

	UNIVERZITET "UNION" RAČUNARSKI FAKULTET Knez Mihailova 6/VI 11000 BEOGRAD	Broj:
		Datum:

UNIVERZITET UNION

RAČUNARSKI FAKULTET

BEOGRAD

RAČUNARSKE MREŽE I KOMUNIKACIJE

DIPLOMSKI RAD

Kandidat: Aleksandar Nikolić

Broj indeksa: 85/05

Tema rada: IMPLEMENTACIJA IPV6 U MPLS MREŽI U ISP OKRUŽENJU

Mentor rada: prof. Mirjana Radivojević

Beograd, 2010.

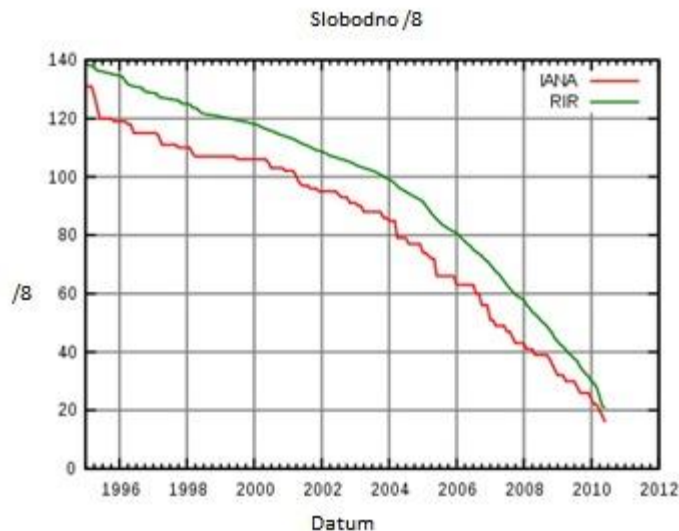
Sadržaj

1.	Uvod	2
2.	IPv6 protokol	6
2.1.	Promene u IPv6 protokolu	7
2.2.	Opis IPv6 zaglavlja	7
2.3.	Predstavljanje IPv6 adresa	9
2.4.	IPv6 tipovi adresa	10
2.5.	Specifični tipovi adresa	12
2.6.	Otkrivanje suseda i autokonfiguracija	14
2.7.	Adresna autokonfiguracija	15
2.8.	IPv6 rutiranje	16
2.9.	IPv6 QoS	17
2.10.	IPv6 sigurnost	17
2.10.1.	IPv4 pitanja sigurnosti	18
2.10.2.	IPv6 sigurnosna poboljšanja	19
2.10.2.1.	Veći adresni prostor	20
2.10.2.2.	IPsec	20
2.10.2.3.	Zaglavlje provere (Authentication Header – AH)	21
2.10.2.4.	Enkapsulacija sigurnosnog tovara	22
2.10.2.5.	Transport i tunel modovi	23
2.10.2.6.	Pregovaranje protokola i upravljanje razmenom ključeva	23
2.10.3.	IPv6 pitanja sigurnosti	23
2.10.3.1.	Pitanja u vezi sa dvostrukom konfiguracijom	23
2.10.3.2.	Pitanja manipulacije zaglavlja	24
2.10.3.3.	Pitanje poplava (flooding)	24
2.10.3.4.	Mobilnost	24
2.11.	IPv4/IPv6 tranzicioni mehanizmi	25
2.11.1.	Dvostruka konfiguracija (Dual stack)	25
2.11.2.	IPv4-IPv6 tunelovanje	26

2.11.2.1.	6to4 (IPv6 to IPv4).....	27
2.11.2.2.	ISATAP (<i>Intra Site Automatic Tunneling Addressing Protocol</i>)	28
2.11.2.3.	IPv6 preko IPv4 (<i>6over4</i>)	28
2.11.2.4.	Agenti tunela (Tunnel Brokers).....	29
2.11.2.5.	Teredo.....	30
2.11.2.6.	DSTM (<i>Dual stack transition Mechanism</i>).....	33
2.11.3.	Translacija	34
2.11.3.1.	SIIT (<i>Stateless IP/ICMP Translation algorithm</i>).....	34
2.11.3.2.	BIS (<i>Bump in the Stack</i>)	35
2.11.3.3.	BIA (Bump in the API).....	36
2.11.3.4.	NAT-PT (<i>Network Address Translation with Protocol Translation</i>).....	37
3.	MPLS	38
3.1.	MPLS osobine i izgled.....	39
3.1.1.	Izgled MPLS labele (MPLS label stack)	40
3.2.	MPLS komponente	42
3.2.1.	Kontrolni deo i deo za rukovanje paketima (<i>Control /Data plane</i>).....	43
3.2.2.	Nezavisna kontrola.....	44
3.2.3.	Naređena kontrola	44
3.3.	Razmena labela	45
3.3.1.	TDP (<i>Tag Distribution Protocol</i>).....	45
3.3.2.	LDP (<i>Label Distribution Protocol</i>).....	46
3.4.	Primer mreže provajdera.....	47
3.4.1.	Tipovi uređaja za obeležavanje paketa.....	48
3.4.2.	Pojašnjenje primera	48
3.5.	Putanje kroz MPLS	49
3.6.	Isporučivanje IPv6 preko MPLS mreže	50
3.7.	Ivica provajdera – 6 PE (<i>IPv6 Provider Edge</i>)	52
3.8.	6VPE (<i>IPv6 VPN Provider Edge</i>).....	57
4.	Zaključak.....	59
5.	Literatura	61

1. Uvod

Sa razvojem Internet globalne računarske mreže, kao i sa pojavom novih aplikacija i servisa broj korisnika se stalno povećava. Uporedo sa porastom popularnosti Interneta i porastom broja korisnika raste i potreba za javnim *IP (Internet Protocol)* adresama, koje su ustvari numeričke oznake koje se dodeljuju uređajima koji učestvuju u računarskoj mreži, kako bi se omogućila komunikacija između ostalih uređaja. Trenutna verzija IP protokola, na kojem počiva Internet, za adresiranje uređaja koristi IP adresu koja ima veličinu 32 bita. Samim tim očigledno je da je adresni prostor odnosno broj IP adresa ograničen što predstavlja veliki problem za dalji razvoj globalne računarske mreže.



Slika 1: Vremenski prikaz dodele i alokacije

Statistika pokazuje da je već do 2004. godine dve trećine adresnog prostora protokola IPv4 bilo alocirano, a dostizanje granice dodele IPv4 adresa se predviđa pre početka 2012. godine (slika 1). Ovo rapidno alociranje adresnog prostora IPv4 protokola je dovelo do razvoja tehnika očuvanja adresnog prostora kao što su:

1. Korišćenje privatnih adresa – koje se koriste kada rutiranje preko internet nije obavezujuće, tj. kada je potrebno povezati uređaje lokalno u jednu mrežu. Nazvane

su privatne zbog toga što rutiranje paketa iz ovih mreža nije moguće preko Interneta. Ove adrese može koristiti svako bez prethodnog odobrenja od strane regionalnog Internet registra *RIR (Regional Internet Registry)*, čiji je zadatak da upravlja dodelom i registracijom adresa unutar određenog regiona u svetu (slika 2).



Slika 2: Prikaz regionalnih Internet registara u svetu

2. Upotreba *NAT (Network Address Translator)* uređaja – koja omogućava da se dodeli jedna javna adresa organizaciji koju će koristiti više uređaja iz njene privatne mreže. Uređaji na kojima se definiše *NAT* tehnika se postavljaju na krajevima mreža tj. na prelazu između privatnih i javnih mreža. Korisnicima je transparentna promena *IP* adrese koja se dešava na *NAT* uređaju.
3. Upotreba uređaja za dinamičku dodelu adresa *DHCP (Dynamic Host Configuration Protocol)* – koji dodeljuje privremene adrese koje se nakon isteka određenog vremena mogu ponovo koristiti.
4. Virtuelno hostovanje (*virtual hosting*) – metoda hostovanja višestrukih imena domena na računaru koji koristi samo jednu *IP* adresu.
5. Bolja kontrola regionalnih Internet registara – *RIR (Regional Internet Registry)* ima zadatak da upravlja dodelom i registracijom adresa unutar određenog regiona u svetu. Bolja kontrola će omogućiti da se pažljivije dodeljuju adrese, ali i da se dodeljivanje adresa ne čini bez prethodno napravljenog plana.

6. Ponovno numerisanje mreža (*network renumbering*) – ovaj vid očuvanja adresnog prostora se sprovodi kako bi se povratili veliki blokovi adresnog prostora koji su nesvesno, pogrešno alocirani u ranim danima Interneta.

Iako su ove tehnike zaista dovele do očuvanja adresnog prostora, one su postale nekompatibilne sa određenim brojem fundamentalno novih trendovima internet kao što su:

- *Peer to peer* aplikacije – mnoge aplikacije u prošlosti su bile zasnovane na klient/server arhitekturi, ali sada većina aplikacija se oslanja na ovakav vid (*peer-to-peer*) komunikacije, kao što su VoIP (Voice...), otvorene aplikacije za razmenu fajlova (npr. Napster), kao i aplikacije za video igre (*video gaming*). Iako ove aplikacije mogu da se prilagode da rade zajedno sa NAT-om, ovo zahteva značajan operacioni trošak i zahteva napredna podešavanja i rekonfiguraciju unutar NAT uređaja.
- Naglo povećanje broja uređaja koji mogu da se konektuju na internet mrežu – iako su personalni računari ključni uređaji kojima treba IP adresa kako bi komunicirali preko Interneta i drugim tipovima uređaja je potrebna Internet konektivnost, kao što su mobini telefoni, prenosivi uređaji za zabavu, PDA, kao i kućni i industrijski sistemi za kontrolu.
- Stalna konekcija (*Always-on*) – mnogi uređaji imaju stalnu konektivnost zbog napretka tehnologija pristupa kao što su širokopojasni pristup (*broadband*) i bežični pristup (*wireless*).

Sve prethodno navedeno dovodi do velikih zahteva po pitanju korišćenja adresnog prostora. S obzirom da je IPv4 već previše opterećen, rešenje problema je uvođenje nove verzije protokola koji može da pruži dovoljno adresnog prostora za sve ove uređaje. Novi protokol se naziva IPv6 i za adresiranje uređaja koristi IP adresu veličine 128 bita čime se broj IP adresa višestruko povećava.

Međutim, uvođenje novog protokola uslovljeno je nizom problema, počev od podrške na mrežnim uređajima, na čemu proizvođači opreme intenzivno rade, kao i podrškom za IPv6 u MPLS backbone mrežama servis provajdera.

S obzirom da su današnje mreže bazirane na IPv4 protokolu, pri čemu je MPLS protokol u mrežama servis provajdera i operatera inicijalno kreiran za IPv4 mreže, postavlja se pitanje kako je najbolje migrirati na IPv6 protokol, a istovremeno nastaviti sa pružanjem definisanih usluga i servisa.

Pitanja koja se nameću i koja treba razrešiti su sledeća:

- Može li da se uvede IPv6 protokol bez menjanja dosadašnje IPv4 mreže tj. njenog jezgra?
- Može li da se napravi IPv6 VPN (*Virtual Private Network*) bez menjanja dosadašnjeg IPv4 jezgra mreže?
- Da li nam i dalje treba MPLS? Dakle da li nam trebaju LDP i RSVP protokoli i da li treba da imaju podršku za IPv6?
- Možemo li da izgradimo IPv6 mrežu sa MPLS-om?

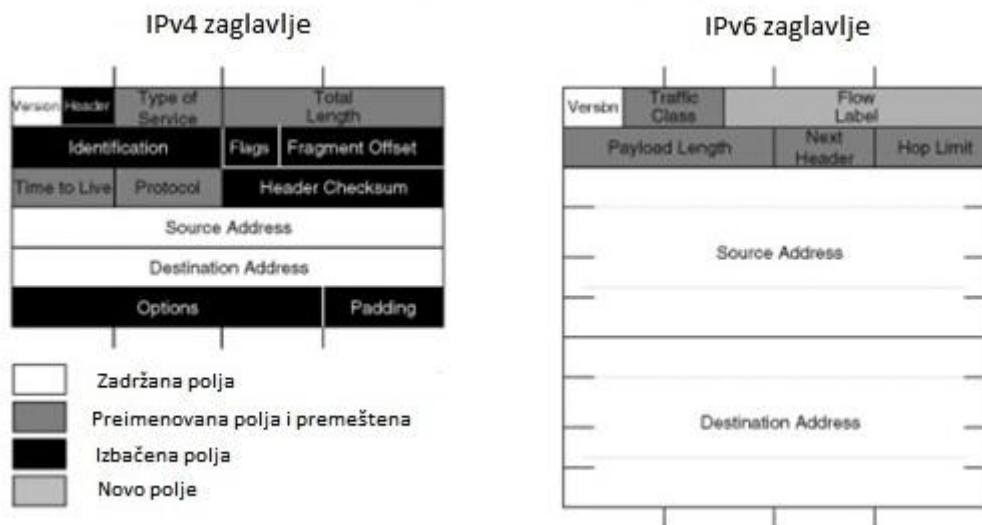
U ovom radu biće predstavljen IPv6 protokol i njegove karakteristike, kao i struktura MPLS okosnice (*backbone*) mreža Internet servis provajdera. Biće analizirane opcije za migraciju sa postojeće IPv4 strukture na novu IPv6 MPLS mrežu.

Takođe će u ovom radu, na primeru jedno core ISP mreže, biti prikazani principi implementacije i integracije IPv6 protokola i MPLS tehnologijom i objašnjeni svi problemi koji se pri tom javljaju.

U zaključku rada izneću svoje mišljenje o navedenom problemu kao i neka od mogućih rešenja.

2. IPv6 protokol

IPv6 protokol je definisan 1998. unutar IETF grupe standarda. Za adresiranje uređaja koristi IP adresu veličina 128 bita, za razliku od IPv4 koja je koristila IP adresu veličine 32 bita, čime se broj IP adresa višestruko povećava. Veličina adresnog polja je učetvorostručena na 16 bajtova, pružajući teoretski 2^{128} ili $3.4 \cdot 10^{38}$ adresabilnih čvorova, koji pružaju više nego dovoljno globalno jedinstvenih IP adresa. IPv6 je dodatno pojednostavljen i za razliku od prethodne verzije tj. IPv4 ima samo 8 od 14 polja koliko je imao IPv4 (slika 3).



Slika 3: Poređenje zaglavlja IPv4 i IPv6.

Format zaglavlja je pojednostavljen korišćenjem polja fiksnih veličina koja su obavezujuća i korišćenjem serijskog povezivanja (*daisy chained*) opcionalnih polja koja su između zaglavlja i polja podataka.

Dodatno, IPv6 obezbeđuje veću sigurnost i integritet podatka (celovitost podatka), kao i autokonfiguraciju, *multicast* i *anycast*. Pored toga, pruža integrisani kvalitet servisa QoS (*Quality of Service*) kao i povećanu efikasnost pri slanju informacija. Za razliku od IPv4

koji koristi NAT, IPv6 formira komunikaciju između dva uređaja tzv. s kraja na kraj bez prisustva NAT-a.

2.1. Promene u IPv6 protokolu

- IP adresa je sa 32 bita povećana na 128 bita znatno uvećavajući adresni prostor
- Pojednostavljeno zaglavlje – neka polja su izbačena, dok su neka napravljena opcionalnim tako da se time omogućava smanjenje troška procesiranja i limitiranje troška propusnog opsega IPv6 zaglavlja.
- Poboljšana podrška za opcije i ekstenzije – promene u smislu IP opcija zaglavlja koje omogućavaju efikasnije prosleđivanje i blaža ograničenja na dužinu opcija, kao i veća fleksibilnost za uvođenje novih opcija u budućnosti.
- Autentifikacija i mogućnost privatnosti – ekstenzije za podršku autentifikaciji, celovitosti podataka i (opcionalno) poverljivim podacima za IPv6.
- Sposobnost označavanja toka podataka – nova mogućnost je dodata, kako bi se omogućilo označavanje pripadnosti paketa, za koju pošiljalac zahteva posebno rukovanje npr. posebno rukovanje za servise u realnom vremenu.

2.2. Opis IPv6 zaglavlja

Version	Traffic Class	Flow Label
Payload length	Next Header	Hop Limit
Source Address (128 bits)		
Destination Address (128 bits)		

Slika 4: Opis polja u zaglavlju

Na slici 4 možemo videti definisana polja u IPv6 zaglavlju. Polja su u daljem tekstu bliže objašnjena.

Version – verzija protokola, sada je IPv6 (6) umesto IPv4 (4). Polje ima 4 bita.

Traffic class – nekada polje *type of service* (IPv4) – koristi se od strane izvorišnih (*originated*) čvorova i/ili prosleđivačkih rutera kako bi prepoznali različite klase ili prioritete IPv6 paketa. Polje ima 8 bitova.

Flow label – novo polje – ovo polje se koristi kako bi se naznačilo da li je potrebno specijalno rukovanje od izvorišta do odredišta za sekvencu pateka. Polje ima 20 bitova. Može da se koristi za višeslojno prebacivanje (*multilayer switching*) i brzo prebacivanje paketa (*fast packet swiching*), a može da se koristi i za pružanje kvaliteta servisa - QoS.

Payload length – nekada polje *total length* – polje koje nam pokazuje ukupnu dužinu IPv6 podataka u paketu. Polje ima 16 bitova.

Next header – nekada polje *protocol* – polje koje identifikuje sledeći enkapsulirani protokol (TCP,UDP). Polje ima 8 bitova.

Hop limit – nekada polje *TTL* – polje koje označava životni vek paketa. Kada paket prođe ruter, vrednost upisana u ovom polju se smanjuje. Kada vrednost dođe do nule paket se odbacuje. Polje ima 8 bitova.

Source address – IPv6 izvorišna adresa. Polje ima 16 bajtova = 128 bita

Destination address – IPv6 odredišna adresa . Polje ima 16 bajtova =128 bita.

Extension headers – Nakon prvog obaveznog zaglavlja, može slediti određeni broj dodatnih zaglavlja pre enkapsuliranih podataka. Oni su kreirani kako bi se postigla efikasnost i fleksibilnost pri kreiranju IPv6 datagrama. Polja koja su potrebna u specijane svrhe se mogu smestiti u ova dodatna zaglavlja. Ovo omogućava da veličina glavnog zaglavlja ostane mala i da sadrži samo ona polja koja su zaista bitna u tom trenutku. Broj ovih polja nije fiksna, ali ukoliko ih postoji više, definisan je tačno utvrđeni redosled i zbog toga moraju biti procesuirana u striktno definisanom redosledu u kome se pojavljuju u paketu.

2.3. Predstavljanje IPv6 adresa

Adresiranje u IPv6 se dosta razlikuje od načina adresiranja koje se koristilo u IPv4. U IPv4 adresiranju, IP adrese su se prikazivale razdvajane tačkama i pisane u decimalnom prikazu (npr. 192.186.100.1), dok se adrese u novom IPv6 protokolu 128 bita duge i prezentuju se kao serija od osam šesnaestobitnih polja odvojenih dvotačkama (heksadecimalni prikaz). Postoje tri forme za prezentovanje IPv6 adrese:

Prva forma je predstavljanje u obliku heksadecimalne vrednosti kao osam šesnaestobitnih delova 128 bitne adrese:

Primer – 200A:1234:00CD:0000:0000:005C:7F3C:E34B

Druga forma je slična prethodnoj, međutim postoje olakšice pri zapisivanju, npr. ukoliko su uzastopna polja od 16 bitova prezentovana nulom, možemo ih preskočiti i predstaviti samo kao dve dvotačke:

Primer – 200A:1234:00CD:0000:0000:005C:7F3C:E34B

Olakšano predstavljanje – 200A:1234:CD::5C:7F3C:E34B

Treća forma se koristi pri mešovitoj arhitekturi IPv4 i IPv6 adresa. Pri ovom predstavljanju bitovi u gornjem delu predstavljaju IPv6 adresu, a bitovi u donjem delu predstavljaju IPv4 adresu.

Primer - 0:0:0:0:0:0:10.1.2.3

Sraćeni oblik - ::10.11.3.123

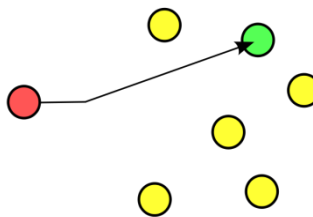
Za prefiksnu reprezentaciju, slična je *CIDR (Classless Inter Domain Routing namenjen alociranju IP adresa i rutiranju Internet protokol paketa)* u IPv4 pa čak je i zapisana u istoj formi. Predstavljanje kod IPv4 je bila u obliku *IPadresa/prefiksna dužina*, npr. 192.186.100.0/24. Kod IPv6 je kao što je rečeno zadržana isto predstavljanje:

primer – 200A:1234:00CD::/48 gde je 48 prefiks.

2.4. IPv6 tipovi adresa

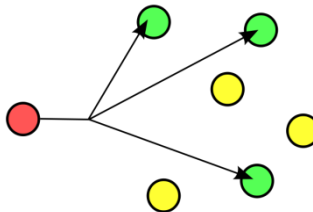
IPv6 definiše tipove adresa prema bitovima koji se nalaze na početku tzv. *leading bits*. Na osnovu tih bitova postoje sledeće podele adresa:

- *Unicast* (slika 5) – Slanje podataka jednom uređaju na mreži. Ova adresa se odnosi na jedan interfejs. IPv6 ima više tipova *unicast* adresa, a to su globalna (*global*) i IPv4 mapirana (*IPv4 mapped*). *Unicast* adrese su hijerarhijski uređene i sastoje se od:
 - Javna topologija (*Public topology*) – koju dobijamo od *ISP*-a ili *RIR*-a (*Regional Internet Registry*)
 - Topologija sajta (*Site topology*) – koja je ekvivalentna IPv4 privatnom opsegu adresa.
 - *Interface ID* – identifikuje interfejs specifičnog čvora.



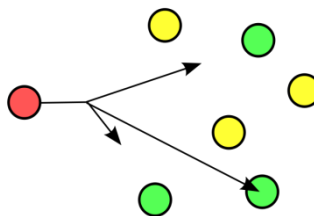
Slika 5: *Unicast* primer

- *Multicast* (slika 6) – identifikuje grupu interfejsa obično različitih čvorova. *Multicast* adrese se koriste za slanje informacija ili servisa svim interfejsima koji pripadaju određenoj grupi. Na primer jedna *multicast* adresa za komunikaciju sa svim IPv6 čvorovima na jednom linku. *Multicast* je dosta efikasniji kroz mrežu, zbog toga što šalje pakete samo jednom ruteru, koji dalje prepoznaje da je paket namenjen udaljenoj lokaciji i određenoj grupi interfejsa kojoj prosleđuje pakete. Nema pravljenja kopija paketa tako da se štedi propusni opseg, a i koristi se veći adresni opseg nego u slučaju IPv4 protokola. Prepoznaju se po 8 bitnom prefiksu 11111111 ili FF::/8.



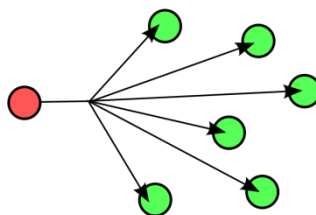
Slika 6: *Multicast* primer

- *Anycast* (slika 7) – mrežna adresna i ruting šema u kojoj se identifikuje lista uređaja ili čvorova. Paket poslat ovoj adresi je isporučen jednom interfejsu, najbližem interfejsu. Najbliži interfejs je definisan kao onaj koji ima najmanju udaljenost pri rutiranju (npr. najmanje skokova).



Slika 7: *Anycast* primer

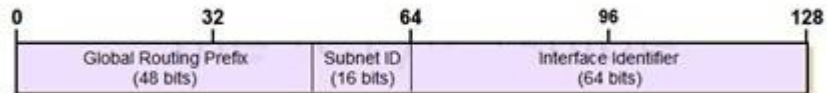
Broadcast (slika 8) – tip adrese koji je postojao u IPv4 izbačen je iz IPv6 i zamenjen *multicastom* ili *anycastom*, jedan od razloga je i taj što je rezultovao brojnim problemima, jedan od njih je npr. *broadcast storm*.



Slika 8: *Broadcast* primer

2.5. Specifični tipovi adresa

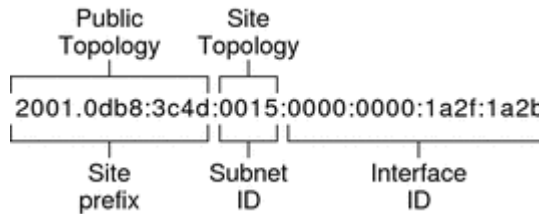
Postoje različiti tipovi adresa, koji će se koristiti za različite vidove komunikacije. Adresa je podeljena u tri dela (slika 8) i svaki tip adrese je konfigurisan specifično kako bi omogućio komunikaciju uređaja na različitom nivou (lokalno ili preko interneta)



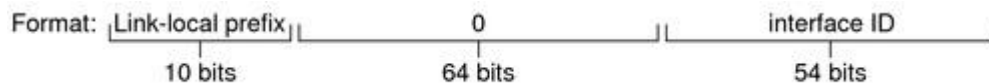
Slika 8: Globalna *unicast* adresa

- Globalna *unicast* adresa – javna rutabilna adresa koja može da se koristi na Internetu ili u bilo kom javnom domenu koji je povezan sa jednim čvorom. Pretpostavlja se da će se ovaj tip najviše koristiti za Internet saobraćaj i zbog toga je alociran najveći deo adresnog prostora. Ovo je *unicast* adresa koja je globalno dodeljena, predstavljena je sa 3 binim prefiksom 001 i prosleđuje se naviše tj. prema ISP. Globalna *unicast* adresa je podeljena u tri dela (slika 8).

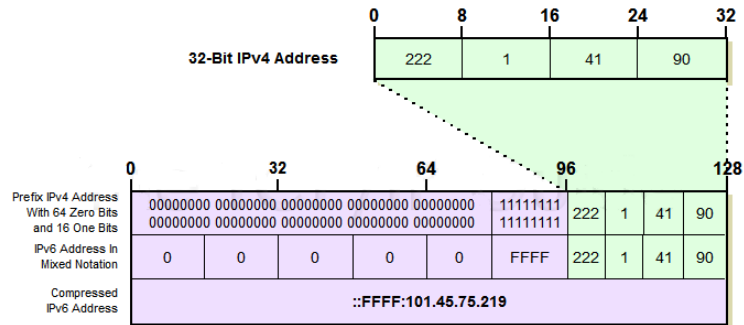
Ovaj tip adrese se sastoji od globalnog prefiksa rutiranja koji je mrežni identifikator (*network ID*) ili prefiksa adrese (001) koji se koristi za rutiranje, identifikatora pod mreže (*subnet ID*) koji identifikuje pod mrežu unutar istog sajta i identifikatora interfejsa (*interface ID*) koji jedinstveno identifikuje određeni interfejs. Globalni prefix rutiranja i identifikator pod mreže predstavljaju dva osnovna nivoa na kojima adrese moraju da budu hijerarhijski konstruisane: globalne ili sajt specifične. Prefiks rutiranja se sastoji od određenog broja bitova koji se dalje može podeliti prema RIR ili ISP potrebama, kako bi opisali topografiju Interneta kao celine. Identifikator interfejsa je jedinstven unutar određenog prefiksa i pod mreže. Primer se može videti na slici 9.

Slika 9: Primer globalne *unicast* adrese

- Link lokalna *unicast* adresa – je *unicast* adresa za lokalnu komunikaciju između uređaja na istom linku. Ova adresa se koristi u toku autokonfiguracije, otkrivanja suseda, otkrivanja rutera. Format je sličan globalnoj *unicast* adresi (slika 8), sa tim što se za prefiks koristi 10 bitni prefiks 1111111010 ili FE80::/10, identifikator pod mreže su sve nule, a interfejs identifikator je jedinstveni interfejs na linku (slika 10). Kada se komunicira sa link lokalnom adresom mora se označiti odlazni interfejs zato što je svaki interfejs konektovan na link FE80::/10

Slika 10: Link lokalna *unicast* adresa

- IPv4 mapirane IPv6 adrese – predstavljaju regularnu adresu IPv4 čvora koja je mapirana u IPv6 adresni prostor i koristi se na uređaju koji prepoznaje samo IPv4 protokol (slika 11) kako bi prezentovala IPv4 adresu čvora kao IPv6 adresu. Prefiks koji identifikuje je ::FFFF/96



Slika 11: IPv4 mapirana adresa

- Jedinstvena lokalna IPv6 *unicast* adresa – ovo je *unicast* adresa koja je globalno jedinstvena i namenjena je lokalnoj komunikaciji, obično unutar sajta. Od njih se ne očekuje da budu rutabilne na globalnom IPv6 Internetu. Ove adrese imaju prefiks FC00::/7 koji identifikuje lokalnu IPv6 *unicast* adresu.

2.6. Otkrivanje suseda i autokonfiguracija

Dosta je truda uloženo u dizajn IPv6 kako bi se vremenski skratio postupak konfiguracije i kako bi se uprostila automatska (*plug and play*) operacija. Protokol za otkrivanje suseda (*Neighbor discovery protocol – NDP*) dozvoljava IPv6 hostu i ruteru na linku da dinamički otkriju ili razmenjuju relevantne informacije lokalne za link. *NDP* izvršava devet specifičnih zadataka koji su podeljeni u tri grupe:

I grupa – funkcije otkrivanja host – ruter:

- Otkrivanje rutera – mogućnost otkrivanja suseda
- Otkrivanje prefiksa – IPv6 hostovi otkrivaju set adresnih prefiksa za link
- Otkrivanje parametara – IPv6 čvorovi uče parameter linka npr. najveća jedinica prenosa *MTU (Maximum Transmission Unit)*
- Adresna autokonfiguracija – IPv6 automatski alociraju IPv6 adresu za interfejs

II grupa – funkcije komunikacije host – host

- Rezolucija adrese – IPv6 čvorovi utvrđuju na link nivou adresu destinacije (kao protocol otkrivanja hardverske adrese *ARP (Address Resolution Protocol)* kod IPv4)
- Utvrđivanje sledećeg skoka – IPv6 hostovi se mogu instrukcijama podesiti da redirektuju saobraćaj boljem ruteru (ruteru sledećeg skoka)
- Otkrivanje nedostupnih suseda – IPv6 čvorovi mogu da detektuju nedostupnost suseda
- Detekcija duplih adresa – IPv6 čvorovi mogu da detektuju duple adrese

III grupa – funkcije redirekcije – ugrađena podrška redirekciji saobraćaja.

2.7. Adresna autokonfiguracija

Nova funkcija u IPv6 je autokonfiguracija IP adrese. Ovo je mehanizam za automatsko dodeljivanje IP adrese. Pri projektovanju IPv6 protokola se vodilo računa i o uređajima koji nisu samo personalni računari, a koji treba da budu povezani na mrežu (npr. terminali, mobilni). Projektovani mehanizmi omogućavaju automatsko (*plug and play*) umrežavanje uređaja smanjujući trošak administrativnog održavanja. Interfejs koji koristi IPv6 obično dobija dve adrese, a to su link lokalnu, koja se koristi za kontrolne funkcije i globalnu adresu, koja se koristi za klasičnu komunikaciju razmene informacija (*data communication*). Još jedna novina kod IPv6 je to što nema ograničenja na broj adresa koji može da se dodeli jednom interfejsu, za razliku od IPv4 koji je mogao da ima samo jednu IP adresu za jedan interfejs.

Sledeći mehanizmi se koriste za konfiguraciju IP adrese čvora:

- Ručno konfigurisanje – pojedinačno unošenje IP adrese za interfejs uređaja.
- *Stateless* – Obično se autokonfiguracija odnosi na ovaj mehanizam IPv6 konfiguracije IP adrese uređaja. Uređaj može automatski da generiše svoju jedinstvenu interfejs adresu na osnovu informacija iz mreže. Zove se *stateless* zato što sam uređaj i niko drugi ne upravlja njegovom IP adresom. IPv6 adresa se dobija

kombinacijom identifikatora interfejsa koji predstavlja interfejs na linku (npr. fizička adresa(MAC)) i prefiksa objavljenog od strane rutera na istom linku.

- *Statefull* – IPv6 uređaj može da komunicira sa *DHCP* serverom kako bi dobio sve konfiguracione informacije.

Prednosti *stateless* autokonfiguracije su :

- Ne zahteva se *DHCP* server dodele adresa
- Omogućavanje vrućeg uključivanja (*hot plug-in*) – interfejs se automatski konfiguriše po uključivanju mrežnog kabla u interfejs, bez dodatne administracije.
- Odgovara aplikacijama (npr. telekonferencije) koje ne zahtevaju posrednike koji mogu da uspore vezu, kao što su *DHCP* server ili proksi
- Smanjuje se ukupni trošak, koji je potreban ukoliko autokonfiguracije ne bi bilo.
- Odgovarajući je za bežične mreže

2.8. IPv6 rutiranje

IPv6 rutiranje je slično IPv4 rutiranju, ali kako bi ruteri rukovali podacima koji prolaze kroz mrežu, oni moraju znati kako da rukuju sa povećanim adresama u IPv6. Svi protokoli koji se danas koriste u IPv4 su prepravljeni kako bi podržali novi IPv6 protokol. Neki od protokola su: *RIPv6* – *RIPng* (*Routing Information Protocol new generation*) kome je dodata podrška IPv6, *IS-IS v6* (*Intermediate System to Intermediate System*) koji može da rukuje i sa IPv4 i sa IPv6, *OSPFv3* (*Open Shortest Path First – IPv6*) rukovanje samo sa IPv6, tamo gde se zahtevaju obe IP verzije protokola, ruter mora da ima i *OSPFv2* i *OSPFv3* omogućenu podršku, *BGP – IPv6*, odnosno *MP-BGP* (*Multi Protocol BGP*) koristi se za podršku IPv6 protokolu u rutiranju između domena.

2.9. IPv6 QoS

Nekoliko osobina je dodato u IPv6 kao što su nivoi osiguranog servisa, povećana sigurnost i poboljšana pouzdanost. Mreže u kojima je konfigurisan IPv4 obično rukuju svakim paketom na osnovu algoritma najboljeg pokušaja (*best effort*), bez obzira na to što su neki paketi osetljivi na vreme isporuke. Ovakve mreže (IPv4 sistemi), nemaju način da utvrde da li je podatak unutar paketa vremenski osetljiv ili ne (npr. video strimovanje (*video streaming*) ili fajl transfer izveštaji (*file transfer reports*)). U IPv6 protokolu, postoji način kako da se rukuje paketima, definisanjem različitih nivoa prioriteta. Neki od nivoa prioriteta su:

- Level 0 – nema naznačenog prioriteta
- Level 1 – pozadinski saobraćaj (*vesti*)
- Level 2 – tihi transfer podataka (*e-mail*)
- Level 3 – rezevisano
- Level 4 – transfer podataka uz učešće (*FTP*)
- Level 5 – rezervisano
- Level 6 – interaktivni saobraćaj (*Telnet*)
- Level 7 – kontrola saobraćaja (rutiranje, upravljanje mrežom)

2.10. IPv6 sigurnost

Kada je IPv4 dizajniran, Internet je bilo „prijateljsko okruženje”, sa ne toliko velikim brojem korisnika. Zbog toga same komponente sigurnosti nisu bile predefinisane pri projektovanju, već kasnije, kada se shvatilo da je sigurnost ipak jedna od bitnijih delova, dodavane. Sa godinama Internet se razvio i postao milionska mreža, koja je sa „prijateljske mreže” prešla na „neprijateljsku”, tj. u takvo okruženje u kome private informacije lako mogu da dođu u pogrešne ruke. Iako su predstavljene nove sigurnosne tehnike (*SSL*, *IPsec*, itd.) kako bi se prevazišle neke od najpoznatijih Internet sigurnosnih mana, sve to i dalje nije bilo dovoljno. Nažalost, uprkos svim skorašnjim unapređivanjima, samoj arhitekturi

Interneta i dalje nedostaju odgovarajuće sigurnosne mere tj. odgovarajući sigurnosni okvir. Svesni ograničenja trenutne infrastrukture Interneta, koja se temelji na IPv4 protokolu, *IETF* (internacionalno udruženje mrežnih dizajnera, operatera, proizvođača i istraživača zaduženih za evoluciju Internet arhitekture i održavanje, kako bi se osigurao besprekoran rad Interneta) je 1998. godine predložio novi protokol IPv6, koji rešava nekoliko pitanja koja utiču na mreže zasnovane na IPv4 protokolu. Neka od rešenja su:

2.10.1. IPv4 pitanja sigurnosti

Pre utvrđivanja sigurnosnih pitanja IPv6, moramo upoznati IPv4 sigurnost i njegove mane. Kao što je već rečeno IPv4 je projektovan, bez razmišljanja o sigurnosti. Za svoj s kraja na kraj (*end-to-end*) model IPv4 pretpostavlja da će se sigurnost konfigurisati od strane krajnjih čvorova. Na primer ukoliko aplikacija zahteva sigurnost, kao što je mejl (*e-mail*) enkripcija, to bi trebalo da bude odgovornost aplikacija na kraju čvorova, kako bi se pružila ova usluga. Današnji Internet nastavlja da bude transparentan i bez sigurnosnog okvira za neke pretnje kao što su:

- DOS napadi (*Denial of Service attack*) – u ovoj vrsti napada određene usluge su potpuno preplavljene sa ogromnim brojem nelegitimnih zahteva, koji rezultuje time da ciljani sistem zbog ogromnog broja zahteva koje ne stiže da obradi, postaje neupotrebljiv za legalne korisnike. Ova ranjivost proizilazi iz arhitektonske mane IPv4 protokola na višesmerne poplave (*broadcast flooding*) ili napadi „štrumfova” (*Smurf attack*).
- Distribucija malicioznog koda: virusi i crvi, mogu koristiti kompromitovani sistem kako bi zarazili udaljene sisteme. S obzirom da je kod IPv4 adresa mali adresni prostor, ovo može olakšati distribuciju zlonamernog koda.
- Napadi čovek u sredini (*man-in-the-middle attack*): kod IPv4 protokola postoji nedostatak odgovarajućeg mehanizma autentifikacije, koji mogu olakšati ovakve napade. Napadač iskorišćava ovu ozbiljnu manu, presreće pakete i kontroliše čitavu komunikaciju između dva uređaja, koji nisu svesni da su paketi presretnuti i da čitava komunikacija ustvari ide preko napadača.

- Napadi fragmentacije (*fragmentation attack*): ovaj napad iskorišćava način na koji operacioni sistem rukuje sa velikim IPv4 paketima. Primer ovog napada je ping smrti (*ping of death*), pri kome ciljani sistem je potopljen sa fragmentovanim *ICMP* ping paketima. Uz svaki paket, veličina rastavljenog ping paketa raste izvan limita IPv4 za veličinu paketa, samim tim se ciljani sistem ruši.
- Skeniranje portova i izviđački napadi (*Port scanning and reconnaissance*): pri ovim napadima ceo deo mreže je skeniran kako bi se našle potencijalne mete. Nažalost IPv4 adresni prostor je toliko mali da skeniranje kompletne C klase može potrajati nešto više od 4 minute.
- *ARP* trovanje i *ICMP* redirekcija: u IPv4 mrežama *ARP* je odgovoran za mapiranje IP adrese uređaja sa fizičkom ili MAC adresom. Ova informacija se nalazi na svakom uređaju u određenoj memorijskoj lokaciji – *ARP* tabeli. Svaki put se za nepoznati uređaj u mreži šalje *ARP* zahtev u mrežu. Zatim, nepoznati uređaj odgovara višesmernim emitovanjem svoje IP adrese sa odgovarajućim informacijama. *ARP* trovanja se dešavaju kad se falcifikovani *ARP* odgovor emituje sa netačnim informacijama mapiranja, koje mogu da nateraju da se paketi pošalju na lažnu adresu. Sličan pristup se koristi i kod napada *ICMP* redirekcije.

Međutim mnoge tehnike su razvijene kako bi se savladale neke IPv4 sigurnosna ograničenja. Na primer iako *NAT* i *NAPT* (*NAT port translation*) bili uvedeni kako bi se olakšalo ponovno korišćenje i očuvanje IPv4 adresnog prostora, ove tehnike takođe pružaju i određeni nivo zaštite protiv nekih spomenutih pretnji. Isto tako uvođenje *IPsec* omogućilo je korišćenje enkripcije komunikacije, iako je njegoa implementacija opcionalna i dalje nastavlja da bude isključiva odgovornost čvorova na kraju.

2.10.2. IPv6 sigurnosna poboljšanja

Iako IPv6 protokol nije nešto mnogo sigurniji od IPv4 protokola, čak je sigurnost IPv6 protokola samo nešto bolja od sigurnosti IPv4 protokola. Ipak iako sigurnost IPv6 protokola nije radikalno nova, neke od prednosti IPv6 protokola su:

2.10.2.1. Veći adresni prostor

Skeniranje portova je jedan od najpoznatijih izviđačkih tehnika u upotrebi danas. Ovime se kao što je rečeno mogu iskoristiti ranjivosti koje se otkriju skeniranjem. U IPv4 mreži skeniranje je relativno jednostavan zadatak. Većina segmenata su C klase sa 8 bitova dodeljenih za adresiranje hostova. Skeniranje tipične IPv4 podmreže po stopi jedan host po sekundi, prevodi u:

$$2^8 \text{ hostova} \times \frac{1 \text{ sekunda}}{1 \text{ hostu}} \times \frac{1 \text{ minut}}{60 \text{ sekundi}} = 4.267 \text{ minuta}$$

U IPv6 mrežama pogled je znatno drugačiji. IPv6 mreža koristi 64 bita za dodeljenih za adresiranje hostova, shodno tome tipičnoj IPv6 je potrebno:

$$2^{64} \text{ hostova} \times \frac{1 \text{ sekunda}}{1 \text{ hostu}} \times \frac{1 \text{ godina}}{3.1536.000 \text{ sekundi}} = 584.942.417.355 \text{ godina}$$

Iz ovoga proizilazi da je skeniranje ovakvih mreža skoro nemogući zadatak, ali nije baš u potpunosti nemoguće.

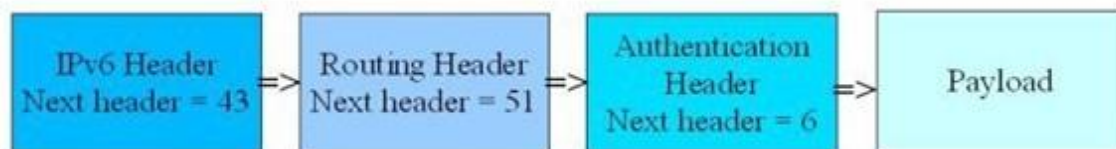
2.10.2.2. IPsec

Kao što je već pomenuto IPv4 protokol nudi, kao opciju, *IPsec*. Nasuprot ovome kod IPv6 protokola, RFC4301 obavezuje IPv6 protokol da koristi *IPsec* na svim čvorovima. *IPsec* se sastoji od skupa kriptografskih protokola koji osiguravaju razmenu podataka i ključeva. *IPsec* koristi dva kriptografska protokola i to: autentifikaciono zaglavlje (*Authentication Header – AH*) i enkapsulaciju sigurnosnog tereta (*Encapsulation Security Payload – ESP*). *AH* nam omogućava autentifikaciju i integritet podataka, dok nam *ESP* osigurava sigurnost,

integritet i poverljivost podataka. U IPv6 protokolu oba kriptografska protokola, *AH* i *ESP* definišu se kao produžetak zaglavlja. Osim toga *IPsec* osigurava i treći protokol, a to je protokol za upravljanje pregovorima i razmenom ključeva poznat pod nazivom *IKE* (*Internet Key Exchange*). Ovaj protokol pruža inicijalnu funkcionalnost, koja je potrebna za uspostavljanje i pregovaranje u vezi sa sigurnosnim parametrima između krajnjih tačaka. Osim toga beleže se i čuvaju ove informacije kako bi garantovala da komunikacija nastavlja da bude sigurna sve do kraja.

2.10.2.3. Zaglavlje provere (*Authentication Header – AH*)

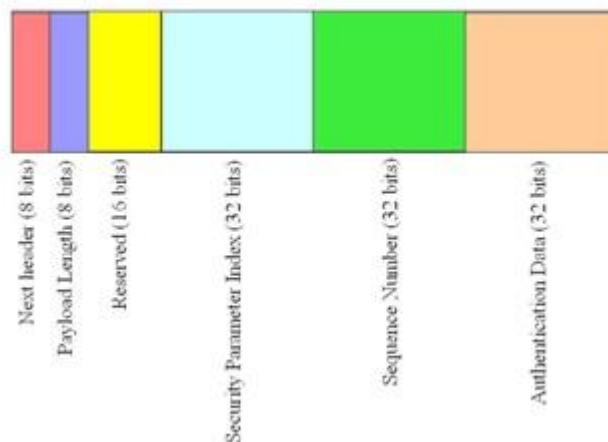
Kao što je već spomenuto, *AH* zaglavlje sprečava da IP paketi budu neovlašćeno promešani ili izmenjeni. U tipičnom IPv4 paketu, *AH* je deo tereta (*payload*). Sledeća slika (slika 12) prikazuje IPv4protokol sa *AH* kao delom tereta.



Slika 12: Redosled IPv6 ekstenzionih zaglavlja

Kada je *AH* implementiran, postojala je briga o tome kako da se integriše u novi IPv6 protokol. Problem je usmeren na činjenicu da IPv6 ekstenziona zaglavlja mogu da se promene u tranzitu. Kako bi se prevazišao ovaj problem, IPv6 *AH* je dizajniran sa fleksibilnošću na umu – protokol autentifikuje i proverava integritet samo za ona polja u zaglavlju IPv6 paketa, koja se ne menjaju u tranzitu. Takođe u IPv6 paketu, *AH* je inteligentno smešten na kraj lanca zaglavlja, ali ispred svih *ESP* zaglavlja za proširenje ili bilo kog protokola višeg nivoa kao što su *TCP* ili *UDP*.

AH zaglavlje takođe pruža i opcionalnu zaštitu od ponovljenih napada. Protokol koristi polje sekvencnog broja (slika 13) kao deo mehanizma kliznog prozora koji sprečava kašnjenja i maliciozna ponavljanja.



Slika 13: Prikaz tipičnog izgleda AH zaglavlja

2.10.2.4. *Enkapsulacija sigurnosnog tovara*

ESP pruža iste funkcionalnosti kao i *AH*, dok dodatno pruža tajnost. U *ESP* zaglavlju za proširenja indeks sigurnosnih parametara (*Security Parameter Index* – *SPI*) polje identifikuje koju grupu sigurnosnih parametara pošiljalac koristi kako bi osigurao komunikaciju. *ESP* ne pruža isti nivo autentifikacije koji je dostupan sa *AH*. Dok *AH* autentifikuje celo IP zaglavlje (ustvari samo one koji se ne menjaju u toku prenosa), *ESP* autentifikuje samo prateće informacije. *ESP* pruža integritet podataka, implementirajući vrednost provere integriteta (*Integrity Check Value* – *ICV*) koji je deo autentifikacionog polja. *ICV* koristi heš funkciju sa *SHA-1* i *MD5* kao podrazumevanom.

2.10.2.5. Transport i tunel modovi

U IP mrežama IPsec pruža dva moda za osiguranje saobraćaja. Transport mod se koristi za pružanje sigurne komunikacije sa kraja na kraj, osiguravajući samo teret paketa. Tunel mod štiti ceo IPv4 paket. Međutim, u IPv6 nema potrebe da ovim modom, zbog toga što oba i AH i ESP protokol pružaju dovoljno funkcionalnosti za osiguravanje IPv6 saobraćaja.

2.10.2.6. Pregovaranje protokola i upravljanje razmenom ključeva

Kao dodatak AH i ESP, IPsec dodatno određuje funkcionalnost za pregovaranje protokola i upravljanje razmenom ključeva. IPsec mogućnosti enkripcije zavise o mogućnošću pregovaranja i razmene ključeva između partija. Kako bi se ovaj zadatak obavio IPsec određuje IKE protokol, koji pruža sledeće mogućnosti:

- pregovaranje oko protokola koji treba da se koristi
- laka razmena ključeva, uključujući i njihovu periodičnu razmenu
- zapisivanje svih dogovora

2.10.3. IPv6 pitanja sigurnosti

Sa strane sigurnosti, IPv6 protokol predstavlja značajan napredak u odnosu na stariji IPv4 protokol. Međutim, uprkos mnogim unapređenjima, IPv6 nastavlja da bude ranjiv. Ovde ćemo opisati neke od oblasti IPv6 gde sigurnost nastavlja da bude važna stavka.

2.10.3.1. Pitanja u vezi sa dvostrukom konfiguracijom

Internet nastavlja da bude IPv4 zasnovan. Međutim, možemo da očekujemo da će se ovo sa vremenom menjati i da će sve više mreža prelaziti na novi protokol. Nažalost ovaj proces će prilično trajati, ali su se već razvile tehnike za upotrebu obe vrste protokola. Te tehnike

su npr. *6to4* dvostruko konfigurisanje. Ovo je sjajno prelazno rešenje, ali je povećava potencijalnu ranjivost mreže. Ovo se dešava kao rezultat postojanja dve infrastrukture sa specifičnim sigurnosnim problemima. Međutim, najviše problema koje mogu rezultirati nesigurnom mrežom, nisu stvar novog IPv6 protokola, već loše konfiguracije sistema.

2.10.3.2. Pitanja manipulacije zaglavlja

Korišćenje zaglavlja ekstenzije i IPsec mogu nas zaštititi od nekih zajedničkih napada zasnovanih na manipulaciji zaglavlja. Činjenica da *EH* (*Encapsulation Header*) mora da bude obrađena od strane svih čvorova može biti korišćena da se preplave određeni čvorovi npr. *firewall*. Podvala (*spoofing*) nastavlja da bude moga pretnja u IPv6 mrežama, ali samo od strane čvorova u istom segmentu mreže. Isto se odnosi i na *6to4* mreže.

2.10.3.3. Pitanje poplava (flooding)

Skeniranje za važeće adrese hostova i usluge znatno je otežano u IPv6. Nove usluge, kao što su multicast adrese i dalje nastavljaju da budu izvor problema. *Smurf* napadi su i dalje mogući u multicast saobraćaju.

2.10.3.4. Mobilnost

Ovo je nova osobina IPv6, koja nije postojala kod prethodnika. Mobilnost je veoma kompleksna funkcija koja povećava sigurnost. Ona koristi dve adrese stvarnu i mobilnu adresu. Stvarna adresa je IPv6 adresa sadržana u ekstenzionom zaglavlju (*EH*), dok je druga privremena adresa sadržana u IP zaglavlju. Zbog karakteristika ovih mreža, privremena komponenta može biti izložena napadima zavaravanja na kućnom agentu. Mobilnost zahteva potpuno nove mere zaštite, koju mrežni administratori moraju dobro poznavati.

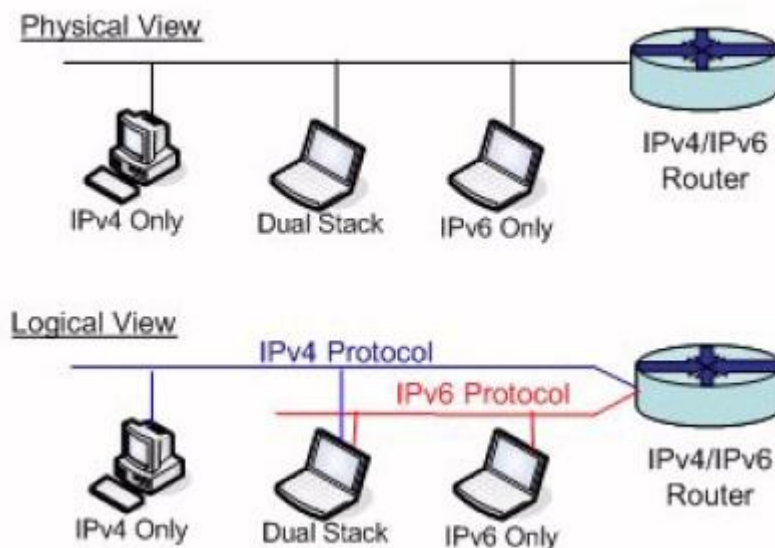
2.11. IPv4/IPv6 tranzicioni mehanizmi

Trenutno Internet se sastoji od prirodnih (native) IPv4, prirodnih IPv6 i IPv4/IPv6 dvostrukih mreža. Nažalost IPv4 i IPv6 nisu kompatibilni protokoli. Kada su oba protokola prisutna, a Internet korisnici žele da se konektuju bez ikakvih restrikcija, potrebni su tranzicioni mehanizmi. Predstavljeni su mnogi mehanizmi kako bi omogućili prelazak bez ometanja sa IPv4 na IPv6 mreže. Ranije su IPv6 mreže bile implementirane kao zasebna ostrva, međutim sada se ta ostrva povezuju u moru IPv4 mreža. Razne vrste tranzicija su podeljejne u tri velike grupe:

- dvostruko konfigurisanje (dual stack)
- tunelovanje (tunneling)
- translacija (translation)

2.11.1. Dvostruka konfiguracija (Dual stack)

Mehanizam koji uključuje oba protokola koji rade paralelno i dozvoljavaju da mrežni čvorovi komuniciraju ili uz pomoć IPv4 ili uz pomoć IPv6 protokola. Ova konfiguracija nam omogućava transport IPv4 i IPv6 paketa. Odluka koji će se protokol koristiti se donosi na osnovu polja verzije i destinacionog tipa IP adresa. Iako je ovo najrasportanjenije tranziciono rešenje, ono omogućava komunikaciju samo sličnih mrežnih čvorova (IPv4 sa IPv4 i IPv6 sa IPv6 čvorovima). Sledeća slika (slika 14) nam je opisana fizička i logička struktura IPv4 i IPv6 uređaja, ali i onih koji su dvostruko konfigurisani sani.



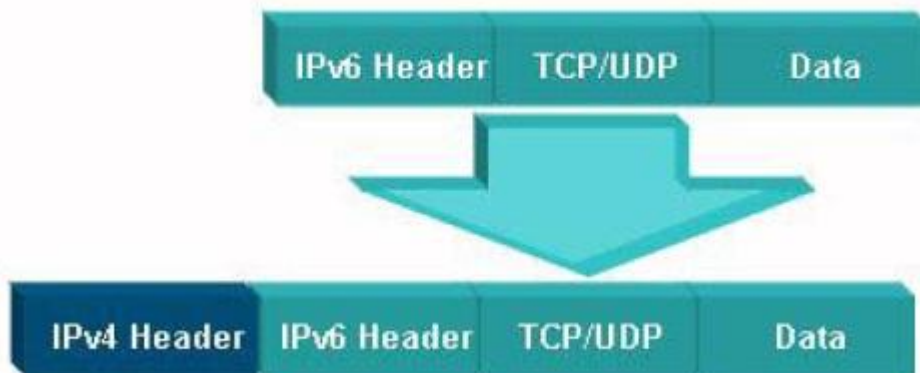
Slika 14: Prikaz fizičke i logičke structure IPv4 i IPv6 uređaja

2.11.2. IPv4-IPv6 tunelovanje

Ovo rešenje nam omogućava da se nekompatibilne mreže premoste (*bridging*). Iako postoji dosta varijanti tunelovanja, oni su generalno podeljeni u dve grupe

- centralizovana – predefinisana, konfigurisana ručno od strane administratora
- automatska – konfigurišu se automatski (*on the fly*), na osnovu informacija sadržanih u IPv6 paketu (npr. izvorišna i odredišna IP adresa). Postoje sledeće automatske tehnike tunelovanja: 6to4, ISATAP, 6over4, Tunnel brokers, Teredo, dual stack transition mehanizam (DSTM).

Uopšteno, IPv6 paketima se dodaje IPv4 zaglavlje, što je na slici 15 predstavljeno.

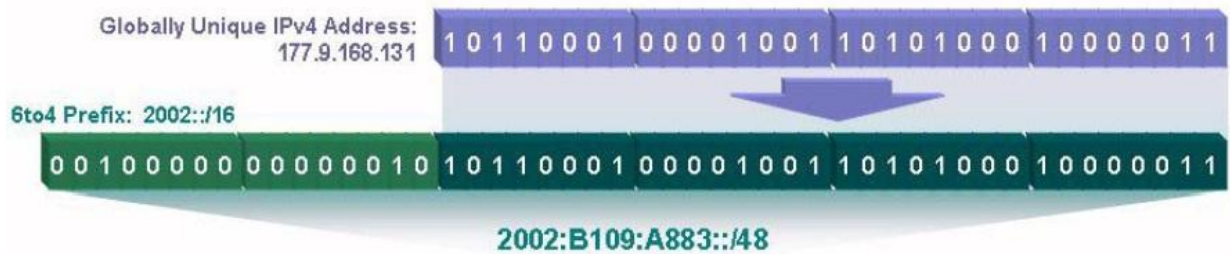


Slika 15: Dodavanje IPv4 zaglavlja

Uz pomoć ovoga moguće je da se IPv6 paket šalje preko IPv4 mreže. Na ruteru ili hostu izvorišta primenjuje se enkapsulacija. Dodaje se IPv4 destinaciona adresa, IPv4 izvorišna adresa, a u polje protokol u IPv4 zaglavlju, ubacuje se vrednost 41 koja nam govori da se radi o enkapsuliranom IPv6 paketu. Na izlaznom hostu izvršava se dekapulacija kako bi se skinulo IPv4 zaglavlje, a potom se paket rutira ka destinacionoj IPv6 adresi.

2.11.2.1. 6to4 (IPv6 to IPv4)

Tehnika tunelovanja koja se oslanja na određeni IPv6 format adrese kako bi se identifikovali 6to4 paketi, a potom tunelovani pravilno. Adresa se sastoji od prefiksa 2002::/16 iza koga sledi globalna jedinstvena IPv4 adresa za ciljani destinacioni sajt. Ta adresa se pretvara u heksadecimalni broj. Ova konkatencija formira /48 prefix koji je prikazan na slici 16.



Slika 16: Formirani /48 bitni prefiks

Jedinstvena IPv4 adresa predstavlja IPv4 adresu *6to4* rutera koji završava (*terminate*) *6to4* tunel. 48 bitni *6to4* prefiks služi kao globalni prefiks rutiranja i identifikator pod mreže se može dodati kao sledećih 16 bitova, iza kojih sledi identifikator interfejsa kako bi se potpuno definisala IPv6 adresa.

2.11.2.2. ISATAP (*Intra Site Automatic Tunneling Addressing Protocol*)

Ovo je eksperimentalni protokol koji pruža automatsko IPv6 preko IPv4 (IPv6-over-IPv4) tunelovanje za host – host, ruter – ruter i host – ruter konfiguracije. On se formira tako što se koristi IPv4 adresa za definisanje identifikatora interfejsa. IPv4 adresa se ne pretvara u heksadecimalni broj. Identifikator interfejsa se sastoji od ::5EFE:a.b.c.d, gde je a.b.c.d IPv4 notacija (npr. IPv4 = 177.9.168.133 što se potom pretvara u ::5EFE:177.9.168.133).

2.11.2.3. IPv6 preko IPv4 (*6over4*)

IPv4 multikast se zahteva i smatra se virtuelnim ethernetom od strane *6over4* šeme. Zbog perspektive virtuelnog eterneta, IPv6 se formira koristeći link lokalni opseg (FE80::/10 prefiks). IPv4 adresa hosta se sastoji od *6over4* dela interfejs identifikatora svoje IPv6 adrese (npr. host sa IPv4 adresom 192.223.16.85 bi formirao interfejs identifikator ::C0DF:1055, pa tako bi *6over4* adresa bila FE80::C0DF:1055). IPv6 paketi se tuneluju u

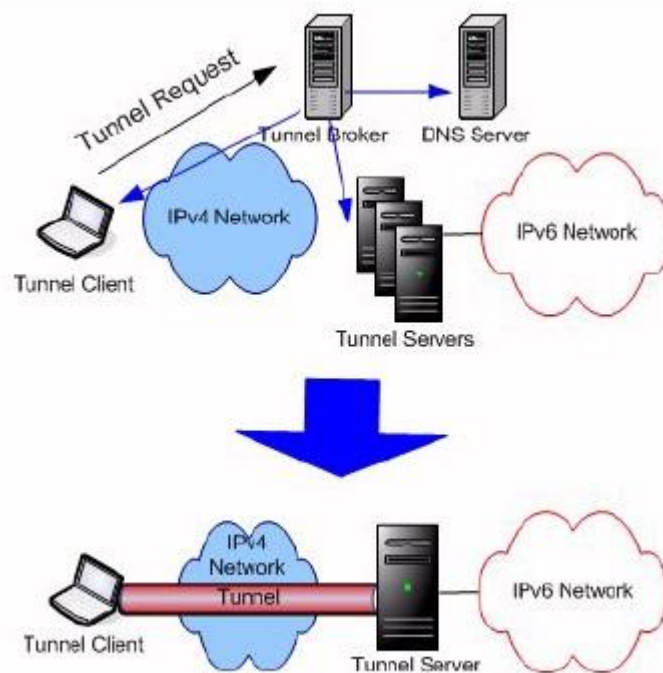
IPv4 zaglavlju koristeći odgovarajuće multikast adrese. Svi pripadnici multikast grupe primaju tunelovane pakete, a potom primalac skida IPv4 zaglavlje i obrađuje IPv6 paket. Dovoljno je da jedan ruter podržava IPv6 i *6over4* tehniku kako bi mogao da se ponaša kao kraj tunela. *6over4* podržava IPv6 multikast i unikast tako da host može izvršavati IPv6 otkrivanje rutera i suseda kako bi locirao IPv6 rutere.

2.11.2.4. Agenti tunela (Tunnel Brokers)

Ovo je još jedna od tehnika automatskog tunelovanja preko IPv4 mreža. Broker upravlja tunelima preko dvostruko konfigurisanih klijenata i broker servera, koji su konektovani na željenu IPv6 mrežu. Klijenti koji žele da se konektuju na IPv6 mrežu se obraćaju brokerima za pomoć (usmereni od strane DNS), klijenti se autentifikuju brokeru, broker može tražiti i sertifikate za autorizaciju, klijent takođe daje i IPv4 adresu kraja tunela zajedno sa željenim FQDN (*Fully Qualified Domain Name*) klijenta, zahtevanu IPv6 adresu kao i podatak da li je to host ili ruter. Kada se klijent autorizuje broker izvršava sledeće korake:

- Dodeljuje i konfiguriše server tunela, a potom ga obaveštava o klijentu.
- Dodeljuje IPv6 adresu ili prefiks klijentu zasnovanu na zahtevanom broju adrese i tipu klijenta (ruter ili host)
- Registruje klijentov FQDN u DNS
- Informiše klijenta o dodeljenom serveru tunela i povezuje tunel sa IPv6 parametrima uključujući adresu/prefiks i DNS ime.

Iz perspektive klijenta uspostavljanje tunela slično je uspostavljanju VPN konekcije. Sledeća slika (slika 17) opisuje prethodno objašnjene korake.

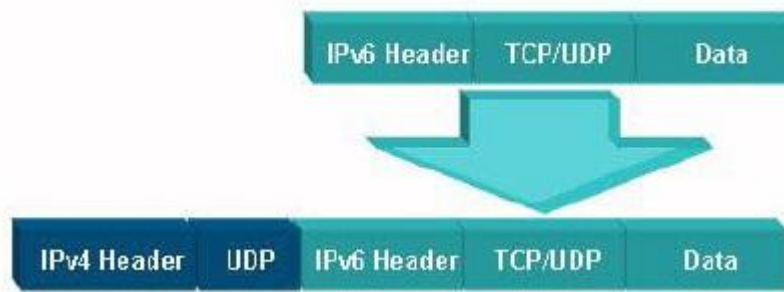


Slika 17: Tunnel Brokers

2.11.2.5. Teredo

Tunelovanje kroz zaštitni zid (*firewall*), koji ima i NAT iskonfigurisan može biti veliki izazov, ako ne i nemoguće. To je zbog toga što NAT neće dozvoliti propuštanje IPv4 protokola sa poljem postavljenim na 41. Teredo omogućava propuštanje IPv6 tunelovanja preko UDP IPv4 za host – host automatsko tunelovanje. Teredo dodaje dodatno UDP zaglavlje (slika 18) kako bi omogućio prolazak kroz NAT. S obzirom da stvara preterano opterećenje (*overhead*) opisan je kao poslednje sredstvo. Teredo zahteva sledeće elemente:

- Teredo klijent
- Teredo server
- Teredo relej (*relay*)



Slika 18: Teredo dodaje UDP zaglavlje

Teredo proces tunelovanja počinje kada Teredo klijent izvrši kvalifikacionu proceduru kako bi otkrio Teredo relej (*relay*), koji je najbliži nameravanoj destinaciji IPv6 hosta i identifikuje tip NAT uređaja. Teredo klijent mora biti prekonfigurisan sa IPv4 adresom Teredo servera koga će koristiti. Utvrđivanje najbližeg Teredo releja se izvršava putem pinga. Mogu se izvršiti dodatni koraci kako bi se identifikovao tip NAT uređaja. Generalno postoje sledeći tipovi NAT uređaja:

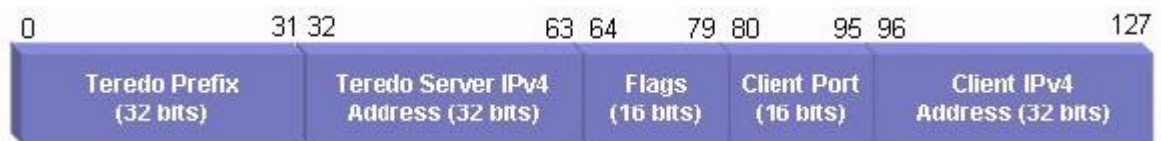
- Potpuni generator (*Full cone*) – svi IP paketi iz iste interne IP adrese i porta se mapiraju od strane NAT uređaja u odgovarajuću eksternu adresu i port. Eksterni hostovi mogu komunicirati sa hostovima prostim slanjem mapirane eksterne adrese i porta.
- Ograničen generator (*Restricted cone*) – svi IP paketi iz iste interne IP adrese i porta se mapiraju od strane NAT uređaja u odgovarajuću eksternu adresu i port. Eksterni hostovi mogu komunicirati sa internim hostovima samo ukoliko su prethodno unutrašnji hostovi njima poslali paket.
- Portom ograničeni generator (*Port restricted cone*) – svi IP paketi iz iste interne IP adrese i porta se mapiraju od strane NAT uređaja u odgovarajuću eksternu adresu i port. Eksterni hostovi mogu komunicirati sa internim hostovima, samo ukoliko su im prethodno interni hostovi poslali paket i to koristeći eksternu host adresu i port.
- Simetrični (*Symmetric*) – svi IP paketi iz iste interne IP adrese i porta za datu eksternu IP adresu i port su mapirani u određenu adresu i port. Paketi koji potiču sa iste interne IP adrese i porta, ali su namenjeni drugačijoj destinacionoj adresi i portu

rezultuju različitom eksternom mapiranjem. Eksterni hostovi mogu da komuniciraju sa internim hostovima, samo ako su prethodno interni hostovi nešto njima slali koristeći eksternu IP adresu i port.

Identifikacija NAT uređaja kao Full cone ne zahteva dalje kvalifikacije, ali restricted zahtevaju. Teredo ne podržava simetrične NAT uređaje. Kako bi završio mapiranje unutar NAT uređaja paket balon (bubble packet) se šalje hostu. Ovaj paket ima samo IPv6 zaglavlje i nema payload. Ovaj metod pomaže NAT uređaju da kompletira mapiranje interne i eksterne IP adrese i porta za Port restricted cone scenario.

Generalno paket „balon” se šalje direktno od izvorišta Teredo klijenta do destinacionog hosta. Problem je ukoliko je host iza zaštitnog zida (*firewall*), tada ovaj paket može da bude odbačen. Opet kao rešenje je da umesto klijenta paket šalje Teredo server, pošto je sa druge strane sigurno klijent, pa samim tim već postoje određena mapiranja koja osiguravaju da će paket od servera doći do klijenta. Nakon što klijent primi paket od servera, on će odgovoriti na poruku, ovime završavajući mapiranje na NAT uređaju.

Format Teredo IPv6 adrese je prikazana na slici 19.

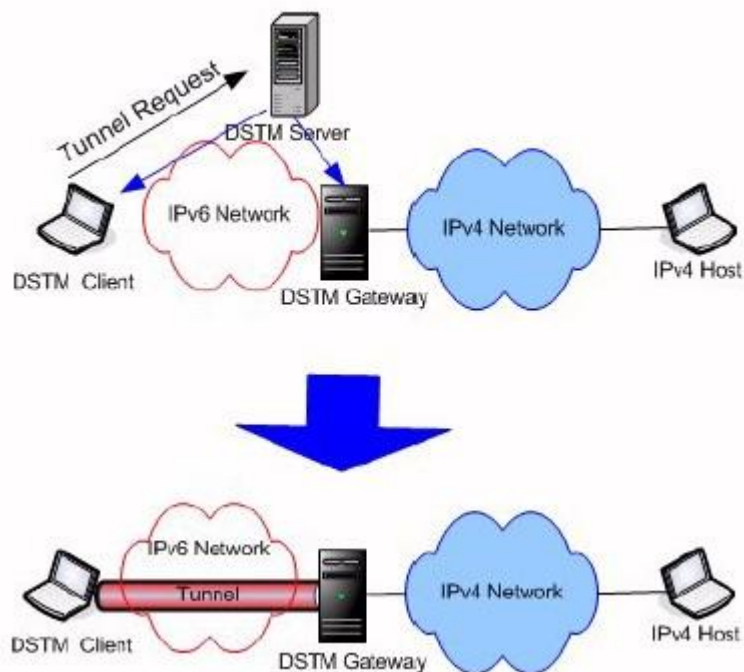


Slika 19: Format Teredo IPv6 adrese

Prefiks je predefinisani IPv6 prefiks 2001::/32, potom sledi Teredo server IPv4 adresa, pa potom „zastavice” koje nam govore o kom NAT uređaju se radi (*full cone*=0x8000, *restricted* ili *port restricted*=0x0000). Polja porta klijenta i IPv4 adrese klijenta se prezentuju maskiranjem vrednosti i to tako što se obrću vrednosti svakog bita.

2.11.2.6. DSTM (*Dual stack transition Mechanism*)

Ovo je metoda automatskog tunelovanja IPv4 paketa preko IPv6 mreže, sve do destinacione IPv4 mreže i hosta. Host koji se nalazi u IPv6 mreži, a želi da komunicira sa IPv4 uređajem zahteva postojanje DSTM klijenta. Nakon rezolucije imena, klijent inicijalizuje DSTM proces koji je vrlo sličan sistemu brokera, kao što je prikazano na sledećoj slici (slika 20).



Slika 20: DSTM proces

Proces počinje kada DSTM klijent kontaktira DSTM server kako bi dobio IPv4 adresu, ali i adresu DSTM izlaza (*default gateway*). Adresa se koristi kao izvorišna adresa. Ova adresa se enkapsulira u IPv6 zaglavlje koristeći DSTM klijent izvorišnu IPv6 adresu i DSTM adresu izlaza, IPv6 adresu, kao adresu destinacije. Sledeće polje u IPv6 zaglavlju indicira enkapsulirani IPv4 paket sa ovim *4over6* pristupom tunelovanja.

Varijanta DSTM podržava VPN zasnovani pristup od DSTM klijenta koji je van prirodne (*native*) mreže (npr. radnik koji radi od kuće). U ovom primeru pretpostavljajući da je

DSTM klijent dobio IPv6 adresu, ali ne i IPv4 adresu, on se može konektovati na DSTM server kako bi dobio IPv4 adresu. Ovaj pristup bi trebao da zahteva autentifikaciju kako bi se uspostavio VPN između DSTM klijenta i DSTM izlaza.

2.11.3. Translacija

Ova metoda koristi translaciju IPv4 u IPv6 na određenom nivou protokola, obično na mrežnom nivou, transportnom ili aplikacionom. Za razliku od tunelovanja, translacija pretvara pakete iz IPv4 u IPv6 i obrnuto.

2.11.3.1. SIIT

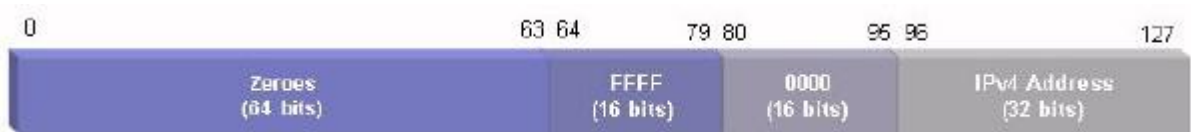
(Stateless IP/ICMP Translation algorithm)

SIIT pruža prevođenje IP zaglavlja paketa između IPv4 i IPv6. SIIT se nalazi na IPv6 uređaju i konvertuje odlazeće IPv6 zaglavlja paketa u IPv4 zaglavlja. Kako bi izvršavao ovaj zadatak, IPv6 uređaj mora da bude iskonfigurisan i sa IPv4 adresom (statički ili dinamički). SIIT algoritam prepoznaje pokušaj komunikacije između IPv6 i IPv4 uređaja ,kada je IPv6 adresa ustvari IPv4 mapirana adresa formatirana kao na slici 21.



Slika 21: Format IPv6 adrese

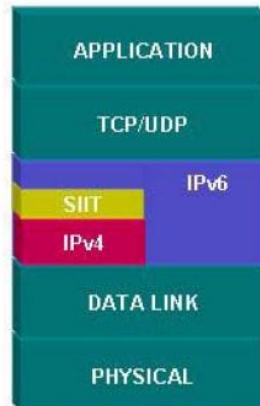
Mehanizam koji prevodi IPv4 u IPv6 adresu je pružen uz pomoć BIS (*bump-in-the-stack*) ili BIA (*bump-in-the-API*) tehnikama. SIIT algoritam obavlja translaciju zaglavlja kako bi doprineo transmisiji preko data link ili fizičkog nivoa. Izvorišna IP adresa koristi drugačiji format od IPv4 transliranog formata (slika 22)



Slika 22: Format izvorišne adrese

IPv4 mapirana adresa nije pogodna za izvorišnu adresu pri tunelovanju, zbog toga što bi to diskvalifikovalo komunikaciju kroz bilo koji tunel koji se nalazi između tunela. Korišćenje IPv4 transliranog formata prevazilazi potencijalno ograničenje.

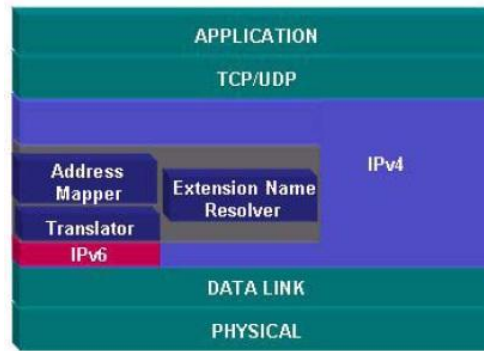
Primer SIIT steka (*stack*) je prikazan na slici 23, iako je obično uokviren sa BIA ili BIS.



Slika 23: Primer SIIT steka

2.11.3.2. BIS (*Bump in the Stack*)

Omogućava uređaju koji koristi IPv4 aplikaciju da komunicira preko IPv6 mreža. On osluškuje tok podataka između TCP/IPv4 modula i uređaja na nivou linka (npr. mrežne kartice) i prevodi IPv4 pakete u IPv6 pakete. Sledeća slika (slika 24) prikazuje BIS.



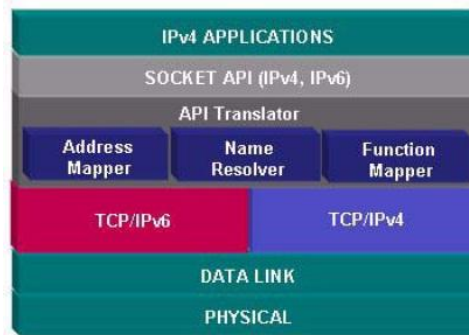
Slika 24: Prikaz BISa

Translator komponenta prevodi IPv4 zaglavlje u IPv6 zaglavlje prema *SIIT* algoritmu. Prošireni prevodilac imena (*Extension Name Resolver – ENR*) osluškuje DNS zahteve za A unos (*IPv4 host name*), a nakon što naiđe na takvu komunikaciju on pravi dodatni upit za oba A i AAAA (*IPv6 host name*) tipa rekorda za isto ime uređaja. Ukoliko se ne dobije odgovor za AAAA, onda se koristi IPv4 komunikacija, u suprotnom *ENR* uz pomoć komponente za mapiranje adresa (*Address Mapper*) povezuje vraćenu IPv4 adresu sa IPv6 adresom. Ukoliko se samo primi AAAA odgovor, onda komponenta za mapiranje adresa dodeljuje IPv4 adresu iz konfigurisanog opsega adresa.

2.11.3.3. BIA (Bump in the API)

Za razliku od modifikacije zaglavlja, koju izvršava BIS, BIA pristup prevodi između IPv4 i IPv6 *API (Application Programming Interface)*. BIA je implementirana između aplikacionog i transportnog (TCP/UDP) nivoa (slika 25) i sastoji se iz:

- API translator
- Komponente za mapiranje adresa (Address mapper)
- Prevodilac imena (Name resolver)
- Komponenta za mapiranje funkcija (Function Mapper)

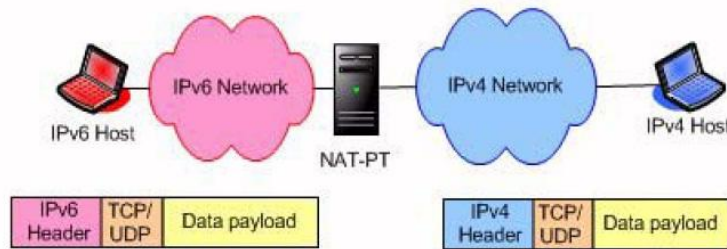


Slika 25: Komponente i implementacija BIA

Kada IPv4 aplikacija pošalje DNS upit, kako bi utvrdila destinacionu adresu, prevodilac imena presreće ovaj upit i kreira novi upit zahtevajući A i AAAA unos. DNS odgovor sa A unosom će pružiti odgovor sa datom IPv4 adresom. Odgovor sa samo AAAA unosom stimulira da prevodilac imena zahteva IPv4 adresu od komponente za mapiranje adresa, kako bi mapirao vraćenu IPv6 adresu. Na kraju prevodilac imena vraća zahtevani A unos aplikaciji koja je tražila taj podatak još na početku. Komponenta za mapiranje adresa održava ovo mapiranje IPv6 adresa sa onima dodeljenim iz internog adresnog opsega koji se sastoji nedodeljenog IPv4 adresnog prostora (0.0.0.0/24). Komponenta za mapiranje funkcija presreće API funkcionalne pozive i mapira IPv4 API pozive sa IPv6 pozivima.

2.11.3.4. NAT-PT (Network Address Translation with Protocol Translation)

Ovaj način translacije radi kao i svaki NAT, samo sa dodatkom translacije protokola. NAT uređaj se nalazi između IPv6 i IPv4 mreža. Ovaj uređaj ima IPv4 i IPv6 opsege adresa koje koristi kako bi osigurao pravilnu translaciju (slika 26).



Slika 26: NAT-PT

2.11.3.5. NAT-PT (Network Address Port Translation with Protocol Translation)

Ovaj vid translacije omogućava komunikaciju između IPv4 i IPv6 čvorova koristeći samo jednu IPv4 adresu. Dakle umesto da održava jedan na jedan mapiranje IPv6 i IPv4 jedinstvenih adresa kao što je slučaj kod NAT-PT, NAT-PT mapira svaku IPv6 adresu sa jednom IPv4 adresom sa jedinstvenim TCP ili UDP portom.

3. MPLS

MPLS (Multiprotocol Label Switching) je metod prosleđivanja paketa (frejmova) visokih performansi kroz mrežu. On omogućava da ruteri na obodu mreže obeležavaju pakete (frejmove), dok oni unutar iste mreže vrše samo prosleđivanje na osnovu oznaka (labela). *MPLS* integriše menadžment saobraćaja i performanse L2 nivoa sa skalabilnošću i fleksibilnošću rutiranja na 3 nivou. Zbog toga na njega se gleda kao na protocol L2.5 nivoa. Posebna prednost je kada se postavi u *ATM* mrežu kako bi ponudio skalabilnost IP preko *ATM* mreže (*IP – over – ATM*).

Konvencionalno rutiranje je zasnovano na razmeni informacija o dostupnosti mreže, kako paket putuje kroz mrežu, svaki ruter izvlači informacije koje su bitne za prosleđivanje iz L3 zaglavlja. Ove informacije se potom koriste za indeksiranje tabele rutiranja kako bi se utvrdio sledeći skok (*next hop*) za paket. Ovo se ponavlja na svakom ruteru u mreži. Pri svakom skoku u mreži, optimalno prosleđivanje paketa se mora ponovo utvrditi. Konvencionalno prosleđivanje IP paketa ima nekoliko ograničenja, kao što je ograničena

sposobnost za rad sa adresnim informacijama izvan destinacione adrese koja se nosi u paketu. Svi IP destinacionog prefiksa isti paketi se tretiraju slično, iz ovoga proizilaze mnogi problemi, kao što je teškoća sa upravljanjem saobraćaja u IP mrežama, takođe prosleđivanje IP paketa ne podnose lako dodatne adresne podatke npr. *VPN*.

Glavni koncept *MPLS*-a je dodavanje labele u svaki paket. Na osnovu ove labele vrši se prosleđivanje kroz mrežu.

Labela sumira esencijalne informacije u rutiranju paketa kao što su:

- Destinaciju
- Prethodnika
- *VPN* članstvo
- QoS informacije iz RSVP
- Rutu paketa izabranu od TE

Sa labelom celokupna analiza L3 zaglavlja se izvršava samo jednom i to na obodu mreže, prilikom dodavanja labele. Dalje unutar mreže vrši se samo prosleđivanje paketa na osnovu zadate labele.

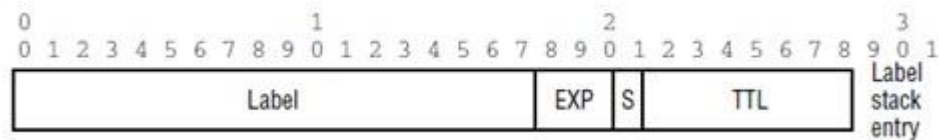
MPLS je tehnologija koja ubrzava i usmerava tok mrežnog saobraćaja i čini ga lakšim za upravljanje.

3.1. MPLS osobine i izgled

U početku pravljen kako bi se spojili IP i ATM zbog brojnih razlika, ali i potrebe da se spoje, a kasnije je postao standard za obeležavanje paketa.

3.1.1. Izgled MPLS labela (MPLS label stack)

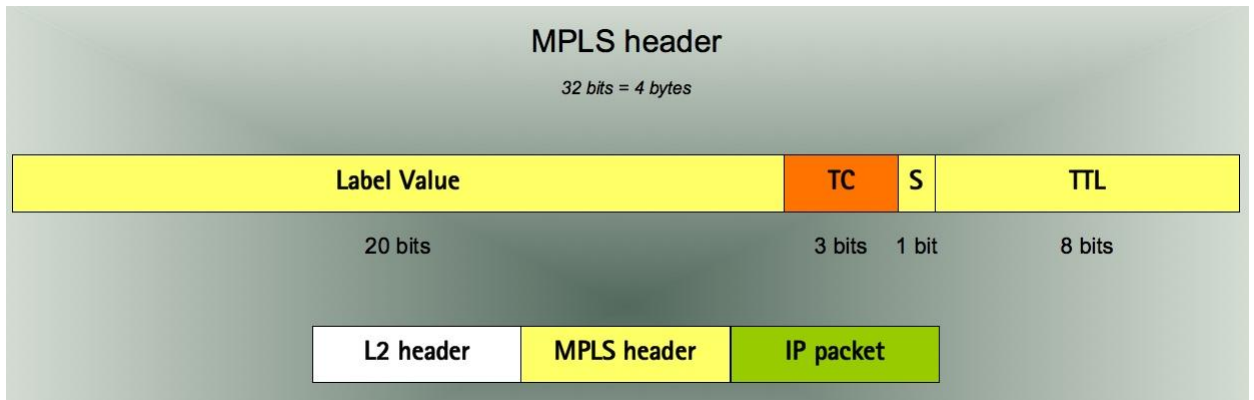
MPLS funkcioniše tako što stavlja prefikse na pakete u MPLS zaglavlje, koje sadrži jednu ili više „labela“. Ovo se zove *label stack*. Sastoji se od 32 bita, a svaki *label stack* unos sadrži 4 polja (okteta). Izgled labela je prikazan na slici 27, dok su polja pojedinačno u daljem tekstu pojašnjena.



Slika 27: Izgled MPLS labela

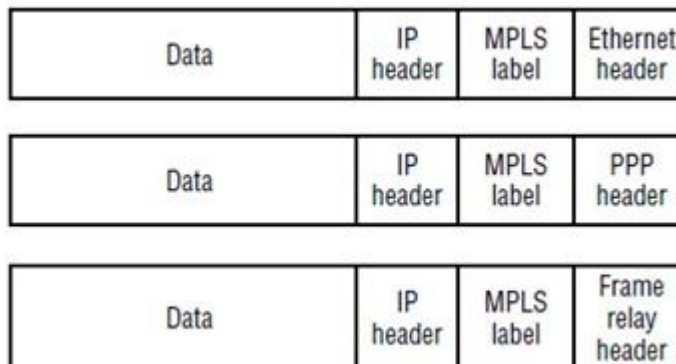
- *Label* – polje veličine 20 bita, koje dovodi do toga da je moguće imati preko 1 miliona labela. Ono nam opisuje putanju koju paket treba da sledi do destinacije.
- *Eksperimental* (exp) – polje veličine 3 bita koji služi za mapiranje standardnog IP *ToS* tipa servisa (*Type of Service*) u eksperimentalno polje MPLS *CoS* – klase servisa (*Class of Service*).
- *S* – Stack – labela se mogu nadovezivati jedna na drugu i ovo polje nam opisuje zadnju labelu u nizu. Veličine je jednog bita.
- *TTL* – polje koje opisuje životni vek MPLS paketa, ovo je kopirana vrednost IP TTL, zatim smanjena za 1 i upisana u ovo polje. Veličine je 8 bita. Kada TTL vrednost dođe do vrednosti 0, paket će biti odbačen.

Sada kad znamo kako izgleda MPLS možemo da opišemo i gde se on postavlja (slika 28).



Slika 28: Mesto MPLS zaglavlja u paketu

MPLS label stack se nalazi nakon *data link layer* zaglavlja (L2), ali i pre mrežnog zaglavlja (L3). Takođe prilikom enkapsulacije MPLS labela, tj. njena pozicija, ostaje ista, ali se koristi drugačija enkapsulacija (slika 29).



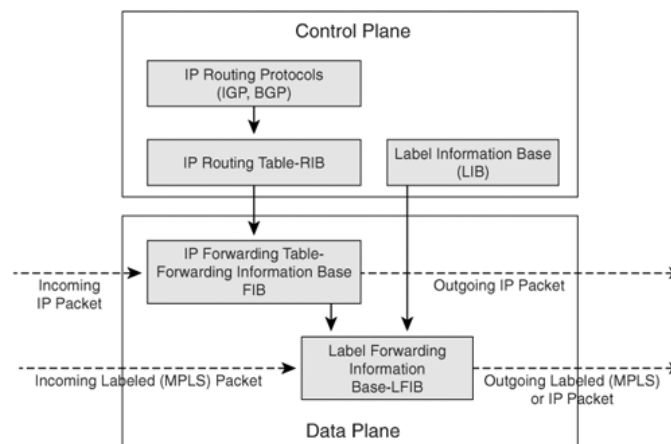
Slika 29: Primeri različitih enkapsulacija MPLS oznaka

3.2. MPLS komponente

Kako bi mogao da ispuni svoje zadatke, *MPLS* mora da bude u mogućnosti da prenosi različite protokole, da obeležava pakete i da ih prebacuje (*switch*) radije nego rutira, s obzirom da je ovaj proces mnogo brži od rutiranja.

CEF (*Cisco Express Forwarding*) je temelj na kome *MPLS* i njegovi servisi rade na ruteru. Zbog toga je *CEF* obavezan preduslov kako bi *MPLS* mogao uopšte da se implementira. *CEF* je CISCO mehanizam koji se koristi na CISCO ruterima i koji pojednostavljuje i povećava IPv4 performanse prosleđivanja na ruteru. *CEF* se sastoji iz dve logičke komponente (slika 30):

1. *Control Plane*
2. *Data Plane (Forwarding)*



Slika 30: Predstavljanje *Control* i *Data Plane*

3.2.1. Kontrolni deo i deo za rukovanje paketima (*Control /Data plane*)

Kontrolni deo (*Control plane*) se odnosi na pravljenje tabele rutiranja i *ARP* (*Address Resolution Protocol*) tabele, na osnovu kojih se pravi IP CEF i tabela CEF suseda. Ona sadrži podmreže, hostove itd. koji mogu vrlo lako i brzo da se pretražuju na osnovu indeksa. Indeks u tabeli suseda pruža adresu sledećeg skoka (*next hop*), interfejs i L3 *rewrite* informacije. Informacije koje smo dobili prosleđujemo logičkom delu za rukovanje paketima (*Data plane*), koji je ustvari hardverski prosleđivački deo.

Sve rute koje su naučene uz pomoć protokola za rutiranje se prosleđuju *FIB* (*Forwarding Information Base*) tabeli koja u delu za rukovanje paketima služi za prosleđivanje na osnovu IP adrese. Postoji direktno 1 na 1 preslikavanje između tabele rutiranja i *FIB* tabele. Na svakom *CEF* omogućenom ruteru se koristi minimum *FIB* tabela, kao dodatak ovome, *MPLS* može da označi neku putanju, ta putanja se beleži u *LIB* (*Label Information Base*) tabeli koju koristi *LDP* (*Label Distribution Protocol*) protokol distribucije labela koji mapira IP adresu u labelu sledećeg skoka. Naučene labele u kontrolnom delu se prosleđuju *LFIB* (*Label FIB*) tabeli koja izvršava uparivanje sledećeg skoka sa odgovarajućim interfejsom. Samo one labele koje se koriste se prosleđuju *LFIB* tabeli za prosleđivanje.

Informacije o dostupnosti destinacione mreže koje se preuzimaju iz baze podataka protokola rutiranja se beleže u *FIB* tabeli. *LIB* tabela se popunjava uz pomoć protokola za distribuciju labela (*LDP*), a koja se zatim koristi zajedno sa *FIB* tabelom kako bi se formirala *LFIB* tabela. *LIB* tabela sadži sledeće informacije : *label-in*, *label-out*, *port-in*, *port-out* i instrukcije. Kada se učitava *LIB* protok informacija može započeti.

LIB tabela se učitava koristeći nezavisnu kontrolu (*Independet control*) ili na osnovu naređenje kontrole (*Ordered control*).

3.2.2. Nezavisna kontrola

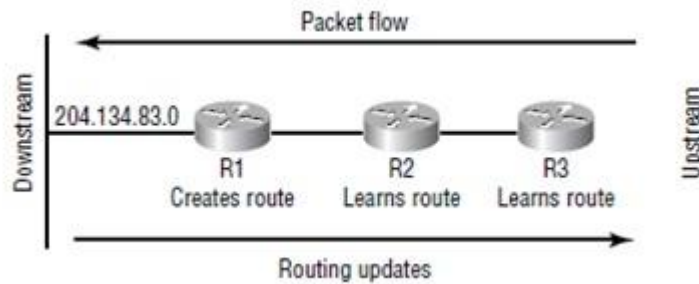
Nezavisna kontrola funkcioniše tako što svaki ruter osluškuje informacije o rutiranju, kojima je preplavljena mreža od strane ostalih rutera i na osnovu tih informacija on pravi svoju *LIB* tabelu. Zatim se ove informacije koje su zabeležene u *LIB* šalju dalje u mrežu ostalim ruterima. Ova metoda ima prednosti zbog toga što ruter može osluškivati bilo koji protokol rutiranja, a potom generisati svoju tabelu i distribuirati je dalje. Konvergencija je brza i dešava se nakon bilo koje promene. Problem je, međutim u tome što ne postoji centralizovano mesto kontrole, kao i to da kontrolu saobraćaja je teže postići.

3.2.3. Naređena kontrola

Naređena kontrola kada se koristi, tada jedan ruter, koji je obično izlazni *LSR* za dinamičku konfiguraciju ili ulazni *LSR* ruter za statičku konfiguraciju, je odgovoran za distribuiranje informacija o označavanju. Tada „uzvodni” *LSR* mora da čeka oznaku od „nizvodnog” *LSR*. Kontrola je tada centralizovana i bolje postignuta. Mana je ta da u slučaju neke promene konvergencija sistema je sporija, kao i to što postoji tačka održivosti (*single point of failure*) koja ukoliko otkaze može doći do zastoja kompletnog sistema.

Kada se koristi nezavisna kontrola oznake su spojene sa *FEC* nezavisno od rutera koji je „nizvodno” (*Down-stream*), u procesu poznatijem kao nizvodni netraženi proces (*Downstream Unsolicited*). Naređena kontrola koristi oznake koje se spajaju sa *FEC* od strane „nizvodnog” rutera kada oni zahtevaju neku vrstu spajanja sa *LMR* (*Label Manager Router*). Ovaj proces je poznat kao nizvodni proces na zahtev (*Down-stream on Demand* (*DOD*)).

Na sledećem primeru (slika 31) pojasnićemo termine „nizvodni” (*Downstream*) i „uzvodni” (*Upstream*) ruter.

Slika 31: Predstavljanje „nizvodnog” (*Downstream*) i „uzvodnog” (*Upstream*) rutera

Ruter R1 objavljuje pod mrežu 204.134.83.0, koju ruter R2 uči. Ruter R3 će naučiti rutu od R2 rutera. Saobraćaj koji je namenjen pod mreži 204.134.83.0 mora proći kroz ruter R1. Termini „uzvodno” i „nizvodno” su u relaciji sa tokom saobraćaja korisnika, a ne sa tokom saobraćaja koji prenosi informacije o rutiranju. Dakle saobraćaj namenjen pod mreži 204.134.83.0 teče od uzvodnih (*Upstream*) rutera R3 i R2 prema nizvodnom (*Downstream*) R1 ruteru.

3.3. Razmena labela

Razmena labela se odvija uz pomoć dva protokola: *TDP* (*Tag Distribution Protocol*) koji je u CISCO vlasništvu. Spaja tagove (što je isto kao i labela) sa rutama u tabeli rutiranja i *LDP* (*Label Distribution Protocol*) koji je IETF standard i ima za cilj distribuciju labela u MPLS okruženju. On se oslanja na osnovne informacije koje dobija od *IGP* (*Interior Gateway Protocol*) kako bi prosledio označene pakete.

3.3.1. TDP (*Tag Distribution Protocol*)

Ovaj protokol je u CISCO vlasništvu. On koristi *TCP* (*Transmission Control Protocol*) za transport i zbog toga je konekciono orijentisan i pruža garantovanu sekvencijalnu dostavu.

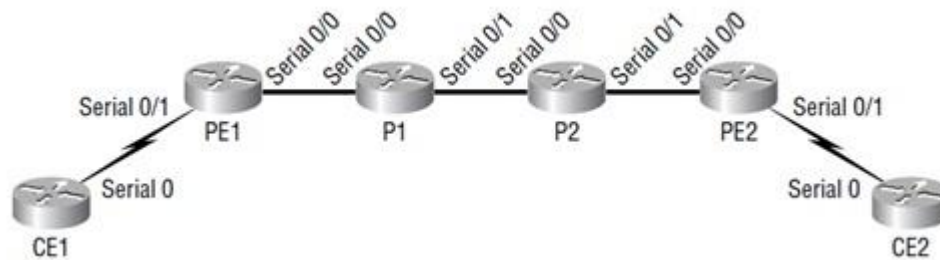
TDP koriste ruteri za označavanje paketa (*Tag Switching Routers – TSR*) kako bi komunicirali o povezivanju oznaka sa čvorovima (*peers*). TDP podržava više mrežnih protokola kao što su IPv4, IPv6, IPX, AppleTalk. TDP omogućava da TSR distribuira, zahteva, otpušta informacije o označavanju (*Tag Binding Information*) za različite mrežne protokole. TDP takođe pruža mogućnost za otvaranje, nadgledanje i zatvaranje TDP sesija na kojima postoje otkrivene greške za vreme sesija.

3.3.2. LDP (*Label Distribution Protocol*)

Predstavlja set procedura i poruka kojim se utvrđuje kako će jedan LSR (*Link State Router*) ruter obavestavati ostale rutere o povezivanju (*bind*) labela koje je uradio. LSR koristi ovaj postupak obavestavanja kako bi se napravila putanja kroz mrežu, mapiranjem informacija o rutiranju mrežnog nivoa direktno na putanju data link nivoa. U jednoj sesiji svaki LSR ruter može da nauči o ostalim mapiranim labelama. Ovo znači da je ovaj protocol bidirekcion. LDP je IETF standard. On omogućava da LSR zahteva, distribuira i otpušta labele ruterima u mreži. LDP omogućava da LSR otkrije potencijalne čvorove (*peers*) i uspostavi sesiju sa tim ruterima zarad razmene labela. Nakon što dva LSR razmene LDP parametre oni će formirati LSP (*Label Switched Path*).

3.4. Primer mreže provajdera

Pre nego što se primer pojasni moramo razjasniti same uređaje na slici 32.



Slika 32: Primer jednostavne mreže provajdera

Objašnjenje uređaja sa slike:

CE – *Customer Equipment* – uređaj koji šalje neobeležene pakete u mrežu provajdera

PE – *Provider Equipment – edge LSR* – uređaj koji se nalazi na obodu mreže provajdera i on obeležava pakete i prosleđuje obeležene pakete ka mreži provajdera, ali takođe isti uređaj skida obeležja sa paketa i šalje neobeležene pakete ka mreži korisnika

P – *Provider Router* – uređaj koji pripada provajderu, on se nalazi u jezgru mreže provajdera i samo prosleđuje obeležene pakete

3.4.1. Tipovi uređaja za obeležavanje paketa

Postoje dva tipa rutera za obeležavanje paketa

- *LSR (Label Switched Router)* – uređaj koji vrši prosleđivanje paketa na osnovu labela (oznaka)
- *Edge LSR* – uređaji koji imaju zadatak da obeležavaju neobeležene IP saobraćaj (*PUSH*) i potom prosleđuju pakete na osnovu labele u jezgro mreže. Nalaze se na granicama (*PE* ruteri). Skidanje labele se naziva *POP(ing)*.

3.4.2. Pojašnjenje primera

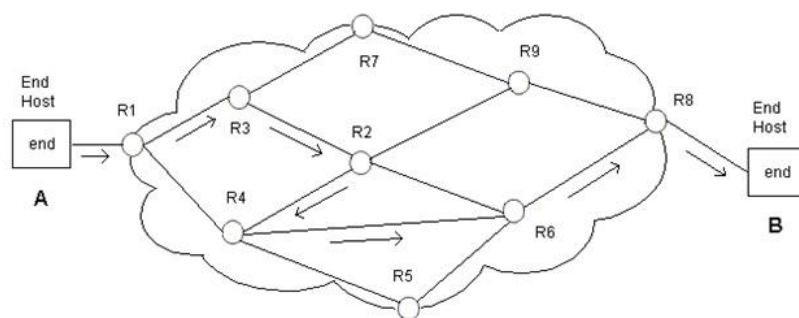
Na datom primeru (slika 19) predstavljena je jednostavna mreža provajdera na kojoj *CE* ruter označava korisničku opremu koja šalje neobeležene, dakle čiste IP pakete prema mreži provajdera. Kod standardnog rutiranja na svakom ruteru, koji se nalazi na putanji paketa, bi se gledala tabela rutiranja i na osnovu tih informacija bi se paket prosledio sledećem ruteru, što bi dovelo do toga da se izvrši dalje rutiranje sa jedne na drugu stranu. Ali kod MPLS pošto neobeleženi paketi stignu na dolazni interfejs graničnog rutera *PE*, taj isti ruter vrši obeležavanje paketa određenim labelama (oznakama), kako bi se dalje prosleđivali paketi na osnovu labela, a ne na osnovu IP zaglavlja. Ovaj ruter koji vrši označavanje paketa se zove ulazni (*Ingress*) ruter. Prenosanje paketa putem labela se postiže brže prenošenje paketa sa kraja na kraj, zbog toga što se ne gleda tabela rutiranja već samo keš verzija tabele rutiranja, što opet dovodi do zaključka da je prenošenje paketa sa kraja na kraj dosta brže. Nakon što se označi paket, oni se prosleđuju jezgru mreže tj. *P* ruterima koji samo vrše prosleđivanje paketa na osnovu labela. Kada paketi stignu do kraja mreže, tamo postoji još jedan *PE* ili *Edge LSR* koji je zadužen za uklanjanje labele kako bi dalje paketi mogli da se rutiraju opet na osnovu IP zaglavlja. U suprotnom uređaji klijenata ne bi znali kako da obrade pakete i oni bi bili odbačeni. Izlazni ruter se zove *Egress* ruter.

POP – Penultimate POP hopping je proces pri kome se uklanja labela sa paketa i on se neoznačen kao IP dalje šalje. Ovo je napravljeno kako bi se uštedelo na vremenu koje treba

poslednjem ruteru, pred odredištem (pred klijentskim uređajima), da ukloni labelu i pošalje paket. Ovde je očigledno kome treba paket da se pošalje i nema potrebe dodatno označavati paket. Dakle kada paket stigne na pretposlednji ruter (*penultimate*) labela se uklanja i paket se šalje na osnovu IP adrese ruteru koji dalje uzima paket sa IP zaglavljem i traži tu pod mrežu. U suprotnom, kao što je već rečeno, poslednji ruter bi morao da prevodi labelu u IP, a tek onda da gleda rutingu tabelu što bi dovelo do usporavanja. U datom primeru na slici 19 ovo bi trebalo omogućiti na P2 ruteru u jezgri koji će skinuti labelu paketu, proslediti IP ELSP ruteru PE2 koji će samo gledati na osnovu IP adrese u tabeli rutiranja gde treba da prosledi paket i potom da ga prosledi na odgovarajući interfejs.

3.5. Putanje kroz MPLS

LSP (Label Switched Path) – predstavlja putanju kroz MPLS mrežu koja je uspostavljena od strane protokola (npr. *LDP*). Putanja je napravljena na osnovu kriterijuma iz *FEC (Forwarding Equivalence Class)* što predstavlja opis paketa sa sličnim ili identičnim karakteristikama (npr. ista destinaciona IP adresa), koji možda mogu da se proslede u istom pravcu tj. primenom iste MPLS oznake. Takođe se često koristi i *QoS (Quality of Service)* klasa za određenu putanju. *LSP* je unidirekcionalan skup *LSR* uređaja koje paket mora proći kako bi stigao do cilja. Na slici 33 možemo videti *LSP* putanju od hosta A do hosta B.



Slika 33: Prikaz *LSP* putanja

3.6. Isporučivanje IPv6 preko MPLS mreže

Mnoge tehnike migracija su pružene provajderima kako bi dodali IPv6 servis u svoj portfolio usluga. Međutim, kada provajder već ima u svojoj mreži IPv4 sa MPLS-om u jezgru mreže ovome se mora prići veoma oprezno. *IP MPLS* jezgro mreže će verovatno biti noseći deo infrastrukture za ostvarivanje prihoda i usluga (kao što su MPLS VPN usluge), kao i za tranzit ključnih slojevitih (*overlay*) mreža (kao što su *PSTN* kanali (*trunking*) i ATM kanali). Dakle od najvećeg značaja je da se prilikom migracije na IPv6 servis izbegne uvođenje određenog rizika ili nestabilnosti jezgra mreže koje već pruža višeservisne usluge brojnim klijentima. Takođe deo naprednih MPLS osobina može da se primeni kao što su npr. upravljanje saobraćajem (*Traffic Engineering*), brzo preusmeravanje (*Fast Reroute*), *MPLS kvalitet usluga QoS (Quality of Service)*.

Pristup migraciji na IPv6 ne sme da ometati poslovanje ovih osobine za IPv4 saobraćaj i trebalo bi da omogući dobit dodavanjem servisa koji idu uz IPv6. Konačno tamo gde oprema ima velike MPLS performanse prosleđivanja (*forwarding*), može biti poželjno za IPv6 pristup migraciji da iskoristi MPLS prosleđivanje za IPv6 saobraćaj.

Najočigledniji pristup za uvođenje IPv6 servisa je, naravno, da se unapredi cela mreža tako da podržava IPv6 rutiranje i prosleđivanje prirodno (*native*). Iako je ovaj pristup najintuitivniji i jasno nudi vrlo skalabilnu podršku globalne IPv6 konektivnosti, provajderi koji već imaju IPv4 MPLS u svojoj mreži često ne zadržavaju ovaj pristup, bar u relativno kratkom vremenskom roku, iz sledećih razloga:

- Oni više vole da izbegnu ili bar barem odlože odgovarajuće uvođenja čistog (*native*) IPv6 u jezgro mreže.
- Ovaj pristup ne dozvoljava MPLS osobine da budu primenjene u jezgru kao što su brzo prosleđivanje, inženjering saobraćaja, MPLS kvalitet usluga (QoS).
- To zahteva IPv6 prosleđivanje visokih performansi na celokupnoj instaliranoj opremi, ali ne dobijaju se prednosti MPLS prosleđivanja.
- Nije lako proširiv da obogati IPV6 usluge konektivnosti u MPLS VPN servis

Postoji još jedan pristup migraciji, a to je da se koristi IPv6 preko IPv4 tunelovanje (*IPv6 over IPv4 tunneling*). Ovo zahteva kreiranje IPv4 tunela na *CE* ruterima i pokretanje IPv6 rutiranja na vrhu tunela. Ovaj pristup koristi jednostavnu i dobro poznatu tehniku i uopšte ne zahteva unapređivnje okosnici (*backbone*) mreže, čak ni na *PE* ruterima. I zato što tunelovan IPv6 saobraćaj automatski pruža VPN izolovanost IPv4 MPLS servisa, ovakav pristup tunelovanju od *CE* do *CE* rutera je primenjen od strane nekih provajdera kao način brzog početka pružanja IPv6 servisa, što omogućava kratko vremene izlaska usluge provajdera na tržište. Međutim, uvođenje IPv6 servisa koristeći ovaj pristup podrazumeva prilagođeni dizajn, konfiguraciju i operacije.

Takođe ovaj pristup pati od uobičajenih izazova skalabilnosti tehnika tunelovanja (kreiranje i upravljanje tunelima, rutiranjem, od svakog *CE* rutera do svih ostalih *CE* rutera). Iz ovih razloga, operatori su u potrazi za ostalim pristupima kako bi podržale razvoj IPv6 servisa.

Pseudo-žična tehnologija (*pseudo-wired technology*), dozvoljava da se emulira prirodni (*native*) servis preko mreže za prosleđivanje paketa (*packet switched network*). Ovaj prirodni servis može biti *SONET*, *ATM*, *Frame Relay* ili *Ethernet* dok mreža za prosleđivanje paketa može biti *Ethernet*, *MPLS* ili *IP* (bilo verzija *IPv4* ili *IPv6*). Ovo može opet biti korišćeno za povezivanje IPv6 rutera eventualno rezultujući IPv6 konektivnošću podržane preko *IPv4 MPLS* mreže. Kao i *IPv6* preko *IPv4* tunelovanje i ovaj pristup izbegava bilo koje *IPv6* unapređivanje u jezgru, ali takođe dolazi do sličnih izazova skalabilnosti zato što odgovarajuća mreža kola mora da bude uslužna preko *IPv6* uređaja. Kao i tunelovanje ovaj pristup su koristili provajderi za rano uvođenje *IPv6* servisa u eksploataciju i time dobijajući grupu klijenata kojoj je to neophodno.

Konačno *6PE* i *6VPE* pristupi, retrospektivno, dozvoljavaju podršku globalne *IPv6* usluge dostupnosti (*IPv6 reachability service*) i *IPv6 VPN* servisa preko *IPv4 MPLS* okosnice. Ovi pristupi su se pokazali veoma atraktivnim operatorima koji koriste *IPv4* okosnicu zato što:

- Ne zahtevaju unapređivanje *P* rutera. Time održavaju stabilnost jezgra i minimiziraju operativne troškove.

- Dozvoljavaju stepenasto uvođenje i to unapređivanjem PE rutera i dobijanjem IPv6 usluga, a gde se koriste reflektori ruta (*Route Reflectors* – koja je komponenta rutiranja i pruža alternativu logičkim *full mesh* zahtevima iBGP) ili unapređivanjem tih ili isporučivanjem odvojene mreže reflektore ruta za IPv6.
- Veoma su skalabilni zato što se oslanjaju na isti model kao i MPLS VPN, gde dodavanje novih sajtova je ustvari konfigurisanje portova za priključak za određeni sajt.
- Iskorišćavaju prednost MPLS prosleđivanja u jezgru mreže i njegove velike performanse.
- Osiguravaju da IPv6 saobraćaj automatski prosperira od naprednih MPLS osobina koje mogu biti isporučene u jezgru mreže: FRR, TE, MPLS QoS.

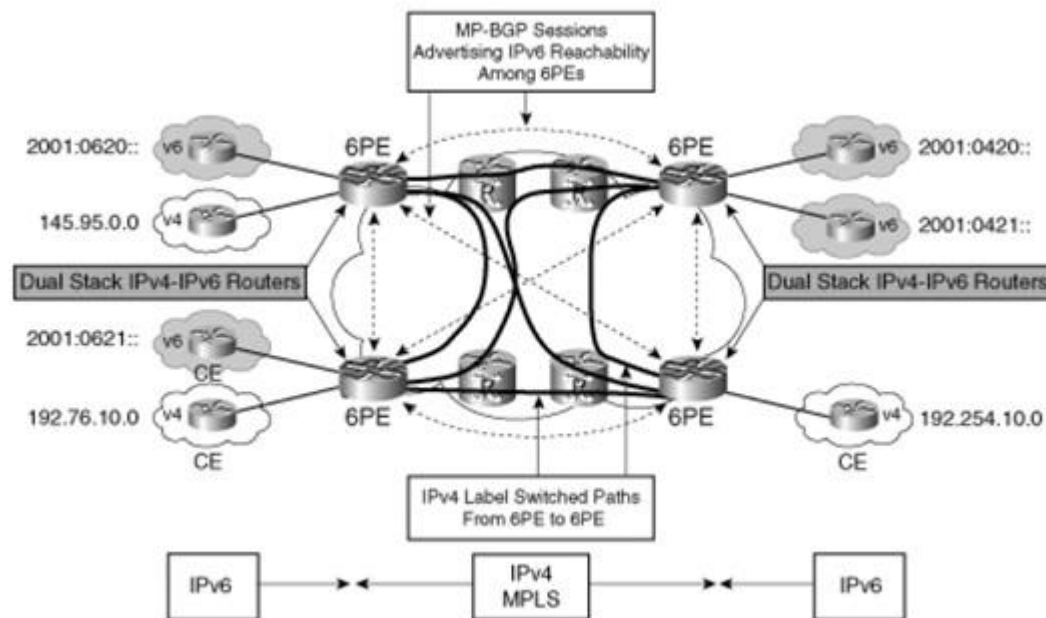
3.7. Ivice provajdera – 6 PE (*IPv6 Provider Edge*)

MPLS VPN arhitektura L3 nivoa uvodi fundamentalnu paradigmu. Ova paradigma je usmeravanje i transport IPv4 VPN saobraćaja transparentno preko IPv4 MPLS jezgra koje ostaje potpuno nesvesno IPv4 VPN ruta i prepoznaje samo interne IPv4 rute. Ovo se postiže kombinovanjem hijerarhijskog rutiranja u kome jezgro mreže uspostavlja IPv4 PE to PE veze, dok IPv4 dostupnost se oglašava samo između PE rutera transparentno preko jezgra.

Tunelovanje IPv4 VPN paketa kroz jezgro od PE rutera do PE rutera u IPv4 MPLS LSP tako da jezgru ni ne treba da bude svesno IPv4 VPN.

6PE rešenje koristi istu paradigmu kako bi se postigla globalna IPv6 dostupnost preko IPv6 nesvesne IPv4 MPLS okosnice. Ključna razlika je da informacije o dostupnosti koje se objavljuju između PE rutera preko MP-BGP nisu više IPv4 VPN prefiksi već IPv6 VPN prefiksi. Dakle PE ruteri postaju dvostruko konfigurisani (*dual stack* – što znači da oni imaju pokrenut i IPv4 i IPv6 na sebi) i nazivaju se 6PE ruteri. Podržavaju i IPv4 i IPv6 na pristupnim interfejsima, ali i dalje podržavaju samo IPv4 i IPv4 MPLS na interfejsima prema jezgru mreže.

P ruteri (ruteri u jezgru mreže) ostaju nesvesni IPv6 i pokrenut je uobičajen IPv4 protokol rutiranja i IPv4 distribucija oznaka. Ovaj vid arhitekture je prikazan na slici 34.



Slika 34: Prikaz 6PE arhitekture

Jedan od načina da se pogleda 6PE rešenje je da se rešenje može razmatrati kao IPv4 MPLS jezgro mreže koje efektivno nosi saobraćaj dodatne virtualne privatne mreže (VPN), čiji je saobraćaj i adrese ustvari IPv6. Baš kao i u IPv4 VPN, ruteri u jezgru ostaju potpuno nesvesni ruta koji pripadaju određenoj virtualnoj privatnoj mreži. Moramo imati na umu, međutim, da taj posebni „VPN“ ne uključuje mehanizme kontrole ruta kao što su *VRF* (*Virtual Routing and Forwarding* – omogućava nam da unutar jednog rutera postoji više instanci tabele rutiranja. Ovo nam opet pomaže tako da može postojati više IP adresa koje se preklapaju, a da ne dovode do konflikta), *RD* (*Route Distinguisher* – koji nam omogućava da pravimo razliku između mreža tj. VPN koji su slično konfigurisane, ali pripadaju različitim klijentima) i *RT* (*Route Targets* – ovo su dodatne BGP vrednosti koje

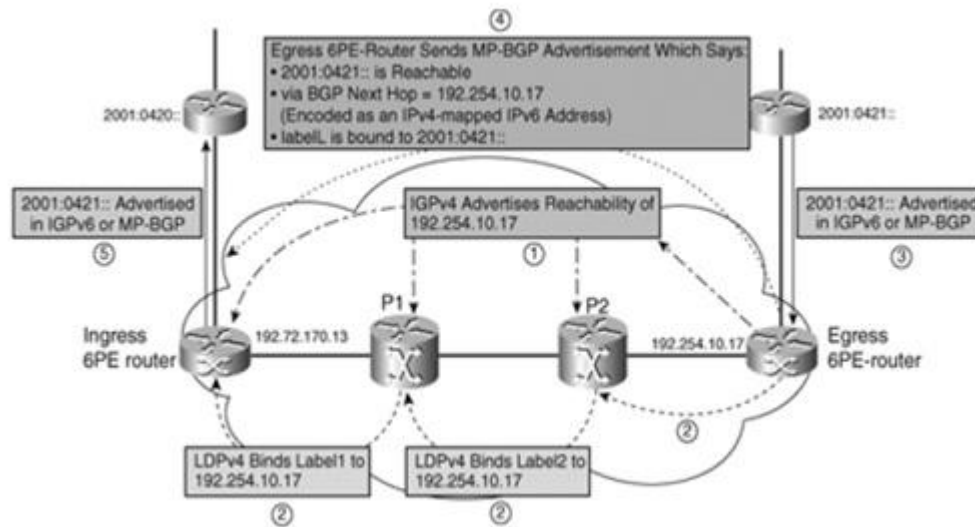
su najbliža aproksimacija VPN identifikatora u MPLS/VPN arhitekturi), zato što rutiranje i prosleđivanje tabela IPv6 prirodno odvojeno od onih u IPv4.

Iz perspektive kontrolnog dela sa slike 35, sledeći koraci su preduzeti pre nego što IPv6 komunikacija može da počne od izvorišnog IPv6 sajta konektovanog na 6PE ruter (*ingress* – ulazni ruter) do IPv6 destinacionog rutera u drugom IPv6 sajtu, takođe 6PE rutera (*engress* – izlazni uređaj)

1. Dostupnost IPv4 adrese na logičkom (*loopback*) interfejsu izlaznog rutera se objavljuje u jezgro mreže preko IPv4 IGP svim P ruterima i svim ostalim 6PE ruterima.
2. Labele se distribuiraju u jezgru mreže svim P i 6PE ruterima za ovu IPv4 logičku adresu, pomoću uobičajenih IPv4 tehnike distribucije labela, kao što je LDP protokol. Ovo rezultuje uspostavljanjem IPv4 konektivnosti od ulaznog do izlaznog rutera u formi IPv4 LSPa.
3. 6PE ruteri pokreću MP-BGP međusobno koristeći označenu IPv6 adresnu familiju. Zbog toga što jezgro podržava samo IPv4, MP-BGP sesije rade na IPv4 steku. Nakon što je naučeno kako se dolazi do destinacionog IPv6 prefiksa (npr. preko IPv6 protokola rutiranja koji je pokrenut na IPv6 CE ruteru ili preko konfigurisane statičke IPv6 rute) izlazni 6PE ruter objavljuje dostupnost ovog prefiksa preko MP-BGP, svim ostalim 6PE ruterima koristeći MP-BGP. Budući da jezgro mreže omogućava IPv4 konektivnost samo između 6PE rutera, prilikom objavljivanja IPv6 dostupnosti, izlazni 6PE ruter mora preneti i na druge 6PE rutere svoju IPv4 adresu kao BGP adresu sledećeg skoka (*next hop*). Međutim BGP pretpostavlja da je BGP polje sledećeg skoka iste adresne familije kao i informacije o dostupnosti (koje su IPv6 u ovom slučaju). Kao što je objašnjeno IPv6 arhitektura adresiranja definiše IPv4 mapirani IPv6 adresni format, upravo za svrhu predstavljanja adrese IPv4 čvora kao IPv6 adrese. Prema tome IPv4 adresa izlaznog rutera je enkodirana u BGP polje sledećeg skoka kao IPv4 mapirana IPv6 adresa. Labela je takođe objavljena od strane izlaznog 6PE rutera za IPv6 prefiks. Konačno, izlazni 6PE ruter popunjava unos u njegovoj LFIB tabeli za ovu labelu/prefiks koji pokazuje kako da se proslede paketi primljeni sa datom labelom. U zavisnosti od politike

dodeljivanja labela ovo može biti skidanje oznake (POP labele) i prosleđivanje interfejsa sledećeg skoka do destinacije, odnosno to može biti skidanje oznake (POP labele) i izvođenje pretraživanje IPv6 baze za prosleđivanje (FIB tabela).

4. Nakon pokretanja uobičajenog BGP algoritma selekcije putanje, ulazni 6PE ruter popunjava svoju IPv6 FIB tabelu sa unosom za objavljeni IPv6 prefiks koji pokazuje da paket namenjen za taj IPv6 prefiks:
 - a. određen je da se enkapsuliran koristeći MPLS labelu čija je krajnja labela, labela objavljena u MP-BGP za IPv6 prefikse i čija je početna labela, labela objavljena u jezgru za IPv4 logičku adresu izlaznog 6PE rutera koji je BGP sledeći skok za taj IPv6 prefiks.
 - b. određen je da se prosledi na interfejs sa IPv4 najkraćom putanjom do 6PE izlaznog uređaja.
5. Ukoliko se koristi protokol za rutiranje između izvorišnog sajta i 6PE ulaznog rutera, ulazni ruter će objaviti dostupnost IPv6 prefiksa u tom protokolu rutiranja.



Slika 35: 6PE operacije kontrolnog dela

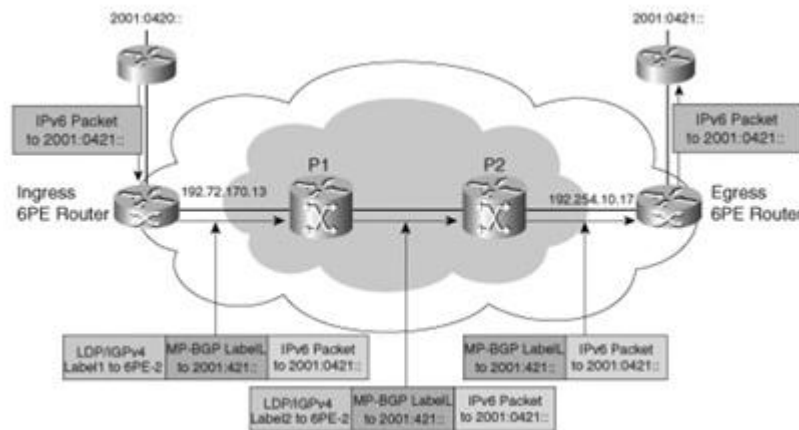
IPv6 komunikacija sada može da se odvija preko IPv4 MPLS okosnice mreže. Kada ulazni 6PE ruter primi IPv6 paket, on će obaviti pogled na destinacionu IPv6 adresu u njegovoj

IPv6 FIB bazi podataka, koji odgovara unosu popunjenog od strane kontrolnog dela, kao što je diskutovano. Dakle ulazni 6PE ruteru postavlja labelu ispred IPv6 paketa sa donjom labelom koja je objavljena od strane MP-BGP za IPv6 prefiks i sa gornjom labelom koja je objavljena od strane protokola distribucije iz jezgra, za BGP sledeći skok IPv4 logičku adresu. 6PE ruter konačno prosleđuje taj označeni paket prema jezgru na interfejsu sledećeg skoka prema izlaznom 6PE ruteru.

P ruteri obavljaju regularne IPv4 operacije zamene labela, koja rezultuje zamenom najviših, prvih labela (ili POP – skidanje labela ukoliko se koristi PHP). Konačno paket je primljen od strane izlaznog rutera, tamo gde se koristi PHP (skidanje labela na pretposlednjem ruteru) u jezgru mreže, paket je primljen sa jednom oznakom, koja je oznaka objavljena od strane MP-BGP za IPv6 prefiks. U suprotnom izlazni 6PE ruter prvo izvršava POP odnosno skidanje IPv4 oznake, tako dobijajući oznaku objavljenu u MP-BGP.

Kada se izvodi pogled u LIB bazi podataka za oznaku izlazni 6PE ruter nalazi unos popunjen od strane kontrolnog dela koji im govori kako treba propisno da se prosledi paket.

Na slici 36 predstavljen je primer prosleđivanja paketa na svakom skoku. U ovom primeru kome se koristi PHP i MP-BGP raspoređuje odvojene oznake za svaki IPv6 prefiks, tako da se paketi mogu presleđivati na osnovu labela od strane izlaznog 6PE rutera bez potrebe za IPv6 pogledom.



Slika 36: 6PE operacije dela za rukovanje podacima

Još jednom ovde možemo videti da su IPv6 uređaji u IPv6 sajtovimu potpuno nesvesni da se u jezgru mreže dešava IPv6 prosleđivanje paketa preko MPLS. Oni rade u njihovom regularnom IPv6 načinu. Slično tome P ruteri u MPLS jezgru su potpuno nesvestni da MPLS paketi koji se prosleđuju ustvari prenose IPv6 saobraćaj.

Kao i u MPLS VPN arhitekturi 6PE ruteri mogu da pokrenu potpunu mešovitu (*full-mesh*) MP-BGP sesiju za razmenu IPv6 informacija o dostupnosti ili se mogu koristiti uobičajene BGP tehnike skaliranja kao što su reflektori ruta (*Route Reflector*).

6PE pristup može funkcionisati na višestrukim autonomnim sistema. Jedan takav senario u kome je 6PE rešenje lokalizovano unutar svakog autonomnog sistema. U ovom slučaju ruteri na ivicama autonomnih sistema (ASBR) razmenjuju regularne IPv6 informacije o dostupnosti (možda čak i da nisu svesni toga da druga mreža koristi 6PE pristup). Drugi senario je gde 6PE rešenje proširuje (*span*) višestruke autonomne sisteme. U tom slučaju označene IPv6 informacije o dostupnosti se razmenjuju između reflektora ruta ili direktno između 6PE rutera, različitih autonomnih sistema. U isto vreme označene IPv4 rute za 6PE logičke adrese rutera se razmenjuju između ASBR rutera različitih autonomnih sistema, kako bi se osiguralo da 6PE ruteri mogu biti dostupni kroz različite autonomne sisteme preko IPv4 LSP.

3.8. 6VPE (IPv6 VPN Provider Edge)

Kao dodatak IPv6 servisa globalne povezanosti koji može da se pruži sa 6PE rešenjem, provajderi biće upitani od svojih klijenata da li pružaju IPv6 VPN usluge. Primarni razlog za ovaj servis je potreba za izolacijom krajnjih korisničkih intraneta. 6VPE kombinuje IPv6 upravljanje 6PE sa VPN rukovanjem IPv4 MPLS VPN-a (opisanog ranije) kako bi podržale ovakav IPv6 VPN servis preko IPv4 MPLS jezgra.

Primetne ekstenzije 6PE pristupa su:

- Korišćenje različitih adresnih familija u MP-BGP definisanog za 6VPE svrhu, koja je IPv6 VPN adresna familija. VPN IPv6 adresa je 24bajtni entitet koji počinje sa 8

bajtnim razjašnjivačem ruta (RD) i završava se sa 16 bajtnom IPv6 adresom. Uloga i kodiranje RD-a je ista kao i sa IPv4 VPN-om.

- Korišćenje VRF (Virtual Routing and Forwarding) koncepta L3 MPLS VPN arhitekture, u kojoj svaki VPN ima zaseban skup tabela za rutiranje i prosleđivanje, zajedno sa svim povezanim mehanizmima za kontrolu unosa i izvoza ruta u i iz VRF-a, uključujući označavanje ruta sa targetima ruta.

6VPE pristup donosi iste prednosti kao i 6PE pristup. Na primer kao i sa 6PE samo PE ruteri koji zaista povezuju IPv6 VPN servis treba da budu unapređeni kako bi podržali IPv6 i 6VPE funkcionalnost. Dakle provajderi mogu da uvedu IPv6 VPN servis bez dodatnih unapređivanja ili konfiguracionih izmena na ruterima unutar jezgra.

Takođe zbog toga što se 6VPE pristup oslanja na iste mehanizme kao i L3 MPLS VPN za IPv4, provajderi mogu ponuditi krajnjim korisnicima iste usluge kao što su mogli sa IPv4, praveći IPv6 VPN servis znatno jednostavnijeg za shvatanje i integraciju istog unutar korisničkih intraneta.

Konačno, provajderi mogu da se oslone na isti set alatki operacija, administracije i održavanja (OAM) – kako bi podržali i IPv4 i IPv6 VPN usluge, čime se drastično smanjuju troškovi poslovanja.

4. Zaključak

Cilj ovog rada bio je da se ukaže na teškoće koje postoje u radu sa trenutnom IPv4 vrstom protokola, predstavi novo rešenje u vidu novog IPv6 protokola, a potom da se ukaže i na problem koje naizgled jednostavna tranzicija sa IPv4 na IPv6 može da dovede u mreži provajdera koji već pruža klijetnima svoje usluge.

S obzirom da je tranzicija neophodna postoji dosta tehnika opisanih ranije (npr. dvostruka konfiguracija (*dual stack*), prevođenje IPv4 u IPv6 (*IPv4 to IPv6 translation*), tunelovanje (*tunneling*)), a sami provajderi imaju dosta opcija, ali samo pojavljivanje novog protokola, IPv6, dovodi do novih izazova koji uključuju skalabilnost, integraciju i sigurnost. Drugim rečima samim uvođenjem IPv6 u postojeću IPv4 mrežu, dovodi do spuštanja nivoa sigurnosti u već ne toliko sigurnu IPv4 mrežu kao što je predstavljeno ranije u radu.

Ova opasnost može biti ublažena faznom pristupu tranzicije. Iako postoji veliki broj tehnika tranzicije, ne može se ni za jedan reći da je najbolji. Svaki ima svoje prednosti i mane, kao i trenutke kada ih je najbolje koristiti.

Integracijom MPLS u mreže provajdera dovodi se do bržeg prenošenja paketa, a dodavanjem VPN u MPLS mreže dovodi do potpunog izolovanja saobraćaja klijenata, bez pravljenja viška saobraćaja (*overhead*). Međutim najveća mana MPLS je ta što mora postojati jaka okosnica koja će omogućiti korisnicima da osete dobit od MPLS. Zbog ovoga se još opreznije mora prići problem tranzicije na novi protokol, kako klijenti ne bi bili oštećeni.

Sve prethodno rečeno dovodi do zaključka da je tranzicija na novi protokol, koji dovodi do višestrukih poboljšanja, više nego poželjna i kao rešenje se najviše ističu dvostruka konfiguracija i tunelovanje.

Iako je finansijska situacija i stanje firme, ali i orkuženja, uvek u prvom planu, ovde se takođe kao rešenje opet ističu dvostruka konfiguracija i tunelovanje zbog toga što provajderi nisu finansijski opterećeni, bar ne onoliko koliko bi bili da se odluče za prirodnu (*native*) IPv6 mrežu. Ovo prelazno rešenje će im omogućiti pružanje usluga koje nudi IPv6

ali i MPLS. Kasnije, kada stasa za sada mladi IPv6 protokol, ali i prateće IPv6 aplikacije, lako će se odlučiti na prelazak na prirodnu (*native*) mrežu.

5. Literatura

1. IPv6 over MPLS networks Cisco press (www.cisco.com)
2. IPv6 security issues, Samuel Sotillo
3. IPv6 Internetworking Technology handbook – CISCO (www.cisco.com)
4. MPLS – CISCO (www.cisco.com)
5. www.tcpipguide.com
6. Stateless autoconfiguration (IPv6.com)
7. MPLS RFC 3031
8. IPv4, IPv6 (technet.microsoft.com)
9. IPv6 RFC 2460
10. IPv4 RFC 791
11. IPv6 over MPLS, Patrick Grossetete – CISCO IOS IPv6 product manager
12. IPv6 Security – CISCO
13. Sybex – CISCO MPLS Study Guide