje czas jaki upłynął pomiędzy wysłaniem zapytania i otrzymaniem pakietu z odpowiedzią pomiędzy stacją roboczą a serwerem *ftp*. Każda reakcja stacji klienckiej zawiera się w tej statystyce. Usługa *ftp* posiada najmniejszy priorytet pola Type of Service i działa w trybie Best Effort, oznaczającą transmisję danych o możliwie najwyższej przepustowości, ale bez żadnej gwarancji utrzymania poziomu jakości usług.

Kolorem czerwonym zaznaczono średni czas reakcji w symulacji sieci MPLS, gdzie usługa uruchomiona jest zgodnie z harmonogramem od setnej sekundy. Na początku różnica między MPLS a IP wynosi około 5 sekund, jednak pod koniec symulacji poziom średniego czasu wyrównuje się z uwagi na rozłożenie ruchu wszystkich usług i odpowiednie przydzielenie klasy przepływu pakietów do danej usługi w kolejce PQ.

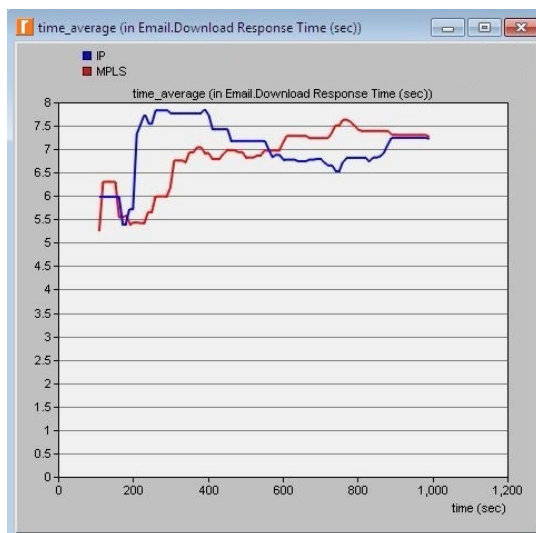


Rysunek 14. Średni czas reakcji odczytu obiektu na stronie

Schemat przedstawia średni czas reakcji odpowiedzi dla każdego obiektu przykładowej strony wykonanego przez stację roboczą na serwerze *http*. Aplikacja obsługiwana jest w trybie standardowym zawierając średni priorytet w polu ToS.

Z wykresu wynika, że MPLS utrzymuje obsługę zapytania na stałym, niskim poziomie z odpowiedzią poniżej 0,5 sekundy. Z kolei niebieski wykres reprezentujący kolejkę PQ, podkreśla fakt, że w pierwszej fazie symulacji, podany mechanizm nie jest w stanie zapewnić określonej jakości usług w porównaniu do MPLS. Podany priorytet uniemożliwia obsługę ruchu na wysokim poziomie, zgodnie z poprzednim wykresem po około 300 sekundach symulacji można zauważyć znaczny spadek średniego czasu reakcji IP, zbliżając się powoli do niewielkiej różnicy w stosunku

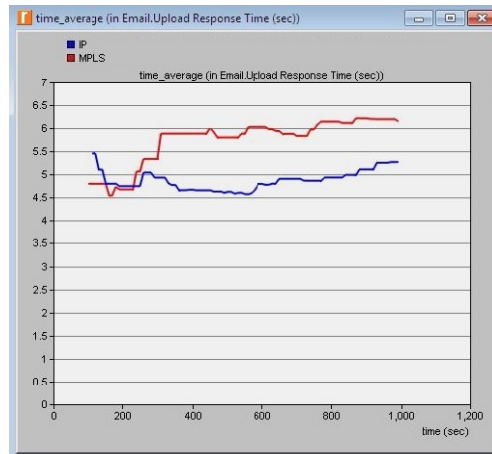
do MPLS. Wpływ na początkowy stan uwidoczniiony na wykresie, ma na pewno wielkość i różnorodność danych, znajdujących się na przykładowej stronie.



Rysunek 15. Średni czas reakcji pobierania wiadomości z serwera *email*

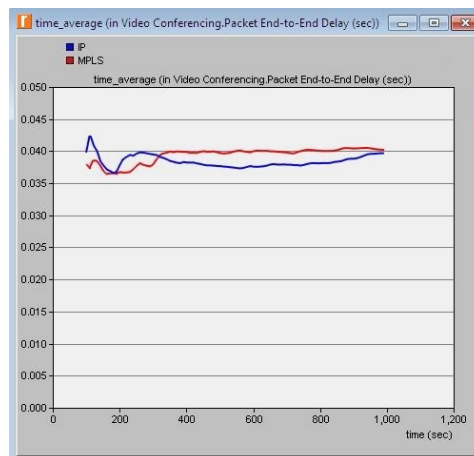
Powyższy wykres przedstawia porównanie zachowania serwera usługi *email* w sieciach IP oraz MPLS. Charakterystyka określa średni czas wysłania zapytania ze stacji do serwera aplikacji oraz otrzymania odpowiedzi zwrotnej z wiadomością. Opóźnienie sygnalizacji dla ustanowienia połączenia również zawiera się w tym czasie, stąd takie wysokie średnie, oscylujące pomiędzy 5 a 8 sekund dla obydwu sieci.

Usługa posiada najwyższy możliwy priorytet, pomijając aplikacje czasu rzeczywistego – Excellent Effort. Przy współpracy w sieci IP z kolejkowaniem PQ, można wywnioskować, że MPLS doskonale sobie radzi z obsługą usługi *email*, praktycznie wyrównując jej stan w końcowej fazie testów, co widać na poniższej charakterystyce.



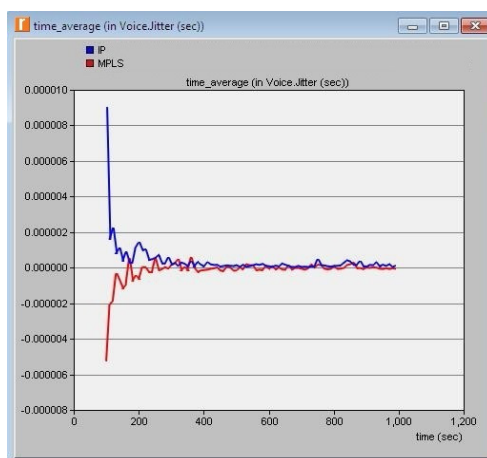
Rysunek 16. Średni czas reakcji wysyłania wiadomości na serwer *email*

Schemat jest bardzo podobny do poprzedniego, gdzie wartości średnie są w podobnym przedziale między 4,5 a prawie 6,5 sekundy. Czas liczony jest na podstawie wysłania wiadomości *email* ze stacji klienckiej do serwera poczty oraz otrzymania potwierdzenia jej wysłania. W tym przypadku również opóźnienie konfiguracyjne powoduje zwiększenie czasu reakcji. MPLS przy wysyłaniu wiadomości posiada jeszcze lepszą charakterystykę niż przy idealnej kompozycji kolejkowania PQ z wysokim priorytetem aplikacji w sieci IP.



Rysunek 17. Średni czas opóźnienia pakietów wideokonferencji

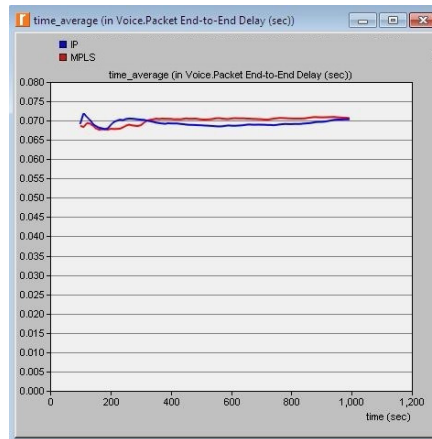
Powyższy wykres przedstawia średni czas opóźnienia pakietu wideo wysłanego do docelowego węzła warstwy aplikacji. Średnia zebrana jest ze wszystkich węzłów sieciowych. Opóźnienie wynosi około 40 milisekund, co jest bardzo dobrym wynikiem, mimo małej rozdzielczości obrazu. Obydwie charakterystyki są na wyrównanym poziomie, jednak w pierwszej fazie można dostrzec lekki wzrost dla sieci IP z uwagi na wypełnianie się kolejki. Wartość pola Type of Service oznaczona jest jako strumieniowanie wideo i posiada jeden z najwyższych priorytetów zaraz za usługami głosowymi, z czego wynika mała wielkość średniego opóźnienia.



Rysunek 18. Średnia zmienność opóźnienia pakietów głosowych

Pokazany na rysunku nr 18 schemat przedstawia największą zaletę sieci MPLS, które bezkonkurencyjnie potrafią obsłużyć aplikacje czasu rzeczywistego, w tym przypadku ruchu głosowego. Usługa posiada najwyższy priorytet pola ToS – Interactive Voice. Sporą zmienność opóźnienia można dostrzec na początku symulacji, w momencie wpływania pakietów do kolejki. Po około 5 minutach proporcje się wyrównują. Zmienność opóźnienia jest liczona na podstawie znaczników czasu odtwarzania głosu. Pierwsze dwa znaczniki są wyznaczane na źródłowym węźle, ostatnie dwa na docelowym. Charakterystyka jest różnicą sumy dwóch zmiennych docelowych i źródłowych.

W przypadku sieci MPLS zmienność opóźnienia jest ujemna z uwagi na mniejszą wartość różnicy na węźle końcowym w stosunku do węzła początkowego. Oznacza to bardzo dobrą charakterystykę pracy dla aplikacji czasu rzeczywistego.



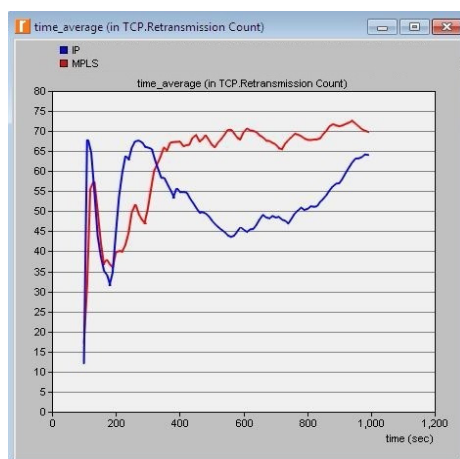
Rysunek 19. Średni czas opóźnienia pakietów usługi *voice*

Początkowy skok w sieci IP można porównać do zmiennej na poprzednim wykresie. Po kolejkowaniu pakietów i obsłudze zgodnie z priorytetem, stany wyrównują się i średnia wynosi około 70 milisekund. Wynik jest prawie dwukrotnie wyższy niż obsługa ruchu wideo, z uwagi na większe wymagania zasobów co do aplikacji głosowej oraz fakt, że videokonferencja przeprowadzana była w niskiej jakości. Opóźnienie pakietów głosowych jest wyliczane na podstawie sumy kilku parametrów: opóźnienie sieciowe – odniesienie wysłania pakietu RTP (czasu rzeczywistego) od odbiorcy oraz otrzymanie jej przez węzeł docelowy, średnia czasu kompresji, dekompresji, kodowania i dekodowania. Statystyka jest uśredniania ze wszystkich węzłów w sieci.



Rysunek 20. Średni czas opóźnienia stosu TCP

Na powyższym rysunku pokazano średni czas opóźnienia pakietów odebranych przez warstwy TCP w całej sieci dla każdego połączenia. Miarą jest czas pomiędzy wysłaniem pakietu danych aplikacji ze źródłowej warstwy TCP a kompletnym odebraniem go przez tę warstwę w docelowym węźle. Średni czas opóźnienia wynosi pomiędzy 1,5 a 4,5 sekundy, gdzie TCP jest wykorzystywany w trzech usługach o najmniejszej wartości priorytetowej, tj. *email*, *ftp*, *http*. Początkowy nagły skok jest powiązany z opcją zaznaczoną przy ustawieniu atrybutów każdej aplikacji, który oznacza włączenie równoległe wszystkich usług w sieci, zgodnie z harmonogramem symulacyjnym.



Rysunek 21. Liczba retransmisji TCP

Widoczna na rysunku nr 21 charakterystyka określa liczbę retransmisji TCP w określonym, uśrednionym przedziale czasowym w trakcie trwania symulacji. Jest to suma całkowita zawierająca tylko retransmisje wykonane z niepotwierdzonego bufora TCP. Wyliczana jest na podstawie otrzymanych segmentów TCP wzdłuż całej sieci, ze wszystkich połączeń. Z zamieszczonego schematu wynika, że największą retransmisję miała sieć MPLS pod koniec symulacji, około 73 jednostki. Pierwszy duży skok powiązany jest z poprzednim wykresem i dotyczy uruchomienia wszystkich usług naraz, co spowodowało sporą liczbę ponownej transmisji danych w krótkim okresie czasu.

Podsumowując przeprowadzoną symulację, MPLS jest w stanie zagwarantować jakość usług dla różnego rodzaju aplikacji, szczególnie czasu rzeczywistego. W porównaniu do sieci IP, która wykorzystywała jeden z optymalnych mechanizmów kolejowania, rozwiązanie budowy sieci routerów LSR, nadających etykiety pakietom

jest o wiele bardziej wydajne i skalowalne. Warto zaznaczyć, że inżynieria ruchu i tunelowanie ruchu nie były wykorzystywane w topologii MPLS, i co za tym idzie, charakterystyki opóźnień i innych wartości mogłyby być jeszcze lepsze. Przy poprawnej konfiguracji urządzeń, MPLS jest w stanie obsłużyć wiele usług uruchomionych jednocześnie, nie ograniczając przy tym zasobów oraz nie zwiększając opóźnień czy jej zmienności, co miałyby katastrofalne skutki dla aplikacji czasu rzeczywistego.

Podsumowanie

Analizując charakterystyki opisanych mechanizmów technologia MPLS może przynieść następujące korzyści:

1. Organizacja ruchu – MPLS dostarcza mechanizmy do zarządzania strumieniami ruchu o różnym natężeniu, między różnymi urządzeniami w sieci. Dodatkowo umożliwia obsługę interfejsu dla istniejących protokołów routingu oraz stosu warstwy drugiej. Udostępnia również narzędzia do odwzorowywania adresacji IP w proste, stałej długości etykiety, wykorzystywane przez różne techniki przełączania i przesyłania dalej.
2. Dystrybucja etykiet – MPLS określa trzy procedury tworzenia etykiet. Pierwsza z nich uwzględnia topologię, która wykorzystuje przetwarzanie połączone z protokołami routingu, takimi jak BGP i OSPF. Druga, zważy na wymagania, korzystając z protokołu RSVP w celu zapewnienia ruchu na odpowiednim poziomie. Ostatnia, uwzględnia ruch, używając rejestrowane pakiety do wywołania procedur przydzielania i dystrybucji etykiet.
3. Quality of Service – jakość usług w sieci MPLS odbywa się w dwóch trybach. Pierwszy z nich gwarantują routery QoS obsługujące MPLS, które wykorzystują kolejkovanie pakietów, na podstawie trzech bitów EXP. Drugi tryb zorganizowany jest podstawie sygnalizacji pasma. Realizacja odbywa się na podstawie klas ruchu przez utworzenie wielu niezależnych połączeń wirtualnych.
4. Traffic Engineering – implementacja TE w sieci MPLS umożliwia sterowanie trasami przesyłu informacji w pewnym stopniu niezależnie od osiągalności sieci określonych w tablicy routingu. Inżynieria ruchu stanowi koncepcję optymalizacji wykorzystania zasobów w rdzeniu sieci.
5. Wirtualne sieci prywatne – MPLS VPN posiada wysokie parametry skalowalności, elastyczność w doborze technologii dostępowych oraz w celu definiowania komunikacji pomiędzy sieciami VPN. Dodatkowo można łatwo zorganizować routing pomiędzy sieciami VPN oraz skonfigurować nowego użytkownika.

**A study of opportunities and limitations of traffic management network
using MPLS protocol in modern ICT networks**

Abstract

The article presents MPLS (Multiprotocol Label Switching), which represents a new level in the development of standards used in combinations of switching technology of the second layer ISO/OSI reference model and routing technology in the third layer. The primary objective of the standardization of MPLS is to create a flexible network structure, which will be much more efficient than existing systems and networks, and will be much more scalable. Having analyzed the results obtained in the simulation tests, we can conclude that the use of MPLS network is able to guarantee optimal quality of service for all kinds of applications.

Keywords – MPLS, traffic engineering, quality of services, delay, delay variatio