

Universitatea „Politehnica” București

Facultatea de Electronică, Telecomunicații și Tehnologia Informației

Combaterea congestiei cu ajutorul algoritmilor RED

Proiect de licență

prezentat ca cerință parțială pentru obținerea titlului de
Inginer în departamentul *Electronică Aplicată și Tehnologia Informației*,
domeniul *Calculatoare și Tehnologia Informației*,
programul de studii *Ingineria Informației*

Conducător științific

Conf. dr. ing. Ștefan Stăncescu

Absolvent

Damaschin Ionuț

Anul 2014

Universitatea "Politehnica" din București
Facultatea de Electronică, Telecomunicații și Tehnologia Informației

Departamentul Electronică Aplicată și Ingineria Informației

Aprobat Director de Departament :


Prof. Dr. Ing. Sever PAȘCA

TEMA PROIECTULUI DE DIPLOMĂ
a studentului Damaschin D. Ionuț George, grupa 443A

1. Titlul temei: **Combaterea congestiei cu ajutorul algoritmilor RED.**

2. Contribuția practică, originală a studentului va consta în *(în afara părții de documentare)*:

Se vor face simulări cu simulatorul Network Simulator prin scripturi TCL; se generează fișiere de trace pe baza cărora se vor realiza reprezentările grafice și tabelele pentru compararea parametrilor de calitate și transfer în funcție de algoritmi folosiți; graficele vor fi realizate cu GNUPLOT; se va modifica sursa NS care implementează algoritmi RED, se vor experimenta transferuri cu algoritmi modificați și se vor compara performanțele.

3. Proiectul se bazează pe cunoștințe dobândite în principal la următoarele 3-4 discipline:
Rețele de calculatoare, Sisteme de operare, Arhitecturi de rețea și Internet

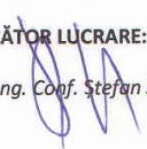
4. Realizarea practică/ proiectul rămân în proprietatea: UPB

5. Proprietatea intelectuală asupra proiectului aparține: UPB

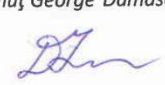
6. Locul de desfășurare a activității: UPB, ETTI

7. Data eliberării temei: 12.12.2013

CONDUCĂTOR LUCRARE:


Prof. Dr. Ing. Conf. Ștefan Stancescu

STUDENT:


Ionuț George Damaschin

**DECLARAȚIE DE ORIGINALITATE A
PROIECTULUI DE DIPLOMĂ/LUCRĂRII DE DISERTAȚIE¹⁾**
(MODEL)

Subsemnatul DAMASCHINDIONUT GEORGE ²⁾, posesor al

B.I./C.I./pașaport seria AS, nr. 643 720, având CNP

1	9	1	0	7	1	6	1	6	0	1	3	4
---	---	---	---	---	---	---	---	---	---	---	---	---

,

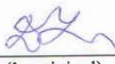
am întocmit proiectul de absolvire/diplomă/lucrarea de disertație cu
tema:³⁾ COMBATAREA CONGESTIEI CU AJUTORUL ALGORITMILOR RED

în vederea susținerii examenului de finalizare a studiilor universitare de licență/masterat,
organizat de către Facultatea de ELECTRONICĂ, TELECOMUNICAȚII ȘI TEHNOLOGIA INFORMAȚIEI
Departamentul ELECTRONICĂ APLICATĂ ȘI ÎNGINERIA INFORMAȚIEI din cadrul Universității
POLITEHNICA din București, sesiunea IULIE, anul universitar 2014.

Luând în considerare conținutul art. 143 din Legea Educației Naționale nr. 1/2011,
al. (4) „Îndrumătorii proiectelor de diplomă/lucrărilor de disertație răspund în solidar cu
autorii acestora de asigurarea originalității conținutului acestora” și al. (5) „Este
interzisă comercializarea de lucrări științifice în vederea facilitării falsificării de către
cumpărător a calității de autor al unui/unei proiect de diplomă/lucrări de disertație”,
declar pe proprie răspundere, că acest/această proiect/lucrare este original(ă), fiind
rezultatul propriei activități intelectuale, nu conține porțiuni plagiate, iar sursele
bibliografice au fost folosite cu respectarea legislației în vigoare.

Cunosc faptul că plagiatul sau prezentarea unui/unei proiect/lucrări, elaborat(ă) de
alt absolvent sau preluată de pe internet, din manuale și cărți, fără precizarea sursei
constituie infracțiune (furt intelectual și nerespectarea dreptului de autor și a proprietății
intelectuale) și atrage după sine anularea examenului de diplomă/disertație, precum și
răspunderea penală.

Data: 23.06.2014

Semnătura: 
(în original)

¹⁾ Declarația se completează „de mână” și reprezintă un document care se depune la înscrierea pentru susținerea examenului de finalizare a studiilor.

²⁾ Numele, inițiala prenumelui tatălui și prenumele, cu majuscule.

³⁾ Denumirea proiectului de diplomă/lucrării de disertație, cu majuscule.

Cuprins

Introducere.....	6
CAPITOLUL 1	
Elemente de bază ce definesc congestia	
1.1 Introducere.....	7
1.2 Factori ce pot influența congestia.....	7
1.3 Diferențe între controlul fluxului și controlul congestiei.....	8
1.4 Caracteristicile de bază ale controlului congestiei.....	10
1.5 Politici folosite pentru combaterea congestiei.....	12
CAPITOLUL 2	
Tehnici cunoscute de prevenire a pierderii pachetelor în rețele	
2.1 Introducere.....	14
2.1.1 Caracteristici ale gateway-ului.....	14
2.2 Protocolul TCP.....	15
2.2.1 Generalități ale protocolului TCP.....	15
2.2.2 Funcții oferite rețelei.....	16
2.3 Algoritmi TCP pentru evitarea congestiei	19
2.3.1 Algoritmul Reno.....	20
2.3.2 Algoritmul Tahoe.....	20
2.3.3 Algoritmul TCP Vegas.....	20
2.3.4 Algoritmul TCP New Reno.....	20
2.3.5 Algoritmul TCP Hybla	20
2.3.6 Algoritmul TCP BIC	21
2.3.7 Algoritmul TCP Cubic	21
2.3.8 Algoritmul de prevenire a congestiei DECbit	21
2.4 Diferențe între gateway-urile DECbit și RED.....	21
CAPITOLUL 3	
Random Early Drop	
3.1 Introducere.....	23
3.2 Identificarea unui user ce are comportament necorespunzător	28
3.3 Calculul probabilității ca pachetele să fie marcate	28
3.3.1 Metoda 1 de calcul: Variabilă aleatoare geometrică	29
3.3.2 Metoda 2 de calcul: Variabilă aleatoare uniformă	29
3.4 Calculul dimensiunii medii a cozii	30
3.5 Alegerea pragurilor minth și maxth	30
3.6 Alegerea valorilor pentru pondere.....	31
3.7 Moduri de evaluare ale algoritmului RED.....	33
3.8 Senzitivitatea parametrilor	35

3.9 Alte versiuni de implementare ale algoritmului RED.....	36
3.9.1 WRED (Weighted Random Early Detection).....	36
3.9.2 FRED (Fair Random Early Detection)	38
3.9.3 ARED (Adaptive Random Early Detection)	38
3.9.4 RRED (Robust Random Early Detection)	41
CAPITOLUL 4	
Programe si limbaje utilizate	
4.1 Network Simulator.....	42
4.1.1 Istoric.....	42
4.1.2 Instalarea simulatorului	42
4.2 Limbajul OTcl.....	43
4.3 Network Animator	43
4.4 Concluzii despre simulator.....	44
CAPITOLUL 5	
Modificări aduse algoritmului RED	
5.1 Introducere.....	45
5.2 Reprezentarea grafică a funcțiilor.....	47
5.3 Analiza fișierelor de trace.....	50
5.3.1 Structura fișierelor de trace.....	51
5.4 Rezultate și observații.....	53
Concluzii.....	60
Bibliografie.....	61
Anexa 1.....	62
Anexa 2.....	65

Introducere

În aceasta lucrare mi-am propus să studiez combaterea congestiei folosind algoritmul RED. Pe baza unor repere generale se va realiza un studiu al algoritmilor folosiți pentru combaterea congestiei, urmând apoi să discut despre principiile de bază pentru combaterea congestiei. Se vor urmări mai mulți algoritmi pentru abordarea congestiei în momentul în care aceasta a intervenit.

Am ales algoritmul RED (Random Early Detection), cunoscut și ca Random Early Drop, care are avantajele de a controla congestia la nivel de transport de pachete și că acesta funcționează cu actualele protocoale de transport (TCP/IP). De asemenea nu este nevoie ca toate ruterele din rețea să folosească același mecanism de control al congestiei. RED calculează probabilitatea de aruncare a unui packet în două etape, atunci când dimensiunea cozii scade între th_min și th_max , unde th_min reprezintă pragul de minim al cozii, th_max reprezintă pragul de maxim al cozii.

Simulatorul în care se va lucra în decursul acestui proiect este Network Simulator 2. Tot în cadrul acestui proiect, pornind de la algoritmul RED inițial, cel inventat de Sally Floyd și de Van Jacobson și urmând modificările aduse acestui algoritm prin schimbarea codului sursă, am reușit să compar anumite funcții și să obțin soluții optime totodată pentru diferitele cazuri cum ar fi variația cozii medii sau frecvența de aruncare a pachetelor chiar și întreruperile din cadrul rețelei. Totodată am reușit să compar performanțele funcțiilor introduse în codul sursă cu funcția algoritmului RED nemodificat.

În implementarea și dezvoltarea acestui proiect, un rol important l-a avut domnul profesor coordonator Ștefan Stancescu.

CAPITOLUL 1

Elemente de bază ce definesc congestia

1.1 Introducere

Definiție

Performanțele unei subrețele se vor deteriora pentru cazul în care în subrețeaua respectivă sau o parte din ea este existent un număr foarte mare de pachete. O astfel de situație putem spune că se mai numește congestie a rețelei (congestion). În figura 1.1 este prezentat un grafic ce reflectă congestia în rețelele de telecomunicație. [1]

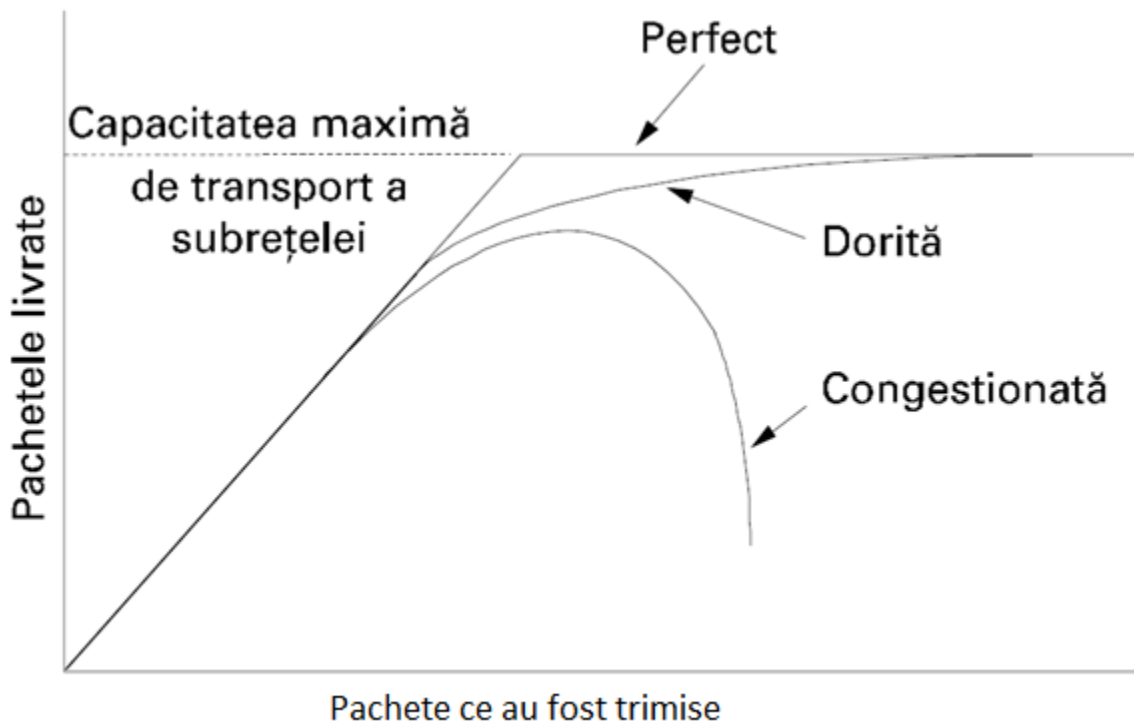


Figura 1.1 : Performanțele și congestia scad radical pentru un trafic agresiv în rețea. [1]

1.2 Factori ce pot influența congestia

Pachetele care sunt trimise de-a lungul subrețelei de calculatoare vor fi expediate într-un mod integral, dar nu și cele care conțin erori ale transmisiei și doar pentru cazul în care numărul lor nu va putea depăși capacitatea de transport din subrețeaua respectivă de calculatoare.

Numărul pachetelor care au fost expediate va trebui să fie direct proporțional cu numărul pachetelor ce sunt primite de către un nod.

Cu toate acestea dacă traficul are momente când devine din ce în ce mai agresiv, în acele momente routerelor vor începe să nu poată să mai dea randament pentru toate pachetele ce se

trimit, unele dintre ele urmând să se piardă, acest lucru putând să fie din ce în ce mai rău, pe măsură ce traficul devine tot mai mare. Cu alte cuvinte, se poate ajunge la pierderea a tuturor pachetelor, nici un pachet neputând fi expediat, iar nivelul performanțelor să devină nul pentru un trafic extrem de agresiv.

În rețelele de calculatoare, există mai mulți factori care pot produce congestia.

O coadă (sau queue) se va crea în momentul în care simultan mai multe grupuri de pachete pe trei ori patru linii de transmisie vor avea nevoie de o singură linie de ieșire.

În cazul în care nu va exista îndeajuns spațiu de memorie pentru a le păstra pe toate în coadă, atunci anumite pachete vor fi aruncate.

Neagle în anul 1987 a descoperit prin mai multe experimente că dacă se dorește a adăuga spațiu memorie suplimentară, acest lucru se poate face însă între anumite limite. Adăugarea de spațiu memorie suplimentară (în cantități enorme), infinită în routere nu ar duce la avantaje ale reducerii congestiei, ci dimpotrivă, congestia în rețeaua de calculatoare s-ar face din ce în ce mai mare ajungând să se deterioreze complet. Acest lucru se datorează faptului că pachetele au fost declarate ca fiind pierdute mai devreme ca ele să intre în coadă, iar din cauza unui timeout repetat, acestea au trimis de două ori aceeași informație (duplicat).

Datorită dimensiunii duble ce va fi trimisă către următorul router (sau nod), acest lucru va încărca rețeaua respectivă din ce în ce mai mult. Cu alte cuvinte, pe măsură ce un pachet dorește a fi transferat de-a lungul a unui anumit număr de routere cu memorie infinită, atunci când trece doar prin unul singur își va duplica dimensiunea; pentru următorul router își va duplica din nou dimensiunea, și tot așa până când rețeaua va fi încărcată total.

Un alt factor care poate duce la congestie sunt procesoarele lente care îngreunează transferul de pachete.

Pentru cazul în care CPU sau unitatea centrală a unui router prezintă performanțe scăzute, atunci aceasta va îngreuna viteza de rulare a target-urilor pe care trebuie să le îndeplinească, de exemplu actualizarea tabelului sau introducerea pachetelor în interiorul cozilor etc.

Cu toate că dacă liniile de comunicație nu vor fi utilizate pentru capacitate, se pot totodată crea cozi în rețea.

Un cu totul alt factor care determină congestia este reprezentat de lărgimea benzii scăzute a liniilor rețelei. Acest lucru se poate remedia, dar nu foarte mult, și doar dacă link-urile (sau liniile) vor fi schimbate cu altele mai performante iar procesorul să rămână același, sau pentru aceleași linii și schimbarea procesorului cu unul performant. Aceste modalități de cele mai multe ori ajută rețeaua, însă fac doar ca punctul critic să se deplaseze. Pentru a îmbunătăți într-un mod incomplet sistemul (și nu în totalitate), va trebui ca punctul critic să fie mutat într-o altă parte, însă acest lucru nu este posibil uneori, deoarece deși problema mutării punctului critic este des întâlnită, nu de fiecare dată părțile componente ale sistemului vor fi compatibile. Această problemă va rămâne neschimbată până când toate componentele vor avea un oarecare echilibru.

1.3 Diferențe între controlul fluxului și controlul congestiei

Între controlul fluxului și controlul congestiei există o oarecare diferență pe care este important să o precizăm de-a lungul discuției.

- controlul fluxului face referire la transportul și mai precis la traficul de pachete dintre locul de unde au fost trimise și locul unde au ajuns. Principala acțiune ce trebuie îndeplinită de către

controlul fluxului este a nu da voie unui expeditor să trimită pachete în continuu cu viteze mai mari decât cele cu care cel care primește poate utiliza acele date.

De cele mai multe ori, destinatarul va trimite un mesaj (de feedback) către expeditor pentru a-l înștiința despre modul cum se desfășoară transmisia în acel loc.

- controlul congestiei are rolul de a garanta ca rețeaua poate transporta tot traficul de pachete care o afectează. Aceasta putem spune că este o situație ce se desfășoară la un nivel global și totodată urmărește cum se comportă toate routerele, computerele gazdă, procesarea datelor ce sunt memorate și apoi retransmise (modul store and forward) de către routerele din rețea, și în același timp toți factorii care minimizează capacitatea de transfer a pachetelor în rețea.

Diferența dintre controlul congestiei și controlul fluxului nu este atât de ușor sesizabilă, însă putem da un exemplu în așa fel încât acest lucru să poată fi evidențiat.

Putem astfel considera o rețea formată din fibre optice cu o capacitate de transmisie de un TerraBit/secunda și un supercalculator care transmite un pachet către un simplu calculator cu viteza de 1 GB/s.

Pentru a vedea diferența dintre aceste două concepte, să considerăm o rețea cu fibre optice cu o capacitate de 1000 gigabiți/sec pe care un supercalculator încearcă să transfere un fișier către un calculator personal la 1 Gbps.

Controlul fluxului este cel care va trebui să forțeze acel supercalculator să facă opriri dese cu scopul de a lăsa calculatorul obișnuit să se poată odihni chiar dacă în rețea nu există congestive și nu este nici un fel de problemă în acest sens. Mai precis, dispozitivele nu sunt compatibile din punct de vedere al performanțelor și acest lucru îngreunează transmisia.

Pentru un alt caz ce poate evidenția diferențele între flux și trafic, putem exemplifica o rețea de tipul store and forward (sau memorează și retransmite) cu un număr de linii ce au viteze care ajung la 1 MB/secundă și un număr de 1000 de calculatoare performante dintre care 500 trimit pachete cu viteza de 100 KB/secundă către celelalte 500 de calculatoare. În acest caz, problema se datorează traficului, deoarece acesta va depăși limitele rețelei, nefiind vorba despre cazul unui destinatar lent și a unui expeditor rapid.

De cele mai multe ori controlul fluxului și controlul congestiei sunt confundate, iar acest lucru se face datorită faptului că există algoritmi ce sunt rulați pentru a controla congestia, care trimit mesaje de feed-back pentru diferiți emițători, pe care îi mai informează să nu mai trimită pachete, sau să trimită mai puține atunci când rețeaua este congestionată. Mai putem spune că datorită faptului că un destinatar nu poate suporta încărcarea informațiilor sau datorită faptului că rețeaua respectivă este încărcată la maxim, atunci un calculator care primește acele informații va primi și un mesaj pentru a încetini transmiterea. Se va mai reveni la acest aspect în cele ce vor urma.

În continuare, ne vom ocupa pe baza unor repere generale de studiul algoritmilor folosiți pentru combaterea congestiei, urmând apoi să discutăm despre modurile de bază pentru combaterea congestiei. În final, se vor urmări mai mulți algoritmi pentru abordarea congestiei în momentul în care aceasta a intervenit.

1.4 Caracteristicile de bază ale controlului congestiei

Majoritatea situațiilor ce sunt întâlnite la (sisteme) rețelele complexe de calculatoare se pot aborda urmărind teoria controlului. Urmărind acest mod de a vedea lucrurile, atunci toate soluțiile de rezolvare a situațiilor respective vor fi împărțite în două cazuri:

- În buclă închisă
- În buclă deschisă.

Urmărind aceste soluții, se constată că cele în buclă deschisă au rolul de rezolvare a situației urmărind ca proiectarea să se faca riguros, în acest fel va garanta că nu se va mai ivi nici o problemă.

Nu se vor mai putea face modificări din momentul în care rețeaua se va porni și va fi funcțională.

Pentru ca funcționarea rețelei să fie continuă, controlul congestiei se va realiza într-o buclă deschisă unde instrumentele utilizate vor seta când pachetele vor trebui să fie distruse și care dintre ele, când traficul nou va trebui să fie acceptat cât și celelalte decizii pe care vor trebui să le ia pentru diferite routere din rețeaua respectivă. Pentru realizarea acestor lucruri nu se va urmări starea prezentă pe care o are rețeaua, acest lucru reprezentând un factor comun pentru acele decizii.

Pentru cazul soluțiilor ce vor fi realizate în buclă închisă, acestea folosesc termenul de feed-back sau de reacție inversă. Acest gen de tratare a problemelor conține trei pași importanți care vor fi urmați succesiv:

1. Rețeaua este examinată pentru a vedea în ce loc și timpul când congestia se va produce;
2. Informații care au fost obținute la primul pas vor fi expediate ca mesaje către nodurile unde vor avea loc evenimente.
3. Pentru ca situația să se remedieze, funcționarea sistemului va fi minimizată.

Se poate utiliza o numeroasă gamă de metrici pentru a putea examina sistemul de legături. Procentul din totalul de pachete care au fost aruncate datorită insuficienței spațiului de memorie din router, variația întârzierii pachetelor, delay-ul (întârzierea) medie a pachetelor, dimensiunea medie a cozii sau numărul de pachete care vor urma să fie retransmise din cauza că timpul de transmisie a fost încheiat (timeout) fac parte din cele mai folosite metrici de monitorizare. Pentru oricare dintre situații, atunci când valorile unui parametru cresc, atunci va crește și congestia din rețea.

Pentru cazul pasului cu numărul doi, în bucla de reacție se va găsi transferul de informații referitoare la congestia din rețea de la depistarea sa la momentul în care se pot realiza acțiuni pentru remediere.

Așadar, există varianta de a trimite pachetele de la un router în care se găsește congestie, la alte surse în care se poate afla trafic de pachete, astfel putând face un raport pentru acea situație. Se mai poate spune că aceste pachete încarcă rețeaua, astfel acest lucru nefiind dorit, deoarece congestia va fi prezentă de-a lungul rețelei.

Mai există și alte variante în afară de cele prezentate mai sus, ca exemplu util, pentru a se completa de către noduri (routere) atunci când congestia va deveni agresivă, se poate rezerva un câmp în fiecare dintre pachete, sau mai simplu doar un bit. Se vor completa acele câmpuri pentru

toate pachetele care au fost trimise cu scopul de a informa și nodurile vecine, atunci când un router va găsi congestie pe legături.

Mai există o modalitate prin care routerele sau alte noduri în care se găsesc pachete să poată trimite din când în când pachete de test pentru a primi un timp de răspuns și în funcție de acesta să își poată da seama dacă legăturile sunt congestionate, totodată datele de acest tip sunt utilizate pentru a se evita zonele în care există congestie.

Astfel putem exemplifica elicopterele care dețin stații radio și care zboară de-a lungul orașelor cu scopul de a depista congestiile pe drumuri și a informa ascultătorii semnalului radio să evite călătoriile (sau trimiterea pachetelor) spre zonele congestionate.

Pentru toate cazurile în care se urmărește primirea unui mesaj de feed-back, de fapt se așteaptă ca informațiile primite legate de congestive vor face ca nodurile respective care primesc aceste mesaje să ia măsuri utile pentru a nu se produce congestie.

Reglarea atentă a duratelor reprezintă un factor ce trebuie luat în considerare, deoarece acest lucru face ca rețeaua să funcționeze fără probleme.

Rețeaua nu converge niciodată și oscilează într-un mod foarte radical pentru cazul în care un router va anunța mesajul de Stop în momentele în care vor ajunge consecutive câte două pachete, și apoi anunță mesajul Go (pleacă) în momentul în care coada este liberă mai mult de 20 de microsecunde.

Procesul de control al congestiei va reacționa foarte greu și nu va fi util în timp real pentru cazul în care va trebui să se aștepte de exemplu o jumătate de ora pentru a garanta înainte de a da un mesaj de Start sau de Go. Astfel, pentru a funcționa într-o manieră normală, vor fi utile anumite medieri, dar totuși, va fi un lucru greu a identifica acele constante de timp care pot fi cele mai potrivite pentru rețea.

Pentru controlul congestiei există foarte mulți algoritmi ce pot trata această problemă.

Reddy și Yang în anul 1995 au stabilit numeroase reguli de clasificare pentru algoritmi de acest tip, astfel putând realiza o organizare.

Pentru realizarea acestor reguli, au împărțit algoritmi așa cum am prezentat mai sus, în cei cu buclă închisă și cei cu buclă deschisă.

Totodată, algoritmi în buclă deschisă i-au împărțit în algoritmi care au rol asupra recepționării și algoritmi ce au rol asupra emiterii.

Algoritmi în buclă închisă au fost și ei separați la rândul lor tot în două tipuri și anume: cu feed-back explicit și cei cu feed-back implicit. Pentru algoritmi cu feed-back implicit, așa cum este timpul suficient pentru a transmite mesajele de confirmare, emițătorul poate deduce existența congestiei doar din informațiile venite de la noduri vecine.

Pentru algoritmi cu feed-back explicit, pentru a avertiza sursa respectivă despre congestie, pachetele vor fi expediate înapoi către sursa din momentul în care congestia se va produce.

Se mai poate spune că încărcarea rețelei pentru un moment, va fi relativ mare în raport cu încărcarea suportată de către componentele unui nod al rețelei, acest lucru ducând tot la congestie.

Pentru această rezolvare a acestei probleme există două posibilități:

- încărcarea pachetelor în interiorul rețelei va trebui redusă;
- adăugarea unei mai mari capacități privind resursele

Putem folosi ca exemplu o rețea care să utilizeze linii de telefonie care să ajute la o urcare progresivă a lărgimii benzii dintre anumite noduri. Creșterea lărgimii de bandă este mărită de către creșterea puterii transmisiei în cazul sistemelor în care se folosesc sateliți.

Atunci când traficul va fi separat în mai multe căi de legătură, în loc să se utilizeze cea mai bună cale, acest lucru va duce la o creștere a lărgimii de bandă. Routerile ce sunt utilizate pentru backup (folosite ca rezervă atunci când altele cedează), pentru ca rețeaua să fie mai rezistentă la defecțiuni, se vor utiliza, asigurând astfel o capacitate ridicată atunci când apare congestia agresivă.

Nu întotdeauna se poate realiza o creștere a capacității, deoarece capacitatea a crescut până la valoarea sa maximă, pentru aceasta situație, există decât o modalitate și anume reducerea încărcării cu pachete.

Pentru a putea reduce încărcarea de pachete din rețea există numeroase modalități, cum putem enumera: scăderea serviciilor pentru o parte din useri sau doar pentru anumiți; userii care refuză serviciile; sau userii care planifică cererile într-un mod care este foarte ușor previzibil.

Dintre aceste metode enumerate, doar pe anumite dintre ele le vom explica mai pe scurt. Ele se pot utiliza mult mai bine pe rețelele virtuale, iar pentru unele subrețe ce utilizează legături virtuale, acestea se vor folosi pentru nivelul rețelei, după care se va urmări nivelul de transport pentru a se realiza astfel controlul congestiei.

1.5 Politici folosite pentru combaterea congestiei

În cele ce vor urma se va face un studiu pe baza combaterii congestiei în rețele. În primul rând se vor examina rețelele care folosesc sisteme cu buclă deschisă, putând astfel să determinăm modul prin care se poate realiza controlul congestiei. Aceste sisteme cu buclă închisă sunt realizate cu scopul de a micșora congestionarea traficului. Scopul principal este de a nu da voie congestiei să se producă, aceste sisteme acționând înaintea producerii. Acest obiectiv pe care îl au se poate realiza utilizând anumite politici pentru numeroase niveluri.

În tabelul 1.1 vor fi enumerate numeroase politici utilizate pentru fiecare nivel, transport, rețea cât și legături de date care au rolul influențării congestiei în sens pozitiv sau negativ. Aceste politici au fost realizate în anul 1990 de către Raj Jain.

Vom începe cu prezentarea nivelului de legătură dintre nodurile rețelei, care este o politică inferioară, apoi vom continua cu prezentarea celorlalte niveluri.

Politica retransmisiei are rolul de a determina intervalul de timp în care se produce un timeout la sursă, și determină în acel interval de timp informații despre pachetele ce vor fi trimise. O sursă rapidă ce realizează un timeout foarte repede și apoi pune toate pachetele în coada de așteptare utilizând retransmiterea lor, va realiza supraîncărcarea în comparație cu o sursă lentă care folosește selecția atunci când are loc retransmisia pachetelor.

Politica ce stă la baza acestora este cea de memorare din zona buffer-ului. În momentul în care destinatarul va arunca toate pachetele în afara de secvența respectivă, atunci pachetele vor fi nevoie să fie retrimise utilizând o nouă încărcare.

Repetarea selectivă este o procedură mult mai sigură decât retransmiterea ultimelor pachete. Congestia mai poate fi afectată și de către politica utilizată pentru confirmare. În cazul în care oricare din pachete vor fi confirmate, atunci acestea vor produce un trafic mai mare. Se mai poate spune că dacă aceste confirmări vor fi preluate de către traficul pentru răspuns, atunci se pot realiza suplimentar noi intervale de timeout și retransmisii de pachete. Volumul de informație poate fi redus atunci când este folosită pentru controlul fluxului o schemă mult prea strânsă (ex : o fereastră prea mică).

Nivelul	Politică
Transportul	Politica retransmisiei de pachete Politica pentru memorarea temporară a pachetelor în afară de secvența respectiva sau out-of-order caching Politica utilizata pentru confirmare Politica pentru controlul fluxului Determinarea intervalului de timeout
Rețeaua	Circuitele virtuale contra datagrame în interiorul rețelei Punerea pachetelor în cozi de așteptare și politici de servire Politica distrugerii pachetelor Algoritmii de dirijare a controlului Gestiunea timpului de viață al pachetelor
Legăturile de date	Politica retransmiterii de pachete Politica pentru memorarea temporară a pachetelor în afară de secvență sau out-of-order caching Politica pentru confirmare Politica pentru controlul fluxului

Tabel 1.1 : Politici ce influențează congestia. [1]

Datorită faptului că un număr mare de algoritmi de control al congestiei rulează doar în rețele cu circuite virtuale, atunci la nivelul de rețea, alegerea dintre folosirea circuitelor virtuale și datagrame are rol de a influența congestia.

Totodată, mai este precizată ordinea prin care pachetele pot fi prelucrate (ex : metoda Round Robin). Metoda prin care se va alege care pachet să fie aruncat în momentul în care nu mai există memorie în nod, folosește politica prin care pachetele sunt distruse. Atunci când este folosită o politică bună, congestia va fi pe cale de a fi eliminată din rețea, iar pentru caz contrar, congestia va fi accentuată.

Evitarea congestiei datorată creșterii traficului de-a lungul link-urilor mai poate fi realizată și utilizând un algoritm bun ce dirijează procesele din rețea. Un algoritm mai puțin bun ar putea trimite pachetele pe aceeași linie ce este congestionată la rândul ei. Întreținerea timpului de viață a unui pachet va putea să determine cât va trăi un pachet până să fie aruncat. Dacă timpul respectiv este unul mare, atunci pachetele ce au fost pierdute vor putea încurca celelalte procese din rețea, iar în cazul în care timpul este foarte mic se vor produce timeout-uri mai devreme ca pachetele să ajungă la receptor, astfel urmând să fie retransmise.

În cazul nivelului de transport vor fi aceleași situații complicate ca la nivelul de legătură de date dar în plus la nivelul de transport aflarea intervalelor de timeout-uri va fi cu mult mai dificilă, din cauza faptului că timpul în care are loc tranzitul de pachete va fi mai greu de estimat decât timpul în care are loc tranzitul pe un link ce leagă două routere.

Atunci când intervalul de timeout va fi mult prea mic, pachete suplimentare se vor trimite fără nici un rezultat. Atunci când intervalul de timeout va fi mult prea mare, congestia din rețea va fi redusă, dar pierderea de pachete nu va afecta timpul de răspuns.

CAPITOLUL 2

Tehnici cunoscute de prevenire a pierderii pachetelor în rețele

2.1 Introducere

Prevenirea pierderilor de pachete este principalul scop al asigurării calității serviciilor (Quality of Service). [3]

Aceasta se realizează printr-un management adecvat al cozilor.

Pierderea pachetelor este cauzată din mai multe motive, dar motivele principale sunt:

- Din cauza erorilor de transmitere a pachetelor, care este un fenomen relativ rar;
- Din cauza congestiilor care apar în rețea.

Pentru evitarea și controlul congestiilor ce apar în rețea trebuie utilizate anumite mecanisme în punctele de acces și în ruterele de legătură din rețea.

TCP-ul, sau aplicațiile client se află la capătul rețelei și trebuie să utilizeze un mecanism de adaptare la condițiile rețelei. Ruterele din rețea sunt cele care asigură, cu ajutorul cozilor, un management al traficului. În arhitectura rutelor, cozile sunt programate să absoarbă vârfurile de trafic. Atunci când dimensiunea cozilor este limitată au loc scăderi și întârzieri ale pachetelor prin rețeaua respectivă. Atunci când cozile sunt pline, pachetele sunt eliminate. Tot ce este important într-o rețea este să nu se ajungă la încărcarea totală a cozilor.

Gateway-urile din rețele de mare viteză sunt proiectate cu cozi de lungime foarte mare pentru controlul congestiilor care apar. Protocolul TCP de transport va detecta o congestie numai dacă un pachet a fost aruncat de către gateway. Deși cozile au dimensiuni foarte mari, care ating limite maxime din punct de vedere al arhitecturii lor, nu este necesar ca acestea să fie pline în majoritatea timpului.

Dacă am avea de-a face cu cozi pline, atunci întârzierile medii din rețea ar crește direct proporțional. Ca urmare, atunci când vitezei rețelei crește, este din ce în ce mai utilă folosirea unor mecanisme care păstrează throughput-ul ridicat, dar dimensiunea cozilor este mică pentru aceste mecanisme, ceea ce reprezintă un oarecare dezavantaj.

Dacă nu este detectat nici un mesaj de răspuns explicit din partea gateway-ului, protocolul de transport, de nivel 4 ar putea afla despre apariția congestiei prin apariția link-urilor bottlenecks sau gâturi, cu ajutorul aruncării de pachete, cu ajutorul modificărilor care apar în throughput dar și prin modificările întârzierilor point to point.

Modificările care apar în traficul conexiunii, lipsa cunoașterii numărului de gateway-uri ce sunt afectate de congestii, schimbările de rutare posibile, modelul traficului, dar și propagarea întârzierilor limitează perspectiva pe care o are conexiunea individuală.

2.1.1 Caracteristici ale gateway-ului:

Un gateway este un nod de rețea echipat pentru interfațare cu o altă rețea care utilizează diferite protocoale de comunicare.

Gateway-ul este cel care garantează cea mai eficientă detecție a apariției congestiei. Doar gateway-ul deține o vedere unică și individuală a modului în care are loc evoluția cozilor în timp.

Gateway-ul este structurat în mai multe conexiuni cu timpi de răspuns care variază, având toleranțele față de întârzieri diferite și cerințe diferite legate de throughput.

Deciziile legate de durata și dimensiunea congestiei permise la nivelul gateway-ului sunt luate cel mai bine chiar de către gateway.

Tot gateway-ul este cel care ia cel mai bine deciziile legate de dimensiunea și de durata congestiei, astfel oferind un support în acest caz.

2.2 Protocolul TCP

2.2.1 Generalități ale protocolului TCP:

Protocolul de control al transmisiei (sau TCP - Transmission Control Protocol) este un protocol utilizat în mare parte de aplicații ce au nevoie de confirmare de primire a informațiilor. TCP creează o conectare virtuală full duplex între două puncte sau noduri terminale, fiecare punct fiind definit de o adresă IP și de către un port (TCP). [2]

TCP folosește noțiunea de numere de port sau altfel cunoscute ca port numbers pentru a identifica recepția și transmisia aplicațiilor punctelor finale (end-points) de la o gazdă, sau a prizelor de Internet sau internet sockets.

Fiecare parte dintr-o conexiune TCP are asociată cu un număr de port de 16 biți, fără semn rezervat pentru a expedia sau a primi date referitoare la aplicație. Sosirea pachetelor de date TCP este identificată ca făcând parte dintr-o anumită conexiune TCP, după priză sau *socket*, care este o combinație între adresa gazdei sursă sau portul sursă și adresa gazdei destinație sau portul de destinație. Acest lucru înseamnă că un calculator server poate susține pentru o mulțime de clienți, în același timp și mai multe servicii per client, atât timp cât fiecare client sau utilizator are în vedere a iniția în același timp mai multe conexiuni către un port de destinație din porturi sursă diferite.

Transmission Control Protocol (TCP) este unul dintre protocoalele de bază din cadrul protocoalelor folosite în Internet. TCP este unul dintre cele două componente originale ale suitei sau totalitatea protocoalelor (celalalt fiind Protocolul Internet, sau IP), astfel încât întreaga suită mai este numită și stiva TCP/IP. În special, TCP oferind încredere suficientă, asigură livrare ordonată a unui flux de octeți de la un program de pe un calculator la un alt program de pe un alt calculator aflat în rețea. Pe lângă îndatoririle sale de oferire a siguranței traficului, TCP mai controlează și mărimea segmentului de date, rata la care se face schimbul de date, evitarea congestionării traficului de rețea precum și debitul de informație. Printre aplicațiile cele mai folosite, ce utilizează TCP putem enumera transferul de fișiere (FTP), World Wide Web(WWW) și poșta electronică (e-mail).

2.2.2 Funcții oferite rețelei

Acest protocol corespunde nivelului de transport din stiva TCP/IP. TCP oferă serviciu de comunicare la un nivel intermediar dintre Protocol Internet (IP) și un program de aplicație. În momentul în care un program de aplicare dorește să trimită o bucată mai mare de date în Internet, software-ul poate emite o cerere unică pentru TCP și să lase protocolul TCP să se ocupe de detaliile de manipulare în loc să fragmenteze datele în pachete IP de dimensiuni mici și să emită o serie de cereri pentru protocolul IP.

Protocolul Internet (IP) realizează schimbul de blocuri de informații numite pachete de date. Un pachet este reprezentat de o secvență de octeți și mai este formată dintr-un antet, urmată de o secțiune de date propriu-zise. Antetul furnizează informații despre destinația unde va trebui să ajunga pachetul și, opțional, dacă este nevoie, de informații de rutare folosite pentru transmitere, până când acesta ajunge la destinația finală. Secțiunea de date conține informația de care au nevoie pachetele, pe care IP trebuie să o transmită.

Din cauza congestiilor din cadrul rețelei, încărcarea traficului, sau alte comportamente imprevizibile ale pachetelor IP, acestea pot fi pierdute, duplicate, sau livrate în altă ordine până când ajung la destinație. TCP-ul detectează problemele de aceste tipuri, solicită retransmiterea pachetelor care au fost pierdute, rearanjează pachetele într-o anumită ordine și ajută la minimizarea traficului ce are loc în rețea, în vederea reducerii apariției altor probleme.

Odată ce protocolul TCP la recepție termină de reasamblat secvența de octeți ce a fost transmisă inițial, o va păstra mai sus către programul de aplicație. Prin urmare, protocolul TCP va ascunde nivelului aplicație detaliile legate de transmise specifice nivelului inferior, din rețea.

Protocolul TCP este optimizat, în așa fel să livreze exact, realizând o livrare la timp a datelor, și totodată, TCP înregistrează rareori, întârzieri relativ mari de timp (de ordinul secundelor, ceea ce este foarte mult față de cazurile în care întârzierile sunt de ordinul milisecundelor cum ar fi normal), iar intervalul de timp de așteptare pentru unele mesaje care sosesc într-o altă ordine sau pentru retransmisia mesajelor care au fost pierdute este oarecum mare. Acest lucru nu se potrivește întotdeauna pentru aplicații în timp real, cum ar fi Voice over IP (VoIP). Astfel, pentru acest gen de aplicații, pot fi utilizate protocoalele cum ar fi Real-Time Transport Protocol (RTP), ce rulează peste User Datagram Protocol (UDP).

Protocolul TCP este un serviciu de încredere ce garantează transportul unui flux de date trimise de la o gazdă la alta fără multiplicarea sau pierderea de date. Atunci când transferul de pachete de date pe rețeaua respectivă, nu este sigur, o tehnică, care se mai numește și confirmare pozitivă cu retransmitere este utilizată pentru garantarea fiabilității transferurilor dintre pachete. Această tehnică deosebit de importantă, constă în faptul că receptorul răspunde cu un mesaj de confirmare sau ACK (acknowledgement) de fiecare dată când primește un pachet de date de la emițător. Expeditorul va păstra o copie a fiecărui pachet care a fost trimis, așteptând confirmarea destinatarului înainte de a trimite următorul pachet. Expeditorul va păstra, și un timer, atunci când pachetul a fost trimis către destinatar, și va relua retransmiterea pachetului în cazul în care

timer-ul urmează să expire iar confirmarea recepției întârzie să aibe loc. Contorul de timp folosit de timer este necesar pentru cazul în care un pachet se va pierde sau este degradat.

TCP se bazează pe un set de reguli ce constau mai mult în asocieri cu alte protocoale: pentru protocol, el este totdeauna asociat cu Protocolul Internet și pentru Protocolul Internet IP, reprezintă o metodă pentru a trimite datele, folosind un "șir de unități de mesaj" între computere prin intermediul Internetului. În timp ce IP-ul se va ocupa doar de transmiterea efectivă a datelor, TCP-ul verifică dacă un mesaj a fost distribuit în unități individuale de date, numite segmente, și să țină evidența segmentelor ce au fost transmise, pentru dirijarea eficientă cu ajutorul rețelei. De exemplu, atunci când un fișier HTML va fi trimis de la un server web, stratul software TCP din acel server împarte secvența de octeți al fișierului respectiv, în segmente și le distribuie în mod individual către stratul software IP (Nivel Internet).

Nivelul Internet mai are rolul și de a încapsula fiecare segment TCP din rețea într-un pachet IP prin adăugarea unui antet, care va adăuga adresa IP destinație printre alte date disponibile. Chiar dacă fiecare dintre pachete au aceeași adresă destinație, pachetele pot fi rutate prin căi diferite prin intermediul rețelei. În cazul în care programul client pe computerul destinație le primește, protocolul TCP (Nivelul Transport) reassemblează segmentele individuale asigurând în același timp ordonarea corectă și lipsa de erori a acestora, înainte ca pachetele să fie livrate către nivelul aplicație sau aplicație.

Modul de lucru pentru protocolului TCP implică existența a trei faze. În prima fază, conexiunea trebuie să fie stabilită (Connection establishment), urmând apoi un proces în care are loc o confirmare pe baza mai multor pași. Imediat ce conexiunea este realizată, va urma transferul datelor sau data transfer.

Odată ce transferul datelor s-a încheiat, conexiunea trebuie terminată (Connection termination) în ideea de a închide calea virtuală și de a elibera resursele hardware/software implicate în proces. [2]

O sursă TCP este considerată ca fiind saturată pentru cazul unei sesiuni FTP sau File Transfer Protocol, unde toate segmentele au dimensiunile maxime posibile, rezultând pachete de mărimi constante. Termenul de segmente înseamnă același lucru cu termenul pachete. TCP-ul conține o fereastră de congestie și o limită pentru „slow-start” sau început lent, ce sunt descrise de variabilele CWND sau congestion window și Ssthresh sau slow start threshold (pragul începutului lent).

Congestion window (CWND) are rolul de a determina numărul maxim admisibil de pachete ce sunt neconfirmate în rețeaua respectivă, iar slow start threshold (Ssthresh) are rolul de a controla creșterea dimensiunii ferestrei de congestie.

Dacă fereastra de congestie este mai mică decât pragul începutului lent, în faza de „start lent”, fereastra de congestie este incrementată cu câte un pachet pentru fiecare confirmare neduplicată ce a fost primită. Această lucră duce la o creștere exponențială a ferestrei de congestie. Când Ssthresh este mai mic sau egal cu CWND, adică în faza de evitare a

congestiei, CWND crește cu câte un pachet la fiecare RTT (round trip time) sau timp dus-întors, ceea ce corespunde unei creșteri liniare a ferestrei de congestie (Figura 2.1). [3]

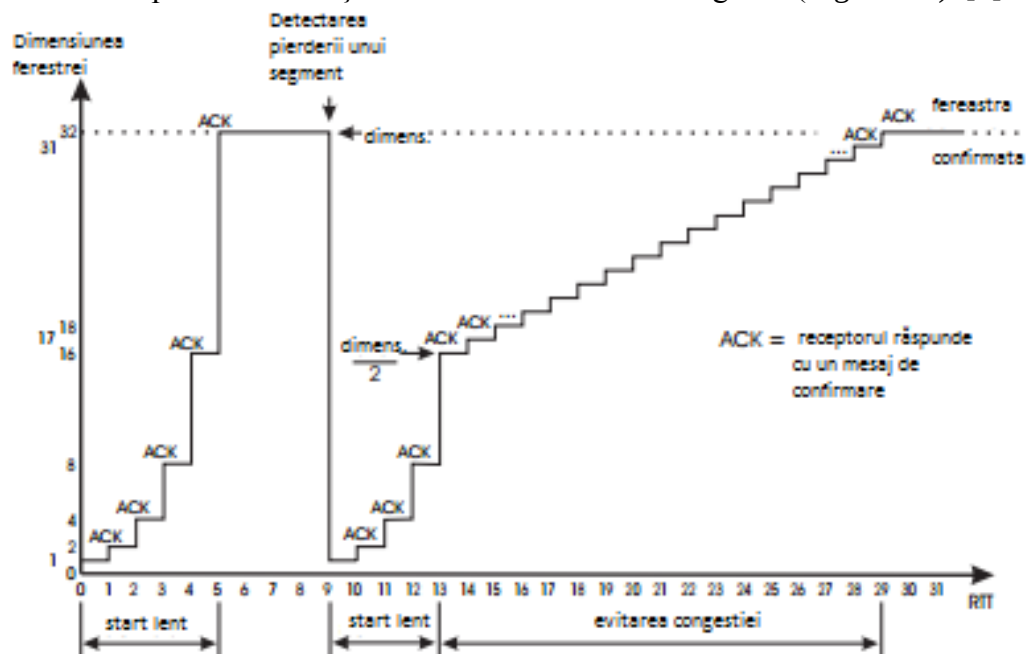


Figura 2.1 : Ilustrare simplificată a startului lent al protocolului TCP pentru algoritmul de evitare a congestiei [3]

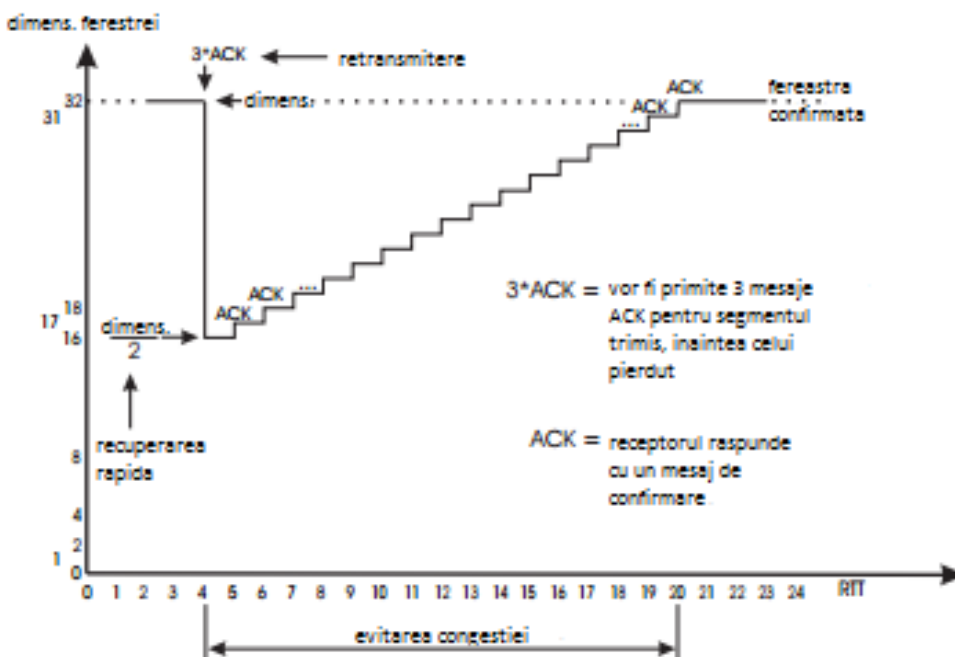


Figura 2.2 Ilustrare simplificată a startului lent al protocolului TCP pentru algoritmul de retransmitere rapidă [3]

Protocolul TCP are doua moduri diferite de a reactiona atunci când au loc pierderi de pachete în rețea. Acesta setează câte un cronometru pentru fiecare pachet trimis. Pentru cazul în care nu va fi primit un mesaj de confirmare pentru acel pachet, înainte de a expira timpul, atunci fereastra de congestie își reduce dimensiunea la mărimea maximă a segmentului (sau pachetului) iar pragul începutului lent este setat la valoarea maximă (adica $SSTHRESH = flight_size / 2, 2 * dimensiunea\ segmentului$). Flight size este cantitatea de informație care a fost neconfirmată în rețea.

Dacă s-ar presupune că emițătorul este saturat, valoarea ferestrei de congestie ar fi egală cu valoarea cantității de informație care a fost neconfirmată.

Atunci când emițătorul recepționează trei ACK-uri ce sunt duplicate, înainte ca timpul să expire, va fi realizată o recuperare rapidă sau fast recovery, mai precis într-un mod mai simplificat, protocolul TCP va trimite pachetul pierdut din nou, iar valoarea pragului începutului lent se va seta la valoare maximă ($flight_size/2, 2 * dimensiunea\ segmentului$) și astfel se va seta și valoarea ferestrei de congestie la noua valoare calculată (Figura 2.2) [3].

2.3 Algoritmi TCP pentru evitarea congestiei

Pentru a evita colapsul congestiei (atunci când buffer-ul este plin și are loc pierderea de pachete în șir), TCP folosește o strategie multi-fețe de control al congestiei. Pentru fiecare conexiune, TCP menține o fereastră de congestie, limitând numărul total de pachete nerecunoscute care pot fi în tranzit end-to-end. Acest lucru este oarecum analog cu fereastră glisantă a TCP-ului folosită pentru controlul fluxului. [4]

Cu alte cuvinte, TCP folosește mecanismul numit start lent pentru a mări fereastra de congestie, după ce o conexiune este inițializată și după un timeout. Acesta începe cu o fereastră de două ori mai mare decât dimensiunea maximă a segmentului. Deși rata inițială este scăzută, rata de creștere este una foarte rapidă: pentru fiecare pachet recunoscut, fereastra de congestie crește cu lungimea maximă a segmentului, astfel încât fereastra de congestie se dublează într-un mod eficient de fiecare dată (dus-întors).

Când fereastra de congestie depășește pragul de început lent, algoritmul intră într-o stare nouă, numită evitare a congestiei. În unele implementări (de exemplu, Linux), pragul de început lent, inițial este mare, și astfel primul start lent, de obicei, se termină după o pierdere. Cu toate acestea, pragul de început lent este actualizat la sfârșitul fiecărui start lent, și va afecta de multe ori start-urile lente ulterioare declanșate de timeout.

Controlul congestiei face referire de fapt la controlarea traficului ce pătrunde într-o rețea de telecomunicații, acest lucru presupunând evitarea supraîncărcării legăturilor dintre nodurile intermediare ale rețelei, numite și link-uri, prin reducerea pachetelor ce vor fi trimise. Așa cum a fost prezentat anterior, controlul congestiei nu trebuie să fie confundat cu controlul fluxului (flow control) ce evită supraîncărcarea receptorului cu date de către emițător.

Există mai mulți algoritmi care previn și în același timp evită congestia, prevăzuți de protocolul TCP. Primii algoritmi apăruiți în domeniul telecomunicațiilor sunt TCP Tahoe, și TCP Reno.

Ambii algoritmi au la bază regulile de „start lent” și „retransmitere rapidă” prezentate mai sus. Cei doi algoritmi, Tahoe și Reno, diferă datorită modului în care reacționează atunci când au loc pierderi de pachete:

2.3.1 Algoritmul Reno

Reno înjumătățește fereastra de congestie, efectuând o retransmitere rapidă (fast retransmit), și intrând totodată în recuperare rapidă (fast recovery) pentru cazul în care sunt recepționate 3 ACK (acknowledged) duplicat.

2.3.2 Algoritmul Tahoe

Atunci când timpul de timeout urmează să expire înainte ca o confirmare să fie recepționată, pierderea este detectată. Din acel moment, algoritmul Tahoe va reduce fereastra de congestie la valoarea maximă a dimensiunii unui segment sau pachet (MSS – maximum segment size) și se resetează la faza de „slow-start” sau început lent.

2.3.3 Algoritmul TCP Vegas

În jurul anilor 1990, toate întârzierile dus-întors și time-outurile erau bazate doar pe ultimul pachet din buffer.

Cercetătorii Lawrence Brakmo și Larry Peterson de la Universitatea din Arizona, au introdus algoritmul TCP Vegas în care întârzierea dus-întors este măsurată pentru fiecare pachet din buffer-ul respectiv și timeout-ul este setat. Totodată, acest algoritm folosește și o creștere aditivă a CWND. Această variantă de algoritm nu a fost atât de mult improvizată în afara laboratorului cercetătorilor.

2.3.4 Algoritmul TCP New Reno

Acest algoritm vine ca un upgrade pentru algoritmul simplu Reno, oferind o îmbunătățire a retransmisiei în timpul fazei de recuperare rapidă sau fast recovery pentru TCP Reno. Pe parcursul fast recovery, pentru fiecare acknowledgement duplicat, va fi transmis un nou pachet care nu a fost trimis, de la sfârșitul ferestrei de congestie CWND, pentru a menține fereastra de congestie plină. Pentru fiecare acknowledgement parțial din spațiul de secvențe numerice, emițătorul va trimite pachetul următor cu o secvență numerică următoare acknowledge-ului.

Datorită faptului că acel cronometru de timeout va fi resetat ori de câte ori va exista un anumit progres în buffer-ul de transmisie, algoritmul New Reno poate umple găuri oarecum mari sau mai multe găuri din spațiul de secvențe.

Datorită faptului că algoritmul New Reno poate trimite pachete noi până la sfârșitul ferestrei de congestie CWND în timpul fast recovery, este menținut un throughput relativ mare pe durata procesului de umplere a găurilor, chiar dacă vor exista mai multe găuri, având mai multe pachete fiecare. Atunci când TCP intră în fast recovery, acesta va înregistra cea mai mare secvență numerică neconfirmată. Dacă această secvență numerică este confirmată, atunci TCP va reveni la starea de evitare a congestiei.

Dacă nu au loc pierderi de pachete, dar cu toate acestea, pachetele vor fi reordonate cu mai mult de trei secvențe numerice, va apare o problemă în algoritmul New Reno. Atunci când acest fenomen are loc, New Reno intră greșit în fast recovery, însă când pachetul reordonat va fi livrat, va avea loc progresul secvenței numerice acknowledged iar din acel moment până la sfârșitul fast recovery, fiecare bit din secvența numerică se va duplica și va oferi retransmisia necesară.

2.3.5 Algoritmul TCP Hybla

Are rolul de eliminare a penalizării conexiunilor TCP care încorporează legături terestre cu latență de dimensiuni foarte mari, sau de legături radio, datorită timpilor de dus-întors care dețin

valori relativ mari. Algoritmul este provenit dintr-o evaluare analitică a dinamicii ferestrei de congestie CWND, care recomandă înlăturarea dependenței performanței față de timpul dus-întors (RTT – round trip time).

2.3.6 Algoritmul TCP BIC (binary increase congestion control)

Acest algoritm este o implementare a TCP-ului cu un algoritm de evitare a congestiei, optimizat pentru rețelele de viteze mari, având latențe mari (numite long fat network sau LFN). Algoritmul BIC este folosit decât în kernel-urile Linux, de la 2.6.8 până la 2.6.18.

2.3.7 Algoritmul TCP Cubic

Algoritmul prezintă o versiune mai puțin agresivă a algoritmului BIC, în care fereastra de congestie este o funcție cubică de timp, de la momentul ultimei congestii, având punctul de inflexiune setat la fereastra anterioară a evenimentului respectiv.

2.3.8 Algoritmul de prevenire a congestiei DECbit

Pentru cazul folosirii algoritmului DECbit, gateway-ul utilizează un bit în antetul pachetului pentru a furniza un mesaj de feedback rețelei legat de apariția congestiei. În momentul în care un pachet ajunge la gateway, aceasta va calcula lungimea medie a cozii pentru ultima perioadă (adică ocupat+inactiv) plus perioada de “ocupat” din momentul respectiv (gateway-ul este ocupat atunci când sunt transmise pachete, dacă nu, el va fi inactiv). [5]

Atunci când lungimea medie a cozii depășește valoarea de unu, gateway-ul va seta bitul de indicare a congestiei în antetul pachetelor ce urmează să sosească.

Sunt folosite ferestre pentru controlul fluxului de către sursă, pe care le updatează o dată pentru fiecare două perioade dus-întors. Pentru cazul în care cel puțin jumătate dintre pachetele din ultima fereastră ar fi avut bitul de identificare al congestiei setat, atunci fereastra ar fi fost scăzută exponențial. În mod contrar, fereastra este crescută liniar.

2.4 Diferențe între gateway-urile DECbit și RED

Între gateway-urile care au implementat algoritmul DECbit și gateway-urile care au implementat algoritmul RED sunt câteva diferențe semnificative [5]:

- Prima diferență face referire la metoda de calcul a dimensiunii medii a cozii. Din cauza faptului că DECbit folosește ultimul ciclu (ocupat+inactiv) plus perioada de “ocupat” curentă pentru a calcula dimensiunii medii a cozii, dimensiunea cozii va putea fi mediată după o perioadă de timp destul de scurtă, ceea ce reprezintă un avantaj.

Pentru rețele de viteză mare, cu buffere foarte mari la nivel de gateway, se dorește ca utilizarea constantei de timp folosite pentru calculul dimensiunii medii a cozii, să fie făcută explicit. Acest lucru nu este posibil pentru algoritmul DECbit, însă este realizat în cazul gateway-urilor RED. În rețele DECbit, sursele urmaresc, astfel monitorizând fracțiunea de pachete ce au fost marcate în ultimul interval, adică cel de dus-întors. Pentru cazul rețelelor cu gateway RED, sursa își va reduce fereastra chiar dacă există un singur pachet care a fost marcat.

- Cea de-a doua diferență dintre cei doi algoritmi, este metoda prin care gateway-ul își alege conexiunile pe care le va înștiința despre apariția congestiei ce va urma. La DECbit nu este

prezentată nici o diferență conceptuală dintre algoritmul pentru setarea bitului de congestie și algoritmul pentru detecția congestiei.

Dacă un pachet ajunge în gateway și dimensiunea medie a cozii este mult prea mare, atunci bitul pentru identificarea congestiei este setat în antetul pachetului respectiv. Datorită acestei metode de marcare a pachetelor, rețele care folosesc algoritmul DECbit pot întâlni probleme în fața traficului fluent, acest lucru fiind evitat pentru cazul gateway-urilor RED, aici marcarea pachetelor făcându-se într-un mod aleator.

Pentru gateway-urile care evită congestia în rețele, destinate pentru folosirea protocolului TCP, motivația în plus pentru ca marcarea să se facă aleator este evitarea sincronizării globale care poate să rezulte în urma unei reduceri în același timp a mai multor conexiuni TCP și a ferestrei de comunicație. Aspectul acesta nu este important decât pentru cazul folosirii algoritmului DECbit, datorită faptului că fiecare sursă își micșorează moderat fereastra ca răspuns la apariția congestiei în rețea.

CAPITOLUL 3

Random Early Drop

3.1 Introducere

Algoritmul RED (Random Early Detection), cunoscut și ca Random Early Drop, Random Early Discard, sau chiar Random Early Delete este un mecanism de prevenire și în același timp de evitare a congestiei. [6]

Acest algoritm deține câteva metode care folosesc pentru alegerea conexiunii care să notifice această congestie și pentru detectarea congestiei. Ruterele folosite de algoritmul RED sunt utilizate în cea mai mare parte cu rețele TCP/IP, fiind proiectate pentru rețele în care un singur pachet aruncat sau marcat este de ajuns pentru a semnaliza congestia la nivel de transport. [5] Sally Floyd și Van Jacobson sunt inventatorii acestui algoritm. Ei au venit în anul 1993 cu o lucrare despre “detectia congestiei în rețele folosind ruterele, folosind un mecanism de detecție aleatoare timpurie”.

Algoritmul Random Early Detection (RED) reduce congestia în cozi aruncând pachetele astfel încât unele conexiuni TCP pe o anumită perioadă limitată de timp vor transmite mai puține pachete prin rețea. RED în mod intenționat va arunca o parte din pachete înainte ca bufferul se umple total, în loc să aștepte până când o coadă se supraîncarcă, cauzând un număr mare de pachete aruncate. Această acțiune este destinată să facă emițătorii care generează traficul să reducă volumul de pachete transmis. Se poate spune că acesta este avantajul de baza al utilizării acestui algoritm. [7]

Numele acestui algoritm Random Early Detection sau detecție aleatoare timpurie descrie modul de utilizare al algoritmului. Atunci când a luat decizia de descărcare de pachete, algoritmul RED selectează într-un mod aleator pachete pe care le descarcă. [5]

Există două aspecte generale care descriu logica de execuție a acestui algoritm:

- RED va trebui în primul rând să detecteze congestia propriu-zisă înainte ca aceasta să se întâmple, sau mai bine spus înainte ca buffer-ul să se umple. Altfel spus, RED trebuie să ia decizia pentru alegerea condițiilor ce duc la aruncarea de pachete.

Atunci când se ia această decizie, va trebui să fie bine definit numărul de pachete care va trebui aruncat. În primul rând RED va trebui să verifice cât de mare este coada dispozitivului.

- algoritmul RED calculează lungimea medie pentru interfața respectivă, urmând apoi să decidă dacă va avea loc fenomenul de congestie sau nu. RED folosește încărcarea medie pentru că încărcarea reală își schimbă valoarea mult mai des. Datorită faptului că RED evită efectele sincronizării dintre pachete, este util ca aruncarea pachetelor să se facă într-un mod proporțional în funcție de cât de încărcat este buffer-ul, evitând astfel aruncarea pachetelor de număr extrem de mare la un moment dat și extrem de mic într-un alt moment.

Avantajul acestui algoritm de control al congestiei la nivelul de transport de pachete este că acesta funcționează cu actualele protocoale de transport (TCP/IP), și în al doilea rând nu este nevoie ca toate ruterele din rețea să folosească același mecanism de control al congestiei. Mecanismul acesta este ușor din punct de vedere al implementării, de evitare a congestiei poate fi folosit în actualele rețele TCP/IP nefiind nevoie de schimbarea protocolului de transport.

Prevenirea apariției congestiei într-o rețea cu ajutorul controlului dimensiunii medii a cozii pachetelor respective reprezintă scopul principal pe care ruterele RED vor trebui să îl îndeplinească.

Algoritmul RED este cel care calculează dimensiunea medie a cozii, utilizând un filtru trece-jos, acesta având o medie ponderată exponențială, dinamică sau mobilă.

Putem compara aceasta dimensiune medie a cozii folosind două praguri și anume un prag care este minim (minimum thresh) și un prag care este maxim (maximum thresh). În acest caz al folosirii pragurilor, când dimensiunea medie a cozii deține o valoare mai mică decât minimum thresh, nu este marcat nici un pachet. Dacă în cazul în care dimensiunea medie a cozii deține o valoare mai mare decât maximum thresh, atunci oricare dintre pachetele care sosesc va fi marcat. Atunci când pachetele care au fost marcate sunt de fapt aruncate, sau cu alte cuvinte dacă toate nodurile principale cooperează între ele, aceasta metodă poate să asigure faptul că dimensiunea medie a cozii să nu poată să dețină o valoare mult mai mare decât maximum thresh.

Probabilitatea de aruncare de pachete

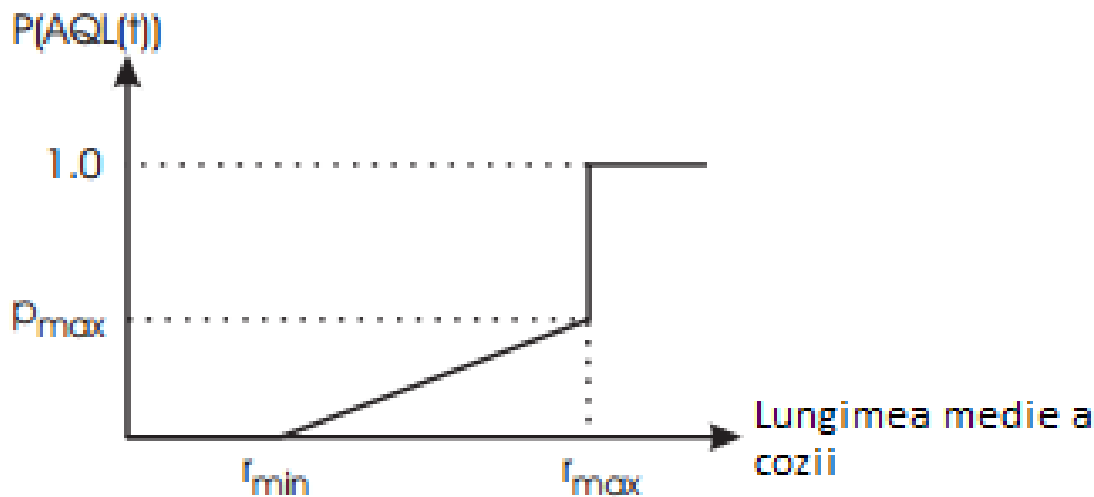


Figura 3.1: Creșterea probabilității de aruncare de pachete folosită de algoritmul RED [3]

Dacă dimensiunea medie a cozii se situează între valoarea de minimum thresh și maximum thresh, atunci fiecare dintre pachetele ce ajung la destinație vor fi marcate cu probabilitatea P_a , acest parametru reprezentând o funcție ce este dependentă de dimensiunea medie a cozii avg .

Ori de câte ori un pachet este marcat, atunci probabilitatea ca pachetul respectiv să facă parte dintr-o conexiune anume, va fi proporțională cu rata pe care conexiunea respectivă o are din lungimea totală de bandă care trece prin gateway. [5]

În cele ce urmează va fi prezentat algoritmul general RED pentru gateway-uri:

Pentru fiecare pachet care a sosit

Se va calcula dimensiunea medie a cozii (avg)

Dacă $th_{max} \leq avg \leq th_{min}$

Calculează probabilitatea P_a

Marchează pachetul care a sosit cu probabilitatea P_a
 Altfel dacă $th_max \leq avg$
 Marchează pachetul care a sosit

Putem spune ca acest algoritm RED pentru gateway-uri este format din 2 algoritmi separați unul față de celălalt:

- Gradul maxim de trafic de pachete sau mai bine spus de pachete care sosesc simultan, care va fi permis în coada gateway-ului este determinat de algoritmul care are rolul de a calcula dimensiunea medie a cozii.
- La un anumit nivel al congestiei algoritmul ce calculează probabilitatea ca un pachet sa fie marcat, va determina cât de des gateway-ul poate să marcheze pachetele . Rolul este ca gateway-ul să poată marca pachetele sosite în perioade de timp egale, cu ajutorul acestui lucru reușind să evite marcarea pachetelor îndeajuns de des și sincronizarea globală, putând astfel să obțină controlul dimensiunii medii a cozii.

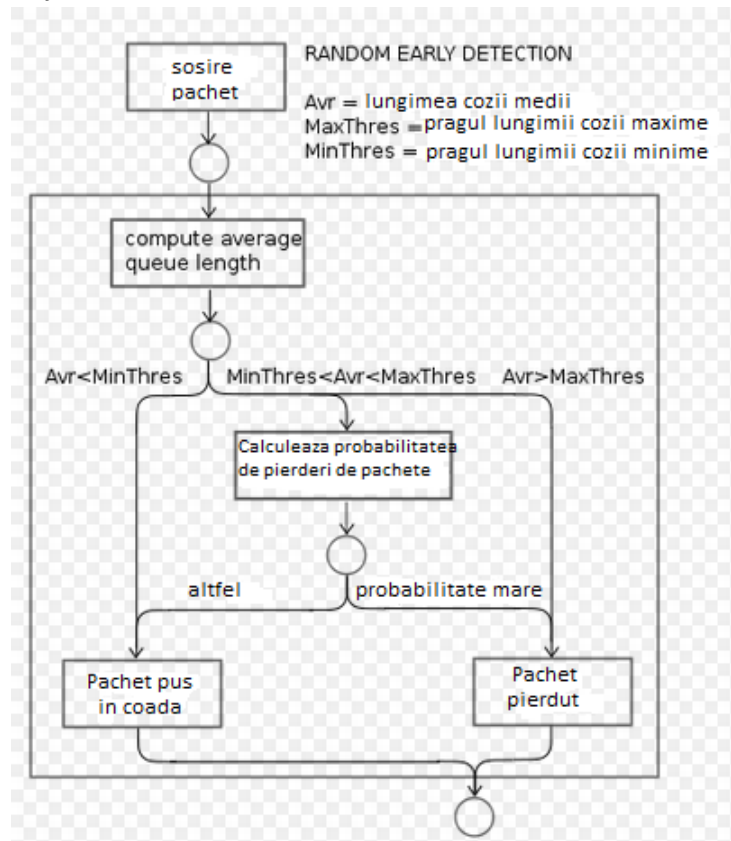


Figura 3.2: Schema algoritmului RED [6]

Intervalul de timp cât coada are valoare nulă (perioadă inactivă) folosește pentru a calcula dimensiunea medie a cozii, totodată estimând numărul de pachete m ce ar fi putut fi transmise de către gateway tot pe durata acestui interval. După perioada inactivă sau idle, gateway-ul calculează dimensiunea medie a cozii presupunând ca m pachete ar fi ajuns pe durata acestei perioade idle într-o coadă nulă. [5]

Valoarea cozii medii variază în funcție de cum sunt alese pragurile de minim (minimum thresh) și de maxim (maximum thresh), iar probabilitatea cu care pachetele sunt marcate (P_b) variază între valoarea nulă (0) și o valoare maximă a probabilității ($\max_p$):

$$P_b = \max_p * \frac{\text{avg} - \text{th}_{\min}}{\text{th}_{\max} - \text{th}_{\min}}$$

P_a , care este probabilitatea finală care marchează pachetele poate crește treptat, pe măsură ce număratoarea pachetelor crește în raport cu pachetul care a fost marcat cel mai devreme:

$$P_a = \frac{P_b}{1 - \text{count} * P_b}$$

Folosind aceasta modalitate, putem garanta că gateway-ul va aștepta puțin timp pentru marcarea unui pachet nou.

În momentul în care dimensiunea medie a cozii va depăși valoarea pragului maxim, atunci gateway-ul marchează fiecare pachet care ajunge în coadă.

Mai exista o altă modalitate pentru algoritmul RED de măsurarea a cozii în biți, renunțând astfel la măsurarea cozii în pachete. Totodata, avg sau dimensiunea medie a cozii va indica într-un mod precis valoarea întârzierii medii care apare în gateway. Algoritmul va fi modificat pentru a asigura că probabilitatea cu care un pachet va fi marcat este proporțională cu dimensiunea pachetului în bytes, acest lucru întâmplându-se doar pentru cazul în care va fi folosită acea modalitate.

$$P_b = \max_p * \frac{\text{avg} - \text{th}_{\min}}{\text{th}_{\max} - \text{th}_{\min}}$$

$$P_a = \frac{P_b}{1 - \text{count} * P_b}$$

$$P_b = P_b * \frac{\text{dimensiune_pachet}}{\text{dimensiune_maxima_pachet}}$$

Inițializare:

avg $\leftarrow$ 0

count $\leftarrow$ -1

Pentru fiecare pachet ce a fost utilizat

Calculează dimensiunea nouă a cozii medii avg:

if avg $\neq$ 0

avg $\leftarrow$ (1- w_q)*avg + w_q *q

else

m $\leftarrow$ f_time – q_time

```

    avg ← 1-wqm * avg
    if thmin ≤ avg < thmax
        count va fi incrementat
        Probabilitatea Pa va fi calculată
        
$$P_b = \max\_p * \frac{avg - th_{min}}{th_{max} - th_{min}}$$

        
$$P_a = \frac{P_b}{1 - count * P_b}$$

        Pachetul care a ajuns la destinație va fi marcat cu probabilitatea Pa
        count ← 0
    else thmax ≤ avg
        Pachetul care a sosit este marcat:
        count ← 0
    else count ← -1
Dacă nu mai sunt pachete în coadă:
    q_time ← time

```

Aceasta este o variantă mai detaliată a algoritmului RED [5]. Totodată putem observa că în acest mod, ca un pachet FTP (File Transfer Protocol) față de un pachet TELNET va putea avea șanse mult mai mari să fie marcat.

Variabilele care au fost salvate:

q_time: intervalul de timp al începerii perioadei inactiv (idle)
 avg: dimensiune medie a cozii
 count: pachetele provenite de la ultimul pachet care a fost marcat

Parametrii fixați:

th_{min}: pragul de minim al dimensiunii cozii
 th_{max}: pragul de maxim al dimensiunii cozii
 max_p: valoarea maximă pe care o poate avea P_b
 w_q: ponderea cozii

Alți parametri:

q: mărimea curentă a cozii
 p_a: probabilitatea curentă de marcare a pachetelor
 f(t): funcție liniară de timp t
 time: timp curent

În funcție de de mărimea traficului pachetelor din coadă ce este permis de către gateway și pentru cazul în care există un trafic agresiv, și de durata acestuia, se poate determina ponderea cozii (w_q). Dimensiunea medie a cozii dorite este cea care determină valorile pentru pragurile de minim al dimensiunii cozii și de maxim al dimensiunii cozii. Dimensiunea medie a cozii care poate stabili

compromisuri dorite (de exemplu cum este compromisul dintre minimizarea intervalului de timp al întârzierii și maximizarea throughput-ului) este dependent de atribuțiile rețelei.

3.2 Identificarea unui user ce are comportament necorespunzător

Detecția congestiei cu ajutorul marcării de pachete este realizată cu ajutorul routerelor RED care pot notifica toate conexiunile dintre noduri. Există un anumit procent din lărgimea benzii pe care îl folosește conexiunea respectivă în momentul respectiv, care este direct proporțional cu probabilitatea ca un pachet dintr-o anumită conexiune să fie marcat, spre deosebire de Tail Drop unde nu sunt respectate aceste reguli.

Routerele care folosesc algoritmul RED pot oferi o implementare eficientă pentru a se putea realiza identificarea legăturilor dintre noduri care folosesc un procentaj mare din lărgimea de bandă în momentul când se produce congestia. Aceste routere au rolul de a identifica foarte ușor care din legăturile disponibile în rețea a primit o mare parte din ultimele pachetele ce au fost marcate, din cauza faptului că ruterele RED vor alege într-un mod random pachetele ce vor urma să fie marcate în momentul în care are loc congestia.

Din momentul în care numărul de pachete ce au fost marcate va fi îndeajuns de mare pentru acea legătură, atunci conexiunea care a acceptat un procent mare de pachete marcate va fi cel mai posibil să fie conexiunea ce a utilizat cel mai mult lungimea de bandă disponibilă. Putem spune ca această indicație ar putea fi folosită de nivele foarte avansate cu scopul restricționării lărgimii de bandă a legăturilor în timpul în care congestia are loc.

Totodată, ruterele RED au posibilitatea păstrării cu ușurință a unei liste ce conține cele n pachete dintre cele care au fost marcate recent. Pentru cazul în care o parte foarte mare din pachetele ce au fost marcate fac parte dintr-o conexiune anume, va fi cel mai probabil ca acea legătură să aibă un procentaj mare din lărgimea medie a benzii.

În momentul în care unele dintre conexiunile TCP vor primi procente relativ mari din lărgimea de bandă, atunci este mult posibil ca acele conexiuni să fie surse ce nu respectă protocoalele TCP actuale, ori mai poate exista o conexiune ce are un interval de dus-întors mai scurt sau o fereastră mult mai mare decât cele folosite pentru alte conexiuni.

Urmărind fiecare dintre cazuri, pentru toate, ruterul folosit de algoritmul RED are posibilitatea de a fi programat astfel încât să poată oferi o prioritate mai mică acelor conexiuni care acceptă un procent mai mare din lărgimea de bandă în intervalul de timp al producerii congestiei în rețea.

3.3 Calculul probabilității ca pachetele să fie marcate

Probabilitatea ca pachetele să fie marcate P_b poate fi calculată ca o funcție liniară de către lărgimea cozii medie. În cele ce vor urma vor fi comparate două metode de calcul al probabilității finale cu care pachetele sunt marcate cât și avantajele folosirii metodei a doua.

Pentru cazul folosirii primei metode, atunci când lărgimea medie a cozii are valori constante, numărul de pachete ajunse la destinație între 2 pachete ce au fost marcate, va fi de fapt o variabilă geometrică aleatoare.

Pentru cazul utilizării celei de-a doua metodă, numărul de pachete care au ajuns la destinație între 2 pachete marcate va fi o variabilă aleatoare uniformă.

Probabilitatea inițială de marcarea a pachetelor este calculată cu ajutorul relației

$$P_b = \max_p * \frac{avg - th_{min}}{th_{max} - th_{min}}$$

Unde $\max_p$ este valoarea maximă pentru probabilitatea cu care pachetele sunt marcate (P_b) când lățimea medie a cozii ajunge la th_{max} .

3.3.1 Metoda 1 de calcul: Variabilă aleatoare geometrică

Dacă s-ar face o presupunere cum că fiecare pachet va fi marcat cu o probabilitate P_b considerând că timpul dintre marcări (notat X) fiind dat de numărul de pachete care ajung la destinație după ce un pachet a fost marcat, până când următorul pachet este marcat. Cu toate acestea fiecare dintre pachete vor fi marcate cu probabilitatea P_b :

$$\text{Prob } X = n = 1 - P_b^{n-1} * P_b$$

Scopul acestei metode este ca pachetele să fie marcate la intervale îndeajuns de regulate folosind o lungime medie a cozii constantă. Nu este necesar ca intervalele de timp dintre marcările ce au avut loc să fie foarte mari și totodată nu este de dorit să avem atât de multe pachete marcate foarte apropiate. Cele două variante pot avea loc atunci când X este o variabilă aleatoare geometrică, iar totodată cele două variante enumerate pot duce la sincronizări globale, având mai multe conexiuni, astfel reducând ferestrele în același interval de timp.

Parametrul X este o variabilă aleatoare geometrică, având la rândul său ca parametrii P_b și $E[X] = \frac{1}{P_b}$.

3.3.2 Metoda 2 de calcul: Variabilă aleatoare uniformă

Folosind această metodă care este o variantă mai bună, X fiind o variabilă aleatoare uniformă din mulțimea $\{1, 2, \dots, \frac{1}{P_b}\}$, presupunând că $\frac{1}{P_b}$ este un întreg.

X va fi o variabilă aleatoare uniformă doar atunci când probabilitatea ca oricare din pachetele ce au sosit să fie marcate să fie $\frac{P_b}{1 - count * P_b}$, $count$ fiind numărul de pachete ce nu au fost marcate care au ajuns la destinație de când ultimul pachet a fost marcat.

Pentru această situație:

$$\text{Prob } X = n = \frac{P_b}{1 - (n-1) * P_b} * \prod_{i=0}^{n-2} (1 - \frac{P_b}{1 - i * P_b})$$

$$\text{Prob } X = P_b \text{ pentru } 1 \leq n \leq \frac{1}{P_b}$$

$$\text{Prob } X = 0 \text{ oentru } n > \frac{1}{P_b}$$

Pentru metoda a doua $E[X] = \frac{1}{2 \cdot P_b} + \frac{1}{2}$

Va fi prezentat un experiment in figura ce va urma in care vor fi comparate cele două metode unde pachetele sunt marcate:

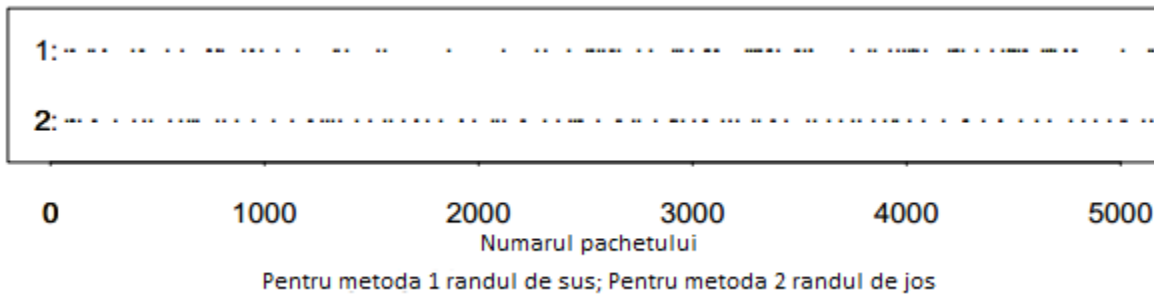


Figura 3.3: Comparație între metodele de marcare (aleator) [5]

Pentru linia de sus ce reprezintă prima metodă, oricare pachet va fi marcat cu o probabilitate P , unde P ia valoarea 0.02. Pentru linia de jos ce reprezintă a doua metodă, oricare pachet va fi marcat cu o probabilitatea $\frac{P}{1+i \cdot P}$, unde P va lua valoarea 0.01 iar numărul de pachete nemarcate i de la ultimul pachet care a fost marcat.

Fiecare din cele doua metode au marcat aproximativ 100 de pachete din 5000 care au sosit la destinație. Pe axa x este reprezentat numărul de pachete, iar pentru ambele metode este reprezentat pe figură câte un punct pentru fiecare pachet ce a fost marcat. Totodată, pachetele marcate sunt mult mai grupate pentru prima metodă, decât pentru cazul folosirii metodei a doua.

3.4 Calculul dimensiunii medii a cozii

Se va folosi un filtru trece jos pentru a calcula dimensiunea medie a cozii. Cu toate acestea, creșterile pe durată scurtă ale lărgimii cozii datorate traficului agresiv nu duc la creșterea foarte mare a dimensiunii medii a cozii.

Mai putem spune că filtrul trece jos are rolul de medie mobilă exponențială ponderată:

$$\text{avg} \leftarrow 1 - w_q \cdot \text{avg} + w_q \cdot q$$

Constanta funcție de timp a filtrului trece jos este determinată de către ponderea w_q . În cele ce vor urma se va discuta despre limitele posibile pentru ponderea w_q .

3.5 Alegerea pragurilor minth și maxth

Pragurile de minim și de maxim au valorile optime care depind în principal de avg, lungimea medie a cozii. Pentru cazul în care traficul este unul agresiv, atunci valoarea minimă a pragului

va trebui să fie suficient de mare pentru a putea permite menținerea folosirii unei legături la un nivel îndeajuns de ridicat. Pentru cazul în care traficul este unul oarecum normal, cu întârzieri destul de mari, atunci utilizarea unui prag minim pachete ar putea duce la o legătură foarte proastă. [5]

Cea mai bună valoare pentru pragul maxim depinde în cea mai mare parte de întârzierea medie maximă permisă de gateway.

Când $th_{min} - th_{max}$ este mai mare decât creșterea lungimii medii a cozii într-o anumită perioadă de timp, atunci algoritmul RED va da un randament maxim. În majoritatea cazurilor, pragul maxim este cel puțin de 2 ori mai mare decât pragul minim.

3.6 Alegerea valorilor pentru pondere

Metoda care realizează medierea nu va reuși să filtreze congestia în urma tranzitului de pachete de la nivelul gateway-ului pentru valori ale ponderii w_q prea mari.

Astfel, se va presupune că în faza inițială coada va fi goală, cu o largime medie a cozii de valoare nulă, urmând apoi o creștere considerabilă a dimensiunii cozii de la 0 la L pachete, după ce L pachete au ajuns la destinație. După ce ultimul pachet a sosit, adică cel numerotat cu litera L, atunci lungimea medie a cozii (avg) va fi [5]:

$$avg_L = \sum_{i=1}^L i * w_q * (1 - w_q)^{L-1}$$

$$avg_L = w_q * (1 - w_q)^L * \sum_{i=1}^L i * \left(\frac{1}{1 - w_q} \right)^i$$

$$avg_L = L + 1 + \frac{(1 - w_q)^{L+1} - 1}{w_q}$$

Lungimea medie a cozii avg_L pentru anumite valori ale ponderii w_q și L vor fi prezentate în figura ce va urma. Pe axa x va fi reprezentată ponderea w_q , care ia valori în intervalul 0.001 și 0.005.

Pe axa y va fi reprezentat L care ia valori cuprinse în intervalul 10 și 100. Pentru $w_q = 0.001$, de exemplu, lungimea medie a cozii avg_{100} va fi de 4.88 de pachete, după o creștere a cozii de la 0 până la 100 de pachete.

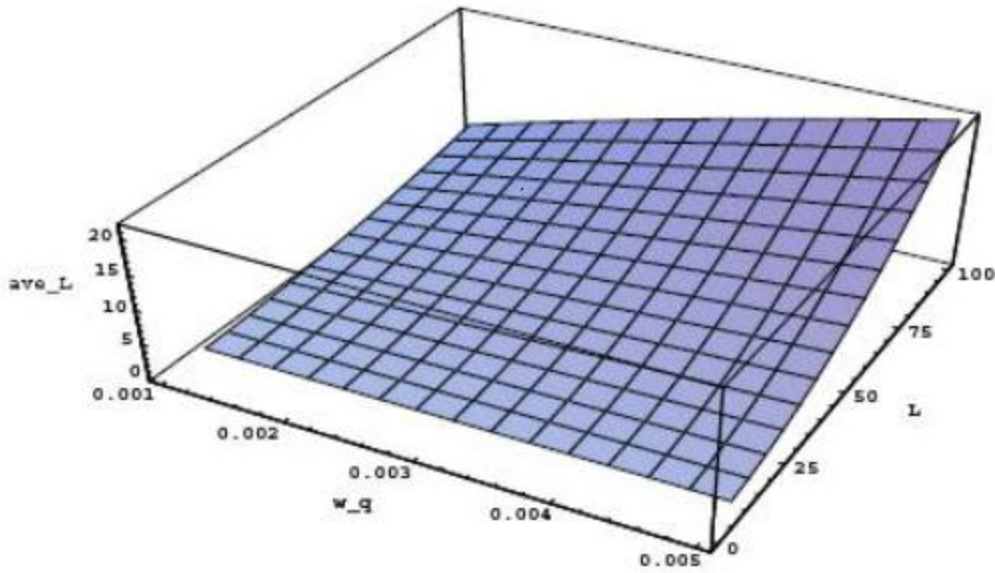


Figura 3.4: Lungimea medie a cozii în funcție de L și de w_q [5]

Avem de asemenea folosit un prag minim th_{min} și se știe că se dorește permiterea traficului agresiv utilizându-se astfel L pachete, atunci va rezulta că ponderea va trebui aleasă în așa fel încât ecuația scrisă mai jos să fie satisfăcută pentru ca lungimea medie a cozii să fie mai mică decât pragul de minim al dimensiunii cozii:

$$L+1+\frac{1-w_q^{L+1}-1}{w_q} < th_{min}$$

Putem exemplifica pentru $L=50$, iar $th_{min}=5$ rezultă că ponderea va trebui aleasă mai mică sau egală cu valoarea de 0,0042.

Valoarea calculată a dimensiunii medii a cozii (avg) cu valori mai mici decât un anumit prag va trebui să fie menținută de către gateway-urile care au implementat algoritmul RED. Totodată, acest lucru nu va avea nici un sens pentru cazul în care dimensiunea medie a cozii calculată (avg) nu va avea valori oarecum asemănătoare cu valorile lungimii medii curente. Pentru cazul în care ponderea este setată la un nivel scăzut, atunci dimensiunea cozii medii va răspunde foarte încet la schimbările lungimii cozii ce au fost făcute. Așadar, gateway-ul nu va putea să observe apariția congestiei.

Dacă s-ar presupune trecerea cozii de la nici un pachet la un pachet, și că în timp ce pachetele sosesc și se întorc având aceeași rată, atunci coada va rămâne la un pachet. Totodată dacă am presupune că în faza inițială lungimea medie a cozii era nulă, în cazul acesta, vor fi

nevoie de sosiri de pachete într-un număr egal cu $\frac{-1}{\ln(1-w_q)}$, cu lungimea cozii ce va rămâne la valoarea 1, până în momentul când aceasta va deveni $0,63 = 1 - \frac{1}{e}$.

De exemplu pentru o valoare a ponderii w_q de 0,001 vor fi necesare 1000 de sosiri de pachete. Dacă am dubla valoarea ponderii, adică la 0,002 vor fi utile 500 de sosiri de pachete, iar pentru valoarea ponderii de 0,003 vor fi utile 333 sosiri de pachete.

3.7 Moduri de evaluare ale algoritmului RED

Există mai multe moduri prin care algoritmul RED poate fi caracterizat, cum ar fi [5]:

Evitarea congestiei

RED face ca dimensiunea medie a cozii ce a fost calculată să nu depășească pragul maxim th_{max} pentru cazul în care pachetele sunt aruncate în momentul în care ajung, atunci când lungimea maximă a cozii va atinge pragul maxim.

Când valoarea ponderii w_q va fi setată într-un mod corespunzător, atunci RED va putea să controleze și largimea medie a cozii curente.

Scalele de timp potrivite

Intervalul de timp suficient pentru ca gateway-ul să poată observa o scădere a ratei de primire de pachete va fi de mai mult de o perioadă dus și întors după ce conexiunea va fi înștiințată despre apariția congestiei prin marcarea unui pachet anume.

Pentru cazul gateway-urilor folosite de algoritmul RED, scalele de timp folosite pentru observarea congestiei sunt aproximativ echivalente cu scalele de timp utilizate de conexiuni pentru a da un răspuns la congestie. Gateway-urile RED au rolul de a înștiința conexiunile pentru a-și reduce fereastra.

Inexistența unei sincronizări globale

Nivelul de congestie este cel de care va depinde rata la care algoritmul RED va realiza marcarea pachetelor. În momentul în care are loc o congestie scăzută, cu alte cuvinte, pentru puține pachete în coadă, probabilitatea de marcarea a oricărui pachet din acea coadă va fi foarte mică și va crește direct proporțional cu creșterea congestiei din rețea.

Algoritmul RED are rolul de a evita sincronizarea globală folosind marcarea de pachete și o rată de marcarea a pachetelor cât se poate de mică.

Simplitatea

Algoritmul RED pentru gateway de marcarea a pachetelor mai poate fi implementat utilizând un overhead moderat în rețele actuale.

Maximizarea puterii globale

Algoritmul RED poate controla într-un mod explicit dimensiunea medie a cozii, iar totodată se mai poate vedea că puterea globală este mai mare decât pentru cazul gateway-urilor de tip Drop Tail.

Pentru noile cercetări ce vor urma va fi nevoie să se determine valoarea optimă a dimensiunii medii a cozii pentru diferite rețele și pentru diferite situații ale traficului.

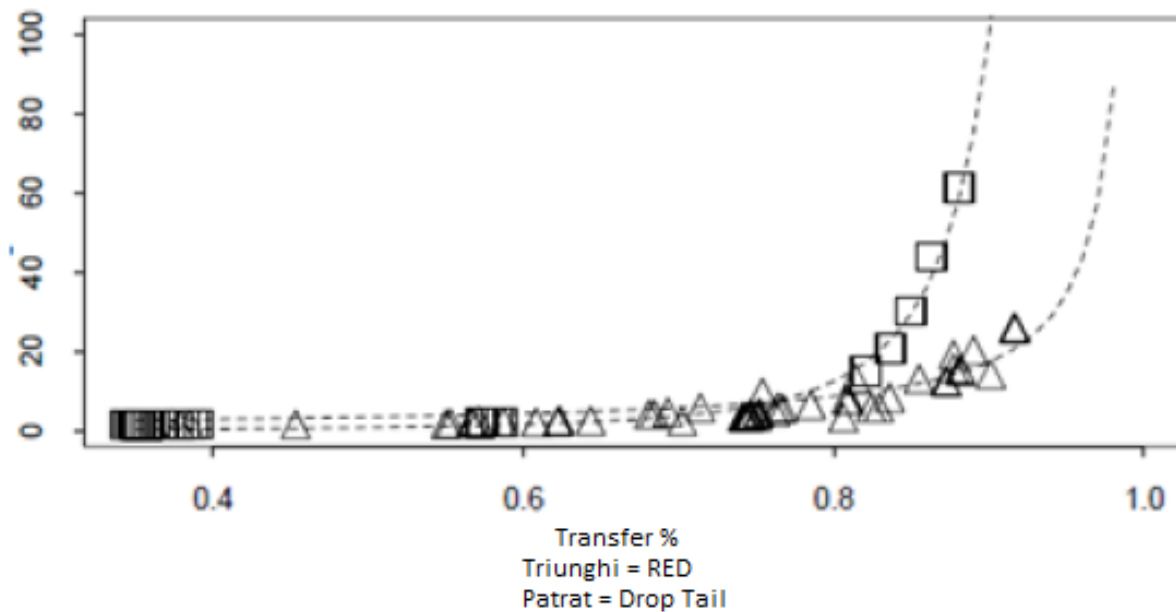


Figura 3.5: Diferențe între transferul făcut de către gateway-urile RED și Drop Tail [5]

Corectitudinea

Corectitudinea reprezintă un scop al mecanismului folosit pentru a evita congestia, însă ea nu este bine definită. Routerule care folosesc algoritmul RED nu pot face diferența între conexiunile private sau alte clase de conexiuni.

Fracția de pachete marcate folosită la algoritmul RED pentru oricare conexiune este direct proporțională cu intervalul de timp când conexiunea respectivă ține banda ocupată.

Potrivit pentru o gamă largă de medii

Procedul aleator de a marca pachetele este utilizat pentru rețele unde conexiunile sunt caracterizate folosind o gamă foarte largă interval de timp (roundtrip sau throughput – transfer); Totodată mai este necesar acest procedeu pentru un număr foarte mare de conexiuni active în același timp. Cu ajutorul modificărilor lungimii medii a cozii pot fi detectate schimbări ale încărcării rețelei, rata de marcare a pachetelor fiind minimizată corespunzător.

Pentru cazul folosirii unei rețele în care routerul RED detectează congestia aruncând pachetele marcate, astfel vor fi mai multe cazuri în care pierderea unui pachet nu va produce o descreștere a încărcăturii de pachete a routerului. Aruncarea unui pachet ce a aparținut unei conexiuni TCP va fi ușor semnalată de către sursă atunci când routerul pierde pachetul respectiv și cel mai posibil după ce timpul de retransmisie expiră.

Pierderea unui pachet ar putea avea loc fără să fie observată de către sursă pentru cazul în care routerul aruncă un pachet ACK sau acknowledged aparținând unei conexiuni TCP sau non-TCP.

Cu toate acestea și chiar pentru cazul utilizării unei rețele congestionate care este caracterizată de trafic alcătuit din conexiuni TCP sau non-TCP, routerul va avea control asupra

lungimii medie a cozii deoarece va arunca toate pachetele pe care le primește din momentul în care pragul de maxim este depășit.

3.8 Senzitivitatea parametrilor

Ruterele folosite de algoritmul RED conțin parametri suplimentari și opționali în același timp, care au rolul de a determina pragul superior sau th_{min} pentru lungimea medie a cozii, totodată determinând intervalul de timp pe care se va calcula această lungime și frecvența maximă cu care se vor marca pachetele. [5] [8]

Tot pentru utilizarea rutelor RED, parametrii pondere (w_q) pragul de minim (th_{min}) și pragul de maxim (th_{max}) sunt utili pentru ca acea persoana care realizează rețeaua să ia cele mai bune decizii în legătură cu lungimea medie a cozii dorită și totodată în legătură cu dimensiunea și durata traficului agresiv care va fi acceptat de către buffer.

Deoarece reprezintă limita superioară pentru probabilitatea de marcarea a pachetelor P_b , parametrul max_p poate să ia valori dintr-o gamă suficient de largă. Atunci când nivelul congestiei este unul foarte ridicat, astfel încât routerul nu va putea să dețină controlul lungimii medie a cozii cu ajutorul marcării la maxim unei fracțiuni de max_p din pachete, va rezulta că lungimea medie a cozii trece peste pragul de maximum, routerul urmând să marcheze toate pachetele până în momentul când congestia va putea să fie controlabilă.

La ruterele de tip Drop Tail, singurul parametru care variază este chiar dimensiunea bufferului. Procedul cu care congestia va fi evitată ar fi normal și necesar să aibă o sensibilitate scăzută pentru parametri, iar acești parametri ar trebui să poată fi utilizați pentru rețelele cu o gamă largă de variație a lungimii de bandă.

În cele ce vor urma au fost descrise câteva reguli care ar trebui urmate pentru ca RED să poată oferi o performanță acceptabilă pentru cazul unei game diversificate de parametri ai rutelor și de numeroase condiții de trafic pe rețea.

Asigurarea calculului corect al lungimii medie a cozii

Încă de la început ponderea w_q va fi aleasă mai mare sau egală cu valoarea de 0,001

Lungimea medie a cozii folosită de router va fi limitată de pragul maxim th_{max} pentru intervalul de timp în care lungimea medie a cozii avg ce a fost calculată va reflecta dimensiunea reală.

Pentru ca dimensiunea medie a cozii să nu aibă întârzieri prea mari atunci când lungimile reale ale cozii sunt determinate, valoarea ponderii sau w_q nu ar fi necesară pentru ca rețeaua să ia valori foarte mici.

Marimea pragului minim va trebui să fie îndeajuns de mare astfel încât să se poată maximiza puterea din rețea

Valorile pragurilor de minim (th_{min}), respectiv de maxim (th_{max}) va fi nevoie să fie îndeajuns de mari pentru ca puterea prezentată în rețea să poată fi maximizată.

Valorile lungimii reale a cozii pot fi totodată destul de variabile datorită faptului că traficul din rețea care va fi de cele mai multe ori agresiv.

Legătura de ieșire sau bottleneck-ul va fi utilizat fără a da un randament foarte bun dacă dimensiunea medie a cozii va fi stabilită la un nivel foarte mic.

Diferența dintre pragul de maxim și cel de minim $th_{max}-th_{min}$ va trebui să fie îndeajuns de mare în așa fel să nu se producă sincronizarea globală

Diferența $th_{max}-th_{min}$ va trebui să fie mult mai mare decât variația uzuală a lungimii medii a cozii pe durata transportului dus-întors, pentru ca sincronizarea globală să nu se poată produce. Lungimea medie calculată a cozii poate să oscileze cu regular până la valoarea maximă a pragului th_{max} , având totodată același comportament ca Drop Tail, acest lucru având loc doar pentru cazul în care acea diferență va fi prea mică.

3.9 Alte versiuni de implementare ale algoritmului RED

3.9.1 WRED (Weighted Random Early Detection)

WRED este utilizat ca o posibilitate pentru modalitatea de aruncare de pachete care se află la finalul unei cozi atunci când coada (sau buffer-ul) este încărcată la maxim (comportament tail-drop). Totodată reprezintă o implementare Cisco pentru algoritmul Random Early Detection (RED) fiind o metodă ce folosește la evitarea congestiei. [9] [11]

Folosind acest algoritm, avem de-a face cu problema sincronizării globale, în momentele în care mai multe transmisii de tip TCP parcurg routerul respectiv ieșind apoi prin aceeași interfață prin care au intrat.

Viteza este oarecum mică atunci când comunicația TCP are loc iar pentru cazul în care banda momentan nu este ocupată cu pachete, atunci oricare dintre comunicațiile din rețea vor avea viteză de transfer mai mare. Pachetele vor începe să fie aruncate din coadă în momentul în care suma tuturor transferurilor va trece peste lățimea de bandă a interfeței utilizată pentru ieșire.

Folosind aceasta modalitate, va fi nevoie ca toate comunicațiile să își minimizeze viteza de transfer simultan, acest lucru ducând apoi la o scădere foarte mare a vitezei globale. Datorită faptului că sesiunile protocolului TCP din link-uri au o viteză relativ mică, vor reveni astfel la cazul inițial atunci când lățimea de bandă este îndeajuns pentru toate transferurile ce au loc și viteza oricărei dintre sesiuni este foarte mică. Așa cum putem observa, acest proces se va relua ori de câte ori va fi nevoie.

Se poate observa că folosind WRED, procentul din banda ce a fost ocupată este cu mult sub 100%. Totodată utilizând acest algoritm, procentul va putea crește constant până la valori oarecum mari. Algoritmul WRED nu dă voie scăderii vitezei de transfer pentru toate sesiunile TCP ce au loc în același timp.

WRED mai poate să utilizeze praguri pentru oricare coadă care a fost creată de către sistem. Aceste praguri vor fi monitorizate până în momentul în care un pachet va fi pus la coadă. Pachetul va fi transferat în coadă atunci când numărul de pachete care se află în coadă este unul mai mic decât pragul inferior utilizat, iar pentru cazul în care depășește limita pragului superior atunci pachetul respectiv va fi aruncat. Aici mai există un caz, pentru care dacă se va afla între cele 2 praguri de minim și de maxim, atunci va fi utilizată o funcție care stabilește dacă pachetul va fi aruncat sau pus în coadă.

Algoritmul WRED poate construi un profil de tip WRED pentru oricare valoare de clasificare a fiecărui transfer, cu alte cuvinte putem spune că acest lucru este de fapt diferența

dintre algoritmul RED și Weighted RED. Un profil de tip WRED este reprezentat de anumite valori pe care le pot lua pragurile maxime și minime. Pragurile pot lua valori care vor fi definite ca numărul pachetelor ce se află în coadă.

Weighted RED are posibilitatea de a trata într-un mod diferit unele tipuri de pachete din rețea utilizând aceste profiluri pentru cazul în care va avea loc congestia în traficul de pachete.

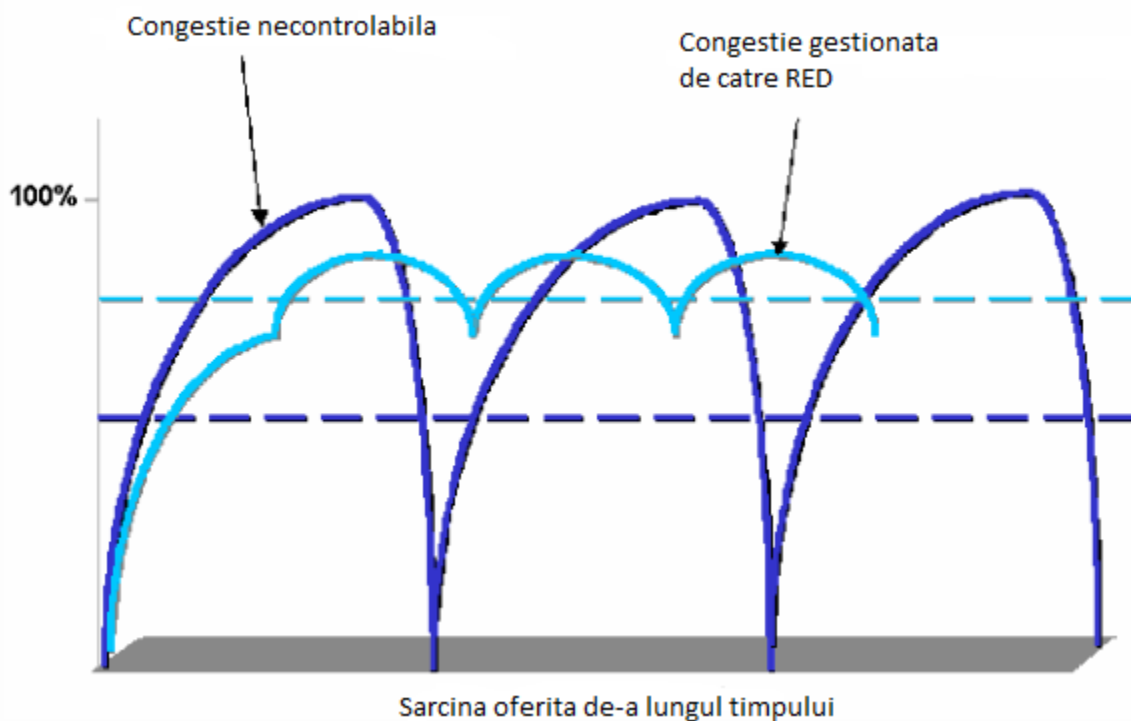


Figura 3.6: Diferențele rezultate în urma implementării algoritmului WRED [10]

Algoritmul RED spre deosebire de algoritmul WRED nu poate realiza o diferențiere QoS (Quality of Service) deoarece în cazul WRED sunt prezentate mai multe avantaje.

Un alt concept de importanță majoră ce trebuie tratat, înainte de a discuta configurația WRED, este referitor la locul unde poate fi folosit WRED și cum interacționează cu uneltele de management al cozilor.

Înainte de implementarea configurației folosite la algoritmul WRED, va trebui să se țină cont de modul cum acesta interacționează cu parametrii care gestionează coada și de locul unde se poate folosi Weighted RED.

Lungimea medie a cozii depinde de lungimea medie anterioară și de lungimea curentă a cozii. Urmărind formula, se va observa că dimensiunea medie a cozii va depinde de lungimea curentă și de lungimea medie precedentă [11] :

$$\text{avg} = o * (1 - 2^{-n}) + c * (2^{-n})$$

Unde o este coada cea veche, iar c este coada curentă.

Algoritmul WRED deține un comportament aproximativ la fel ca și algoritmul RED, astfel calculând lungimea medie a cozii și în același timp poate preciza dacă pachetele vor fi aruncate,

iar dacă acest lucru va fi afirmativ, care este procentul din pachetele disponibile care se vor arunca. Aceasta decizie va fi luată utilizând aceiași parametri ce au fost descriși în capitolele precedente pentru RED.

WRED, tot la fel ca și RED folosește anumite criterii pentru a lua o decizie pentru câte pachete și când acestea să poată fi descărcate, cum sunt următoarele:

- Dimensiunea medie a cozii
- Pragul de maxim
- Pragul de minim

WRED calculează lungimea medie a cozii în prima fază, exact același lucru fiind făcut și de algoritmul RED, urmând în continuare să se compare această lungime cu pragul de maxim și de minim pentru pachetele ce au fost marcate.

Pentru cazul în care valoarea medie a cozii se află între valorile limită ale pragurilor alese, atunci algoritmul WRED va descărca un anumit procent din pachetele disponibile în rețea iar când valoarea medie a cozii va depăși th_{max} , atunci algoritmul va descărca toate pachetele ce au ajuns recent.

3.9.2 FRED (Fair Random Early Detection)

Algoritmul Fair RED utilizează ca baza algoritmul RED unde va mai adăuga noi parametri q_{min} și q_{max} , reprezentând numărul minim și numărul maxim de pachete ce pot fi adăugate în coadă de către oricare din legăturile disponibile.

Totodată, algoritmul Fair RED adaugă o nouă variabilă (globală) notată $avgcq$ și reprezintă o modalitate prin care se poate estima o medie a pachetelor ce vor ajunge în buffer-ul rețelei per link-ul respectiv. Legăturile care vor trimite pachete într-un număr mai mic decât valoarea parametrului $avgcq$, vor avea prioritate.

În variabila $qlen$ va fi reținut numărul de pachete actuale din coadă pentru oricare din conexiunile cu routerul. Totodată, FRED mai poate să rețină și numărul de tentative nereușite ale fiecărei legături de a trimite un mesaj răspuns atunci când congestia va urma. Acest număr poate fi găsit în variabila $strike$.

Cu toate acestea, algoritmul FRED poate atribui legăturilor valori ale variabilei $strike$ foarte mari. [7]

3.9.3 ARED (Adaptive Random Early Detection)

Principalul obiectiv al acestui algoritm este acela de a preveni dacă algoritmul RED este necesar să devină mai puțin sau mai mult agresiv privind modul de examinare a dimensiunii medii a cozii de pachete din rețea. Pentru cazul în care această dimensiune medie a cozii va depăși într-un mod repetat th_{min} sau pragul de minim, algoritmul este prea agresiv sau mai precis examinarea avg va fi una riguroasă.

Algoritmul ARED nu va fi îndeajuns de agresiv pentru cazul în care valorile dimensiunii medii a cozii vor trece într-un mod constant peste th_{max} sau pragul de maxim. [8]

Algoritmul Adaptive Random Early Detect are posibilitatea de a-și putea configura parametrii în funcție de care este nivelul traficului din rețeaua respectivă. Cu toate acestea, dacă dimensiunea medie a cozii va lua valori ce se vor afla între pragurile de minim, respectiv de maxim, atunci valoarea parametrului p_{max} va crește într-un mod constant și în același timp creșterea fiind una scăzută din punct de vedere al valorilor în funcție de cât de încărcat este traficul din rețea.

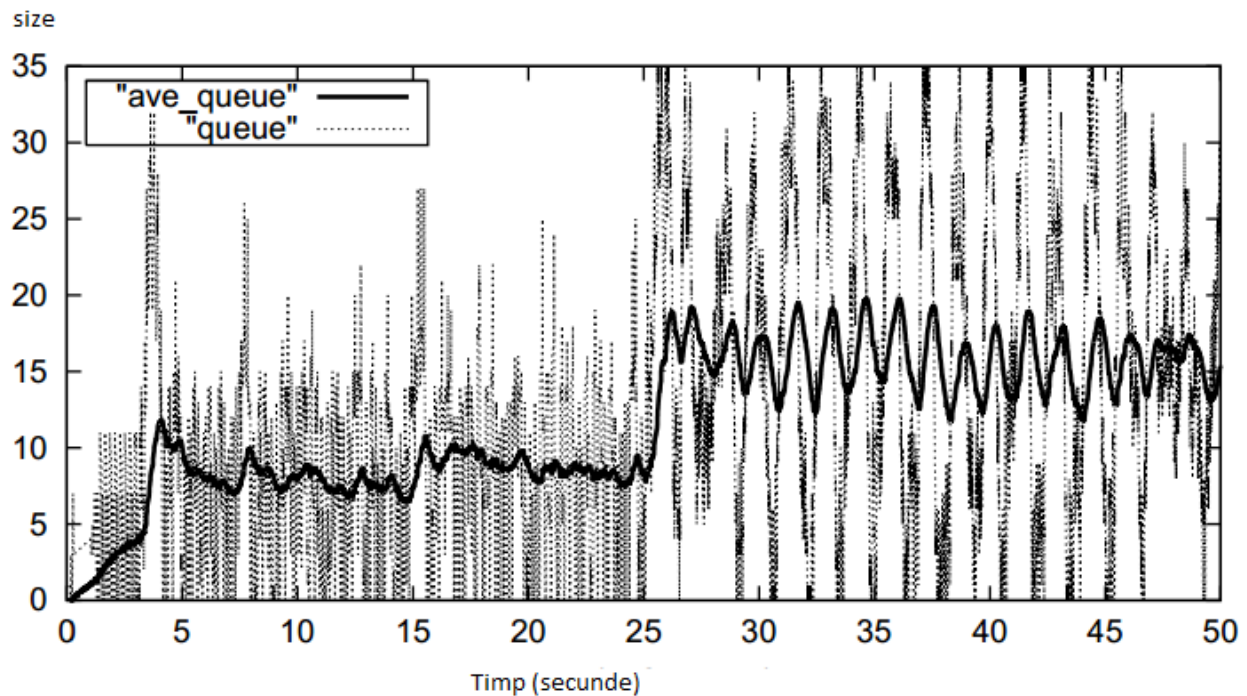


Figura 3.7: Algoritmul RED folosind o congestie ridicată [8]

Am putea spune că principalul dezavantaj la ARED este ca parametrii optimi folosiți pentru a obține rezultate contradictorii nu sunt stabiliți clar în momentul inițial.

Avantajul de bază pentru ARED este că își poate schimba parametrii în funcție de cât de agresiv este traficul.

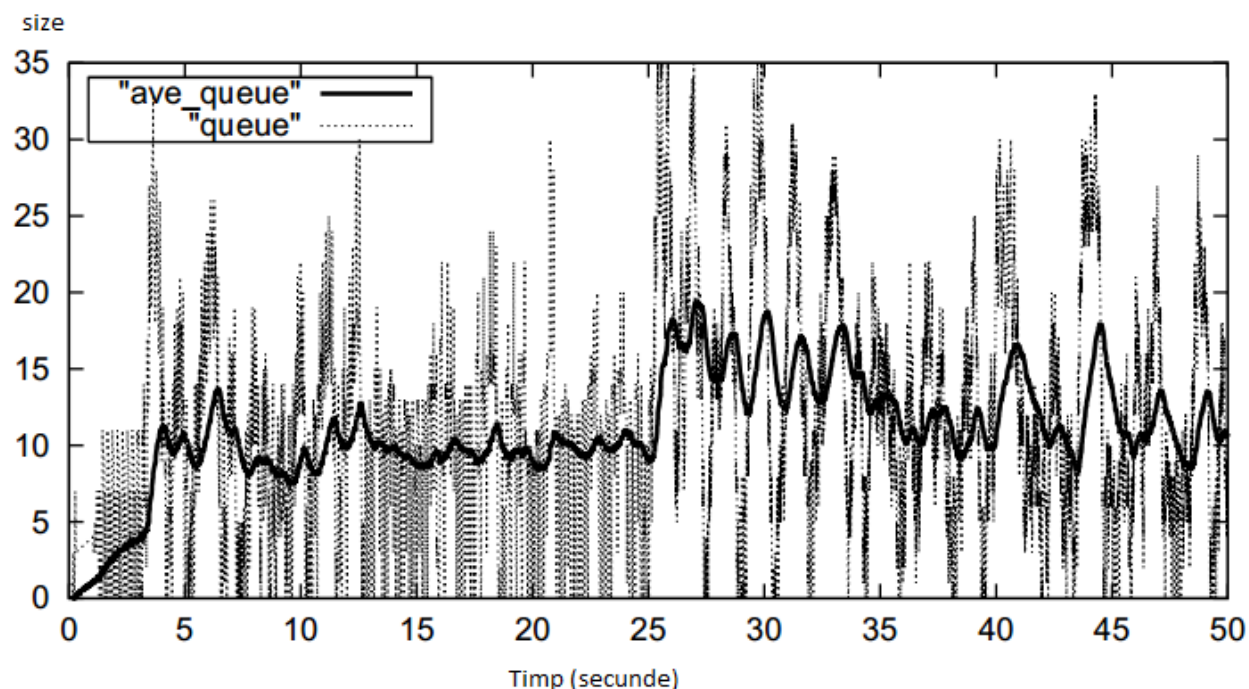


Figura 3.8: Algoritmul ARED folosind o congestie ridicată [8]

În funcție de valorile lungimii medii a cozii, ARED descrește valoarea parametrului p_max . În funcție de pragul ce va fi depășit (th_{min} sau th_{max}), în prezentarea algoritmului ce va urma mai jos, p_max va fi scalat utilizând astfel unul din factorii de creștere (α) sau descreștere (β).

Implementarea algoritmului ARED:

Pentru fiecare interval de timp:

If ($avg > target$ si $p_max \leq 0,5$)

Crește valoarea p_max :

$p_max \leftarrow p_max + \alpha$;

else if ($avg < target$ si $p_max \geq 0,001$)

Micșorează valoarea p_max :

$p_max \leftarrow p_max * \beta$;

Variabile:

avg : lungimea cozii medii

Parametrii fixați, adică am exemplificat valori pentru acești parametri:

$Timp = 0,5$ secunde

Obiectiv ($target$) pentru avg : $[th_{min} + 0,4 * (th_{max} - th_{min}), th_{min} + 0,6 * (th_{max} - th_{min})]$

α : parametru de creștere = $\min(0,01, \frac{p_max}{4})$

β : parametru de descreștere = 0,9

3.9.4 RRED (Robust Random Early Detection)

Algoritmul Robust Random Early Detection constă în adăugarea unui bloc de filtrare și în același timp de detecție, ce au loc pe un router înaintea blocului RED. [12] [13]

În momentul în care un pachet va fi expediat la un interval mic de timp după ce alt pachet din rețea a fost aruncat din coadă, se poate spune că cel de-al doilea pachet ar putea fi un pachet de atac. Astfel se poate realiza o diferențiere între pachetele de tip normal (TCP) și pachetele de atac.

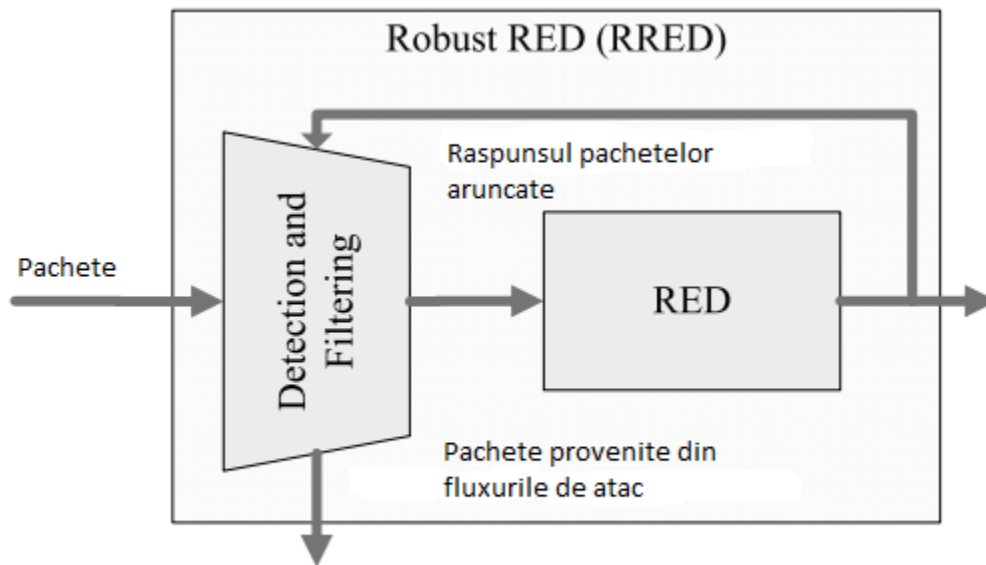


Figura 3.9: Schema arhitecturii algoritmului Robust RED [12]

Scopul acestui algoritm practic este de a realiza o îmbunătățire privind modul de abordare al conexiunilor TCP din rețea împotriva pachetelor de atac LDoS (sau Low-rate Denial-of-Service attacks).

Așadar RRED este o variantă a algoritmului RED simplu, iar modelul pe care îl urmează va fi acela de filtrare și de monitorizare a pachetelor provenite din fluxurile de atac, înaintea aplicării algoritmului RED normal. [13]

CAPITOLUL 4

Programe si limbaje utilizate

4.1 Network Simulator

4.1.1 Istoric

Simulatorul în care se va lucra în decursul acestui proiect este Network Simulator 2.

Unii dintre cei care au ajutat la construirea acestui algoritm au fost chiar inventatorii algoritmului RED Sally Floyd, Kevin Fall, Steve McCanne. NS2 a fost construit între anii 1996-1997, ca o variantă a simulatorului de rețele REAL.

Simulatorul numit REAL a fost inventat de către Srinivasan Keshav în anul 1989 și inițial acesta avea decât datoria de a monitoriza atitudinea dinamică a schemelor pentru controlul congestiei sau a fluxurilor ce apăreau în rețeaua respectivă.[14]

Există numeroase sisteme de operare pe care Network Simulator 2 poate fi disponibil cum putem exemplifica : Linux/GNU, Solaris, FreeBSD, Mac OS X chiar și versiuni Windows care suportă mediul Cygwin.

Singurul domeniu în care poate fi folosit Network Simulator este domeniul cercetării.

Totodată, NS poate oferi un suport foarte mare pentru a realiza simularea comunicației dintre legăturile TCP ale rețelei, protocoalelor multicast și a routerelor ce realizează rețeaua atât pentru rețelele ce utilizează wi-fi sau cele care utilizează cabluri.

4.1.2 Instalarea simulatorului

Simulatorul pe care l-am folosit este Network Simulator 2 cu versiunea v2.35, pe sistemul de operare Ubuntu 14.04, care a fost instalat pe computerul meu HP Compaq Pressario A900 Notebook cu următoarele performanțe:

- memorie RAM de 3GB
- processor Intel Pentium Dual 1,73 GHz
- memoria internă de 180 de GB

Pentru a instala programul NS2, am urmat următorii pași [21] :

Înainte de a instala NS2, trebuie să instalați unele software-uri esențiale:

```
$sudo apt-get install tcl8.5-dev tk8.5-dev  
$sudo apt-get install build-essential autoconf automake  
$sudo apt-get install perl xgraph libxt-dev libx11-dev libxmu-dev
```

Va trebui downloadată sursa NS2 de pe site-ul NS. Așadar fișierul va fi de forma “ns-allinone-2.35.tar.gz”. Odată downloadată această arhivă trebuie despachetată în /home/):

```
$tar -zxvf ns-allinone-2.35.tar.gz -C /home/
```

Următorul pas este instalarea propriu-zisă:

```
$cd /home/ns-allinone-2.35
```

```
$sudo ./install
```

În continuare va trebui modificat fișierul .bahrc:

```
$vi /home/.bashrc
```

Pe ultima linie de comandă din acest fișier va trebui introdus scriptul:

```
export PATH=$PATH:/home/ns-allinone-2.35/bin:/home/ns-allinone-
```

```
2.35/tcl8.5.10/unix:/home/ns-allinone-2.35/tk8.5.10/unix
```

```
export LD_LIBRARY_PATH=$LD_LIBRARY_PATH:/home/ns-allinone-2.35/otcl-
```

```
1.14:/home/ns-allinone-2.35/lib
```

```
export TCL_LIBRARY=$TCL_LIBRARY:/home/ns-allinone-2.35/tcl8.5.10/library
```

Se mai poate face o eventuală verificare a instalării:

```
$cd ns-2.35; ./validate
```

Se poate accesa programul executând comenzile:

```
$nam
```

```
$ns fisier.tcl
```

4.2 Limbajul OTcl

Network Simulator este un simulator scris în limbajul C++ (sau cu alte cuvinte, Network Simulator este compatibil unei ierarhii pentru clasele utilizate de acel interpretor OTcl și unei alte ierarhii pentru clasele utilizate în C++) și este orientat pe obiecte (OO). [14][15][20]

Aceste ierarhii pe care stau la baza simulatorului necesita dependent uneia fata de cealaltă ierarhie.

Între o clasă din cadrul ierarhiei ce a fost compilată și o clasă din cadrul ierarhiei interpretorului există o legătură foarte stransă din punct de vedere al utilizatorilor care folosesc acest simulator. Tcl Object reprezintă o rădăcină pentru oricare dintre aceste ierarhii utilizate. [16] Similarele ce se doresc a fi realizate se fac cu ajutorul interpretorului de către useri.

Interpretorul este cel care face legătura la aceste obiecte (Tcl Object) și în același timp mai are rolul de a reflecta un obiect din cealalta ierarhie a compilării. Ierarhia claselor interpretorului este realizată automat prin metode definite în clasa TclClass.

Cu ajutorul procedurilor ce sunt bine determinate în clasa Tcl Class, se va crea automat ierarhia claselor simulatorului.[17]

Cu alte cuvinte, obiectele care au fost reflectate de catre user vor fi referite prin modalitățile ce se găsesc în clasa Tcl Object. Mai sunt și alte ierarhii care sunt compatibile cu limbajele OTcl, respectiv C++ însă, nu sunt prezentate în clasa TclObject.

4.3 Network Animator

Nam reprezintă un tool al simulatorului NS ce se folosește pentru animațiile rețelelor propriu-zise. Nam se bazează pe Tcl/Tk (TookKit) , utilizat pentru a vizualiza informațiile obținute în urma simulărilor rețelei. Modul de gândire cu care a fost implement acest tool Nam este de fapt inventarea unui animator ce poate citi fragmente imense de informație folosite pentru animațiile respective și în același timp să aibă posibilitatea de a putea observa foarte multe posibilități pentru rețelele propriu-zise. [18][19][20]

Nam a fost realizat cu scopul de a citi comenzile simple, însă cu o încărcătură mare de informație utilizând astfel fișiere mari în care anumiți parametri sunt monitorizați și de a anima rezultatele.

Pentru ca tool-ul să aibă posibilitatea de a utiliza un volum de informație foarte mare, va trebui ca pentru realizarea animațiilor, cantitatea de date ce se va afla în memoria respectivă să ocupe dimensiuni foarte mici. În același timp, acest lucru poate ușura procesul de realizare a animațiilor, deoarece comenzile respective pot fi citite de fiecare dată când este nevoie oferindu-se un acces simplu.

În primul rând, pentru a putea utiliza acest tool, va fi nevoie de un fișier de monitorizare în care vor fi adăugate toate comenzile și informațiile despre parametri de care este nevoie. În fișierul de monitorizare putem găsi detalii legate de nodurile rețelei, legăturile dintre noduri, pachete dar și topologiile dintre legături. Network Simulator este cel care crează automat acest fișier de monitorizare.

Atunci când este pornit, NAM va citi din fișierul din monitorizare informația respectivă, după care va fi creată fereastra de tip pop-up, crează schema simulării apoi se va opri atunci când momentul de timp va atinge valoare nulă.

Cu alte cuvinte, putem spune ca NAM este o interfață pentru useri care are posibilitatea controlului asupra animațiilor.

După ce a fost creat acel fișier de monitorizare, atunci simulatorul va fi pregătit pentru transmiterea informațiilor și totodată vor fi realizate și animațiile.

4.4 Concluzii despre simulator

Datorita faptului ca simulatorul NS2 va trebui să realizeze două obiective pentru realizarea simulării (pentru obținerea datelor simulării și pentru orientarea pe obiecte), atunci acesta va utiliza două limbaje pentru realizarea acestui lucru și anume OTcl și C++.

Totusi, există cazuri în care se va avea de-a face cu simulări ce necesită aspecte pentru protocoalele utilizate în detaliu, acest lucru ducând la folosirea unui limbaj de programare ce deține avantajul de a controla biții ai pachetele într-un mod mai eficient, totodată putând astfel să încarce seturi de informație cu mult mai mari.

Pentru astfel de limbaje, nu este un lucru foarte important determinarea și remedierea defecțiunilor atunci când simularea este în decurs de rulare, însă o importanță majoră este pusă pe viteza de lucru a simulării.

Limbajele folosite de Network simulator dețin avantaje în funcție de scopul în care sunt folosite.

De exemplu limbajul OTcl prezintă un mod de a rula mai greoi, pe când acesta prezintă o flexibilitate mai mare, putând fi realizate schimbările într-un mod mai rapid, pe când limbajul C++ nu prezintă o flexibilitate atât de bună, însă este mult mai rapid atunci când este stabilită rularea respectivă.

Pentru cazurile unde sunt utilizate scenariile apropiate din punct de vedere al valorilor, sau parametrii rețelei având dimensiunile apropiate, atunci timpul de care are nevoie iterația prezintă o importanță deosebită, ca și pentru viteza de lucru a simulării, deoarece acele valori fiind asemănătoare, va fi mult mai dificil a se realiza iterația respectivă. Dar totodată, acel timp în care are loc simularea nu va fi atât de important pentru cazul în care simularea are loc o singură dată.

CAPITOLUL 5

Modificări aduse algoritmului RED

5.1 Introducere

Așa cum am descris în capitolele precedente, algoritmul RED calculează probabilitatea de aruncare a unui packet în două etape, atunci când dimensiunea cozii scade între th_min și th_max , unde th_min reprezintă pragul de minim al cozii, th_max reprezintă pragul de maxim al cozii.

Prima etapă reprezintă calculul probabilității cu care pachetele sunt marcate (P_b).

Cu ajutorul acestor valori se calculează probabilitatea P_b pe baza unei funcții liniare:

$$P_b = P_{max} * (avg - th_min) / (th_max - th_min)$$

Se observă în ecuația anterioară că P_b este dependent de valoarea dimensiunii medii a cozii.

În cea de-a doua etapă, RED contorizează numărul de packete care au trecut prin gateway până când ultimul pachet va fi aruncat, aplicându-se formula:

$$P_a = P_b / (1 - count * P_b)$$

Așa cum am observat în articolul [22], am modificat primul pas, adică probabilitatea P_b aplicând 4 funcții. Fiecare funcție, în algoritm este apelată de o sintaxă de tipul `edp_.func == x`, unde x poate lua valori de la 1 la 4.

Aceste 4 funcții sunt reprezentate mai jos de valorile $v1, v2, v3, v4$:

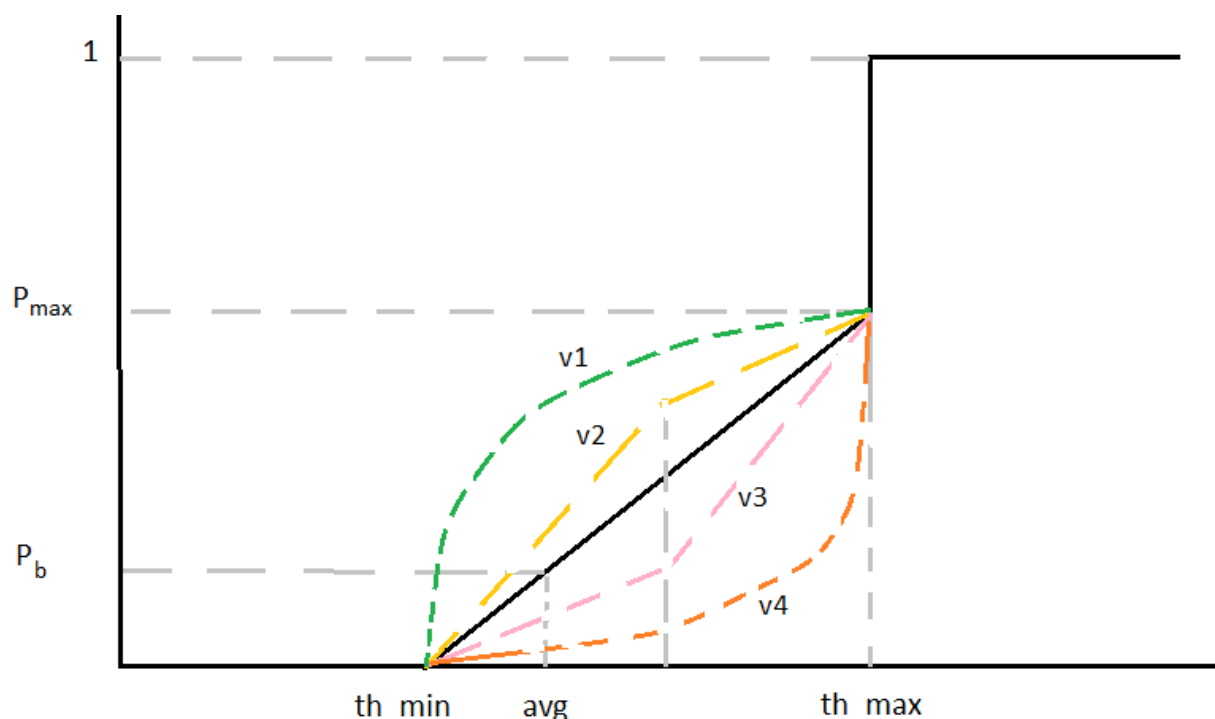


Figura 5.1: Diferite funcții pentru a evalua probabilitatea de aruncare a unui pachet [23]

Putem observa în figura 5.1 de mai sus cum lungimea medie a cozii este dependentă de pragul cozii minime și de pragul cozii maxime.

În implementarea cozii RED vom utiliza formula:

$P_b = P_{max} * (v_a * v_{ave} + v_b)$, unde $v_a = 1.0 / (th_{max} - th_{min})$, $v_b = -th_{min} / (th_{max} - th_{min})$

Această formulă este echivalentă cu prima formulă a P_b

Pentru funcția v_1 , care va fi concavă vom folosi funcția de logaritm:

$P_b = P_{max} * (m * \log(v_{ave}) + n)$

Pentru ca acele funcții curbe să unească punctele de th_{min} și th_{max} avem următoarele ecuații:

$0 = P_{max} * (m * \log(th_{min}) + n)$

$m = 1 / (\log(th_{max}) - \log(-v_b/v_a))$

$n = -m * \log(-v_b/v_a)$

$P_{max} = P_{max} * (m * \log(th_{max}) + n)$

De asemenea, funcția liniară inițială ne dă funcția:

$0 = P_{max} * (v_a * th_{min} + v_b)$

În mod similar, ne putem exprima formula pentru funcția de v_4 funcția convexă (am ales funcție exponențială), folosind th_{max} , v_b și v_a :

$P_b = P_{max} * (m * \exp(v_{ave}) + n)$, unde pentru acest caz:

$m = 1 / (\exp(th_{max}) - \exp(-v_b/v_a))$

$n = -m * \exp(-v_b/v_a)$

În ceea ce privesc funcțiile v_2 și v_3 , care sunt funcții liniare pe porțiuni:

$P_b = P_{max} * (m * v_{ave} + p)$ when v_{ave} in $[th_{min}, (th_{min} + th_{max})/2]$

$P_b = P_{max} * (n * v_{ave} + q)$ when v_{ave} in $[(th_{min} + th_{max})/2, th_{max}]$

Mai putem obține următoarele ecuații cu care vom obține relații dintre m și n .

$m * th_{min} + p = 0$

$m * [(th_{min} + th_{max})/2] + p = n * [(th_{min} + th_{max})/2] + q$

$n * th_{max} + q = 1$

Va rezulta:

$n = [m * (th_{max} + th_{min})/2 - m * th_{min}] / [(th_{max} + th_{min})/2 - th_{max}]$

Pentru funcția v_2 am setat $m = 1.4 * v_a$:

$p = -1.4 * v_a * th_{min}$

$n = [1.4 * v_a * (th_{max} + th_{min})/2 - 1.4 * v_a * th_{min}] / [(th_{max} + th_{min})/2 - th_{max}]$

$q = 1 - n * th_{max}$

Pentru funcția v_3 am setat $m = 0.6 * v_a$:

$p = -0.6 * v_a * th_{min}$

$n = [0.6 * v_a * (th_{max} + th_{min})/2 - 0.6 * v_a * th_{min}] / [(th_{max} + th_{min})/2 - th_{max}]$

$q = 1 - n * th_{max}$

Modificarea făcută algoritmului RED este practic făcută în funcția calculată p_{new} unde au fost introduse toate ecuațiile necesare pentru a realiza aceste funcții. Totodată, după cum se observa în anexa 2 și în ecuațiile de mai sus, am modificat funcțiile v_2 și v_3 , cărora le-am introdus valori diferite (respectiv 1,4 pentru v_2 și 0,6 pentru v_3) pentru parametrul m pe care îl conțin.

5.2 Reprezentarea grafică a funcțiilor

Pentru cazul inițial unde nu a fost făcută nici o modificare a codului sursă, atunci graficul unei funcții va arăta ca în figura 5.2

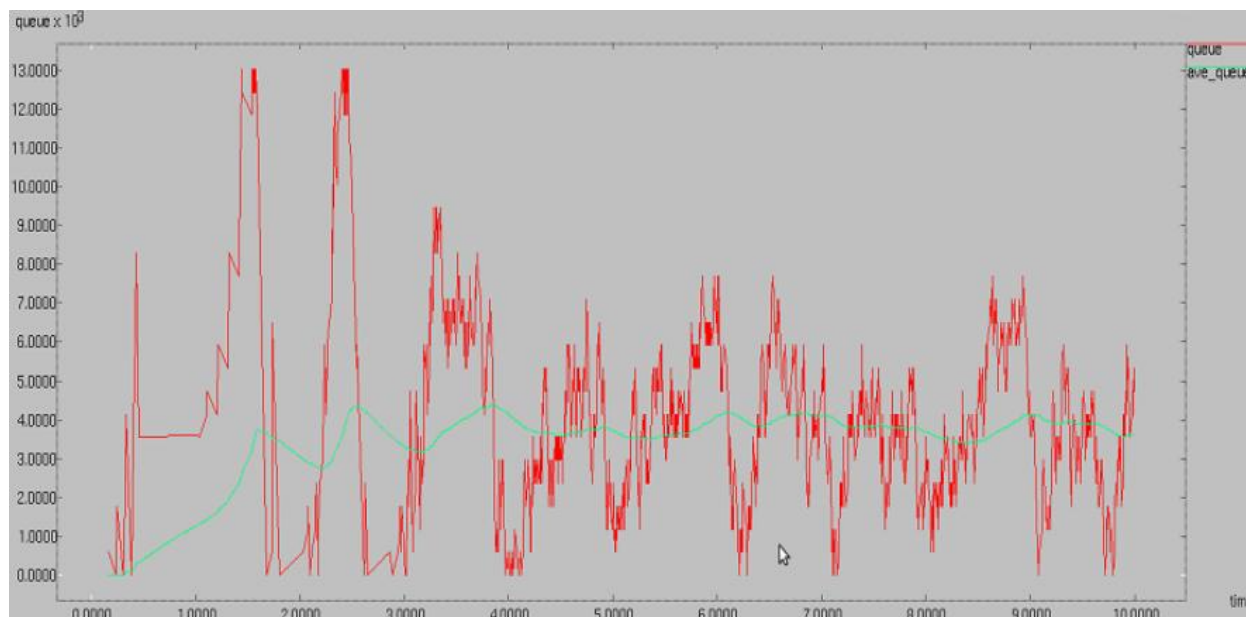


Figura 5.2 : Graficul algoritmului RED fără nici o modificare

Rulând simulările pentru fiecare din cele 4 funcții s-au obținut următoarele grafice, unde `ave_queue` (valoare reprezentată cu verde) este dimensiunea cozii medii și `queue` (valoare reprezentată cu roșu) este valoarea curentă a cozii.

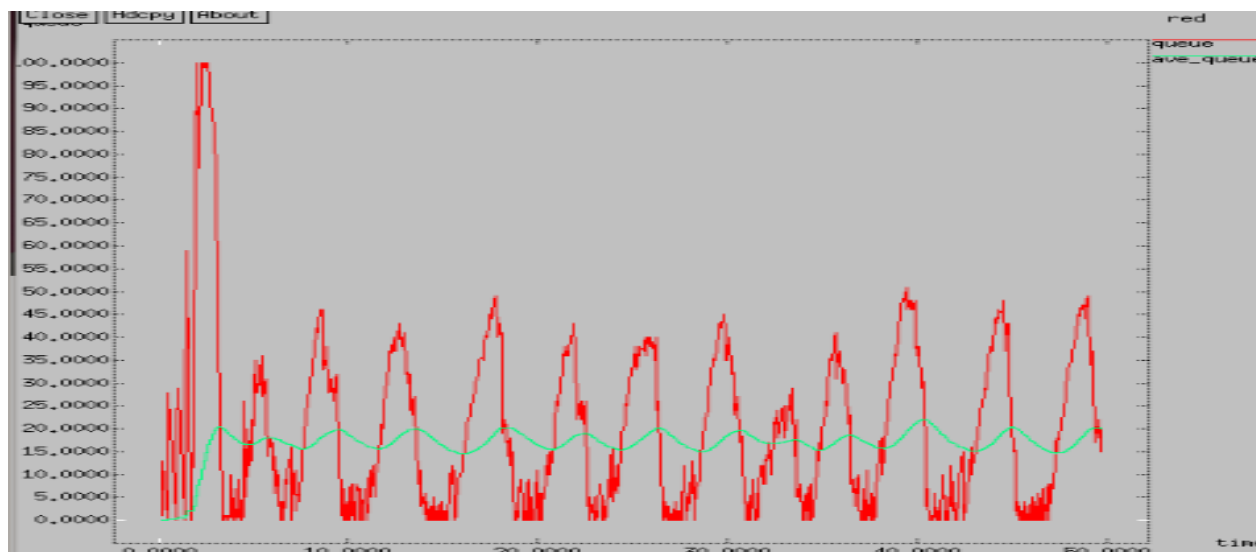


Figura 5.3 : Grafic folosind funcția din algoritm – v1

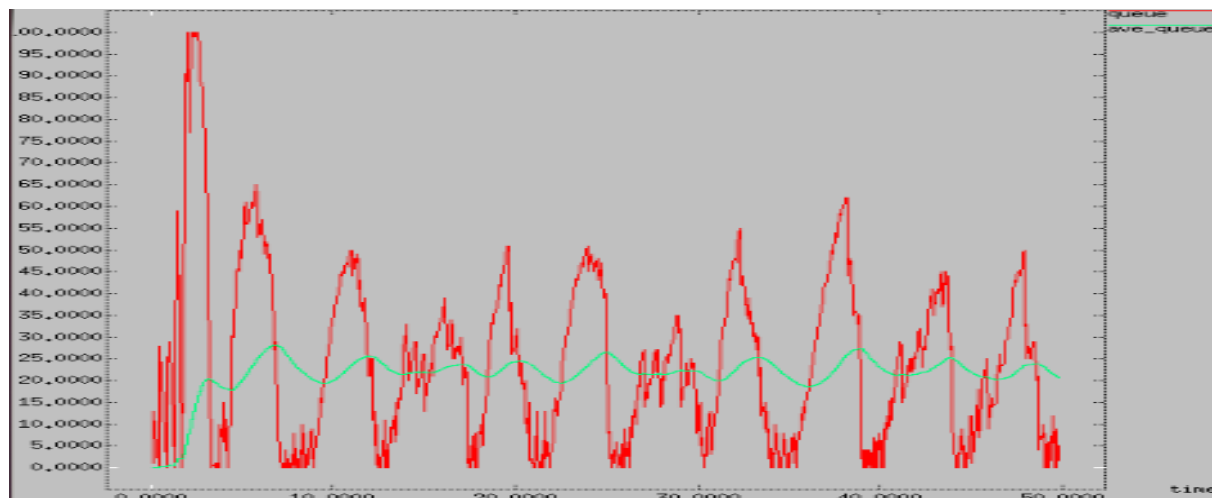


Figura 5.4 : Grafic folosind funcția din algoritm – v2

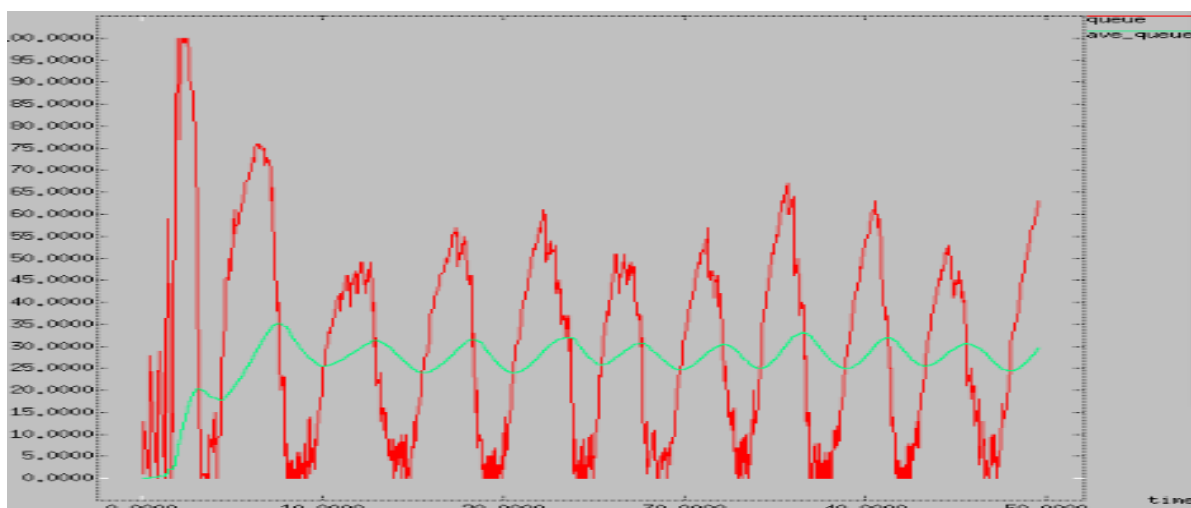


Figura 5.5 : Grafic folosind funcția din algoritm – v3

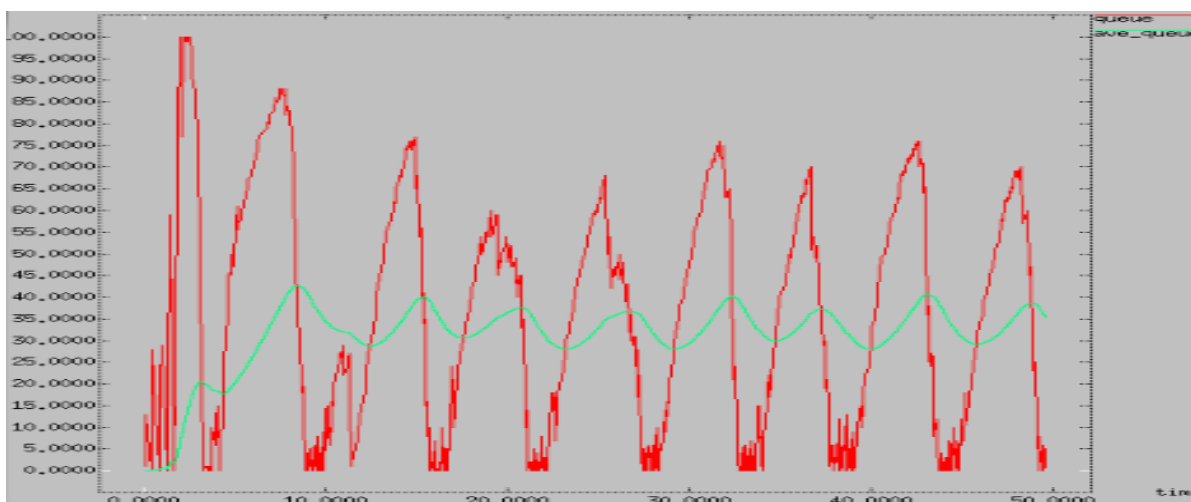


Figura 5.6 : Grafic folosind funcția din algoritm – v4

Din aceste funcții putem observa că valorile cozii medii variază mai mult pentru funcția v4 față de celelalte trei funcții, la v4 valorile cozii curente trecând de la valori foarte mari la valori foarte mici. În cazul primelor 3 funcții, acestea oferă o probabilitate mai mică la aruncarea pachetelor, așa cum vom vedea în simulările ce vor urma. Funcția v4 pare a fi cea mai agresivă din acest punct de vedere.

Pentru aceste 4 funcții am folosit script utilizând limbajul TCL pe care l-am rulat cu comanda :
\$ns RED.tcl x y z

Unde x,y,z sunt parametrii ce descriu funcția:

- x este valoarea uneia din cele 4 funcții;
- y este unul dintre scenariile folosite.

Am utilizat o topologie formată din 16 noduri, acestea fiind legate între ele prin link-uri așa cum urmează:

- Folosirea a 14 link-uri cu lățimea de bandă de 10 Mbps cu delay-ul setat la de 100 ms.

Aceste 14 link-uri practice leagă 15 noduri. 14 dintre noduri trimit pachete către un singur nod, acesta numindu-se nod de acumulare al pachetelor.

Din momentul în care pachetele sosesc în nodul de acumulare, atunci se creează congestia până când o parte din pachetele pe care algoritmul le consideră utile ajung la destinație, adică în nodul terminal.

- Folosirea a unui link bottleneck cu lățimea de bandă de 3 Mbps, cu delay-ul setat la 100 ms.

Am ales ca valoarea lățimii de bandă pentru acest bottleneck să fie una mică, deoarece acesta poate fi luat ca un factor ce duce la congestie, astfel am putut mai ușor evidenția acest aspect.

Dacă aș fi folosit o lățime de bandă mai mare, pachetele sosite în nodul de acumulare ar fi avut suficient spațiu pentru ca buffer-ul să nu se încarce.

Așa cum am prezentat în algoritm, trimiterea pachetelor pe link-ul de bottleneck diferă, transmisia pachetelor către nodul superior fiind diferită față de transmisia pachetelor dinspre nodul superior spre nodul de acumulare al pachetelor.

Cele 14 noduri care trimit pachete către nodul de acumulare folosesc conexiuni FTP care utilizează protocolul TCP prezentat în capitolele anterioare. Nodurile trimit traffic FTP pe parcursul întregii simulări, adică în algoritm am prezentat acest interval de timp ca fiind durata de utilizare a conexiunilor FTP plus o secundă pentru a garanta ca procesele FTP au timp suficient pentru execuție.

Dimensiunea cozii de bottleneck am ales-o ca fiind 100.

Pentru pachetul TCP am ales dimensiunea de 1024 biti, dar practic atunci când pachetul va fi creat și are loc simularea, acesta va fi de 1064 de biți, adică se vor adăuga încă 40 ce sunt utilizați pentru header.

Valorile utilizate pentru parametrii RED în funcție de scenariile folosite la rularea algoritmului din terminal:

Scenariul 1:

minth=15, maxth=45, wq=0,002, unde wq reprezintă încărcarea cozii.

Scenariul 2:

minth=20, maxth=60

Scenariul 3:

minth=25, maxth=75

Scenariul 4:

minth=30, maxth=90

Schema topologiei se găsește în următoarea figură, dar poziția nodurilor este aleatoare atunci când simularea are loc, acest lucru neavând nici o importanță. În exemplul de simulare de mai jos link-ul de bottleneck este cel care leagă nodurile 15 și 16.

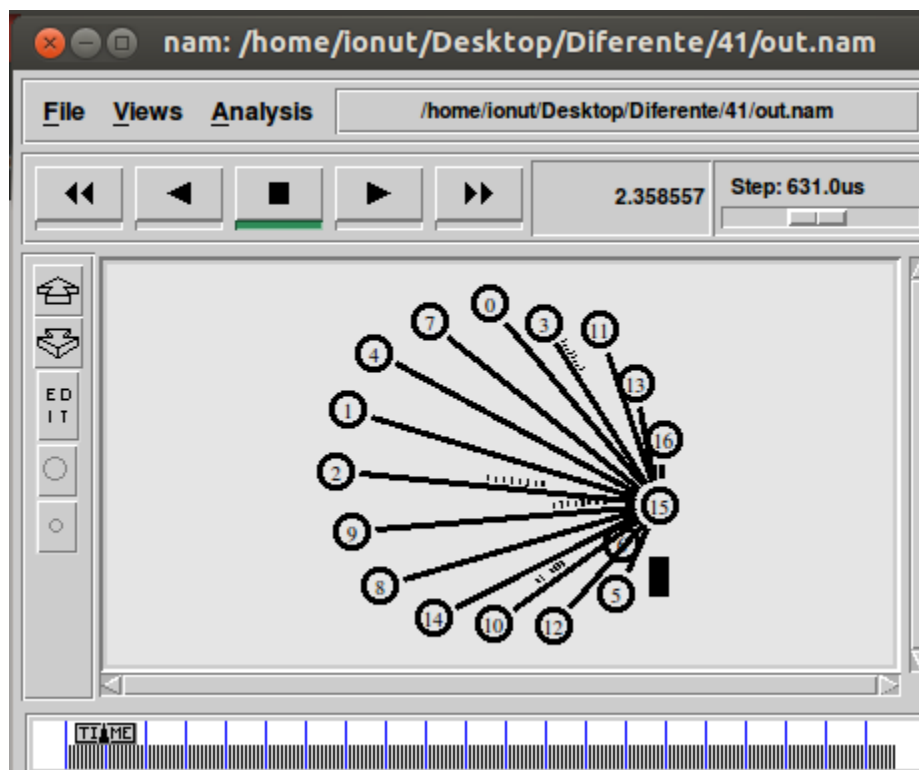


Figura 5.7 : Topologia folosită

Algoritmul TCL poate fi urmărit în Anexa 1.

5.3 Analiza fișierelor de trace

În urma simulărilor acestor funcții utilizând limbajul TCL s-au generat mai multe fișiere de trace, în acestea aflându-se valorile parametrilor ce descriu funcțiile.

Atunci când NS rulează un script tcl, toate evenimentele de tipul pachete primite (receive, r), pachete pierdute (drop, d), pachete puse în capul cozii (enqueue, +) și pachete puse în josul

cozii (dequeue, -) sunt puse într-un fișier de trace (out.tr), care este organizat pe 12 coloane în funcție de anumiți parametri rezultați, ca în figura următoare:

event	time	from node	to node	pkt type	pkt size	flags	fid	src addr	dst addr	seq num	pkt id
Ex: +	160.012933	1	0	tcp	60	-----	2	1.0.1.0	0.0.0.0	0	64

Figura 5.8: Cele 11 detalii despre fiecare tip de eveniment din fișierul principal de trace [23]

5.3.1 Structura fișierelor de trace

Pierderea de pachete (packet loss) apare atunci când unul sau mai multe pachete de date care călătoresc într-o rețea de calculatoare nu reușesc să ajungă la destinația lor. [23] Pierderea de pachete se distinge ca unul dintre cele trei tipuri principale de erori întâlnite în format digital în comunicații; celelalte două fiind bitul de eroare și pachetele false cauzate din cauza zgomotului.

Pachetele pot fi pierdute într-o rețea, deoarece acestea pot fi aruncate atunci când o coadă este supraîncărcată.

Valoarea pierderii de pachete în timpul stării de echilibru este o altă proprietate importantă a unui sistem de control al congestiei. Cu cât valoarea pierderii de pachete este mai mare, cu atât mai dificil este pentru protocoale de transport-layer pentru a menține lățimi de bandă ridicate, sensibilitate la pierderea de pachete individuale, precum și a frecvențelor și modelelor de pierderi în rândul secvențelor de pachete puternic dependente de aplicații.

Debitul de transfer (Throughput)

Aceasta este caracteristica principală de măsurare a performanței, și cel mai des utilizată. În rețelele de comunicații, cum ar fi Ethernet sau pachet radio, debitul de transfer de rețea este rata medie de livrare a mesajului de succes de peste un canal de comunicare. Debitul este de obicei se măsoară în biți pe secundă (bit / s sau bps), și, uneori, în pachete de date pe secundă sau pachete de date pe un interval de timp.

Aceasta măsoară cât de curând receptorul este capabil de a obține o anumită cantitate de date trimise de către expeditor. Este determinat ca raportul dintre totalul datelor primite la un capăt la altul cu întârzieri. Debitul este un factor important care are un impact direct asupra performanței rețelei.

Întârzierea (Delay)

Întârzierea este timpul suplimentar scurs în timp ce un pachet călătorește de la un punct la altul. Cu cât valorile întârzierilor sunt mai mari, cu atât este mult mai dificil pentru protocoalele layer de transport să mențină lățimi de bandă mari.

Lungimea cozii (Queue Length)

Un sistem de așteptare în rețele poate fi descris de pachetele sosesc pentru serviciu, de așteptare pentru serviciu în cazul în care aceasta nu este imediat, iar în cazul în care au așteptat serviciul, părăsesc sistemul după ce au fost servite. Aceasta lungime a cozii este o caracteristică foarte importantă pentru a determina cât de bine lucrează algoritmul controlului congestiei.

Routerele renunță la unul sau mai multe pachete înainte ca buffer-ul să devină complet plin. De fiecare dată când un pachet ajunge în buffer, algoritmul RED calculează lungimea cozii medii (avg).

Dacă lungimea cozii medii este mai mică decât un anumit prag inferior, congestia se presupune a fi minimă sau inexistentă, iar pachetul este pus în așteptare.

Dacă lungimea cozii medii este mai mare decât un anumit prag superior, congestia se presupune a fi de un nivel ridicat și pachetul este aruncat.

Dacă valoarea lungimii cozii medii este între doua praguri, acest lucru ar putea indica debutul de congestie. Probabilitatea de congestie este apoi calculată.

În scriptul TCL realizat, am adăugat în funcția Finish linii de cod care duc la segmentarea fișierului de trace în mai multe fișiere, fiecare reprezentând unul dintre acei parametri care descriu funcția (ex.: ave.tr=lungimea medie a cozii; temp.q=momentele de timp când sunt aruncate pachete etc.).

Cu ajutorul acestei segmentari se poate face mai ușor identificarea și contorizarea valorilor acestor parametri pentru realizarea graficelor.

[24] Cu ajutorul limbajului AWK am extras pe rând evenimentele (+,-,r,d) din fișierul principal de trace (out.tr) folosind următoarea comandă din terminal:

```
$awk -f awk.awk out.tr
```

În urma acestei comenzi am obținut numărul evenimentelor de acel tip.

Algoritmul folosit pentru determinarea numărului de pachete puse în capul cozii (enqueue, +) :

```
BEGIN{
    sum=0;
    }
    {
    if ($1=="+") {
        sum++;
    }
    }
END
{
printf("The number of enqueue is: %d\n",sum);
}
```

5.4 Rezultate și observații

Folosind cele 4 funcții disponibile se vor realiza comparații între ele privind anumiți factori cum ar fi: variația pachetelor aruncate, frecvența de aruncare a pachetelor, numărul pachetelor aruncate, întreruperi, cât și jitter.

În continuare am constatat care este numărul evenimentelor de fiecare tip și am salvat rezultatele într-un fișier de tip .txt pentru fiecare eveniment.

Totodată am constatat că atunci când folosesc aceeași funcție pentru două scenarii diferite obțin același număr de evenimente, ceea ce înseamnă că tipul scenariului (sau alegerea pragurilor de minim și de maxim) folosit nu influențează evenimentele.

În urma simulărilor am introdus rezultatele în tabelul 1:

Funcție\Eveniment	Enqueue (+)	Dequeue (-)	Reveived (r)	Drop (d)
1	65904	65713	65583	191
2	67311	67124	66981	172
3	67584	67350	67207	166
4	66855	66690	66562	324

Tabelul 5.1 : Tabel al funcțiilor în raport cu evenimentele simulării

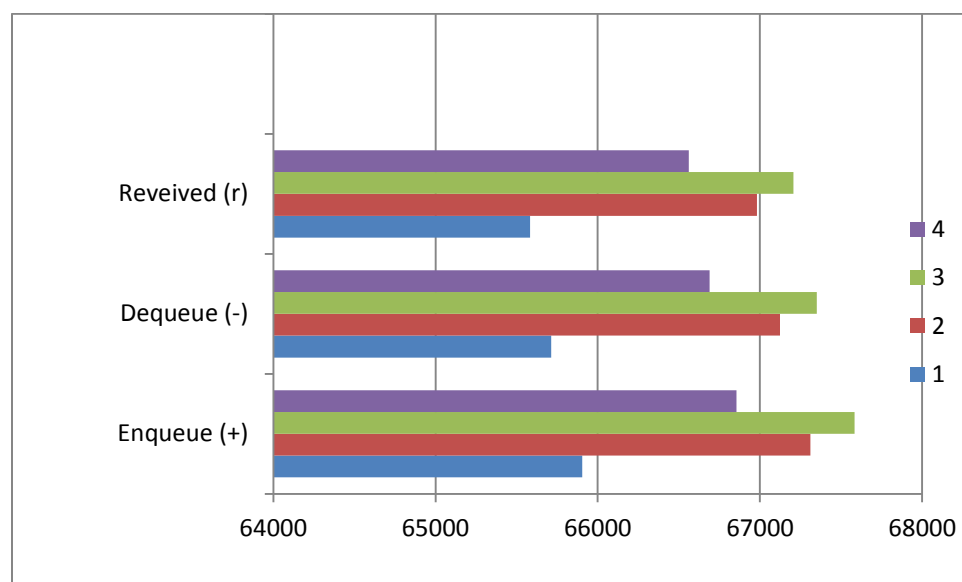


Figura 5.9: Grafic al evenimentelor în funcție de activitatea pachetelor

Se observă din tabel că numărul pachetelor pierdute este mai mic pentru funcțiile v1,v2,v3, iar la funcția v4 este mult mai mare (aproape dublu). Acest fapt se datorează și valorilor cozii medii care variază mai mult pentru funcția v4, aceasta fiind cea mai agresivă din acest punct de vedere.

Se mai poate observa că funcțiile v2 și v3 deși au cel mai mic număr de pachete pierdute, în același timp au și cel mai mare număr de acțiuni de alt tip ale pachetelor. O altă observație este că numărul de acțiuni ale pachetelor este mult mai mic pentru v1 (pentru oricare din evenimente), deoarece funcția v1 nu prezintă o variație atât de mare precum v4.

În continuare se vor realiza comparații privind variația pachetelor aruncate, cât și frecvența de aruncare a pachetelor.

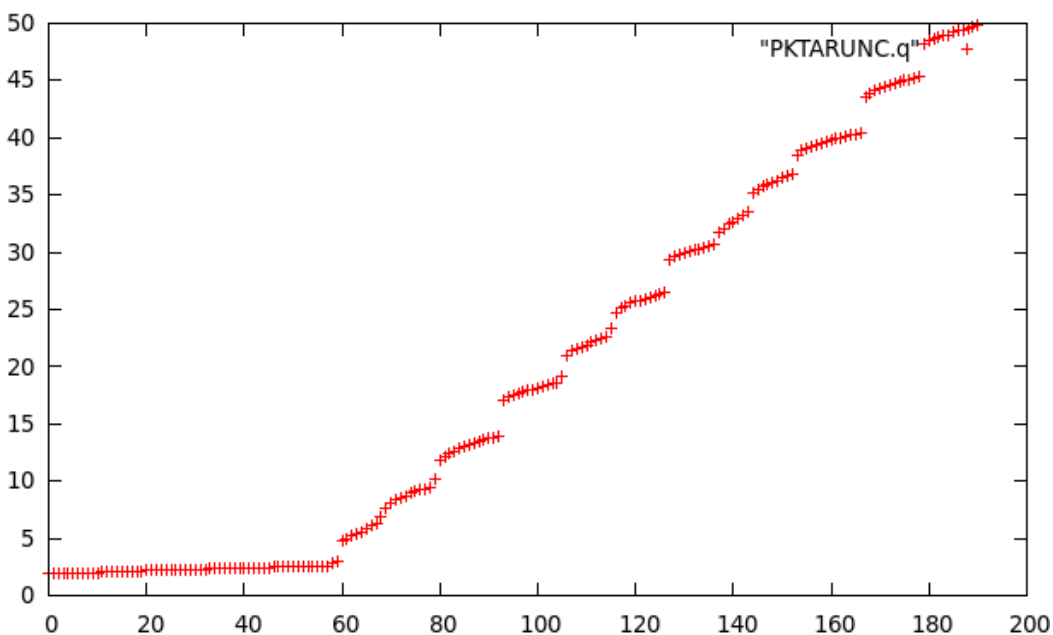


Figura 5.10: Variația pachetelor aruncate în funcție de timp pentru funcția v1

Din grafic se observă că în faza inițială pachetele nu se pierd, urmând după un moment să crească numărul pachetelor aruncate, din cauza încărcării cozii.

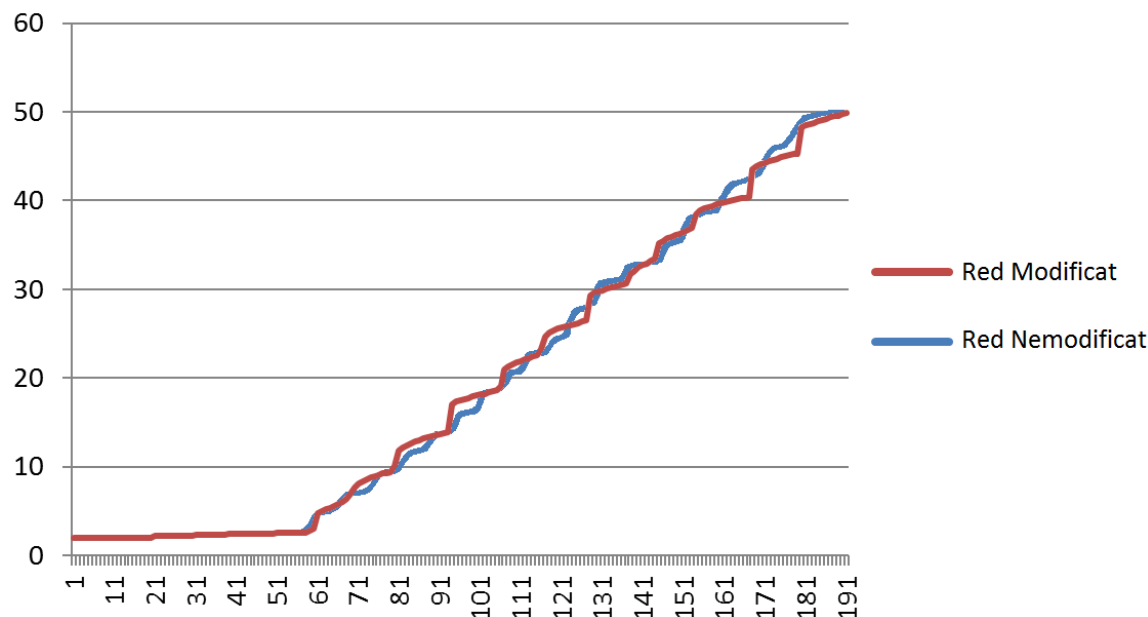


Figura 5.11: Variația pachetelor aruncate pentru algoritmul RED modificat și cel nemodificat

Putem observa din aceasta comparație în figura 5.11 că aruncarea pachetelor este mult mai frecventă pentru RED nemodificat, lucru constatat și în figura 5.13.

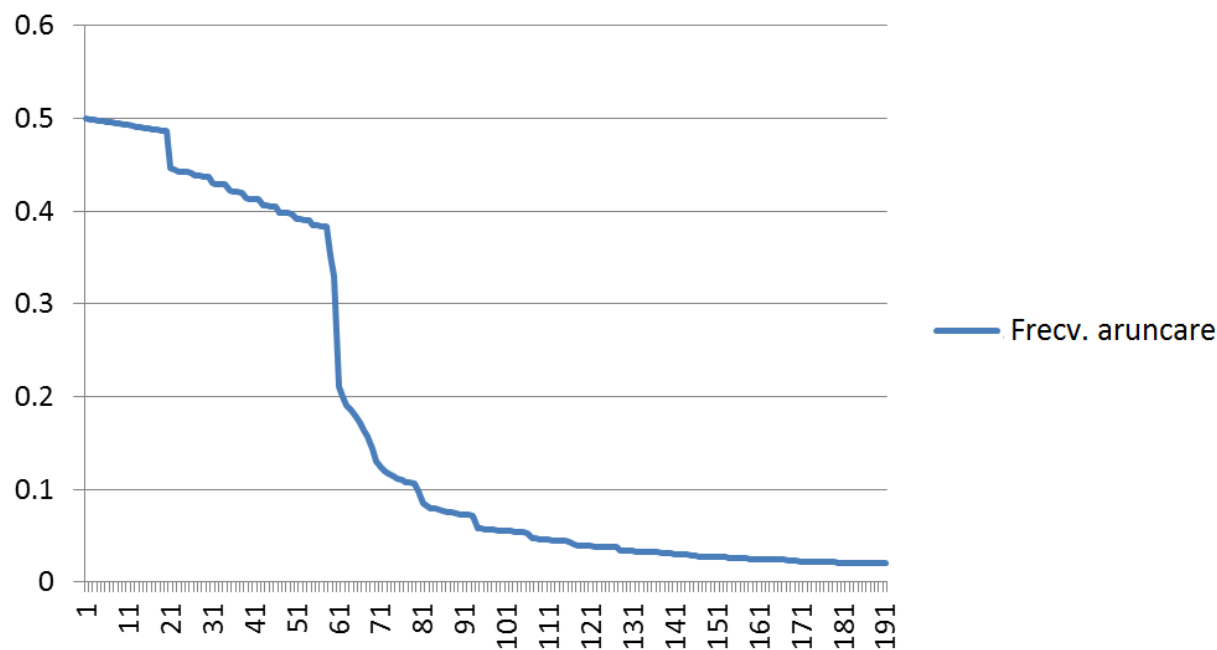


Figura 5.12: Frecvența de aruncare a pachetelor pentru funcția v1

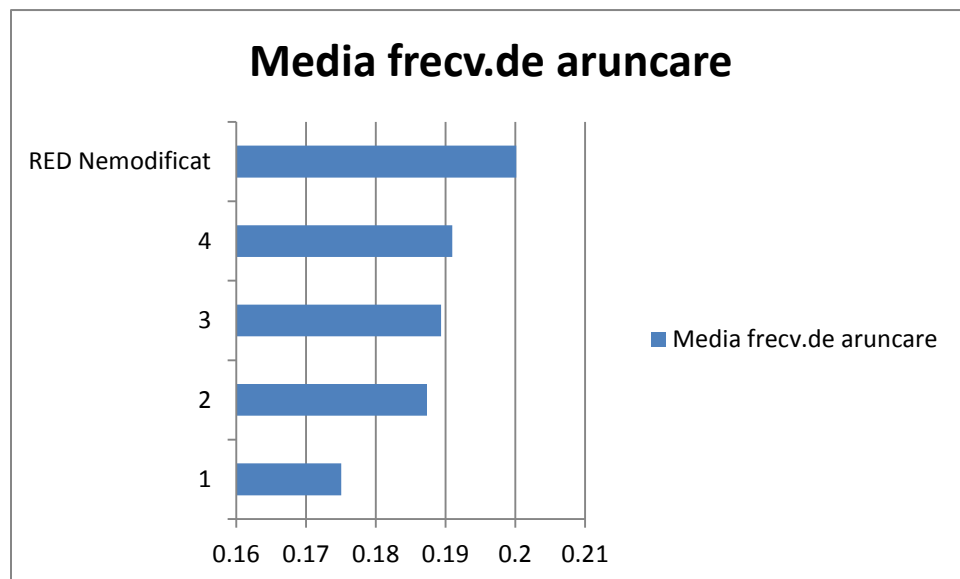


Figura 5.13: Grafic în funcție de media frecvenței de aruncare pentru cele 4 funcții modificate și pentru RED nemodificat

Se observă din figura 5.13 că dintre cele 4 funcții, valoarea cea mai mare a mediei frecvenței de aruncare o deține funcția v4. Mai putem preciza că media frecvenței de aruncare a pachetelor pentru v4 este cu 4,8 % mai mică (daca procentajul este mic, atunci funcția nu este

potrivită) decât cea a algoritmului nemodificat. Funcția optimă pentru media frecvenței de aruncare a pachetelor este funcția v1, deoarece este cu 8% mai mică decât cea pentru algoritmul RED nemodificat.

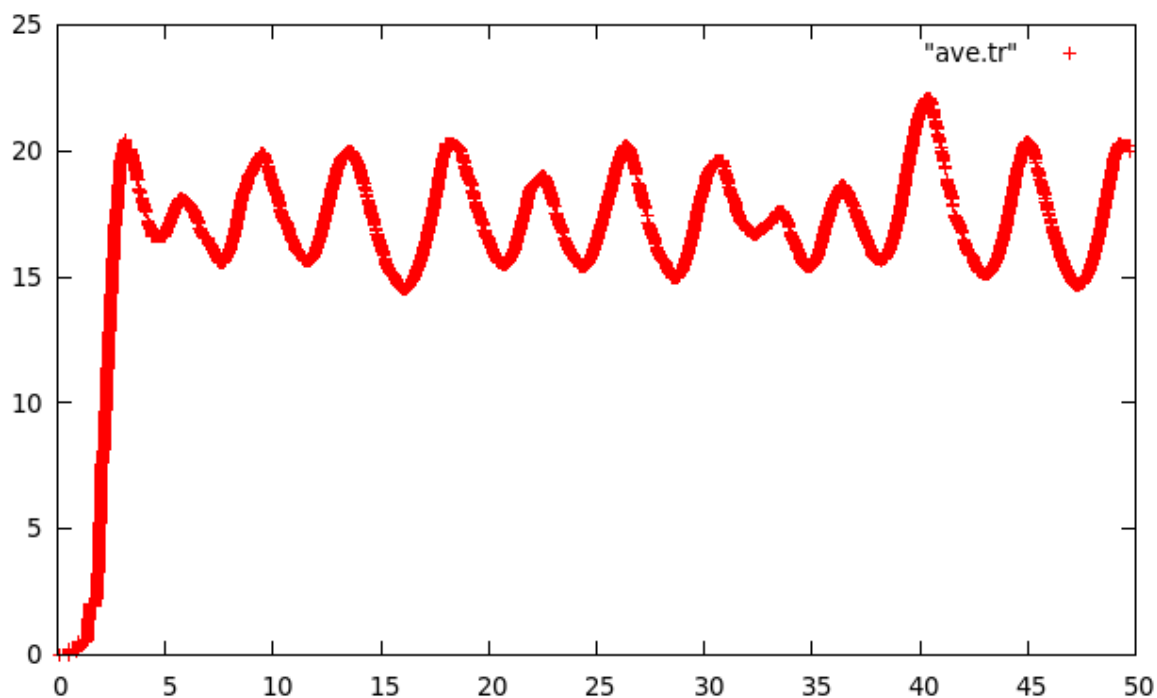


Figura 5.14: Variația cozii medii în funcție de timp pentru funcția v1.

Se observă că valorile cozii medii sunt minime în faza inițială, urmând o creștere bruscă. Acest lucru se datorează timpului până când pachetele vor sosi simultan pe bottleneck.

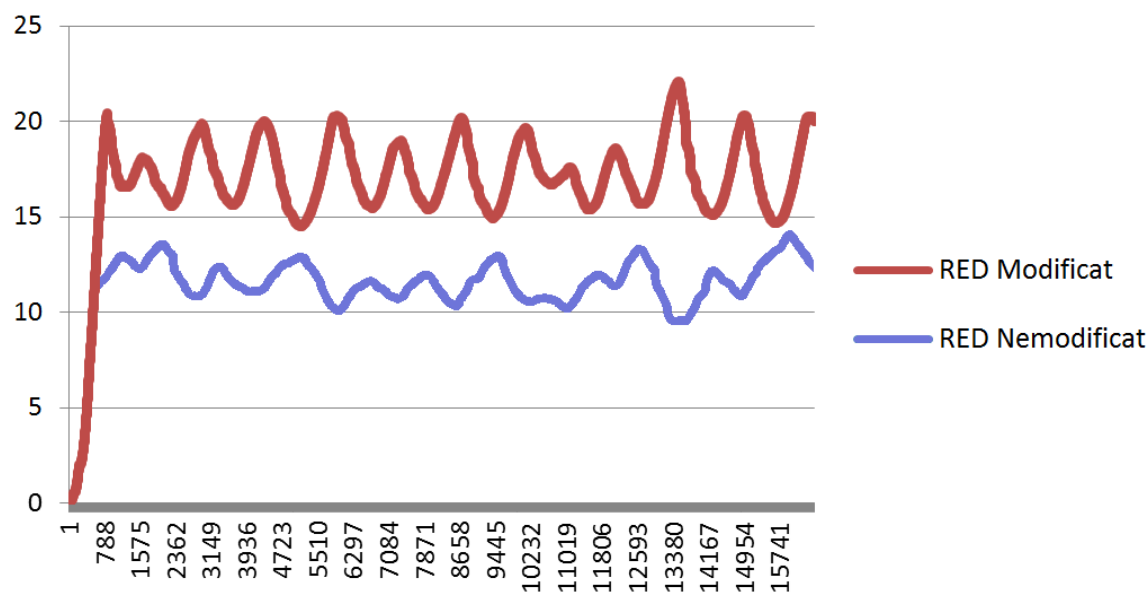


Figura 5.15: Variația cozii medii în funcție de timp pentru RED modificat (v1) și nemodificat

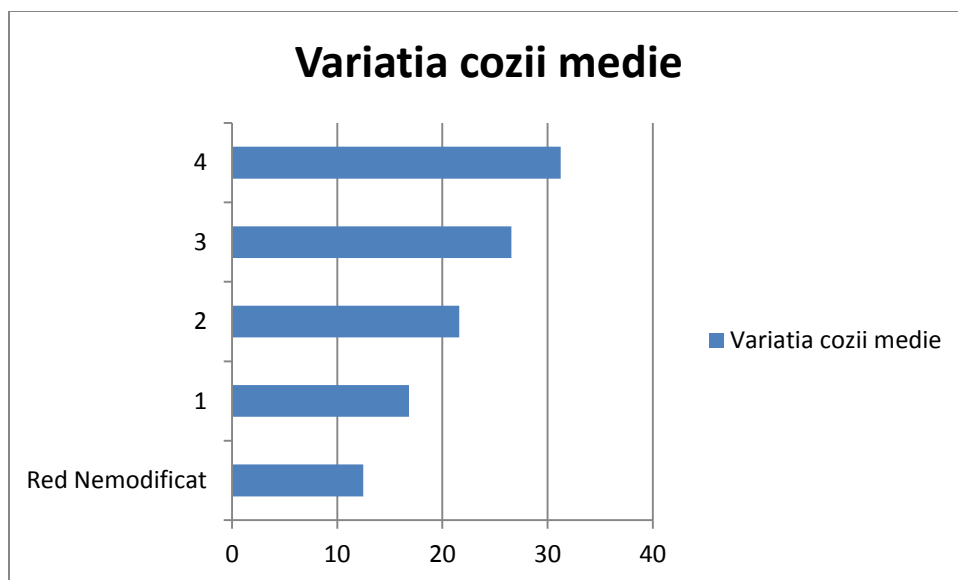


Figura 5.16: Grafic pentru variația medie a cozii (pentru cele 4 funcții)

Urmărind figura 5.16, se observă că funcția v4 variază în amplitudini mult mai mari (de peste două ori) decât cele ale algoritmului nemodificat, chiar și mult mai mari decât celelalte funcții. Pot spune că v4 este funcția cea mai optimă din punct de vedere al mărimilor amplitudinii la care variază.

Am executat cu ajutorul GNU PLOT următoarele grafice folosind fișierele de trace obținute [24] :

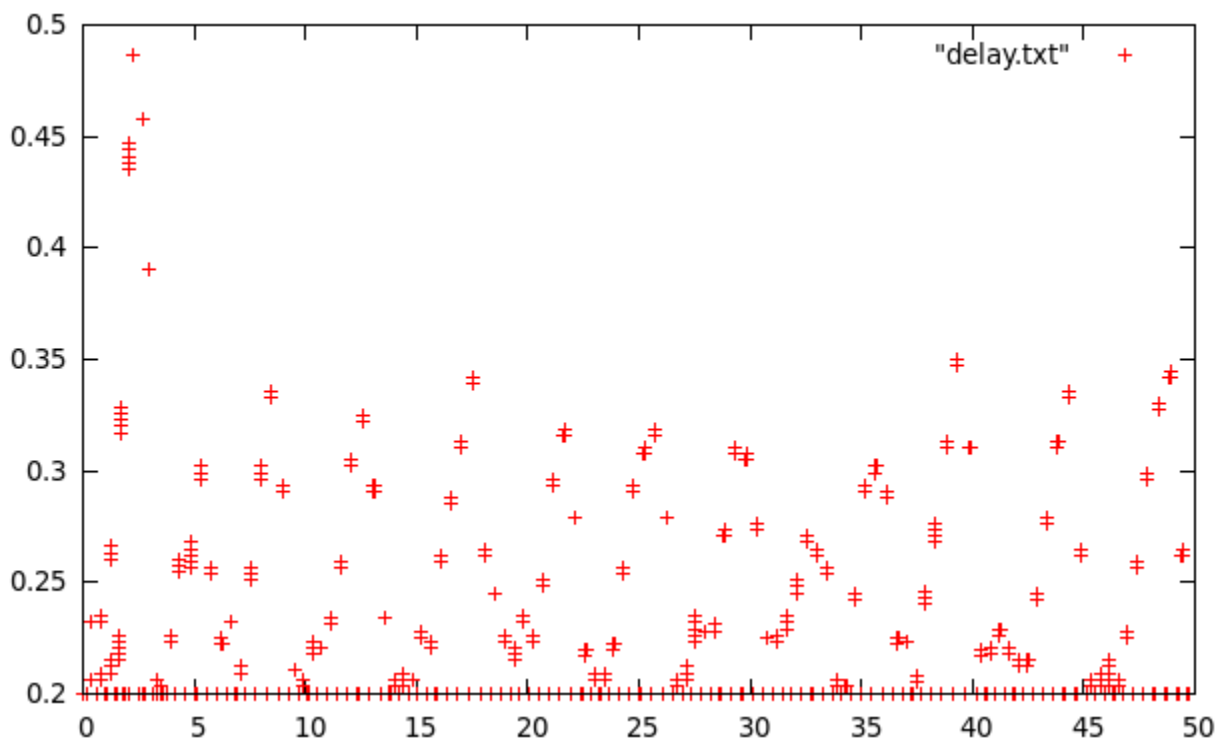


Figura 5.17: Variația întreruperilor (delay) în funcție de timp pentru funcția v1

	Numar Intreruperi	Timpul intreruperii (s)
1	1902	0.234738
2	2838	0.238602
3	2494	0.247863
4	2376	0.256565

Tabel 5.2: Detalii privind delay-ul în toate cele 4 funcții

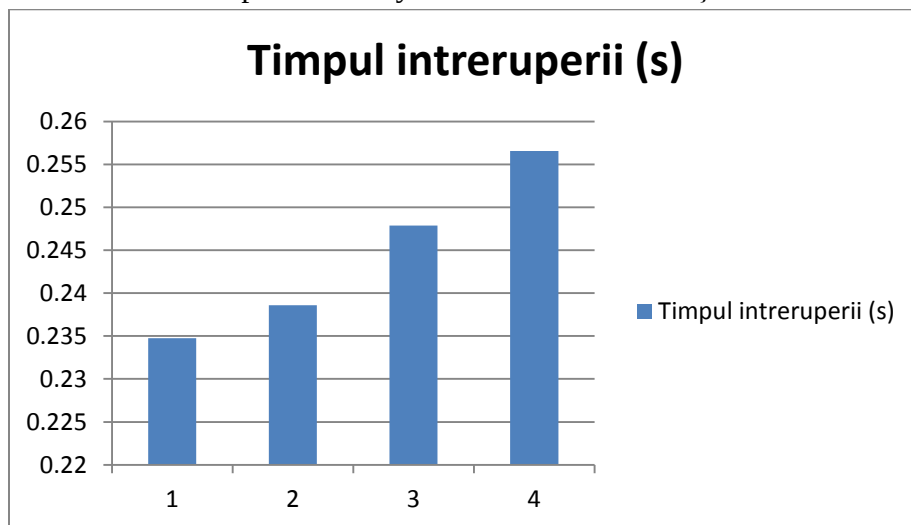


Figura 5.18: Media intervalelor de timp când au loc întreruperile pentru cele 4 funcții

Din tabelul 5.2, cât și din figura 5.18 se observă că funcția v4 deține cel mai mare număr de întreruperi, cât și cele mai mari intervale de timp pentru acestea, ceea ce nu este un lucru bun pentru rețea. Cea mai optimă funcție este v1 privind întreruperile.

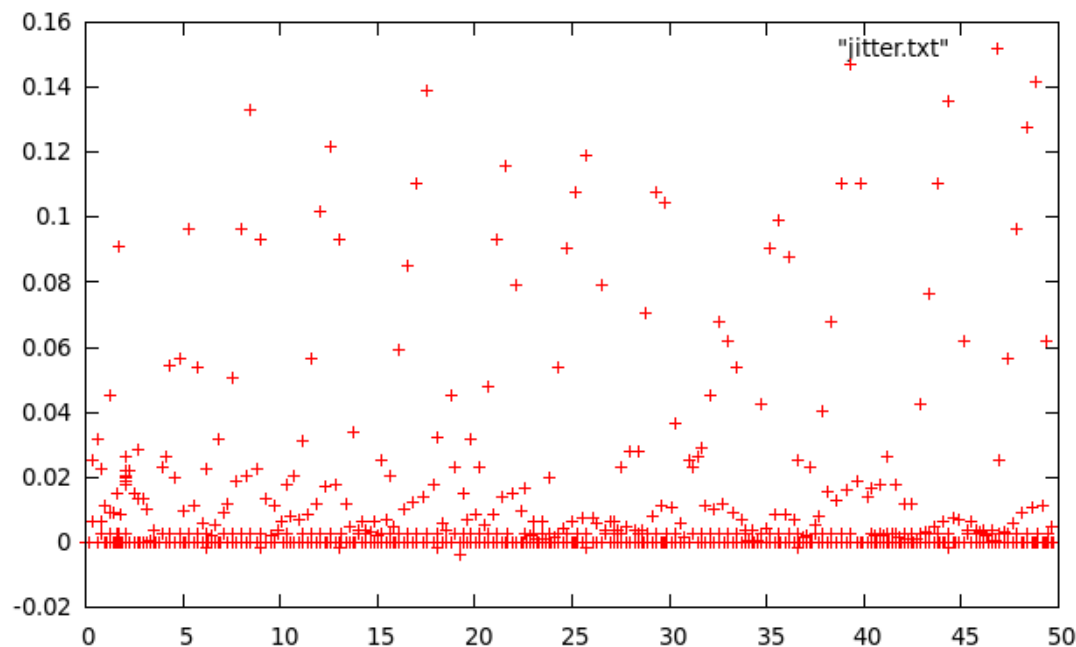


Figura 5.19: Jitter-ul (variația de întârziere) în funcție de timp (pentru funcția v1)

$$\text{Jitter} = ((\text{EndTime}(j) - \text{StartTime}(j)) - (\text{EndTime}(i) - \text{StartTime}(i))) / (j - i) \quad [24]$$

În acest exemplu, am calculat bruiajul pentru traficul CBR prin transmisie UDP

	Valoare Jitter	Numar Jittere
1	0.003747	1902
2	0.002689	2838
3	0.00371	2494
4	0.003853	2376

Tabel 5.3 Detalii privind Jitter-ul în toate cele 4 funcții

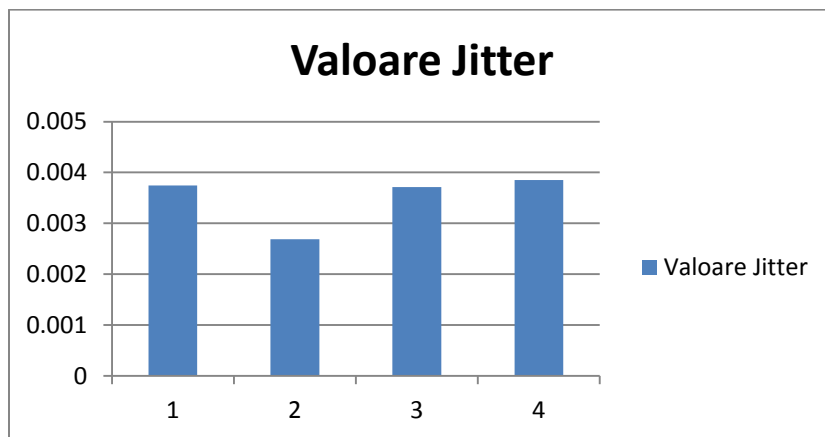


Figura 5.20: Media valorilor jitter-elor celor 4 funcții

Din figura 5.20 se observă că valoarea medie cea mai mică a jitter-elor dintre cele 4 funcții este cea a funcției v2, iar numărul de jittere cel mai mare este tot al funcției v2.

Concluzii

În concluzie, algoritmul Random Early Detection sau RED deține o bună modalitate cu care congestia ce apare la nivel de gateway poate fi evitată prin intermediul protocoalelor de transport [5][3].

În momentul în care valoarea pragului maxim th_{max} va fi depășită de către lungimea cozii medii, atunci gateway-ul va începe să șterga pachetele din coadă. Cu alte cuvinte, folosind acest algoritm, la nivel de gateway congestia poate fi prevenită cu ajutorul informațiilor care circulă în rețea referitoare la valoarea maximă a pragului respectiv sau a lungimii medii a cozii. Aceste lucruri se pot realiza utilizând marcarea pachetelor, aceasta făcându-se în funcție de starea congestiei în acel moment, sau se mai pot realiza primind informații despre lărgimea de bandă din toată rețeaua. Cu ajutorul acestor informații routerele vor putea reduce numărul de pachete pe care vor să îl trimită, dacă este cazul.

Tot în cadrul acestui proiect, pornind de la algoritmul RED inițial, cel inventat de Sally Floyd și de Van Jacobson și urmând modificările aduse acestui algoritm, am reușit să compar anumite funcții și să obțin soluții optime totodată pentru diferitele cazuri. Modificările au fost făcute în codul sursă al simulatorului Network Simulator și totodată am modificat un script .tcl care avea rolul de a genera funcțiile respective, cât și fișierele de trace cu valorile parametrilor.

Totodată, comparând funcțiile respective s-a ajuns la obținerea performanțelor în ceea ce privește media de aruncare de pachete, variațiile cozii medii sau întreruperile.

Fiecare din parametrii au fost examinați în detaliu, iar pentru cazul variației cozii medii am obținut performanțe de peste două ori mai bune decât algoritmul RED inițial, iar pentru media de aruncare de pachete am obținut performanțe cu 8% mai bune.

Bibliografie:

- [1] Computer Networks 4th Edition - Andrew S. Tanenbaum
- [2] http://ro.wikipedia.org/wiki/Protocol_de_control_al_transmisiei
- [3] Analytic Performance Evaluation of the RED Algorithm for QoS in TCP/IP Networks - Stefan Kohler, Michael Menth and Norbert Vicari
<http://www-info3.informatik.uni-wuerzburg.de/TR/tr259.pdf>
- [4] http://en.wikipedia.org/wiki/TCP_congestion-avoidance_algorithm
- [5] Random early detection gateways for congestion avoidance - Sally Floyd and Van Jacobson
<http://www.icir.org/floyd/papers/early.twocolumn.pdf>
- [6] http://en.wikipedia.org/wiki/Random_early_detection
- [7] Quality of Service Control in High-Speed Networks - H. Jonathan Chao, Xiaolei Guo pag 215
- [8] Adaptive RED: An Algorithm for Increasing the Robustness of RED's Active Queue Management - Sally Floyd, Ramakrishna Gummadi, Scott Shenker
- [9] http://www.cisco.com/c/en/us/td/docs/ios/12_2/qos/configuration/guide/fqos_c/qcfconav.html#wp1003370
- [10] <http://networkengineering.stackexchange.com/questions/5316/how-are-weighted-fair-queuing-and-weighted-random-early-detection-related>
- [11] http://en.wikipedia.org/wiki/Weighted_random_early_detection
- [12] RRED: Robust RED Algorithm to Counter Low-Rate Denial-of-Service Attacks - Changwang Zhang, Jianping Yin, Weifeng Chen, Zhiping Cai
- [13] http://en.wikipedia.org/wiki/Robust_random_early_detection
- [14] [http://en.wikipedia.org/wiki/Ns_\(simulator\)](http://en.wikipedia.org/wiki/Ns_(simulator))
- [15] <http://en.wikipedia.org/wiki/OTcl>
- [16] Introduction to Network Simulator - Giovanni Neglia, Mouhamad Ibrahim
- [17] Introduction to Network Simulator 2 - Noun Choi, Hai Vu, Ryan Burchfield, Sara Arbab Yazd
- [18] <http://www.isi.edu/nsnam/nam/index.html>
- [19] <http://wiki.tcl.tk/16279>
- [20] <http://www.isi.edu/nsnam/ns/tutorial/>
- [21] <http://ajlinx.wordpress.com/2013/06/19/install-ns2-ns-allinone-2-35-on-ubuntu-12-04-for-beginners/>
- [22] <http://www.roman10.net/modification-of-red-aqm-in-ns2/>
- [23] Congestion Control in Communication Network Using RED, SFQ and REM Algorithm - Dolly Kalav, Sudha Gupta
- [24] School of Computer Science and Engineering Seoul National University - Min Chen

Pentru cazul referințelor de tip link, acestea au fost verificate in data de 28.06.2014 și toate sunt funcționale.

Anexa 1

```
set ns [new Simulator]
set nf [open out.nam w]
$ns namtrace-all $nf

set NumSenders 15
set NumReceivers 1

#read scenario, seed and bottleneck bandwidth from the command
line input arguments
set Scenario [lindex $argv 0]
set function [lindex $argv 1]
set seed [lindex $argv 2]
puts "scenario: $Scenario; function: $function; seed: $seed"
ns-random $seed
set BufSize 100
set PktSize 1024
#set winSize 200
#in seconds
set Duration 50
#create all nodes: note that the order of creating the nodes
matter
for {set i 1} {$i <= $NumSenders} {incr i} {
    set s($i) [$ns node]
}
set r1 [$ns node]
set d1 [$ns node]
#open the nam trace file
set nf [open out.nam w]
$ns namtrace-all $nf
#open the traffic trace file to record all events
set nd [open out.tr w]
$ns trace-all $nd
#define a finish procedure
proc finish {} {
    global ns , nf , nd , qtf
    $ns flush-trace
    close $nf
    close $nd

    #puts "running nam..."
    exec nam out.nam &
    #exec xgraph out.tr -geometry 800x400 &
    exec grep "^d" out.tr > aruncm.tr
    exec grep "^+" out.tr > enqm.tr
    exec grep "^-" out.tr > dequ.tr
```

```

set awkCode {
    {
        if ($1 == "Q" && NF>2) {
            print $2, $3 >> "temp.q";
            set end $2
        }
        else if ($1 == "a" && NF>2)
            print $2, $3 >> "ave.tr";
    }
}

set awkCode3 {
    {
        print $2 >> "PKTARUNC.q";
        set end $2
    }
}

set f [open temp.queue w]
puts $f "TitleText: red"
puts $f "Device: Postscript"
if { [info exists tchan_] } {
    close $qtf
}
exec rm -f temp.q ave.tr
exec touch ave.tr temp.q
exec awk $awkCode queue.txt
exec awk $awkCode3 aruncm.tr
puts $f "\"queue
exec cat temp.q >@ $f
puts $f \"\\n\\\"ave_queue
exec cat ave.tr >@ $f
close $f
exec xgraph -bb -tk -x time -y queue temp.queue &
exit 0
}

#link the nodes
if { $Scenario == 1 } {
    Queue/RED set thresh_ 15
    Queue/RED set maxthresh_ 45
} elseif { $Scenario == 2 } {
    Queue/RED set thresh_ 20
    Queue/RED set maxthresh_ 60
} elseif { $Scenario == 3 } {
    Queue/RED set thresh_ 25
    Queue/RED set maxthresh_ 75
} elseif { $Scenario == 4 } {
    Queue/RED set thresh_ 30

```

```

    Queue/RED set maxthresh_ 90
}

Queue/RED set queue_in_bytes_ false
Queue/RED set gentle_ false
Queue/RED set function_ $function

for {set i 1} {$i <= $NumSenders} {incr i} {
    $ns duplex-link $s($i) $r1 10Mb 100ms DropTail
    $ns queue-limit $s($i) $r1 $BufSize
}
#r1 d1 and d1 r1 are different
$ns duplex-link $r1 $d1 3Mb 100ms RED
$ns queue-limit $r1 $d1 $BufSize
#trace the queue: note that link r1 d1 is different from d1 r1
set redq [$ns link $r1 $d1 queue]
set qtf [open queue.txt w]

$redq trace curq_
$redq trace ave_
$redq attach $qtf

#set up TCP connections
for {set i 1} {$i <= $NumSenders} {incr i} {
    set tcp($i) [new Agent/TCP]
    $ns attach-agent $s($i) $tcp($i)
    set sink($i) [new Agent/TCPSink]
    $ns attach-agent $d1 $sink($i)
    $ns connect $tcp($i) $sink($i)
    $tcp($i) set fid_ $i
    $tcp($i) set packetSize_ $PktSize

    #$tcp($i) set window_ $winSize
    #set up FTP over TCP connection as traffic source
    set ftp($i) [new Application/FTP]
    $ftp($i) attach-agent $tcp($i)
    $ftp($i) set type_ FTP
}
#schedule events for the FTP agents
set StartTime [expr [ns-random] / 2147483647.0 / 100]
puts "starttime $StartTime"
#temporarily set to 2
for {set i 1} {$i <= $NumSenders} {incr i} {
    $ns at $StartTime "$ftp($i) start"
    $ns at $Duration+$StartTime "$ftp($i) stop"
}

```



```

#ensure the ftp application have enough time to finish, so we +1
$ns at $Duration+$StartTime+1 "finish"
#run the simulation
$ns run

```

Anexa 2:

```

REDQueue::calculate_p_new(double v_ave, double th_max, int
gentle, double v_a,
    double v_b, double v_c, double v_d, double max_p)
{
    double p;
    double lm, ln, lp, lq, Xmin;
    if (gentle && v_ave >= th_max) {
        // p ranges from max_p to 1 as the average queue
        // size ranges from th_max to twice th_max
        p = v_c * v_ave + v_d;
    } else if (!gentle && v_ave >= th_max) {
        // OLD: p continues to range linearly above
max_p as
        // the average queue size ranges above th_max.
        // NEW: p is set to 1.0
        p = 1.0;
    } else {
        // p ranges from 0 to max_p as the average queue
        // size ranges from th_min to th_max
        if (edp_.func == 1) {
            lm = 1.0/(log(th_max) - log(-v_b/v_a));
            ln = -lm*log(-v_b/v_a);
            p = lm*log(v_ave) + ln;
        } else if (edp_.func == 2) {
            Xmin = -v_b/v_a;
            if (v_ave < (th_max + Xmin)/2) {
                lm = 1.4*v_a;
                lp = -lm*Xmin;
                p = lm * v_ave + lp;
            } else {
                ln = (1.4*v_a*(th_max+Xmin)/2-
1.4*v_a*Xmin)/((th_max+Xmin)/2-th_max);
                lq = 1.0 - ln*th_max;
                p = ln * v_ave + lq;
            }
        } else if (edp_.func == 3) {
            if (v_ave < (th_max - v_b/v_a)/2) {
                lm = 0.6*v_a;
                lp = v_b*lm/v_a;
                p = lm*v_ave + lp;
            }
        }
    }
}

```

```

        } else {
            ln = (0.6*v_a*(th_max+Xmin)/2-
0.6*v_a*Xmin)/((th_max+Xmin)/2-th_max);
            lq = 1.0 - ln*th_max;
            p = ln * v_ave + lq;
        }
    } else if (edp_.func == 4) {
        lm = 1.0/(exp(th_max) - exp(-v_b/v_a));
        ln = -lm*exp(-v_b/v_a);
        p = lm*exp(v_ave) + ln;
        //printf("%f, %f, %f, %f\n", exp(th_max), exp(-
v_b/v_a), lm, v_ave);
    } else {
        p = v_a * v_ave + v_b;
    }

    // p = (v_ave - th_min) / (th_max - th_min)
    p *= max_p;
}
if (p > 1.0)
    p = 1.0;
return p;
}

```