

Universitatea “Politehnica” din București

Facultatea de Electronică, Telecomunicații și Tehnologia Informației

Algoritmi AQM cu detecția/semnalarea congestiei

Proiect de Diplomă

prezentat ca cerință parțială pentru obținerea titlului de
Inginer în domeniul Calculatoare și Tehnologia Informației
programul de studii de licență *Ingineria Informației*

Conducător științific
Conf. dr. ing. Ștefan Stăncescu

Absolvent
Muntean Andrada

2014

ACRONYMS

ACK Acknowledgment
AQM Active Queue Management
BDP Bandwidth Delay Product
CPU central processing unit
CRC cyclic redundancy check
CWND congestion window size
DCE Distributed Computing Environment
DNS Domain Name System
FTP File Transfer Protocol
ICMP Internet Control Message Protocol
IEEE Institute of Electrical and Electronics Engineers
IP Internet Protocol
FTP File Transfer Protocol
LAN local area network
LCP link control protocol
HTTP Hypertext Transfer Protocol
MSL maximum segment lifetime
MSS maximum segment size
NS2 Network Simulator 2
OSI open systems interconnection
PDU protocol data unit
PPP Point-to-Point Protocol
PSH push flag, TCP header
RED Random Early Detection
RFC Request for Comment
RTT round-trip time
SMTP Simple Mail Transfer Protocol
SMS Systems Management Server
SSTHRESH Slow-start threshold Value
SACK selective acknowledgment
SYN synchronize packet
TCP Transmission Control Protocol
TTL time-to-live
UDP User Datagram Protocol
VOIP Voice over Internet Protocol.
QoS Quality of Service

Cuprins

Introducere.....	7
1. Congestia traficului	9
1.1 Definirea Congestiei.....	9
1.2 Principii generale ale controlului congestiei.....	15
1.3 Politici pentru prevenirea congestiei	17
2. Tehnici cunoscute de prevenire a pierderii pachetelor în rețele	21
2.1 Introducere.....	21
2.2 TCP	22
2.3 Algoritmi de management al cozii la receptor (AQM).....	27
3 TCP NewReno.....	33
3.1 Introducere.....	33
3.2. TCP Reno	35
3.2 Caracteristici New Reno	39
4. Modificare Algoritm.....	43
4.1 Introducere.....	43
4.2 Soluția propusă pentru modificarea algoritmului	44
5. Simulări	45
Modificarea simulatorului	47
Scenariul Simulat :	49
6. Prezentarea rezultatelor	49
7. Concluzii.....	58
8. Bibliografie și referințe.....	61

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

Subsemnatul MUNTEAN FV ANDRADA²⁾, posesor al
B.I./C.I./pașaport seria UX, nr. 216665, având CNP

2	9	1	0	7	2	7	3	8	5	5	7	5
---	---	---	---	---	---	---	---	---	---	---	---	---

,
am întocmit proiectul de absolvire/diplomă/lucrarea de disertație cu
tema:³⁾ ALGORITMI AQM DE DETECTIE / SEMNALARE A CONGESTIEI

în vederea susținerii examenului de finalizare a studiilor universitare de licență/masterat,
organizat de către Facultatea de ELECTRONICĂ, TELECOMUNICAȚII ȘI TEHNOLOGIA INFORMATICII
Departamentul CALCULATOARE ȘI TEHNOLOGIA INFORMATICII 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: 26.06.2014

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.

Introducere

În această lucrare propun o modificare a algoritmului New Reno în etapa de “Fast recovery” a prafului de Slow Start threshold . Această modificare urmărește să aducă o îmbunătățire a ferestrei de congestive și a lățimii de bandă. Urmăresc comportamentul algoritmilor New Reno, Reno și Tahoe în comparație cu algoritmul New Reno modificat propus de mine. Deoarece New-Reno are la bază algoritmul Reno, modificat la partea de Fast-recovery, iar Reno la rândul său are la bază algoritmul Tahoe, îmbunătățind partea de fast Recovery, iar Tahoe este alcătuit din Slow start, congestion avoidance and fast retransmit, dacă se îmbunătățește una din cele patru părți ce ajută la funcționarea algoritmului (Slow start, congestion avoidance and fast retransmit, fast-recovery), acesta va reprezenta o versiune îmbunătățită.

Problema atât cu Reno cât și cu NewReno este că în algoritmul de recuperare rapidă, TCP înjumătățește fereastra de congestie, indiferent de starea rețelei. O altă problemă cu NewReno este că, atunci când nu există pierderi de pachete, dar pachetele sunt reordonate cu mai mult de trei confirmări duplicate, NewReno intră în mod eronat în recuperare rapidă, iar acest aspect duce iar la înjumătățirea ferestrei de congestie. Deoarece fereastra TCP de congestie controlează numărul de pachete pe care un expeditor TCP le poate trimite în rețea, în orice moment, prin urmare, procesul de stabilire a dimensiunii ferestrei de congestie la jumătate din valoarea face ca algoritmul NewReno TCP să fie inefficient în ceea ce privește utilizarea capacității de legătură a algoritmului ce stă la baza lui .

În lucrarea de față, propun o modificare a fazei de recuperare rapidă, pe care o să o compar pe aceeași topologie de rețea cu performanțele algoritmilor New Reno, Reno și Tahoe. Urmăresc comportamentul ferestrei de congestive , întârzierea pachetelor, lățimii de bandă , numărul de pachete pierdute , Numarul de pachete pierdute , întârzierea între pachete și întârzierea end-to-end.

Scopul lucrării este cel de a îmbunătăți modul de funcționare a algoritmului New Reno .

1. Congestia traficului

1.1 Definirea Congestiei

Congestia se poate defini în raport cu supraîncărcarea rețelei sau reacția rețelelor la încărcări mari sau foarte mari. Congestia apare atunci când pe o legătură se dorește a fi transmis un volum de date mai mare decât capacitatea sa. Sarcina de a transmite un volum de date mai mare decât capacitatea legăturii pe care se va transmite, fenomenul cauzând pierderi de pachete transmise prin rețea, se numește congestie.

Deoarece se folosește un singur mijloc de comunicare, performanțele rețelei scad o dată cu creșterea pachetelor transferate în rețea, numărul de utilizatori și a dimensiunilor aplicațiilor.

Factorii care duc la producerea congestiei într-o rețea sunt :

- cantitatea traficului;
- dimensiunea pachetelor;
- dimensiunile fizice ale rețelei.

Timpul de răspuns al rețelei depinde de încărcarea acesteia. În funcție de numărul de coliziuni, care cresc o dată cu încărcarea, determinând retransmișii ce încarcă și mai mult rețeaua, va descrește sau va crește performanța rețelei.

Ceea ce duce ca o rețea să fie congestionată poate să fie atât supraîncărcarea bufferelor echipamentelor de rețea și implicit a cozilor rezultând pierderi de pachete), ca și cazul în care bufferele nu sunt supra-încărcate ci doar utilizatorul nu are o disponibilitatea a serviciilor rețelei ce îi poate satisface nevoile aplicației folosite.

O schemă de control al congestiei trebuie să îndeplinească anumite cerințe fundamentale.

Acestea sunt:

- 1) Eficiență
- 2) Heterogenitate
- 3) Capacitatea de a funcționa cu surse defecte
- 4) Stabilitate
- 5) Scalabilitate

6) Simplitate

7) Echitate

Eficiența

Sunt două aspecte de menționat în ceea ce privește eficiența unei scheme de control a congestiei. În primul rând trebuie să se țină cont de cantitatea de trafic suplimentar pe care schema de control a congestiei o aduce în rețea. Am dori desigur că o asemenea schemă să presupună o povară minimă pentru rețea. În al doilea rând trebuie să ne gândim dacă schema de control nu conduce la o subutilizare a resurselor critice ale rețelei. Scheme de control ineficiente pot “gâtui” sursele chiar dacă nu există pericol de congestie, conducând la o subutilizare a resurselor. Nu se dorește să se lucreze suboptimal într-o rețea.

Heterogenitatea

Odată cu creșterea dimensiunilor rețelelor cât și a acoperirii lor, acestea tind să înglobeze o varietate din ce în ce mai mare de arhitecturi hardware și software. O schemă de control care presupune o singură dimensiune de pachet, un singur tip de protocol la nivelul transport sau un singur tip de serviciu nu poate funcționa în regulă într-un așa mediu. Așadar, ne dorim o schemă de control implementabilă peste o largă varietate de arhitecturi de rețea.

Capacitatea de a funcționa cu surse defecte

În rețelele actuale, administratorul de rețea nu deține mijloace administrative de control asupra surselor de mesaje. Deci, sursele (de exemplu, utilizatorii de PC-uri) sunt liberi să manipuleze protocoalele de rețea pentru a maximiza utilitatea obținută de la rețea. Atunci când un ruter informează o sursă să-și reducă rata de trimitere, o sursă funcțională va face acest lucru. Totuși, este posibil că o sursă „defectă” să ignore aceste semnale, în speranța că acest lucru îi va permite să trimită mai multe date. O schemă de control a congestiei nu trebuie să eșueze în prezența unor asemenea surse (o modalitate este “pedepsirea” surselor defecte).

Stabilitatea

Controlul congestiei poate fi privit că o problemă clasică de feed-back negativ. Complexitatea suplimentară se datorează faptului că semnalele de control sunt întârziate. Cu alte cuvinte, există o întârziere finită între momentul detectării congestiei și momentul recepționării acestui semnal de către sursă. Mai mult, sistemul prezintă zgomot iar observările parametrilor sistemului pot fi corupte. Această complexitate poate induce instabilitate în rețea. Astfel, vrem ca schema de control să fie robustă și dacă este posibil, dovedit stabilă.

Scalabilitatea

Este însuși natura sistemelor distribuite să crească odată cu trecerea timpului. Cheia succesului într-un asemenea mediu este scalabilitatea. O schemă de control performantă a congestiei ar trebui să fie scalabilă în ceea ce privește doi parametri: lățimea de bandă și dimensiunea rețelei.

Transmisiunile pe fibră optică au făcut posibilă dezvoltarea de rețele cu lățimi de bandă cu trei ordine de mărime mai mari decât predecesorii lor nu foarte îndepărtați (45000 kbps față de 56 kbps). În altă ordine de idei, sute de mii de rețele locale au fost interconectate într-o enormă inter-rețea ce acoperă întregul glob. Această expansiune face că scalabilitatea unei scheme de control a congestiei să fie imperativă.

Simplitatea

Simplitatea este întotdeauna un bun de preț. Protocoale mai simple sunt mai ușor de implementat și pot suporta mai ușor creșterile de lățime de bandă. De asemenea, protocoalele simple sunt mai probabil acceptate ca standarde. Așadar, o schemă de control al congestiei ideală trebuie să fie ușor de implementat.

Echitatea

Poate fi necesar ca anumite surse să-și reducă transmisiile datorită congestiei. Modalitatea de a alege care surse vor fi restricționate (fie cerându-le să-și reducă transmisiile, fie renunțând la unele dintre pachetele trimise de acestea) determină cât de echitabil își alocă rețeaua resursele. De exemplu, dacă o cerere excesivă din partea sursei A „sufocă” sursa B, este clar că rețeaua nu îl tratează echitabil pe B. Nu este de dorit o schemă de control care să genereze o alocare neechitabilă a resurselor.

Când se vorbește despre echitate, sunt două aspecte de care trebuie ținut cont: definirea și implementarea. Numeroase criterii de definire a echității au fost propuse de-a lungul timpului dar nici unul nu este absolut. Implementarea unui criteriu de echitate generează probleme care uneori depășesc domeniul controlului congestiei. De exemplu, s-ar putea decide că este echitabil să se avantajeze anumite surse care sunt dispuse să plătească pentru acest lucru sau care trimit un anumit tip de pachete (mai importante - email). [1]

Toți algoritmi de gestionare a cozii active inventați până acum au fost proiectați pentru a asigura alocarea de lățime de bandă corect între fluxurile TCP, care sunt proiectate să coopereze cu clasicul TCP de control al congestiei.

Este recunoscut faptul că combinația/combinarea mecanismului de combatere a congestiei împreună cu algoritmul de routare prin cozi nu suportă o împărțire corectă a benzii între cele două linkuri. În particular, conexiunile cu RTT-uri nu primesc rația cuvenită în ceea ce privește lungimea de bandă. Mai mult, mecanismele menționate mai sus nu se supun normelor de control al congestiei. Problema va fi abordată în două perspective. Prima, a fost propus un algoritm de tratare a congestiei cu ajutorul mai multor variante corecte de TCP (Hybla TCP – ce realizează alocarea de lățime de bandă justă prin intermediul supracompensării ferestrei de congestie pentru conexiunile cu RTT-uri lungi.)

Dacă debitul cu care intră pachete într-un nod depășește fie capacitatea de prelucrare a nodului, fie capacitatea legăturii prin care pachetele trebuie să iasă, nodul memorează pachetele într-o structură de coadă, de unde le extrage pe măsură ce pot transmise prin legătura de ieșire. Un efect imediat este creșterea timpului de propagare, din cauza staționării pachetelor în coada de așteptare. Dacă excesul de debit de intrare se păstrează mai mult timp, coada crește până când memoria alocabilă cozii de așteptare este epuizată; în acel moment, nodul va trebui fie să sacrifice pachete, fie să solicite, prin intermediul mecanismului de control al fuxului de la nivelul legăturii de date, micșorarea debitului de intrare. De notat că reducerea și, în extremis, blocarea fuxului de date la intrarea într-un nod poate duce la acumularea de date de transmis în nodul vecin dinspre care vine acel ux, ducând mai departe la blocarea reciprocă a unui grup de noduri.

Principalele probleme ce apar în cazul în care capacitatea nodurilor sau legăturilor directe este depășită sunt următoarele:

- 1) Într-o rețea aglomerată, pachetele sau datagramele întârzie mult sau chiar se pierd, lucru care poate declanșa retrimiteri intempestive de pachete, ducând la aglomerare și mai mare a rețelei și la performanțe și mai scăzute. O astfel de situație, de degradare suplimentară a performanțelor în urma creșterii încărcării, se numește *congestie* și trebuie evitată sau, cel puțin, ținută sub control.
- 2) Capacitatea disponibilă a rețelei trebuie împărțită în mod echitabil între utilizatori.
- 3) Diferite aplicații au diferite priorități cu privire la caracteristicile necesare ale serviciului oferit de rețea. Reacția unei rețele aglomerate trebuie să țină cont de aceste priorități. De exemplu, un nod supraaglomerat poate să fie nevoit să arunce o parte dintre datagramele aflate în tranzit.

Dacă datagramele aparțin unei aplicații de transfer de fișiere, este preferabil să fie aruncate cele mai recente (acestea fiind retransmise mai târziu; dacă se aruncă datagramele mai vechi, este posibil ca destinatarul să nu aibă ce face cu cele mai noi și să trebuiască retransmise toate).

Pentru exemplificare, să considerăm o rețea în care, dintr-un anumit motiv, rata de sosire a pachetelor la un ruter depășește rata de deservire a acestuia. În acest moment, pachetele sunt introduse într-o coadă de așteptare, ceea ce conduce la întârzieri. Întârzierea adițională poate determina sursele să retransmită, crescând astfel și mai mult încărcarea rețelei. Acest tip de feedback conduce la o deteriorare rapidă a situației, unde traficul este dominat de retransmisii și productivitatea efectivă – banda efectivă (throughput) scade brusc. Mai departe, dacă este implementat un mecanism de control al traficului „router-to-router”, noi pachete pot să nu mai

fie admise în coadă, așa că acestea pot inunda un ruter precedent de asemenea. Se poate ajunge astfel la o situație de tipul „**deadlock**”, în care întreg traficul este înghețat.

Avem de-a face cu trei evenimente ce apar simultan. În primul rând, întârzierea datorată cozilor de așteptare crește. În al doilea rând, pot apărea pierderi de pachete. În final, în starea congestionată, traficul este dominat de retransmisii, astfel încât eficiența scade. Definițiile standard ale congestiei sunt de forma următoare: „O rețea este congestionată dacă, datorită supraîncărcării, condiția X se îndeplinește”, unde X reprezintă creșterea excesivă a întârzierii, pierderi de pachete sau scăderi ale benzii efective.

Aceste definiții nu sunt satisfăcătoare din mai multe motive. În primul rând, întârzierile și pierderile de pachete sunt indicatori de performanță ce sunt impropriu folosiți drept indicatori pentru congestie, dat fiind că modificarea acestor indicatori se poate datora unor fenomene diferite de congestie. În al doilea rând, definiția nu specifică în mod exact punctul de la care o rețea poate fi considerată congestionată. De exemplu, în timp ce o rețea care are timpi medii de întârziere în fiecare ruter de 1 până la 10 durate de servire nu este cu siguranță congestionată, nu este clar dacă o rețea ce are timpi medii de întârziere de 1000 durate de servire este sau nu congestionată. Nu pare posibil să se stabilească o valoare de prag care să determine congestia! Nu în ultimul rând, o rețea congestionată din punctul de vedere al unui utilizator nu este neapărat congestionată din punctul de vedere al altuia. De exemplu, dacă utilizatorul A poate tolera o rată de pierdere a pachetelor de 1 la 1000, și utilizatorul B poate tolera o rată de pierdere de 1 la 100, și rata de pierdere efectivă este de 1 la 500, atunci A va considera că rețeaua este congestionată, în timp ce pentru B rețeaua va funcționa în parametri normali. O rețea trebuie considerată necongestionată numai dacă toți utilizatorii săi sunt de acord cu presupusa stare.

Din cele prezentate mai sus, deduce faptul că starea de congestionare a unei rețele depinde de perspectiva utilizatorului. Un utilizator care cere puțin de la o rețea poate tolera o pierdere de performanță mult mai bine decât un utilizator cu cerințe mari. De exemplu, un utilizator care folosește rețeaua numai pentru e-mail va fi mulțumit cu o întârziere de o oră, pe când acesta „performanță” nu este acceptabilă pentru un utilizator care folosește rețeaua pentru transmisiuni de voce în timp real. Punctul cheie este noțiunea de **folos**, folos pe care un utilizator îl obține de la rețea, și modul în care acest folos se degradează odată cu creșterea încărcării rețelei.

Definiție

O rețea este congestionată din perspectiva utilizatorului i dacă utilitatea lui I scade ca urmare a unei creșteri a încărcării rețelei.

Observații:

1. O rețea poate fi congestionată din perspectiva unui utilizator și necongestionată din perspectiva altuia.
2. O rețea poate fi considerată strict necongestionată numai dacă niciun utilizator nu o percepe că fiind congestionată.

O rețea care controlează congestia trebuie să fie sensibilă la funcțiile de utilitate ale utilizatorilor și să fie capabilă să își direcționeze resursele astfel încât să nu existe pierderi de

utilitate odată cu creșterea încărcării. Așadar, rețeaua trebuie să fie capabilă să facă diferența între conversații și să prioritizeze conversațiile în funcție de stringența utilității participanților la o anumită conversație.

Congestia poate fi produsă de mai mulți factori. Dacă dintr-o dată încep să sosească șiruri de pachete pe trei sau patru linii de intrare și toate necesită aceeași linie de ieșire, atunci se va forma o coadă. Dacă nu există suficientă memorie pentru a le păstra pe toate, unele se vor pierde. Adăugarea de memorie poate fi folositoare până la un punct, dar Nagle a descoperit în 1987 că dacă ruterele ar avea o cantitate infinită de memorie, congestia s-ar înrăutăți în loc să se amelioreze, deoarece până să ajungă la începutul cozii pachetele au fost deja considerate pierdute (**timeout** repetat) și s-au trimis duplicate. Toate aceste pachete vor fi trimise cu conștiinciozitate către următorul ruter, crescând încărcarea de-a lungul căii către destinație.

În figura 1.1 se prezintă simptomele unei rețele congestionate.

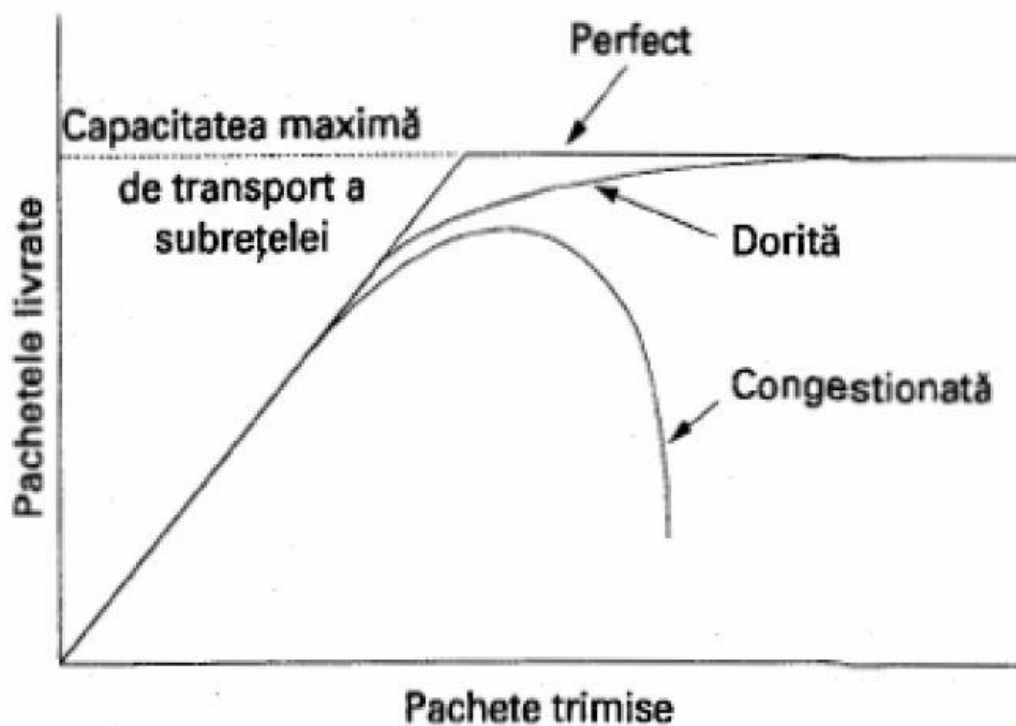


Figura 1.1: Dacă traficul este prea intens, apare congestia și performanțele se degradează puternic.

O altă cauză sunt procesoarele lente care pot cauza congestia. Dacă unitatea centrală a ruter-ului (CPU) este lentă în execuția funcțiilor sale (introducerea în cozi, actualizarea tabelor etc.), se pot forma cozi, chiar dacă linia de comunicație nu este folosită la capacitate maximă. Similar și liniile cu lățime de bandă scăzută pot provoca congestia. Schimbarea liniilor cu unele mai performante și păstrarea aceluiași procesor sau vice-versa de obicei ajută puțin, însă de cele

mai multe ori nu fac decât să deplaseze punctul critic. Adevărata problemă este de multe ori o incompatibilitate între părți ale sistemului. Ea va persista până ce toate componentele sunt în echilibru. [1]

Este important de subliniat diferența dintre controlul congestiei și controlul fluxului, deoarece relația este subtilă.

Controlul congestiei trebuie să asigure că subrețeaua este capabilă să transporte întreg traficul implicat. Este o problemă globală, implicând comportamentul tuturor calculatoarelor gazdă, al tuturor rutelor, prelucrarea de tip store-and-forward din rutere și toți ceilalți factori care tind să diminueze capacitatea de transport a subrețelei.

Controlul fluxului, prin contrast, se referă la traficul capăt la capăt între un expeditor și un destinatar. Rolul său este de a împiedica un expeditor rapid să trimită date continuu, la o viteză mai mare decât cea cu care destinatarul poate consuma datele. Controlul fluxului implică frecvent existența unui feed-back de la receptor către emițător, pentru a spune emițătorului cum se desfășoară lucrurile la celălalt capăt.

Pentru a vedea diferența dintre aceste două concepte, să considerăm o rețea cu fibre optice cu o capacitate de 1000 Gbps pe care un calculator încearcă să transfere un fișier către un calculator personal la 1 Gbps. Deși nu există nici o congestie (rețeaua nu are nici un fel de probleme), controlul fluxului este necesar pentru a forța supercalculatorul să se oprească des, pentru a permite calculatorului personal să primească datele.

La cealaltă extremă, să considerăm o rețea de tip store-and-forward cu linii de 1 Mbps și 1000 de calculatoare mari, din care jumătate încearcă să transfere la 100 kbps către cealaltă jumătate. Aici problema nu apare datorită unui emițător rapid care surclasează un receptor lent, ci pentru că traficul cerut depășește posibilitățile rețelei. Motivul pentru care controlul congestiei și controlul fluxului sunt adesea confundate este acela că unii algoritmi pentru controlul congestiei funcționează trimițând mesaje înapoi către diferitele surse, spunându-le să încetinească atunci când rețeaua are probleme. Astfel, un calculator gazdă poate primi un mesaj de încetinire fie din cauză că receptorul nu suportă încărcarea, fie pentru că rețeaua este depășită.[3]

1.2 Principii generale ale controlului congestiei

Controlul congestiei este un proces foarte complex care nu și-a găsit o rezolvare cât de cât mulțumitoare nici în ziua de azi. Chiar dimpotrivă, acest subiect este din ce în ce mai des abordat în cercetările din informatică din cauza aglomerării din ce în ce mai mari a rețelelor de calculatoare. Putem extinde problema congestiei chiar la nivelul general al unei rețele de telecomunicații. Aspectul care complică și mai mult mecanismele de control al congestiei este acela că nu se poate localiza pe un nivel particular al stivei de protocoale. Algoritmii de control al congestiei pot fi implementați atât la nivelul routerelor de rețea cât și la nivelul protocoalelor de transport implementate în softul de rețea al dispozitivelor comunicante. O soluție exclusivă, bazată doar pe controlul unui singur nivel nu duce la o rezolvare mulțumitoare. De aceea se caută soluții de compromis între cele două niveluri.

Multe din problemele care apar în sistemele complexe, cum ar fi rețelele de calculatoare, pot fi privite din punctul de vedere al unei teorii a controlului. Această abordare conduce la

împărțirea tuturor soluțiilor în două grupe: în buclă deschisă și în buclă închisă. Soluțiile în buclă deschisă încearcă să rezolve problema printr-o proiectare atentă, în esență să se asigure că problema nu apare. După ce sistemul este pornit și funcționează, nu se mai fac nici un fel de corecții.

Instrumentele pentru realizarea controlului în buclă deschisă decid când să se accepte trafic nou, când să se distrugă pachete și care să fie acestea, realizează planificarea deciziilor în diferite puncte din rețea. Toate acestea au ca numitor comun faptul că iau decizii fără a ține cont de starea curentă a rețelei.

Prin contrast, soluțiile în buclă închisă se bazează pe conceptul de reacție inversă (feedback loop). Această abordare are trei părți, atunci când se folosește pentru controlul congestiei:

1. Monitorizează sistemul pentru a detecta când și unde se produce congestia.
2. Trimite aceste informații către locurile unde se pot executa acțiuni.
3. Ajustează funcționarea sistemului pentru a corecta problema.

În vederea monitorizării subrețelei pentru congestie se pot folosi diverse de metrice. Cele mai utilizate sunt procentul din totalul pachetelor care au fost distruse din cauza lipsei spațiului temporar de memorare, lungimea medie a cozilor de așteptare, numărul de pachete care sunt retransmise pe motiv de timeout, întârzierea medie a unui pachet, deviația standard a întârzierii unui pachet. În toate cazurile, valorile crescătoare indică creșterea congestiei.

Al doilea pas în bucla de reacție este transferul informației legate de congestie de la punctul în care a fost depistată la punctul în care se poate face ceva. Varianta imediată presupune trimiterea unor pachete de la ruterul care a detectat congestia către sursa sau sursele de trafic, pentru a raporta problema. Evident, aceste pachete suplimentare cresc încărcarea rețelei exact la momentul în care acest lucru era cel mai puțin dorit, subrețeaua fiind congestionată.

Există însă și alte posibilități. De exemplu, poate fi rezervat un bit sau un câmp în fiecare pachet, pentru a fi completat de rutere dacă congestia depășește o anumită valoare de prag. Când un ruter detectează congestie, el completează câmpurile tuturor pachetelor expediate, pentru a-și preveni vecinii.

O altă abordare este ca ruterele sau calculatoarele gazdă să trimită periodic pachete de probă pentru a întreba explicit despre congestie. Aceste informații pot fi apoi folosite pentru a ocoli zonele cu probleme. Unele stații de radio au elicoptere care zboară deasupra orașelor pentru a raporta congestiile de pe drumuri și a permite ascultătorilor mobili să își dirijeze pachetele (mașinile) astfel încât să ocolească zonele “fierbinți”.

În toate schemele cu feedback se speră că informarea asupra producerii congestiei va determina calculatoarele gazdă să ia măsurile necesare pentru a reduce congestia. Pentru ca o schemă să funcționeze corect, duratele trebuie reglate foarte atent. Dacă un ruter strigă STOP de fiecare dată când sosesc două pachete succesive și PLEACĂ (eng.: GO) de fiecare dată când este liber mai mult de 20 μ sec, atunci sistemul va oscila puternic și nu va converge niciodată. Pe de altă parte, dacă va aștepta 30 de minute pentru a fi sigur înainte de a spune ceva, mecanismul pentru controlul congestiei reacționează prea lent pentru a fi de vreun folos real. Pentru a funcționa corect sunt necesare unele medieri, dar aflarea celor mai potrivite constante de timp nu este o treabă tocmai ușoară.

Se cunosc numeroși algoritmi pentru controlul congestiei. Pentru a oferi o modalitate de organizare a lor, Yang și Reddy (1995) au dezvoltat o taxonomie pentru algoritmii de control al congestiei. Ei încep prin a împărți algoritmii în cei în buclă deschisă și cei în buclă închisă, așa

cum s-a precizat anterior. În continuare împart algoritmi cu buclă deschisă în unii care acționează asupra sursei și alții care acționează asupra destinației. Algoritmi în buclă închisă sunt de asemenea împărțiți în două subcategorii, cu feedback implicit și cu feedback explicit. În algoritmi cu feedback explicit, pachetele sunt trimise înapoi de la punctul unde s-a produs congestia către sursă, pentru a o avertiza. În algoritmi implicați, sursa deduce existența congestiei din observații locale, cum ar fi timpul necesar pentru întoarcerea confirmărilor.

Prezența congestiei înseamnă că încărcarea (momentană) a sistemului este mai mare decât cea pe care o pot suporta resursele (unei părți a sistemului). Pentru rezolvare vin imediat în minte două soluții: sporirea resurselor sau reducerea încărcării. De exemplu subrețeaua poate începe să folosească linii telefonice pentru a crește temporar lățimea de bandă între anumite puncte. În sistemele bazate pe sateliți, creșterea puterii de transmisie asigură de regulă creșterea lățimii de bandă. Împărțirea traficului pe mai multe căi în locul folosirii doar a celei mai bune poate duce efectiv la creșterea lățimii de bandă. În fine, ruterele suplimentare, folosite de obicei doar ca rezerve pentru copii de siguranță (backups) (pentru a face sistemul tolerant la defecte), pot fi folosite pentru a asigura o capacitate sporită atunci când apar congestii serioase.

Totuși, uneori nu este posibilă creșterea capacității sau aceasta a fost deja crescută la limită. Atunci singura cale de a rezolva congestia este reducerea încărcării. Sunt posibile mai multe metode pentru reducerea încărcării, cum ar fi refuzul servirii anumitor utilizatori, degradarea serviciilor pentru o parte sau pentru toți utilizatorii și planificarea cererilor utilizatorilor într-o manieră mai previzibilă.

1.3 Politici pentru prevenirea congestiei

Voi începe studiul asupra metodelor de control al congestiei prin analiza sistemelor cu buclă deschisă. Aceste sisteme sunt proiectate astfel încât să minimizeze congestia în loc să o lase să se producă și apoi să reacționeze. Ele încearcă să-și atingă scopul folosind politici corespunzătoare, la diferite niveluri.

Prevenirea congestiei se bazează pe un set de mecanisme ce ajută cozile echipamentelor să evite congestia. Cozile se supraîncărcă atunci când traficul sosit pe interfețele de intrare depășește capacitatea interfețelor de ieșire. Atunci când mai mult trafic sosește pe interfețele de ieșire decât acestea pot suporta, cozile (memoriile && bufferele) routerelor se încarcă. Mecanismele de management al cozilor ne pot ajuta să evităm supraîncărcarea cozilor și respectiv congestia. Practic, aceste mecanisme ne oferă posibilitatea de a alege între pierderea de pachete și întârzierea sau jitterul – variațiile în timp ale întârzierilor (parametri critici pe rețea).

Să începem cu nivelul legătură de date și să ne continuăm apoi drumul spre niveluri superioare. Politica de retransmisie stabilește cât de repede se produce timeout la un emițător și ce transmite acesta la producerea timeout-ului. Un emițător foarte activ, care produce repede timeout și retransmite toate pachetele în așteptare folosind retransmiterea ultimelor n , va produce o încărcare mai mare decât un emițător calm, care folosește retransmiterea selectivă. Strâns legată de acestea este politica de memorare în zonă tampon. Dacă receptorii distrug toate pachetele în afara secvenței, acestea vor trebui retransmise ulterior, introducând o încărcare

suplimentară. În ceea ce privește controlul congestiei, repetarea selectivă este, în mod sigur, mai bună decât retransmiterea ultimelor n .

În Tabelul 1.1. sunt prezentate diferite politici pentru nivelul legătură de date, rețea și transport, care pot influența congestia

Nivel	Politică
Transport	Politica de retransmisie Politica de memorare temporară a pachetelor în afară de secvență (out-of-order caching) Politica de confirmare Politica de control al fluxului Determinarea timeout-ului
Rețea	Circuite virtuale contra datagrame în interiorul subrețelei Plasarea pachetelor în cozi de așteptare și politici de servire Politica de distrugere a pachetelor Algoritmi de dirijare Gestiunea timpului de viață al pachetelor
Legătură de date	Politica de retransmitere Politica de memorare temporară a pachetelor în afară de secvență (out-of-order caching) Politica de confirmare Politica de control al fluxului

Tabelul 1.1

Și politica de confirmare afectează congestia. Dacă fiecare pachet este confirmat imediat, pachetele de confirmare vor genera un trafic suplimentar. Totuși, în cazul în care confirmările sunt preluate de traficul de răspuns, se pot produce timeout-uri și retransmisii suplimentare. O schemă prea strânsă pentru controlul fluxului (de exemplu fereastră mică) reduce volumul de date și ajută în lupta cu congestia.

La nivelul rețea, alegerea între folosirea circuitelor virtuale și datagrame influențează congestia, deoarece mulți algoritmi pentru controlul congestiei funcționează doar pe subrețele cu circuite virtuale. Plasarea în cozi de așteptare a pachetelor și politicile de servire specifică dacă ruterele au o coadă pentru fiecare linie de intrare, o coadă pentru fiecare linie de ieșire sau ambele. Mai precizează ordinea în care se prelucrează pachetele (de exemplu round robin sau bazată pe priorități). Politica de distrugere a pachetelor este regula care stabilește ce pachet este distrus dacă nu mai există spațiu. O politică bună va ajuta la eliminarea congestiei, pe când una greșită o va accentua.

Un algoritm de dirijare bun poate ajuta la evitarea congestiei prin răspândirea traficului de-a lungul tuturor liniilor, în timp ce un algoritm neperformant ar putea trimite toate pachetele

pe aceeași linie, care deja este congestionată. Gestiunea timpului de viață asociat pachetelor stabilește cât de mult poate trăi un pachet înainte de a fi distrus. Dacă acest timp este prea mare, pachetele pierdute vor încurca pentru mult timp activitatea, iar dacă este prea mic, este posibil să se producă timeout înainte de a ajunge la destinație, provocând astfel retransmisii.

La nivelul transport apar aceleași probleme ca la nivelul legăturii de date dar, în plus, determinarea intervalului de timeout este mai dificil de realizat, deoarece timpul de tranzit prin rețea este mai greu de prezis decât timpul de tranzit pe un fir între două rutere. Dacă intervalul de timeout este prea mic, vor fi trimise inutile pachete suplimentare. Dacă este prea mare, congestia se va reduce, însă timpul de răspuns va fi afectat de pierderea unui pachet. [4]

2. Tehnici cunoscute de prevenire a pierderii pachetelor în rețele

2.1 Introducere

TCP (Transmission Control Protocol) realizează fragmentarea mesajelor în pachete și asigură transmiterea corectă a mesajelor între utilizatori. Pachetele unui mesaj sunt numerotate, putându-se verifica primirea lor în forma în care au fost transmise și reconstituirea mesajelor lungi, formate din mai multe pachete.

TCP este un protocol orientat pe conexiuni, care permite ca un flux de octeți trimiși de la sursă să ajungă fără erori la destinație .

Principalele caracteristici ale TCP sunt:

- Transfer de date în flux continuu - datele circulă în același timp, în ambele sensuri ale conexiunii.
- Siguranța transmisiei - recuperează pachetele transmise cu erori, pierdute sau cu număr de secvență eronat.
- Controlul fluxului de date – în transferul de date dintre două procese, când aplicația destinație trimite o confirmare către emitent, se indică și numărul permis de octeți ce se pot recepționa, pentru a se asigura că transmiterea rapidă de mesaje de către un emițător, nu face ca un receptor lent să primească mai multe mesaje decât poate prelucra. În urma unui astfel de mesaj, emițătorul își va dimensiona pachetele transmise la lungimea indicată de receptor.
- Multiplexarea - permite mai multor procese, care rulează pe același host, să utilizeze facilitățile protocolului TCP simultan.
- Controlul conexiunii - presupune stabilirea numărului de secvență și a dimensiunii ferestrei, pentru fiecare pachet TCP.

Pentru a realiza transferul de date, TCP fragmentează fluxul de octeți în pachete și transmite fiecare pachet la nivelul Internet. La destinație, procesul TCP receptor reassemblează pachetele primite într-un flux de ieșire. Înainte de a începe transferul datelor, între cele două procese utilizator (programe de aplicație) se stabilește, prin intermediul rețelei, o conexiune logică numită circuit virtual.

Livrarea la destinație a fluxului de octeți în ordinea în care au fost emiși, fără pierderi și fără duplicate, se realizează prin folosirea unei tehnici de confirmare pozitivă cu retransmitere

(dacă confirmarea de primire nu a fost transmisă într-un timp prestabilit, pachetul este transmis din nou) , combinată cu fereastră glisantă (sistem de confirmare în expectativă) . Mecanismul ferestrei glisante, utilizat de TCP, operează la nivelul octeților și nu al pachetelor, permițând o transmisiune eficientă și un control al fluxului.

2.2 TCP

TCP implementează un mecanism de control al fluxului pe bază de fereastră. Propriu zis vorbind, un protocol pe bază de fereastră înseamnă că dimensiunea ferestrei curente definește o limită superioară a cantității de date nerecunoscută, care poate fi în tranzit între o pereche dată expeditor-receptor. Inițial, fluxul de control al TCP-ului a fost guvernat doar de dimensiunea ferestrei maximă permisă anunțate de către receptor și politică care a permis expeditorului să trimită noi pachete numai după ce a primit unidirecționat pachetul anterior.

Principalul scop al TCP este de a asigura circuit logic sigur sau serviciu de conexiune între două procese pereche. El nu se bazează pe siguranța altor protocoale de nivel inferior (cum este IP), așa ca TCP trebuie să garanteze el însuși siguranța transmisiei. TCP poate fi caracterizat prin următoarele facilitati pe care le asigură pentru aplicațiile care îl utilizează:

- Transfer de date în flux (Stream Data Transfer)

Din punctul de vedere al aplicației, TCP transferă un sir (stream) continuu de octeți prin rețea. Aplicația nu trebuie să-și facă griji cu segmentarea datelor în blocuri de bază sau datagrame. O face TCP prin gruparea octetilor în fragmente TCP, care sunt transferați IP pentru a fi transmiși la destinație. De asemenea, însuși TCP decide cum să segmenteze datele și poate trimite datele mai departe într-un mod convenabil. Uneori, o aplicație are nevoie să fie sigură că toate datele transmise către TCP au fost în realitate transmise la destinație. Pentru aceasta, o funcție `push` este definită. Ea va forța transmiterea tuturor segmentelor TCP rămase în memorie către destinație. Funcția de închidere (`close`) normală forțează și ea transmiterea datelor la destinație.

- Siguranța

TCP atribuie un număr de secvență fiecărui octet transmis și așteaptă o confirmare (acknowledgment - ACK) de la aplicația care recepționează datele TCP. Dacă confirmarea nu vine într-un interval de timp prestabilit, datele sunt transmise din nou. Deoarece datele sunt transmise în blocuri (segmente TCP) numai numărul de secvență al primului octet este trimis calculatorului destinație. Aplicația TCP destinație folosește numerele de secvență pentru a le ordona atunci când sosesc neordonate și să elimine segmentele duplicate.

- Controlul transmisiei (Flow Control)

Aplicația TCP destinație, când transmite o confirmare (ACK) către emitent, indică de asemenea numărul de octeți pe care îi poate recepționa pe lângă ultimul segment TCP primit, fără să se

suprancarce sau să apară depasirea memoriilor tampon (internal buffers) ale sale. Aceasta este trimis în ACK sub forma numărului cel mai mare de secvență pe care-l poate receptiona fără probleme. Acest mecanism este cunoscut sub numele de fereastră glisantă (window-mechanism) și va fi discutat în detaliu în acest capitol.

- Multiplexarea

Este realizată prin utilizarea porturilor, la fel ca și la UDP.

- Conexiunile logice

Siguranța și mecanismele de control ale transmisiei descrise mai înainte necesită ca TCP să inițializeze și mențină câteva informații referitoare la starea fiecărui flux de date (data stream). Această informație de stare ce include socketurile, numerele de secvență, dimensiunile ferestrelor se numește conexiune logică. Fiecare conexiune este unic identificată de către o pereche de socketuri utilizate de către procesele emitente și receptoare.

- Full Duplex

TCP asigură două fluxuri concurente de date în ambele sensuri.

În afară de advertisementul ferestrei receptorului, $awnd$, controlul congestiei TCP introduce două noi variabile pentru conexiune: fereastră de congestie "CWND" și pragul de start lent. Dimensiunea ferestrei transmițătorului, w , a fost definită ca:

$$w = \min(cwnd, awnd) \quad (2.1)$$

în loc să fie egală cu $awnd$. Fereastră de congestie poate fi considerată ca fiind o contra-fereastră de avertizare, unde variabila " $awnd$ " este folosită pentru a atenționa pe transmițător să suprafolosească resursele receptorului în timp ce scopul " $cwnd$ -ului" este acela de a preveni transmițătorul de a trimite mai multe pachete decât rețeaua poate suporta în condițiile de încărcare curentă.

Controlul congestiei (reactivă) = se aplica după ce rețeaua este supraîncărcată

Congestie Evitarea (Proactive) = se aplica înainte să devină rețeaua supraîncărcată

TCP (Transport Control Protocol) este un protocol punct la punct, orientat pe conexiune, introdus pentru a proteja rețeaua împotriva supraîncărcării puternice și pentru a realiza o împărțire echitabilă a lățimii de bandă pentru diferite surse TCP. Protocolul TCP detectează erori precum pachete eronate, pierdute sau duplicate. Emițătorul atribuie o secvență numerică fiecărui pachet (segment) și solicită o confirmare pozitivă (ACK) de la receptor. Receptorul detectează

pierderea de pachete verificând secvența numerică și confirmând, la fiecare pachet primit, ultimul pachet recepționat în ordinea corectă. Pentru fiecare pachet care nu este în ordine, o confirmare duplicat este trimisă. Acest ACK duplicat are rolul de a notifica perechea de faptul că a fost detectată o pierdere, și pentru a îi transmite emițătorului următoarea secvență numerică așteptată.

Protocolul TCP asigură un flux de octeți sigur, printr-o conexiune nesigură, cap la cap, a rețelei sau inter-rețelei. Inter-rețeaua este formată prin asocierea unor rețele care diferă ca: topologie, lățime de bandă, întârzieri, dimensiuni de pachete, etc. TCP se adresează dinamic la rețeaua Internet și este robust la mai multe tipuri de defecte. Entitatea de transport TCP gestionează fluxurile TCP și interfețele către IP. Conexiunile TCP sunt:

- duplex;
- punct-la-punct (TCP nu suportă difuzarea);
- orientate pe fluxuri de octeți nu de mesaje.

Fluxurile de date de la procesele utilizator sunt împărțite în segmente de maxim 64 kB, tipic de 1500 de octeți. Apoi fiecare segment este expediat ca o datagramă IP separată. Datagramele IP recepționate sunt furnizate entității TCP care reconstruiește fluxul original de octeți, în ordine, deoarece IP nu asigură livrarea sau transportul ordonat al datagramelor. Serviciul TCP este asigurat prin crearea de către emițător și receptor a unor puncte finale numite socluri sau socket-uri. Portul este un punct de acces la servicii de nivel transport, TSAP (Transport Service Acces Point). Numerele de port mai mici decât 256 sunt alocate porturilor general cunoscute, rezervate serviciilor standard.(telnet-23, http -80).

Serviciile TCP principale asigurate de TCP sunt:

- **expedierea datelor (SEND)**, atunci când e convenabil pentru TCP,
- **urgentarea expedierii (PUSH)**, cere transmiterea imediată dacă e posibil,
- **urgentarea recepției (URGENT DATA)**: transmiterea și recepția imediată pentru a întrerupe o prelucrare distantă, deoarece PUSH nu are efect la distanță, dar URGENT DATA are.

Formatul antetului TCP

Antetul TCP are 20 octeți plus o parte opțională. Segmentul TCP, format din antet și date, are lungimea maximă de 64kB, adică 65.536 B, dar e limitat de **MTU (Maximum Transfer Unit)** al fiecărei rețele traversate. Eventual rețelele din cale mai fragmentează segmentul TCP. Fiecare fragment va avea propriile antete TCP și IP.

Port SursA								Port destinatie
Numar de secventa								
Numar de confirmare								
Lungime Antet (4)	Rezervati (6)	U R G	A C K	P S H	R S T	S Y N	F I N	Dimensiunea ferestrei
Suma de Control								Indicator URGENT
Optiuni (1-n) cuvinte pe 32 biti								
Date (optional)								

2.1 Formatul Antetului TCP

Unde:

- Portul sursa : Numărul portului sursa pe 16 biti folosit de receptor pentru a raspunde.
- Portul destinatie: Numărul portului destinatie pe 16 biti.
- Numarul de secventa: Numărul de secventa al primului octet de date din acest segment. Daca bitul de control SYN este setat, numarul de secventa este numarul de secventa initial (n) si primul octet de date este n+1.
- Numarul de confirmare: Daca bitul de control ACK este setat, acest camp conține valoarea urmatorului număr de secventa pe care receptorul se asteaptă să-l primească.
- Data Offset: Numarul de cuvinte de 32 biti din antetul TCP. El indica unde incep datele.
- Reserved: Șase biți rezervați pentru utilizare in viitor; trebuie sa fie zero.
- URG: Precizează că articolul pointer de urgenta are semnificatie in acest segment.
- ACK: Precizează că articolul de "Numar de secventa confirmat" are semnificatie in acest segment.
- PSH: Functia push.
- RST: Resetează conexiunea.
- SYN: Sincronizează numerele de secventa.
- FIN: Nu mai sunt date de la emitent.
- Fereastra: Folosit in segmenteleACK. El specifica numarul de octeti de date incepand cu acela indicat in "numarul de secventa confirmat" pe care receptorul (adica emitentul acestui segment) doreste (si poate) sa accepte.
- Suma de control: Complementul față de unu al sumei in complement față de unu al tuturor cuvintelor de 16 biți din pseudo- antet, antetul TCP si datele TCP. La calcularea sumei de control, campul "Suma de control" este considerat zero.
- indicatorul de urgență – permite identificarea poziției unor date de urgență, în cadrul protocolului TCP;
- date – datele protocolului nivelului superior

Mecanismul de comunicare TCP

Stabilirea conexiunii:

- înțelegerea în 3 pași (hand-shake) sau comunicarea cu confirmare SYN = 1 arată o cerere de conectare **datelor Terminarea conexiunii**
- fluxul: fiecare octet e numerotat modulo 2
- antetul conține numărul de secvență al primului octet

- controlul fluxului: prin credite (număr de octeți)
- datele sunt transmise când vrea entitatea TCP: - PUSH - transmite acum
- urgent: transmite această dată din fluxul normal care are indicatorul urgent
- dacă se recepționează un TPDU nepotrivit acestei conexiuni, se pune reset =1 în segmentul care pleacă

Terminarea conexiunii

- deconectarea normală cu FIN = 1;
- deconectarea bruscă (abort) cu pierdere de date.

Politici de implementare

- **emisie trebuie aleasa dimensiunea segmentului și a ferestrei de transmisie:**
 - dacă segmentul este prea mic, antetul este o încărcare prea mare
 - dacă segmentul este prea mare, rezultă întârzieri prea mari
 - alegerea ferestrei este o problemă la liniile cu lățime de bandă mare și întârziere mare: cu fereastra de 64 kB și linia T3 de 4,736 Mbps, transmiterea a 64 KB durează 12 msec. Dacă RTT (Round Trip Time), întârzierea dus-întors este de 50 msec, atunci 75 % din timp emițătorul este inactiv, așteptând confirmări.

- **livrarea datelor**
 - datele pot fi memorate sau sunt livrate imediat, ordonat. Cu mențiunea urgent sunt livrate imediat, dacă e posibil
 - segmentele în afara secvenței sunt acceptate sau rejectate
 - la expirarea timpului se retransmite numai primul segment, toate, sau individual, (se mențin cronometre separate per segment)
 - confirmarea se face imediat sau cumulat (se așteaptă date în sens opus sau expirarea timpului)

- **controlul fluxului și al congestiei**
Există trei ferestre, pentru emisie W_t , pentru receptivă W_r , și pentru congestivă W_c . Se alege fereastra de emisie

$$W_t = \min (W_r, W_c) \quad (2.2)$$

O sursă TCP este considerată saturată în cazul unei sesiuni FTP, atunci când unde toate segmentele au dimensiunea maximă posibilă, rezultând pachete de lungimi constante. Termenul de pachet este echivalent cu termenul segment. TCP-ul deține o fereastră de congestie și o limită pentru „slow-start” ce sunt descrise de variabilele CWND (Congestion Window) și Ssthresh (**Slow-start threshold Value**) . CWND determină numărul maxim de pachete ce pot fi neconfirmate în rețea, iar Ssthresh controlează creșterea dimensiunii CWND. Dacă CWND este mai mică decât Ssthresh, în faza de „slow-start”, fereastra de congestie este incrementată cu câte un pachet pentru fiecare confirmare neduplicată primită. Aceasta duce la o creștere exponențială a CWND. Când CWND este mai mare sau egală cu Ssthresh, în faza de evitare a congestiei, CWND crește cu câte un pachet la fiecare RTT (round trip time), ceea ce corespunde unei creșteri liniare a CWND

TCP reacționează la pierderea pachetelor în două feluri diferite. Setează câte un cronometru pentru fiecare pachet trimis. În cazul în care nu este primită o confirmare pentru acel

pachet, înainte de expirarea timpului, CWND își reduce dimensiunea la dimensiunea maximă a segmentului iar Ssthresh este setat la *max*:

$$Ssthresh = \frac{flight_size}{2.2 * dimsegment} \quad (2.3)$$

Flight size reprezintă cantitatea de informație neconfirmată în rețea. Presupunând că emițătorul este saturat, valoarea CWND este egală cu valoarea flight size. Când emițătorul recepționează trei ACK-uri duplicat, înainte de expirarea timpului, este realizată o recuperare rapidă (fast recovery): într-o manieră simplificată, TCP trimite pachetul pierdut din nou, reduce Ssthresh la *max* și setează CWND la noua valoare calculată Ssthresh.

Congestia rețelelor este o cauză a distribuției inegală a resurselor de rețea și a fluxului de rețea, iar algoritmi de control al congestiei pot fi împărțiți în algoritmi „Source” și „link”. Algoritmul de control al congestiei TCP este cel mai utilizat algoritm „Source”, dar utilizarea TCP-ului nu a putut evita posibilitatea de apariție a congestiei în rețele. Algoritmii „Link” ar putea depăși în mod eficient impasul și problemele de sincronizare la nivelul de strategie drop-coadă, iar cercetarea se concentrează pe algoritmi de AQM (Active Queue Management). Există trei tipuri de algoritmi Aqm : Red (Random early Decetion), Algoritm AQM bazat pe teoria controlului și Algoritm AQM bazate pe teoria optimizării.

2.3 Algoritmi de management al cozii la receptor (AQM)

Principalele caracteristici de performanță ale AQM includ :

- Utilizarea eficientă a cozii: coada ar trebui să evite supraîncărcarea rezultată în pierderea de pachete și retransmisii nedorite sau pachete goale= goliciuni au rezultat într-un link subîncărcat
- Întârzierea cozii: este de dorit să se păstreze întârzierea și micile sale variații.
- Robustețea: Schema AQM are nevoie de a menține un comportament robust, în ciuda diferitelor condiții de rețea, cum ar fi variații în numărul de sesiuni de TCP sau variații în întârzierea propagării și capacitatea legăturii.

TCP a fost optimizat pentru rețele cu fir. Orice pierdere de pachete este considerată a fi rezultatul unei congestii a rețelei și dimensiunea ferestrei de congestie este dramatic redusă ca o măsură de precauție. O serie de algoritmi alternativi de control al congestiei au fost propuși pentru a ajuta la rezolvarea problemei fara fir : cum ar fi Las Vegas, Westwood, Veno și Santa Cruz.

TCP îndeplinește două sarcini de control fundamentale:

- 1) De control end-to-end al debitului, pentru a evita supraincercării buffer-ului receptorului ;
- 2) Adaptarea controlului congestiei la conectare, adaptarea cantitatii de informații transferate în rețeaua la starea congestionarea rețelei. Această funcție este bazată pe recuperarea congestiei, mai degrabă decât de prevenire. Crește traficul TCP oferit rețelei până când rețeaua este obligată să renunțe la unele pachete; apoi se strânge pentru a da timp rețelei de a reveni din congestie, și începe din nou să crească traficul oferit pentru a forța un alt episod de congestie. Controlul fluxului TCP este pus în aplicare prin intermediul receptorului, „*advertised window*“ (*rcvwnd*) W_{re} , stabilit de către receptorul în antetul de confirmare a segmentelor (ACK). Controlul congestiei este pus în aplicare prin intermediul ferestrei de congestie (CNWD) W_c , calculată de către transmițător după algoritmul de control al congestiei TCP. În final cantitatea de date care îi este permis TCP-ului să o transmită la un moment dat este calculată ca

$$W_{tr} = \min\{ W_c, W_{re} \} - W_v(2.4)$$

W_v este cantitatea de date restante, adică, a ferestrei de transmisie de alunecare, care este umplută cu datele transmise deja, dar încă nerecunoscute.

Controlul congestiei se referă la controlarea traficului ce pătrunde într-o rețea (de telecomunicații sau de calculatoare), ceea ce presupune evitarea supraîncărcării legăturilor dintre nodurile intermediare ale rețelei, prin reducerea numărului de pachete trimise. Așa cum s-a prezentat anterior, controlul congestiei nu trebuie confundat cu controlul fluxului care protejează receptorul de a fi „copleșit” de către emițător.

Slow-start este parte a strategiei de control al congestiei folosit de TCP, protocolul de transmitere a datelor utilizate de mai multe aplicații pe Internet. Slow Start este folosit în combinație cu alți algoritmi pentru a evita trimiterea de mai multe date decât rețeaua este capabilă să transmită, scopul este de a evita congestionarea rețelei. Algoritmul este specificat de RFC 5681.

Slow-start este unul dintre algoritmi care TCP îi folosește pentru a controla congestia în interiorul rețelei. De asemenea, este cunoscut sub numele de fază de creștere exponențială. În timpul fazei de creștere exponențială, începe creșterea în regim încet a creșterii ferestrei de congestie TCP de fiecare dată când o confirmare este primită. Aceasta crește dimensiunea ferestrei cu numărul de segmente recunoscute. Acest lucru se întâmplă până când o confirmare nu este primită pentru se ajunge la unele segment sau o valoare de prag predeterminată. În cazul în care are loc un eveniment de pierdere, TCP presupune că este din cauza congestiei rețelei și ia măsuri pentru a reduce sarcina oferite pe rețea. Odată ce pragul a fost atins, TCP intră în fază de creștere liniară (de evitare a congestiei). În acest moment, fereastra crește cu un segment pentru fiecare RTT. Acest lucru se întâmplă până când se produce un eveniment de pierdere.

Cu toate că strategia este denumită "Slow-Start", creșterea ferestrei de congestie este destul de agresivă, mai agresivă decât faza de evitare a congestiei. [29] Înainte de „Slow Start” a fost introdus în TCP, faza inițială de evitare pre-congestie fost chiar mai rapid.

Algoritmul de bază slow-start : Algoritmul începe în faza de creștere exponențială, inițial cu o dimensiune a ferestrei de congestie (*cwnd*) de 1, 2 sau 10 [30] segmente și crește prin

dimensiunea unui segment (SS) pentru fiecare nou ACK primit. În cazul în care receptorul trimite un ACK pentru fiecare segment, acest comportament se dublează în mod eficient dimensiunea ferestrei fiecare călătorie dus-întors a rețelei. În cazul în care receptorul acceptă un ACK întârziat, rata de creștere este mai mică, dar totuși crește cu un minimum de un SMS de fiecare dată pentru un drum dus-întors. Acest comportament continuă până când dimensiunea ferestrei congestiei (cwnd) atinge dimensiunea ferestrei de avertizare a receptorului sau până când se produce o pierdere.

Atunci când are loc o pierdere, jumătate din cwnd curent este salvat ca un prag de pornire lentă (SSThresh) și Slow Start începe din nou de la dimensiunea ferestrei cwnd inițială. Odată ce cwnd ajunge la SSThresh, TCP trece în modul de evitare a congestiei în cazul în care fiecare nou ACK crește cwnd de $SS \times SS / cwnd$. Acest lucru duce la o creștere lineară a cwnd.

Fast Retransmit : Un expeditor TCP folosește un cronometru pentru a recunoaște segmente pierdute. În cazul în care o confirmare nu este primită pentru un anumit segment într-un timp specificat (în funcție de estimat timpul de întârziere tur-retur), expeditorul va asuma segmentul a fost pierdut în rețea, și va retransmite segmentul. Confirmarea Duplicate este baza pentru mecanismul Fast Retransmit , care funcționează după cum urmează: după ce a primit un pachet (de exemplu, cu numărul de ordine 1), receptorul trimite o confirmare prin adăugarea de la 1 la numărul de ordine (de exemplu, numărul de ordine 2), ceea ce înseamnă că receptor primește numărul de pachete de 1 și se așteaptă ca numărul de pachete de 2 de la expeditor. Să presupunem că cele trei pachete ulterioare s-au pierdut. Între timp receptorul primește un număr de pachete de 5 și 6. După ce a primit numărul de pachete de 5, receptorul trimite o confirmare, dar numai pentru numărul de ordine 2. Când receptorul primește numărul de pachete de 6, acesta trimite un alt valoare confirmare de 2. Deoarece expeditorul primește mai mult de o confirmare cu același număr de ordine (2 în acest exemplu), aceasta se numește duplicat confirmare. [31]

Recuperarea rapidă : Există o variație a algoritmului slow-start, cunoscut sub numele de recuperare rapidă, care utilizează retransmiterea rapidă, urmată de evitarea congestiei . În algoritmul de recuperare rapidă, în modul de evitare a congestiei, atunci când pachetele nu sunt primite (ceea ce ar însemna detectarea a trei ACK duplicat), dimensiunea ferestrei de congestie este redusă la pragul de Slow Start , mai degrabă decât valoarea inițială.

Există mai mulți algoritmi de prevenire a congestiei prevăzuți de TCP. Primii algoritmi apăruiți sunt TCP Tahoe și TCP Reno, ambii având la bază regulile de „slow-start” și „fast retransmit”. Cei doi algoritmi, Tahoe și Reno, diferă prin modul în care reacționează la pierderea de pachete:

- **Tahoe**: pierderea este detectată atunci când timpul de timeout expiră înainte ca o confirmare să fie recepționată. În acel moment, Tahoe reduce fereastra de congestie la valoarea maximă a dimensiunii unui segment (MSS – maximum segment size) și se resetează la faza de „slow-start”.
- **Reno**: în cazul în care sunt recepționate trei confirmări duplicate (ACK duplicat), Reno înjumătățește fereastra de congestie, efectuează o retransmitere rapidă (fast retransmit) și intră în recuperare rapidă (fast recovery). Până la mijlocul anilor 1990, toate timeout-urile și întârzierile dus-întors erau bazate numai pe ultimul pachet din buffer. Cercetătorii Larry Peterson și Lawrence Brakmo de la Universitatea din Arizona, au introdus **TCP Vegas** în care timeout-ul este setat iar întârzierea dus-întors este măsurată pentru fiecare pachet din buffer. De asemenea, acest algoritm folosește o creștere aditivă a ferestrei de

congestie. Această varianta nu a cunoscut o răspândire foarte mare în afara laboratorului cercetătorilor.

- **TCP New Reno** îmbunătățește retransmisia în timpul fazei de recuperare rapidă (fast recovery) a TCP Reno. Pe parcursul recuperării rapide, pentru fiecare ACK duplicat, este transmis un nou pachet netrimis de la sfârșitul ferestrei de congestie, pentru a menține fereastra plină. Pentru fiecare ACK parțial din spațiul de secvențe numerice, emițătorul trimite următorul pachet cu secvența numerică următoare ACK-ului. Întrucât cronometrul de timeout este resetat ori de câte ori există un progres în buffer-ul de transmisie, New Reno poate umple găuri mari sau mai multe găuri în spațiul de secvențe. Deoarece New Reno poate trimite pachete noi la sfârșitul ferestrei de congestie în timpul recuperării rapide, este menținut un throughput mare pe durata procesului de umplere a găurilor, chiar și atunci când există mai multe găuri, de mai multe pachete fiecare. Când TCP intră în recuperare rapidă, acesta înregistrează cea mai mare secvență numerică neconfirmată. Când această secvență numerică este confirmată, TCP revine la stare de evitare a congestiei. Atunci când nu au loc pierderi de pachete, dar în schimb pachetele sunt reordonate cu mai mult de 3 secvențe numerice, apare o problemă în New Reno. Când acest fenomen are loc, New Reno intră greșit în recuperare rapidă, însă când pachetul reordonat este livrat, are loc progresul secvenței numerice ACK și din acel moment până la sfârșitul recuperării rapide, fiecare bit al secvenței numerice produce un duplicat și retransmisia nu este necesară.
- **TCP Hybla** are rolul de a elimina penalizarea conexiunilor TCP ce încorporează legături terestre cu latențe foarte mari sau legături radio, datorită timpilor dus-întors mai mari. Provine dintr-o evaluare analitică a dinamicii ferestrei de congestie, ce sugerează înlăturarea dependenței performanței de timpul dus-întors (RTT – round trip time).
- **TCP BIC** (binary increase congestion control) este o implementare a TCP cu un algoritm de control al congestiei optimizat pentru rețele de viteze mari, cu latențe mari (numite LFN – long fat networks, în RFC 1072). BIC este folosit implicit în kernel-urile Linux, de la 2.6.8 la 2.6.18.
- **TCP Cubic** este o versiune mai puțin agresivă a BIC, în care fereastra este o funcție cubică de timp, de la momentul ultimei congestii, cu punctul de inflexiune setat la fereastra anterioară evenimentului.

În rutare, management activ al cozii (AQM) reordonează arbitrar sau pierde pachetele de rețea în interiorul bufferului de transmisie a interfaței unui controler de rețea. Sarcina este efectuată de către designerul de rețea. Un router de Internet menține de obicei un set de cozi, câte unul pe fiecare interfață. Sarcina activă a cozii este de a pierde sau marca pachetele înainte ca ea să fie plină. De obicei, acestea operează prin menținerea uneia sau mai multor probabilități de pierdere sau de marcare.

Pentru a evita colapsul congestiei, TCP folosește o strategie multi-fațete de control a 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 TCP folosită pentru controlul fluxului. TCP folosește un mecanism numit start lent [29] pentru a mări fereastra de congestie, după o conexiune este inițializată și după un timeout. Acesta începe cu o fereastră de două ori dimensiunea maximă a segmentului (MSS). Deși rata inițială este scăzută, rata de creștere este foarte rapidă: pentru fiecare pachet recunoscut, fereastra de congestie crește cu 1 MSS, astfel încât fereastra de congestie dublează în mod eficient de fiecare dată dus-întors (RTT). Când fereastra de congestie depășește un prag ssthresh algoritmul intră într-o stare nouă, numită evitare a congestiei. În unele implementari (de exemplu, Linux), ssthresh inițială este mare, și astfel primul start lent, de obicei, se termină după o pierdere. Cu toate acestea, ssthresh este actualizat la sfârșitul fiecărui start lent, și va afecta de multe ori începe lent ulterioare declanșate de timeout.

Evitare a congestiei. Atâta timp cât ACK-uri non-duplicat sunt primite, fereastra de congestie este aditiv a crescut cu un MSS de fiecare dată dus-întors. Când un pachet este pierdut, probabilitatea de a fi primit ACK duplicat este foarte mare (este posibil, deși puțin probabil ca fluxul tocmai a suferit reordonare pachete de extremă, care ar putea, de asemenea, ACK duplicat prompte). Comportamentul de Tahoe și Reno diferă în modul în care detectează și reacționează la pierderea de pachete:

- Tahoe: ACK triple duplicat sunt tratate la fel ca un timeout. Tahoe va efectua "Fast Retransmit", a stabilit pragul de start lent la jumătate ferestrei de congestie curent, reduce fereastra de congestie la 1 MSS, și a reseta pentru faza de Slow Start. [29].
- Reno: În cazul în care sunt primite trei ACK duplicat (de exemplu, patru ACK recunoscând în același pachet, care nu sunt piggybacked pe date, și nu se schimbă fereastra de publicitate receptorului), Reno va înjumătăți fereastra de congestie (în loc de a stabili o la 1 MSS ca Tahoe), a stabilit pragul de pornire lent egal cu noul ferestrei de congestie, efectua o retransmite rapid, și intră într-o fază numită de recuperare rapidă. În cazul în care un ACK din, început lent este utilizat ca acesta este cu Tahoe. [28]

Recuperarea rapidă. (În cazul algoritmului New Reno) În această stare, TCP retransmite pachetul lipsește, care a fost semnalat prin trei ACK duplicat, și așteaptă o confirmare a întregii ferestrei de transmisie înainte de a reveni la evitarea congestiei. Dacă nu există nici o confirmare, TCP Reno experimentează un timeout și intră în starea de slow-start. Ambii algoritmi reduc fereastra de congestie la 1 MSS pe un eveniment de timeout.

TCP Cubic îmbunătățește corectitudinea împărțirii prin utilizarea în timp real pentru evoluția din fereastră de congestie. S-au făcut multe mai multe cercetări pentru a se găsi o metodă de tratare a congestiei și pentru împărțirea corectă a lățimii de bandă la nivel de rețea. Au fost propuși mai mulți algoritmi AQM bazați pe echitate. Ei folosesc o gamă largă de tehnici pentru a detecta care link mai multă de lățime de bandă decât ar trebuie și aruncă în rețea mai multe pachete ca urmare a ocupării acestui flux mai mare. În unii algoritmi (de exemplu CHOKe), nu există o detectare explicită a ocupării de prea multă lățime de bandă - echilibrul între consumul de bandă este menținut în mod automat și datorită algoritmului ingenios de gestionare a cozii.

Acum trebuie subliniat faptul că cele mai multe studii anterioare au fost concentrate pe un singur tip de soluție și anume, fie pe nivelul rețea (AQM) sau nivelul transport (controlul congestiei TCP). Când noul AQM a fost propus, a fost testat pentru controlul congestiei clasice (New Reno). Similar, după ce a fost inventată noua variantă de TCP, a fost evaluată cu managementul clasic al cozii.

Schemele pentru gestionarea cozii pot fi clasificate în mai multe moduri. Din punct de vedere al performanței se iau în considerare mai mulți parametrii (de exemplu debit, coadă, dimensiune, corectitudine). Diferitele tipuri de algoritmi AQM sunt axate pentru optimizarea unuia sau poate chiar mai multora dintre acești parametrii.

3 TCP NewReno

3.1 Introducere

Performanța TCP este afectată de congestia în rețea și de selecția unui format de date adecvat pentru capacitatea canalului. Congestie apare mai ales atunci când o cantitate mare de flux (tranzacției FTP) este trimis. Scopul acestei lucrări este de a rezolva acest tip de problemă a congestiei, folosind abordări de simulare în algoritmi deja definiți. Traficul de Internet de astăzi este în mare parte realizat de către protocolul de control al transmisiei (TCP). TCP în co-relație cu UDP reprezintă nucleul stratului curent de transport internet. Studii recente arată mai multe eforturi de eficientizare în direcția de control a congestiei și de selecție a ratei de trimitere potrivită pentru tranzații în TCP. Implementări moderne, TCP, sugerează, de asemenea, diferite scheme de control a congestiei folosind fereastră de congestie (cwnd) și tehnici de întârziere. Cu toate acestea, scopul este de a crește dimensiunea maxim posibil a cwnd, fără pierderi de pachete și selecție a ratei de transmisie adecvată înainte de a apărea congestia. Mai multe variante de TCP sunt disponibile, dintre care s-a recunoscut că standardul de transfer TCP New Reno se deteriorează în rețele de mare viteză cu întârzieri de lungimi de bandă mari (BDP). Au fost propuși pentru a aborda această deteriorare noi algoritmi de control a congestie.

Cererile utilizatorilor de internet legate de calitatea rețelei au crescut ca urmare a serviciilor ce au deveni treptat diversificate și mai sofisticate, din cauza gradului remarcabil în care internetul a crescut, ceea ce se datorează în parte de accesul și tehnologii de rețea. Aplicații care implică timp real servicii de livrare mass-media, cum ar fi VoIP, sisteme video streaming și de conferință TV, toate au experimentat un nivel dramatic de dezvoltare care necesită cantități mari și stabile de resurse de rețea, în scopul de a menține calitatea serviciului (QoS). De exemplu; calitatea în timp real a aplicațiilor de streaming este foarte dependentă de întârziere de propagare și întârzierea bruiaj. Lățimea de bandă disponibilă pe calea rețelei end-to-end este de asemenea un factor important pentru asigurarea lin conținut bogat, inclusiv de voce și video. Există o serie de tehnologii de layer de rețea, cum ar fi IntServ și DiffServ, care furnizează astfel de servicii de rețea de înaltă calitate pe Internet. Cu toate acestea, punerea în aplicare a IntServ sau arhitecturi DiffServ ar avea nevoie de mecanisme suplimentare pentru a fi utilizate la toate routerele prin care fluxurile de trafic traversează în scopul de a beneficia în mod suficient de introducerea IntServ sau DiffServ în rețea. Pe de altă parte, un număr de aplicații video streaming folosesc UDP ca protocol de transport, și UDP controlează viteza de transmitere a datelor în funcție de starea rețelei. Cu toate acestea, aceste mecanisme au un cost mare atunci când se modifică programul de aplicație pentru realizarea cerințelor QoS-aplicații specifice, precum și setarea parametrilor sunt foarte sensibile la diverși factori de rețea.

Mai mult decât atât, atunci când co-exista astfel de aplicații în rețea și împărtășesc resursele în rețeaua congestionată, nu putem estima performanța rețelei sau a cererilor, pentru că

mecanismele de control pentru astfel de aplicații sunt concepute și puse în aplicare în mod independent, fără a ține seama de efectul de interacțiuni cu alte aplicații. Deoarece TCP controlează viteza de transmitere a datelor în funcție de starea rețelei.

- Controlul congestiei TCP este guvernat de doi parametrii:
 - Congestion Window (cwnd)
 - Slow-start threshold Value (ssthresh)

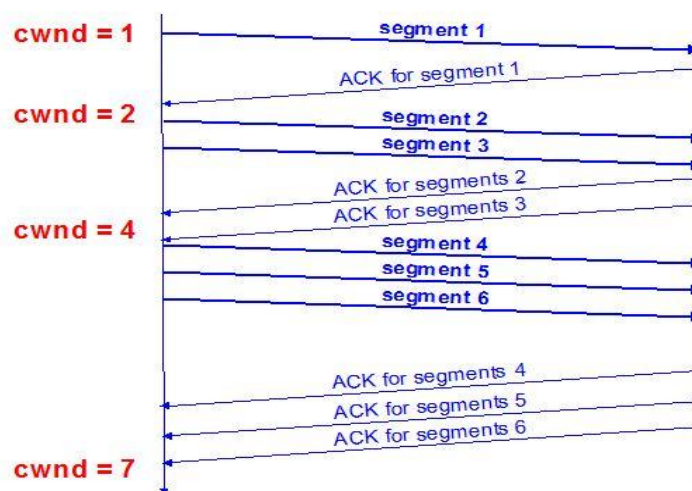
Valorile inițiale sunt : $2^{16} - 1$

- Controlul congestiei are două etape:
 - slow start (cwnd < ssthresh)
 - congestion avoidance (cwnd ≥ ssthresh)

În cazul „Slow Start” valoarea inițială a cwnd este 1 . (Notă: Unitatea este o dimensiune segment. TCP de fapt se bazează pe bytes și suplimente de 1 MSS (dimensiunea maximă a segmentului) . Receptorul trimite o confirmare (ACK) pentru fiecare segment (Notă: În general, un receptor de transport trimite un ACK pentru fiecare alt segment.) De fiecare dată când un ACK este recepționat de către expeditor, fereastra de congestie este majorat cu 1 segment:

$$cwnd = cwnd + 1 \quad (3.1)$$

Dacă un ACK recunoaște două segmente, cwnd este incrementat cu doar 1 segment. Chiar dacă ACK recunoaște un segment care este mai mic decât un MSS bytes lung, cwnd este crescut cu 1. Dimensiunea ferestrei congestie crește foarte rapid. Pentru fiecare ACK, vom incrementa cwnd cu 1 , indiferent de numărul de segmente ACK-uri. TCP încetinește creșterea cwnd când $Cwnd > ssthresh$.



Faza de evitare a congestiei este pornită dacă cwnd a atins valoarea de prag a slow-start-ului.

Dacă $cwnd \geq ssthresh$ apoi de fiecare dată când este recepționat un ACK, $cwnd$ va crește după cum urmează:

$$cwnd = cwnd + \frac{1}{cwnd} \quad (3.2)$$

Astfel $cwnd$ va fi incrementat numai dacă toate segmentele $cwnd$ au fost recunoscute.

3.2.TCP Reno

Performanța la starea de echilibru a unui flux TCP de transfer (de exemplu, un flux cu o cantitate mare de date de trimitere, cum ar fi transferurile FTP) poate fi caracterizată prin rata de trimitere, care reprezintă cantitatea de date transmise de către expeditor în unitatea de timp. O cantitate semnificativă de trafic pe Internet, inclusiv WWW (HTTP), transfer de fișiere (FTP), e-mail (SMTP), și accesul de la distanță (Telnet) de trafic, se efectuează de către protocol de transport TCP. TCP împreună cu UDP formează miezul de layer de transport pe Internet de astăzi. În mod tradițional, simularea și punerea în aplicare și măsurarea au fost instrumentele de alegere pentru examinarea performanțelor diferitelor aspecte ale TCP-ului. Recent, atenția a fost îndreptată către caracterizarea rata de trimitere a unui flux TCP ca o funcție de pierdere de pachete și de întârziere dus-întors.

Pentru a evita colapsul rețelei cauzat de congestie, TCP folosește o strategie de control al congestiei multi-fațete. 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 TCP folosită pentru controlul fluxului. TCP folosește un mecanism numit „slow start” (start lent) pentru a mări fereastra de congestie, dacă o conexiune este inițializată și după un timeout. Acesta începe cu o fereastră de două ori mai mare decât dimensiunea maximă a segmentului (MSS). Deși rata inițială este scăzută, rata de creștere este foarte rapid: pentru fiecare pachet recunoscut, fereastra de congestie crește cu 1 MSS, astfel încât fereastra de congestie se dublează în mod eficient de fiecare dată dus-întors (RTT). Când fereastra de congestie depășește un prag „ssthresh” algoritmul intră într-o stare nouă, numită evitare a congestiei. În unele implementări (de exemplu, Linux), „ssthresh” inițial este mare, și astfel primul start lent, de obicei, se termină după o pierdere. Cu toate acestea, ssthresh este actualizat la sfârșitul fiecărui start lent, și va afecta de cele mai multe ori un început lent ulterior declanșat de timeout.

Evitare a congestiei. Atâta timp cât ACK-uri non-duplicate sunt primite, fereastra de congestie este aditivă, crescând cu un MSS de fiecare dată când este efectuat un ciclu dus-întors. Când un pachet este pierdut, probabilitatea de a fi primit ACK duplicat este foarte mare (este posibil, deși puțin probabil ca fluxul să fi suferit reordonarea pachetelor). Comportamentul algoritmilor Tahoe și Reno diferă în modul în care detectează și reacționează la pierderea de pachete:

Tahoe: ACK triple duplicat sunt tratate la fel ca un timeout. Tahoe va efectua ” fast

retransmit", a stabilit pragul de start lent la jumătate ferestrei de congestie curentă, reduce fereastra de congestie la 1 MSS, și a reseta pentru slow-start state.

Reno: În cazul în care sunt primite trei ACK duplicat (de exemplu, patru ACK recunoscute în același pachet, dar ale caror date nu sunt întorase, și nu se schimbă fereastra de avertizare a receptorului), Reno va înjumătăți fereastra de congestie (în loc de a stabili-o la 1 MSS ca Tahoe), a stabilit pragul de pornire lent egal cu noua fereastră de congestie, efectuând o retransmitere rapidă, și intră într-o fază de recuperare rapidă.

Recuperarea rapidă. Reno în această stare, TCP retransmite pachetul care lipsește, semnalat în urma cu trei ACK duplicat, și așteaptă o confirmare a întregii ferestrei de transmisie înainte de a reveni la evitarea congestiei. Dacă nu există nici o confirmare, TCP Reno experimentează un timeout și intră în starea de slow-start.

Ca o definiție putem spune că:

Pentru controlul congestiei TCP se definesc următorii parametri:

- Dimensiunea maximă a segmentului expeditor (SMS-uri) reprezintă suma maximă de date care pot fi trimise într-un singur segment TCP, fără a include antetul.
- Fereastra expeditorului (swnd) reprezintă numărul maxim de bytes care poate fi trimisi. Valoarea sa are valoarea cea mai mică dintre fereastra receptor și fereastra de congestie.
- Fereastra destinatarului (rwnd) este cel mai nou fereastra de publicitate de către receptor.
- fereastra de congestie (cwnd) este o variabilă de stare TCP, care limitează cantitatea de date care pot fi trimise.
- Fereastra pentru pierderi (lw) este valoarea ferestrei de congestie după o pierdere de pachete a fost detectată.
- prag *slow-start* (lent-start) (sssthresh) este o altă variabilă TCP, care determină ca algoritmul de control al congestiei să fie declansat sau nu : fie de pornirea a unui *slow-start* (dacă $cwnd \leq sssthresh$), fie de evitarea congestiei (dacă $cwnd \geq sssthresh$).

Slow-Start și evitarea congestiei : Slow-Start, denumit la propriu, crește exponențial dimensiunea ferestrei de congestie. Acesta este utilizat de către o entitate TCP la începutul unei transmisiuni sau după detectarea unei pierderi de pachete. Scopul Slow-Start este de a umple cât mai curând posibil, un canal de transmisie. După ce fereastra de congestie a atins valoarea de prag, algoritmul de evitarea a congestiei este activat. Fereastra de congestie continuă să crească liniar, adăugând până la un MSS, dar nu mai puțin de un octet. În ambele cazuri un timer de retransmisie este utilizat pentru fiecare pachet. Timeout-ul semnalează pierderea de pachete. Acest lucru duce la retransmiterea și înjumătățirea pragului Slow-Start. Fereastra de congestie este setată la valoarea pierderii ferestrei.

Expeditorul păstrează două variabile pentru controlul congestiei: o fereastră *slow-start/congestion*, cwnd, și o dimensiune prag, sssthresh, pentru a comuta între cele două algoritmi.

Expeditorul trimite întotdeauna un minim de cwnd și advertisementul ferestrei receptorului. Pe un timeout, jumătate din dimensiunea ferestrei curente este înregistrată în sssthresh, apoi cwnd este setat la 1 pachet (acesta inițiază Slow-Start-ul). Când date noi sunt acked, expeditorul va face următorul algoritm:

If (cwnd<ssthresh)

/* daca in continuare se face procesul de Slow-start , se va deschide fereastra exponential*/ Cwnd+=1;

Else /* altfel se va face congestia evitand cresterea cu -1*/

Cwnd+=1/cwnd;

Astfel, Slow-Start deschide repede fereastra la ceea ce avertismentul de congestiei crede că este un punct de funcționare în condiții de siguranță (jumătate fereastra care ne-a băgat în bucluc), apoi avertismentul de congestie este eliminat și se încet crește dimensiunii ferestrei pentru a sonda pentru mai multă lățime de bandă disponibilă în rețea.

Recuperarea rapidă și retransmiterea rapidă: După punerea în aplicare a algoritmului , noi probleme apar. Primul este legat la detectarea pierderii de pachete. În mod normal, o pierdere de pachete se deduce pe baza de expirare a timer-ului de retransmisie. Aceasta poate duce la întârzieri semnificative în transmisii de date, astfel încât un alt mod de a determina pierderea de pachete a fost adăugat la TCP. În condiții normale, o entitate TCP trebuie să trimită un ACK duplicat pentru fiecare pachet care vine din secvență. Un pachet poate fi primit din secvență ca urmare a dublării de pachete de rețea, întârzieri de pachete sau pierdere. Algoritmul de retransmisie rapida consideră că un pachet a fost pierdut atunci când primește 3 ACK duplicat, înainte de expirare a timer de retransmisie. În acest fel timp prețios este salvat.

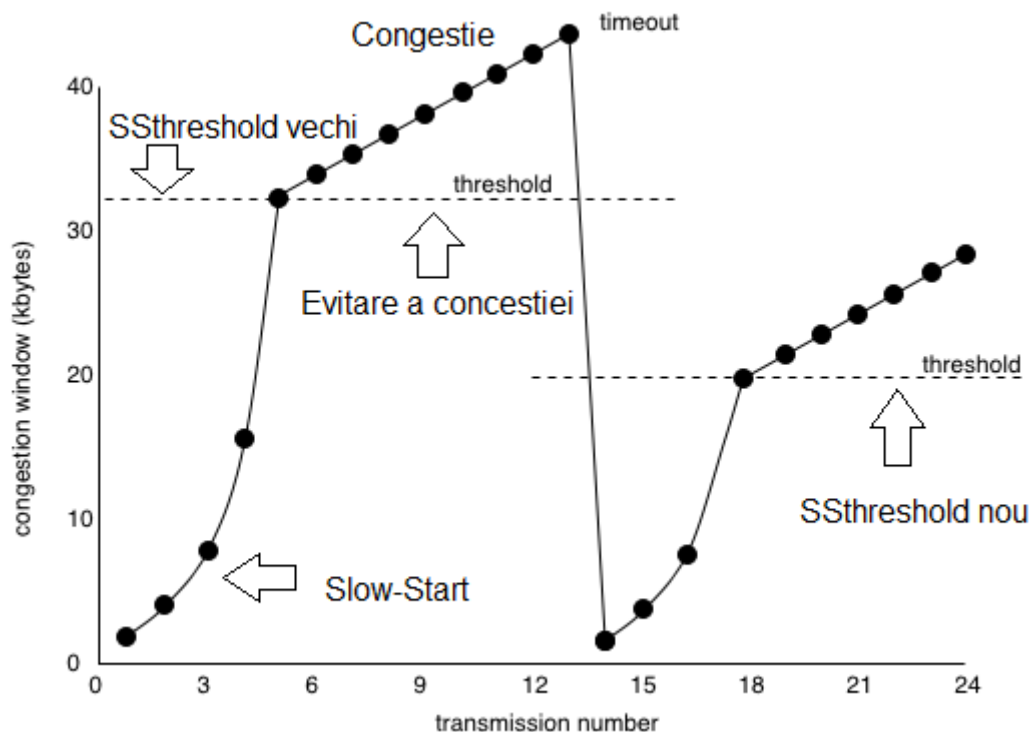


Figura 3.1 : Slow-Start și evitarea congestiei

A doua problemă este legată de scăderea drastică a ferestrei de congestie după o detectare a pierderilor de pachete. Dacă un pachet este pierdut în timpul Slow-Start sau evitarea congestiei, valoarea ferestrei de congestie este setată la valoarea pierderii ferestrei (1 SMSS). Algoritmul de recuperare rapidă abordează această problemă. Noua valoare a ferestrei de congestie, după o pierdere de pachete este detectată de retransmisia rapidă și este setată la $ssthresh + 3 \text{ SMSS}$. Aceasta se numește "inflație artificială" a ferestrei de congestie. Pe lângă faptul că, pentru fiecare nou duplicat ACK primit, fereastra de congestie va crește cu un SMSS. Mecanismul de evitare a congestiei, în cazul în care dimensiunea ferestrei de control a congestiei este W , atunci dimensiunea sa crește cu $1/W$ de fiecare dată când este recepționat un ACK. Invers, fereastra este redusă de fiecare dată când un pachet pierdut este detectat.

Există un comportament de evitare a congestiei TCP în termeni de "runde". O rundă începe cu transmisia de pachete de W , în care W este dimensiunea actuală a ferestrei de congestie TCP. După ce toate pachetele care intră în fereastra de congestie au fost trimise, nu mai este trimis nici un alt pachet până când primul ACK este recepționat de unul dintre aceste pachete W . Acest ACK marchează sfârșitul etapei actuale și începutul runda următoare. În acest model, durata unei runde este egal cu RTT și este considerată independentă față de dimensiunea

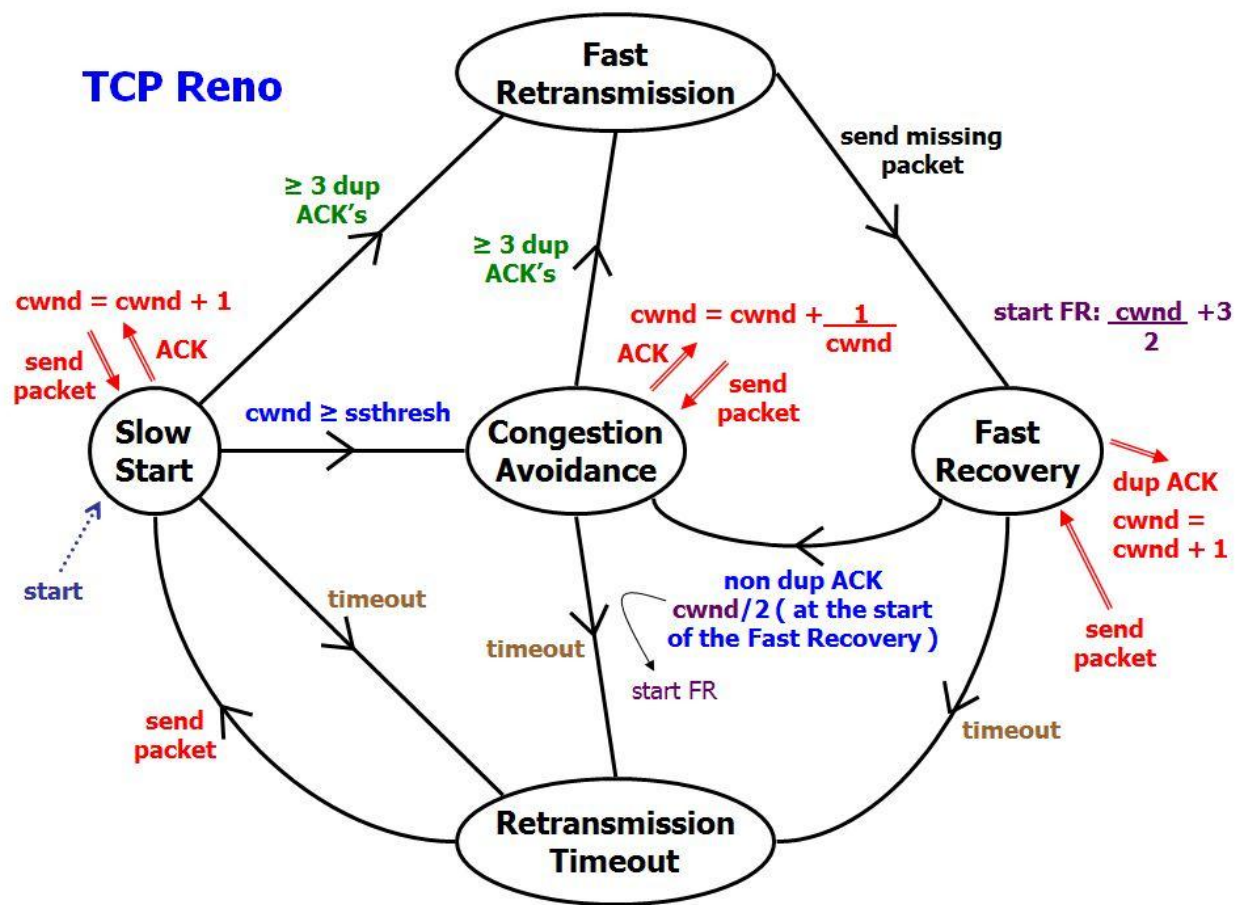


Figura 3.2 : Diagrama Algoritm TCP Reno

ferestrei. Conceptul nostru de runde este similar cu conceptul de "mini-cicluri.

3.2 Caracteristici New Reno

New Reno, definit de RFC 6582 (care a devenit comun definițiile anterioare în RFC 3782 și RFC 2582), îmbunătățește retransmisia în timpul fazei de Fast recovery a algoritmului TCP Reno. Timpul de recuperare rapidă, pentru fiecare ACK duplicat care este returnat la TCP New Reno, un nou pachet netrimis de la sfârșitul ferestrei de congestie este trimis, pentru a menține fereastra de transmisie integrală. Pentru fiecare ACK care face progrese parțiale în spațiul secvență, expeditorul presupune că punctele de ACK la o nou nod , iar următorul pachet aflat după numărul de ordine al ultimului ACK este trimis.

Nou RENO reprezintă o ușoară modificare a algoritmului AQM TCP-RENO. Acesta este capabil să detecteze multiple pierderi de pachete și este mult mai eficient decât RENO în ceea ce privește pierderea de pachete. Ca Reno, New-Reno intră de asemenea în rapidă retransmisie în care primește mai multe pachete duplicate, însă în acest caz , New Reno diferă de RENO prin faptul că nu iese din Fast-recovery cât atunci toate datele care au fost pierdute în momentul în intrării în faza de recuperare rapidă sunt acknowleg. Astfel, se depășește problema cu care se confruntă Reno , de reducere a CWD de mai multe ori .

Faza de fast-transmit are aceleași etape ca la Reno. Diferența constă ca în faza de fast-recovery trebuie să se retransmită toate pachetele pierdute. Ori de câte ori New-Reno intră în faza de recuperare rapidă se constată segmentul maxime, care este foarte important . Deși aceasta fază se desfășoară ca la Reno, cu toate acestea, atunci când un ACK este proaspăt primit, atunci există două cazuri:

- Dacă sunt ACK toate segmentele care au fost retransmise , am intrat în faza de fast-recovery , atunci se iese din această și se setează CWD la ssthresh și se continuă transmisia
- Dacă ACK este un ACK parțială atunci se deduce că următorul segment în linie a fost pierdut și se re-transmite segmentul și se stabilește numărul de ACK duplicat primite de la zero. Se iese de recuperare rapidă, atunci când toate datele din fereastra sunt recunoscute.

Deoarece cronometrul timeout este resetat ori de câte ori există un progres în bufferul de transmisie, aceasta permite New Reno să umple arii mari, sau mai multe arii , în spațiul de secvență - la fel ca TCP Sack . Deoarece New Reno poate trimite noi pachete de la sfârșitul ferestrei de congestie în timpul de recuperare rapidă, randamentul ridicat se menține în timpul procesului de umplere a ariei , chiar și atunci când există mai multe arii, de mai multe pachete fiecare. Când TCP intră în recuperare rapidă se înregistrează cel mai mare număr de secvențe de pachete nerecunoscut. Când acest număr de ordine este recunoscut, TCP revine la starea de evitare a congestiei.

O problemă apare la New Reno, atunci când nu există pierderi de pachete, dar în schimb, pachetele sunt reordonate cu mai mult de 3 numere de secvență de pachete. Când se întâmplă acest lucru, New Reno intră din greșeală în recuperare rapidă, dar în cazul în care pachetul reordonat este livrat, ACK-secvență de numărare a progres apare și de acolo până la sfârșitul recuperării rapide, fiecare bit de secvență a numărului de progres produce un duplicat care este imediat ACKed.

În TCP Reno, dimensiunea ferestrei este schimbată continuu într-o situație tipică. Dimensiunea ferestrei continua să crească până se produc pierderi de pachete. TCP Reno are două faze în creșterea dimensiunii de fereastră: Etapa de start lent și de evitare a congestiei. Când un pachet ACK (confirmare) este primit de TCP la expeditor la momentul $t + t'$ [sec], actuala dimensiune a ferestrei $cwnd(t + t')$ este actualizată de la $cwnd(t)$, după cum urmează:

- Faza de Slow-Start :

If $cwnd(t) < ssthresh(t)$
 $Cwnd(t+t') = cwnd(t) + 1;$

- Faza de evitare a congestiei:

If $cwnd(t) > ssthresh(t)$
 $Cwnd(t+t') = cwnd(t) + 1 / cwnd(t)$

În cazul în care, $ssthresh(t)$ este o valoare de prag la care TCP se schimbă de la faza de început lent la faza de evitare a congestiei. Când sunt detectate pierderi de pachete prin expirarea retransmisie timeout, $cwnd(t)$ și $ssthresh(t)$ sunt actualizate ca:

$$Cwnd(t) = 1; \quad (3.3)$$

$$ssthresh(t) = cwnd(t) / 2; \quad (3.4)$$

Pe de altă parte, atunci când se detectează pierderea de pachete TCP de către un algoritm de retransmisie rapid, se schimbă $cwnd(t)$ și $ssthresh(t)$ ca;

$$Cwnd(t) = ssthresh(t); \quad (3.5)$$

$$ssthresh(t) = cwnd(t) / 2; \quad (3.6)$$

TCP Reno intră apoi într-o fază de recuperare rapidă în cazul în care pierderea de pachete este găsită de către algoritmul retransmis rapid. În această fază, dimensiunea ferestrei este mărită cu un pachet atunci când un pachet ACK duplicat este primit. Pe de altă parte, $cwnd(t)$ este readus la $SSTH(t)$ atunci când pachetul ACK non-duplicat corespunzând pachetul retransmis este primit.

Conexiunea TCP pornește în modul de pornire lent și crește exponențial fereastra de congestie ($cwnd$) până când se trece pragul start lent ($ssthresh$). Odată ce $cwnd$ este mai mare decât $ssthresh$, TCP se mută la faza de evitare a congestiei. În acest mod, obiectivul principal este de a menține debitul mare, fără a provoca congestie. Dacă TCP detectează orice pierdere a

unui segment, acest lucru este un indiciu puternic că rețeaua este saturată, astfel ca o acțiune corectivă TCP reduce rata de flux de date prin reducerea cwnd după care din nou se întoarce pentru a încetini faza de început.

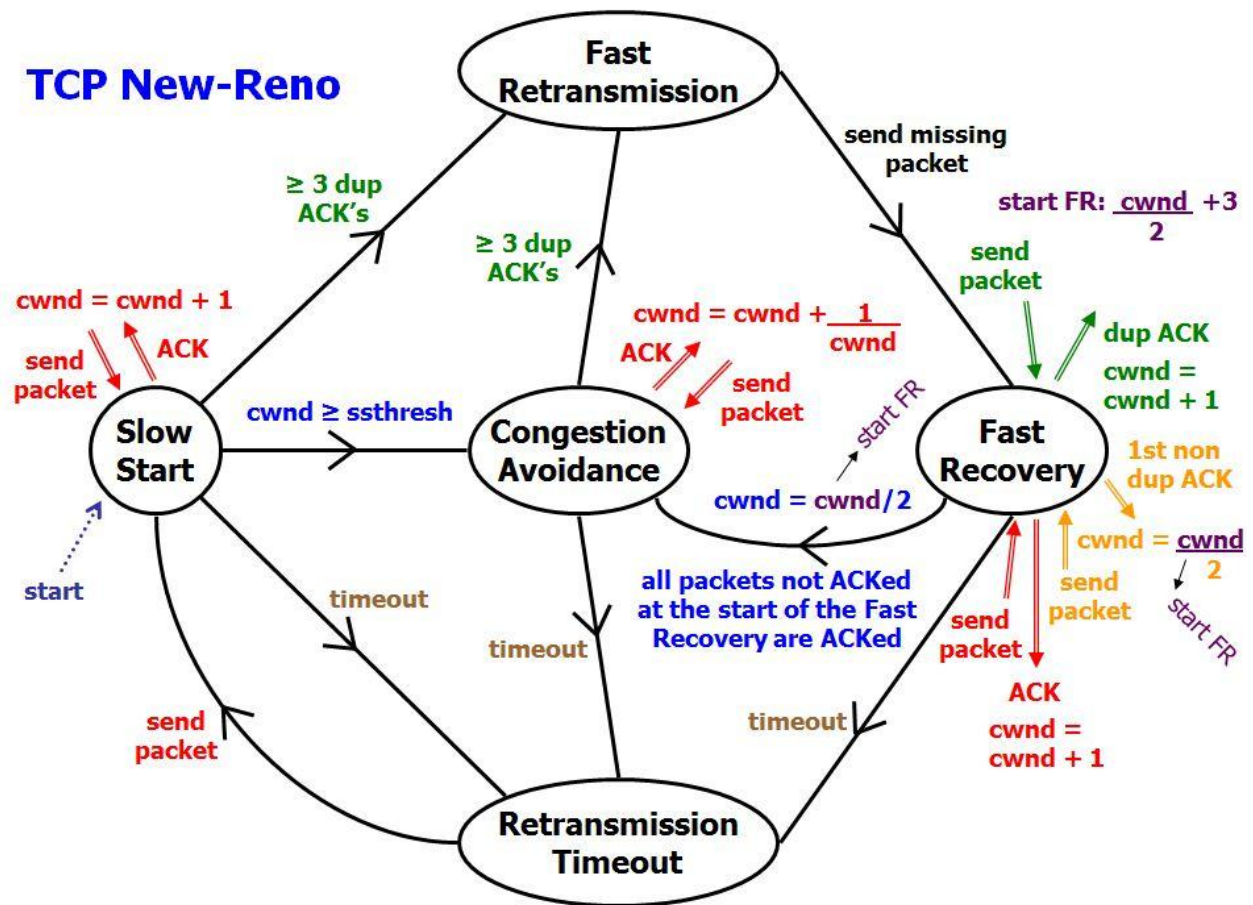


Figura 3.3 Diagrama Algoritm TCP New Reno

Principiu de funcționare a algoritmul New Reno este urmatorul:

- Algoritmul Slow Start
Cwnd=1;
Executa pentru fiecare pachet în parte : $cwnd = cwnd + 1$
Pana cand (evenimentul de congestie sau $cwnd > ssthresh$
- Algoritmul de evitare a congestiei:

// atunci cand se iese din Slow Start si $cwnd > ssthresh$

Executa pentru fiecare Ack: $cwnd = cwnd + (1/cwnd)$

Pana cand expira timpul (Timeout) sau trec 3 Dupacks

- Algoritmul de retransmisie rapidă:
// după ce s-au transmis 3 dupacks
Retrimite pachetele pierdute;
Apelează algoritmul de recuperare rapidă
- Algoritmul de recuperare rapidă
// după retransmisia rapidă nu se intra în Slow Start
 $Ssthresh = cwnd/2$
 $Cwnd = ssthresh + 3$;
Pentru fiecare DACK primit;
 $Cwnd = cwnd + 1$
Trimite un packet nou daca este permis;
Dup ace s-a primit un ACK:
 Daca ACK este partial :
 Rămâi în recuperare rapida;
 Retransmite următorul pachet pierdut ;
 Daca ACK este complet:
 $Cwnd = ssthresh$;
 Iesi din recuperare rapida
 Apelează algoritmul de evitare a congestiei;
- Cand se termina timpul de Timeout:
 $Ssthresh = cwnd/2$;
 $Cwnd = 1$;
 Apeleaza algoritmul de Slow Start

4 . Modificare Algoritm

4.1 Introducere

Problema atât cu Reno cât și cu NewReno este că în algoritmul de recuperare rapidă, TCP înjumătățește fereastra de congestie, indiferent de starea rețelei. O altă problemă cu NewReno este că, atunci când nu există pierderi de pachete, dar pachetele sunt reordonate cu mai mult de trei confirmări duplicat, NewReno intră în mod eronat în recuperare rapidă, iar acest aspect duce iar la înjumătățirea ferestrei de congestie. Deoarece fereastra TCP de congestie controlează numărul de pachete pe care un expeditor TCP le poate trimite în rețea, în orice moment, prin urmare, procesul de stabilire a dimensiunii ferestrei de congestie la jumătate din valoarea face ca algoritmul NewReno TCP să fie inefficient în ceea ce privește utilizarea capacității de legătură.

Reno și New Reno retransmite cel mult un pachet pierdut în unitatea de timp pentru un drum dus-întors. Deoarece se urmărește următoarele seturi de metrici: rata de transmisie, network fairness, comportamentul ferestrei de transmisie și rata de pierdere a pachetelor analizăm comportamentul algoritmului TCP NEW Reno cu Tahoe, Reno. Toți algoritmi AQM menționați, în momentul transmisiei, au fereastra de transmisie cu o valoare scăzută, ulterior crescând în mod constant până la un vârf, unde după un interval de timp fereastra va scade , rămânând constantă până la stabilirea unei noi conexiuni. Pentru fiecare conexiune nouă a existat o scădere a ratei de transmisie, urmată de o fază de stabilizare . Algoritmul New Reno își păstrează comportamentul oscilant la fel și Reno și Tahoe.

În ceea ce privește împărțirea corectă a fluxurilor în rețea (Network Fairness) TCP New Reno are performanțe foarte bune, fluxurile cu debit mai mare nu consumă întreaga bandă , lăsând fluxurile cu debit mai mic, au o transmisie echitabilă .

Numărul de pachete pierdute în rețea este un factor important atunci când se analizează calitatea metodei de control al congestiei, deoarece pierderea de pachete, poate avea ca efect creșterea congestie. Chiar și timeout, primirea de pachete triple este considerat un semn de congestionarea a rețelei. În funcție de numărul de pachete triple, New Reno a înregistrat performanțele cele mai bune conform studiilor .

Deoarece New-Reno are la bază algoritmul Reno, modificat la partea de Fast-recovery, iar Reno la rândul său are la bază algoritmul Tahoe, îmbunătățind partea de fast Recovery, iar Tahoe este alcătuit din Slow start, congestion avoidance and fast retransmit, dacă se îmbunătățește una din cele patru părți ce ajută la funcționarea algoritmului (Slow start, congestion avoidance and fast retransmit, fast-recovery), acesta va reprezenta o versiune îmbunătățită a algoritmului ce stă sub el. Am ales să modific partea de recuperare rapidă .

4.2 Soluția propusă pentru modificarea algoritmului

Propun o modificare a algoritmul de recuperare rapidă . În timpul evitării congestiei , în cazul în care expeditorul TCP primește trei ACK duplicat, se trece în modul Fast Retransmit să Propun o modificare a algoritmul de recuperare rapidă . În timpul evitării congestiei , în cazul în care expeditorul TCP primește trei ACK duplicat, se trece în modul Fast Retransmit să retransmită pachete ce pare a fi pierdute fără a aștepta să expire cronometrul retransmisiei. Apoi se calculează fereastra nouă de congestie , se stabilește ssthresh la maxim două segmente de cwnd , iar cwnd se stabilește la valoarea ssthresh plus numărul de recunoașteri duplicate primite și continuă cu faza de recuperare rapidă. Expeditorul TCP incrementează cwnd cu unul pentru fiecare confirmare duplicat primită, și trimite segmentul nou dacă este permis. Cu ACK parțială, se retransmit segmentele confirmate și se procesează. Cu ACK complet, se stabilește cwnd la valoarea ssthresh și se apelează algoritmul de evitare a congestiei. În cazul în care expeditorul TCP detectează pierderile de expirare a timeout-ului , se stabilește ssthresh la maxim de două segmente și cwnd, și seturile de cwnd la un segment, și intră în algoritmul Slow Start.

Mecanismul propus evită aceste probleme prin modificarea algoritmului de revenire rapidă a TCP - NewReno. Ideea de bază este de a ajusta ferestrei de congestie (cwnd) a expeditorului TCP în funcție de nivelul de congestie în rețea, pentru a permite transferul mai multor pachete la destinație. Ajustarea ferestrei de congestie are trei obiective principale. Prima dintre ele este de a utiliza resursele de rețea disponibile, al doilea este de a minimiza probabilitatea de congestie, iar al treilea este de a oferi o lățime de bandă echitabil între mai multe conexiuni. Aceste obiective ar putea fi realizat prin adoptarea dimensiunii ferestrei de congestiei bazat pe starea rețelei.

5. Simulări

NS2 este un simulator de evenimente, orientat pe simularea în rețea care să permită imitarea într-o varietate de rețele locale și largi. El pune în aplicare diferite protocoale de rețea (TCP, UDP), surse de trafic (FTP, web, CBR, exponențiale on / off), mecanisme de gestionare coadă (RED, DropTail), protocoale de rutare (Dijkstra), etc. NS2 este scris în C++ și Otcl la separa de control și implementări cale de date. Simulatorul are o ierarhie de clase în C++ (ierarhia compilat) și o ierarhie corespunzătoare în interpretul Otcl (ierarhie interpretat). Motivul pentru care ns2 folosește două limbi este că sarcini diferite au cerințe diferite: Un exemplu de simulare de protocoale necesită manipularea eficientă de bytes și antete de pachete de a avea viteză. Pe de altă parte, în studiile de rețea în cazul în care obiectivul este de a varia parametri și a examina rapid o serie de scenarii de timp pentru a schimba modelul și rulați-l din nou, este mult mai important. În NS2, C++ este utilizat pentru a aplica detaliat protocol și, în general, pentru astfel de cazuri în care fiecare pachet de un flux trebuie să fie prelucrate. De exemplu, dacă doriți să pună în aplicare o nouă disciplină de așteptare, apoi C++ este limbajul de alegere. Otcl, pe de altă parte, este adecvat pentru configurarea și configurare. Otcl rulează destul de încet, dar poate fi schimbat foarte repede face construcția simulări mai ușor. În NS2, a compilate C++ a obiectelor poate fi pusă la dispoziția interpretului Otcl. În acest fel, obiectele gata făcute C++ obiecte pot fi controlate de la nivelul OTcl. Există destul de multe tutoriale de înțeles disponibile pentru utilizatorii NS-noi. Trecând prin, pentru exemplu, următoarele tutoriale ar trebui să vă dau o imagine destul de bună a modului de a crea scenarii simple, de simulare cu NS2: <http://nile.wpi.edu/NS/> sau <http://www.isi.edu/nsnam/ns/tutorial/index.html>

OTcl

Network Simulator este un simulator orientat pe obiecte, scris în C++, cu o interfață făcută de un interpretor OTcl. Simulatorul suportă o ierarhie a claselor în C++, și o ierarhie a claselor asemănătoare pentru interpretorul OTcl. Cele două ierarhii sunt strâns legate una de cealaltă. Din perspectiva utilizatorului, există o corespondență unu la unu, între o clasă din cadrul ierarhiei interpretorului și una din cadrul ierarhiei compilate. Rădăcina acestei ierarhii este clasa TclObject. Utilizatorii creează simulări noi prin intermediul interpretorului. Aceste obiecte sunt instanțiate în cadrul interpretorului și sunt oglindite de un obiect corespondent în ierarhia compilată. Ierarhia claselor interpretorului este realizată automat prin metode definite în clasa TclClass. Obiectele instanțiate de utilizator sunt oglindite prin metode definite în clasa TclObject. Există și alte ierarhii în codul C++ și scripturile OTcl. Aceste alte ierarhii nu sunt oglindite în maniera TclObject.

De ce două limbaje?

Network Simulator folosește două limbaje deoarece simulatorul are două tipuri de

obiective pe care trebuie să le realizeze. Pe de o parte, simulări detaliate ale protocoalelor necesită un limbaj de programare al sistemului, care poate manipula eficient bytes, pachete, anteturi și care poate implementa algoritmi ce rulează cu seturi foarte mari de date. Pentru aceste obiective, viteza la rulare este importantă, în timp ce rularea simulării, găsirea bug-urilor și repararea acestora este mai puțin importantă. Pe de altă parte, o mare parte a cercetărilor în rețele implică variații mici ale parametrilor sau configurațiilor sau numărului de scenarii. În aceste cazuri, timpul necesar iterației (schimbarea modelului și rularea acestui din nou) este mai importantă. Din moment ce configurația rulează o singură dată (la începutul simulării), timpul de rulare pentru această parte este mai puțin important.

NS împacă cele două cerințe cu cele două limbaje, C++ și OTcl. C++ este rapid la rulare, însă mai greu la schimbare, făcându-l potrivit pentru implementările detaliate ale protocoalelor. OTcl rulează mult mai greu, însă poate fi schimbat foarte rapid (și interactiv), făcându-l ideal pentru configurarea simulărilor.

Network Animator

Nam este o unealtă pentru animații, bazată pe Tcl/Tk, folosită pentru vizualizarea datelor simulărilor de rețele. Teoria de implementare din spatele Nam a fost crearea unui animator care să poată citi seturi foarte mari de date pentru animații și care să poată fi îndeajuns de extensibil pentru a putea vizualiza o varietate foarte mare de situații de rețele. Având această constrângere, **nam** a fost creat să poată citi comenzi simple de animație din interiorul unor fișiere foarte mari de urmărire. Pentru a putea lucra cu seturi de date foarte mari pentru animații, cantitatea de informație reținută în memorie este foarte mică. Comenzile sunt reținute în fișier și sunt recitite ori de câte ori este necesar.

Primul pas pentru a putea folosi **nam** este crearea unui fișier de urmărire. Acest fișier conține informații referitoare la topologii, noduri, legături dar și pachete. De obicei, acest fișier este creat de către Network Simulator.

După ce a fost generat fișierul pentru urmărire, acesta este pregătit pentru a putea fi animat de către nam. La pornire, nam citește fișierul, crează topologia, creează o fereastră pop-up, realizează layout-ul necesar și se oprește la momentul de timp 0. Network animator are o interfață pentru utilizator însă permite controlul asupra multor aspecte ale animației.

Instalarea simulatorului

Simulările au fost realizate cu ajutorul simulatorului Network Simulator v2.35 rulând pe un sistem de operare Ubuntu 10 instalat pe o mașină virtuală creată pe un Notebook HP 4520s, având următoarele specificații: processor Intel Core i3, memorie RAM 4GB, memorie externă 400GB.

Pentru instalarea simulatorului am urmat următorii pași:

-descărcarea și instalarea fișierului ns-2 all-in-one cu următoarele comenzi:

```
$ wget http://nchc.dl.sourceforge.net/sourceforge/nsnam/ns-allinone-2.34.tar.gz
```

```
$ tar -xvzf ns-allinone-2.34.tar.gz
```

```
$ cd ns-allinone-2.34
```

```
$ sudo apt-get install build-essential autoconf automake libxmu-dev
```

```
$ ./install
```

-setarea variabilelor [6]

```
$ gedit ~/.bashrc
$ source ~/.bashrc
-validare
$ cd ns-2.34
$ ./validate
```

Modificarea simulatorului

Pentru implementarea modificarilor propuse in capitolul anterior, evitarea trecerii în modul Fast Retransmit care retransmite pachete ce par a fi pierdute fără a aștepta să expire cronometrul retransmisiei a trebuit sa modifice algoritmului de revenire rapidă a TCP - NewReno. Aceasta modificare se realizeaza prin inlocuirea algoritmului sursa din NewReno.cc si NewReno.h ale simulatorului NS2.

În fișierul sursă a algoritmului NewReno, NewReno.cc, am făcut următoarea modificare: în algoritmul de recuperare rapidă , Ssthresh nu va mai lua valoarea cwnd/2 , ci va ii va fi atribuita jumătate din cwnd știuta ulima data plus 3 (numărul de ACK neconfirmate în faza de congestie , dar recuperate în urma algoritmului de retransmitere rapidă. Modificarea in codul sursa este urmatoarea:

```
reno_action:

    recover_ = maxseq_;

    reset_rtx_timer(1,0);

    if (!lossQuickStart()) {

        trace_event("NEWRENO_FAST_RETX");

        last_cwnd_action_ = CWND_ACTION_DUPACK;

        last_cwnd_action_ = CWND_ACTION_ECN;

        slowdown(CLOSE_SSTHRESH_HALF|CLOSE_CWND_HALF);

        output(last_ack_ + 1, TCP_REASON_DUPACK);    // from top

        dupwnd_ = numdupacks_;

    }

    return;

}
```

```

void NewRenoTcpAgent::recv(Packet *pkt, Handler*)
{
    hdr_tcp *tcph = hdr_tcp::access(pkt);
    int valid_ack = 0;

    /* Use first packet to calculate the RTT --contributed by Allman */

    if (qs_approved_ == 1 && tcph->seqno() > last_ack_)
        endQuickStart();
    if (qs_requested_ == 1)
        processQuickStart(pkt);
    if (++acked_ == 1)
        basertt_ = Scheduler::instance().clock() - firstsent_;

    /* Estimate ssthresh based on the calculated RTT and the estimated
       bandwidth (using ACKs 2 and 3). */

    else if (acked_ == 2)
        ack2_ = Scheduler::instance().clock();
    else if (acked_ == 3) {
        ack3_ = Scheduler::instance().clock();
        new_ssthresh_ = int((basertt_ * (size_ / (ack3_ - ack2_))) / size_);
        if (newreno_changes_ > 0 && new_ssthresh_ < ssthresh_)
            ssthresh_ = new_ssthresh_;
    }
}

```


Scenariul Simulat :

Topologia utilizată este alcătuită din:

- 7 link-uri cu o lățime de bandă de 10 Mbps
- 3 link bottleneck cu o lățime de bandă de 1 Mbps și un delay de 50 ms

Considerăm că avem șapte conexiuni FTP ce utilizează TCP, și setăm dimensiunea maximă a ferestrei la 30. Fiecare dintre cele șapte noduri trimite trafic FTP pe parcursul întregii perioade de simulare (50 de secunde). Dimensiunea cozii de bottleneck este 25. Se alege dimensiunea de 552 de bytes pentru pachetul TCP (pachetul efectiv creat în timpul simulării are 592 bytes deoarece se adaugă 40 bytes ce corespund headerului).

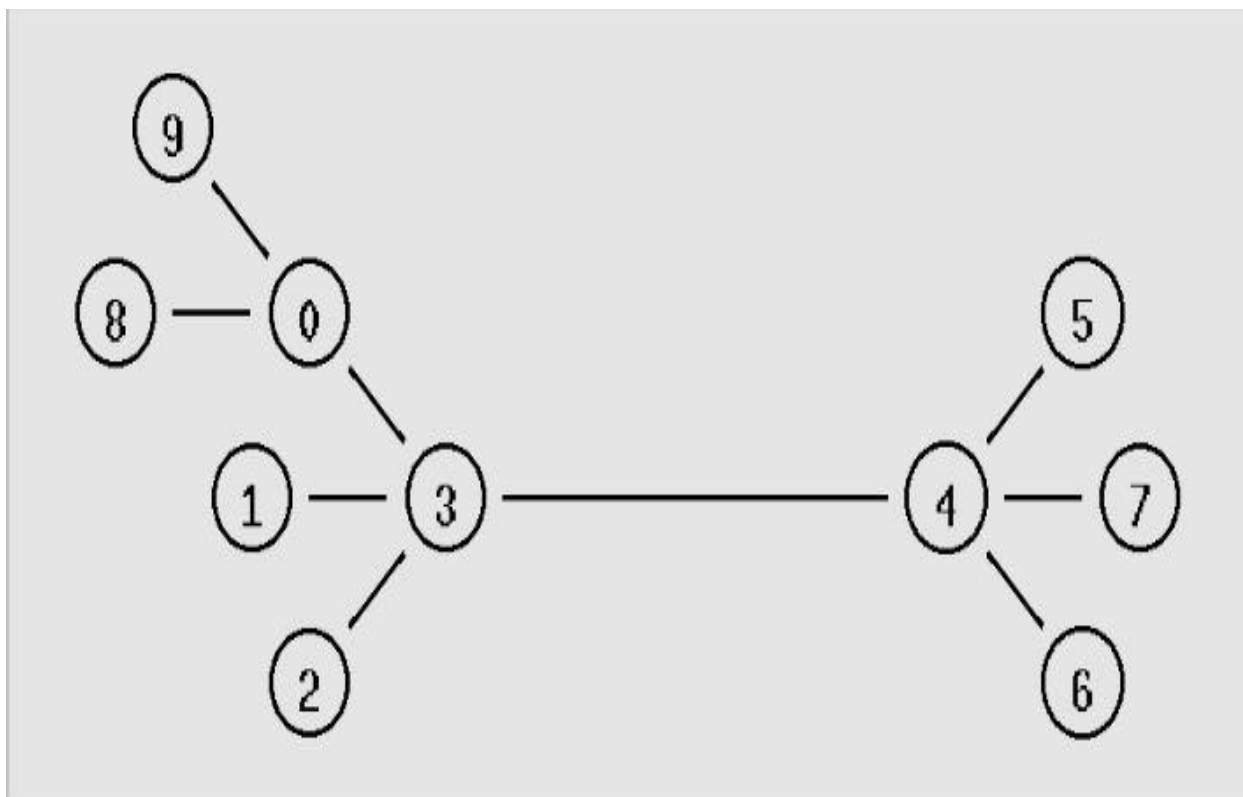


Figura 5.1 Topologia rețelei simulate

6. Prezentarea rezultatelor

Am rulat simulările pentru topologia descrisă în capitolul anterior , iar în continuare vom prezenta graficele care descriu variația ferestrei de congestie în cazul algoritmului modificat New Reno, New Reno, Reno și Tahoe , precum și graficele rezultate în urma comparației lățimii de bandă , numărul de pachete retransmise, pierdute și întârzierea între pachete în cazul fiecărui algoritm din cei menționați.

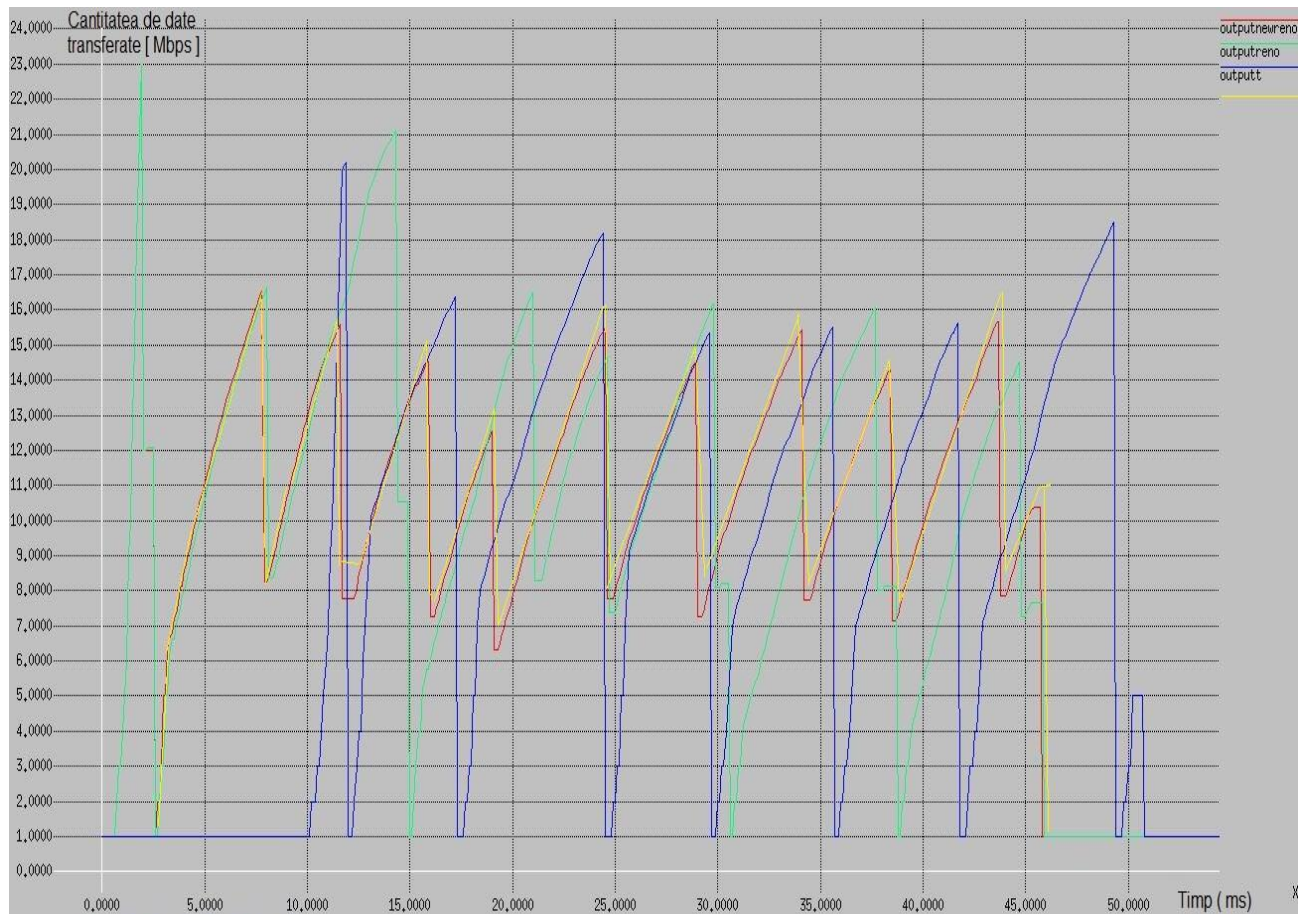
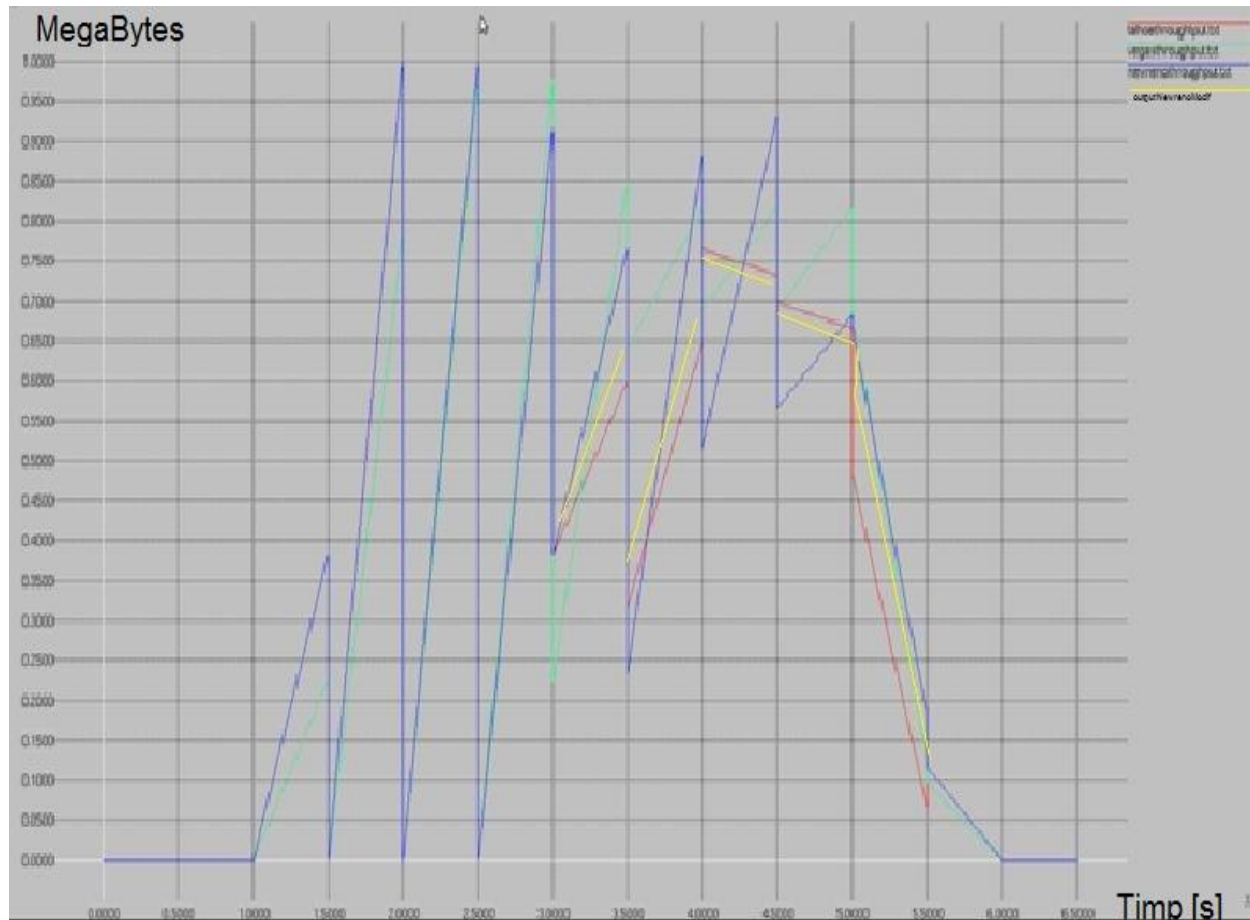


Figura 6.1. Variația ferestrei de congestie în cazul algoritmilor New Reno, Reno și Tahoe

Figura 6.1 prezintă comportamentul ferestrei de congestie (cwnd) pentru mecanismele de control al congestiei; TCP Reno, NewReno, Tahoe și NewRenoModificat . Figura arată schimbarea ferestrei de congestie cu timpul. În TCP Reno și NewReno, atunci când se constată pierderi de pachete; aceste protocoale selectează dimensiunea ferestrei la 50% din dimensiunea curentă. În TCP NewReno Modificat, de fiecare dată când pierderile de pachete sunt detectate, se modifică dimensiunea ferestrei de valori diferite în funcție de starea rețelei

Figura 6.2 Variația lățimii de bandă în cazul algoritmilor New Reno, Reno și Tahoe, modificat New Reno



Se observa că dintre cei patru algoritmi propuși, în ceea ce privește lățimea de bandă, performanțele cele mai bune le are algoritmul New Reno, urmat de varianta modificată. Reno și Tahoe oscilează foarte mult, ceea ce va duce la o transmisie destul de ineficientă, deoarece deși uneori lățimea de bandă are valori mari, aceasta nu se menține, trecând de la valori mari la scăzute, aproape de nule chiar.

Variația numărului de pachete pierdute

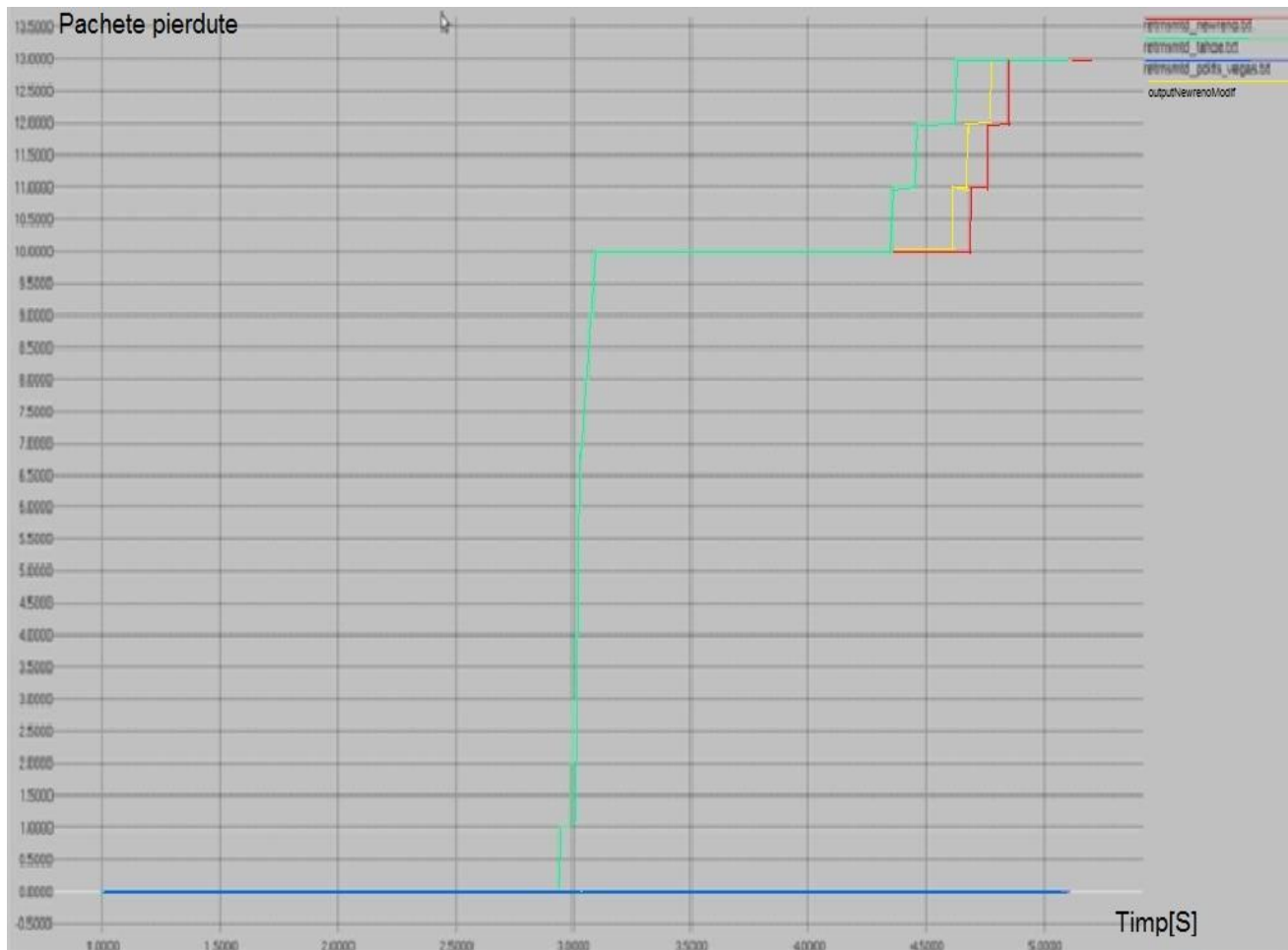


Figura 6.3 Variația numărului de pachete pierdute

Dupa cum se observa, cantitatea cea mai mica de pachete pierdute o are algoritmul New Reno urmat de New Reno modificat. Este de astepta ca algoritmul Tahoe sa fie pe ultimul loc, deoarece pierderea este detectată atunci când timpul de timeout expiră înainte ca o confirmare să fie recepționată. În acel moment, Tahoe reduce fereastra de congestie la valoarea maximă a dimensiunii unui segment (MSS – maximum segment size) și se resetează la faza de „slow-start”.

Delay între pachete

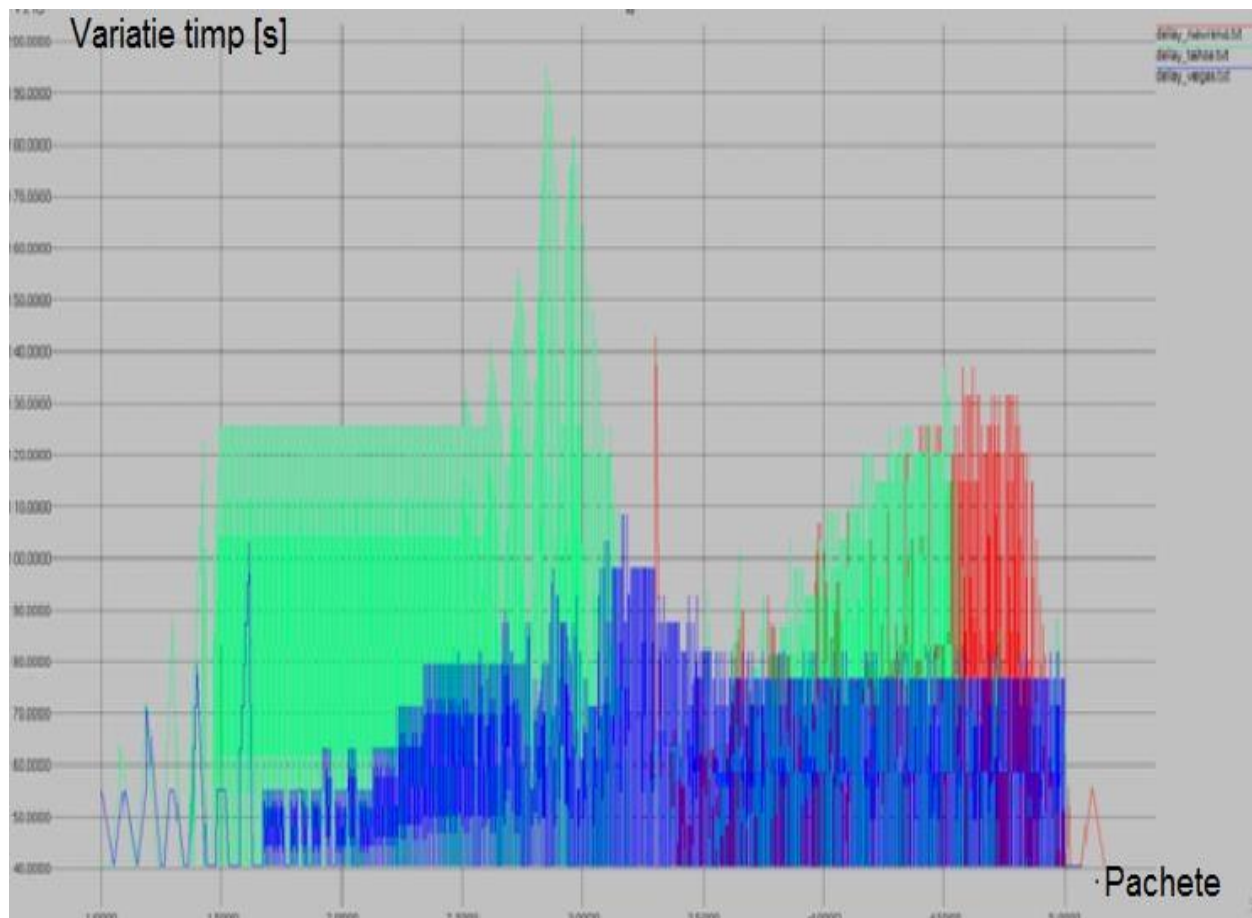


Figura 6.4 : Delay între pachete

Figura prezintă întârziere între pachete în funcție de timp pentru Reno TCP, NewReno, și NewRenoModificat. Așa cum se arată în figură, spre deosebire de NewReno care are întârzieri mai mari decât Reno, NewRenoModificat reduce întârzierea de pachete. Acest lucru se datorează faptului că, în timpul algoritmului de recuperare rapidă, TCP NewReno așteaptă să recupereze toate pachetele pierdute și trimite câteva pachete noi, dar cu NewRenoModificat modificările ferestre de congestie pe baza stării rețelei și ar putea trimite mai multe pachete noi, astfel se reduce întârzierea de pachete.

Tabelul 6.5 arată pierderile de pachete în funcție de timp pentru Reno TCP, NewReno, Tahoe și NewRenoModificat. Este clar din figură că cele patru protocoale au același comportament în timpul început lent și fazele de evitare a congestiei. Cu toate acestea, cu timpul, rata pierderilor de pachete de NewRenoModificat crește. Acest lucru se datorează faptului că, dimensiunea ferestrei congestie crește mai mult decât cea a Reno și NewReno, și mai multe pachete sunt transmise, astfel încât probabilitatea ca pierderile au crescut de asemenea. Acest lucru apare în figură, NewRenoModificat asigură pierderi mai mari decât NewReno. Acest comportament este observat, de asemenea, în NewReno TCP comparativ cu Reno.

Numarul de pachete pierdute in retea sunt:

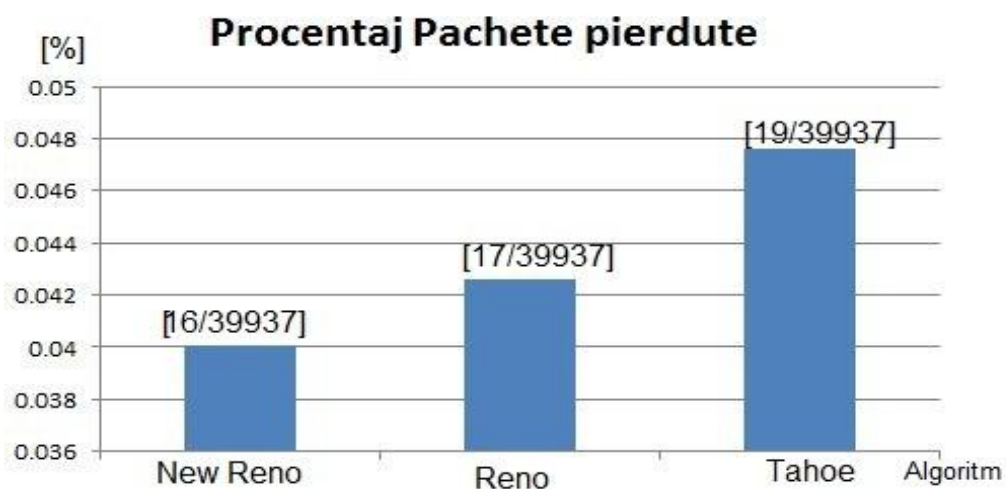


Figura 6.5 : Numărul pachetelor pierdute in procente

Tip Algoritm	New Reno	Reno	Tahoe	NewrenoModificat
Numar pachete pierdute in retea	16	17	18	18

Tabelul 6.1 Numărul pachetelor pierdute

Din totalul de 39937 de pachete trimise in retea.

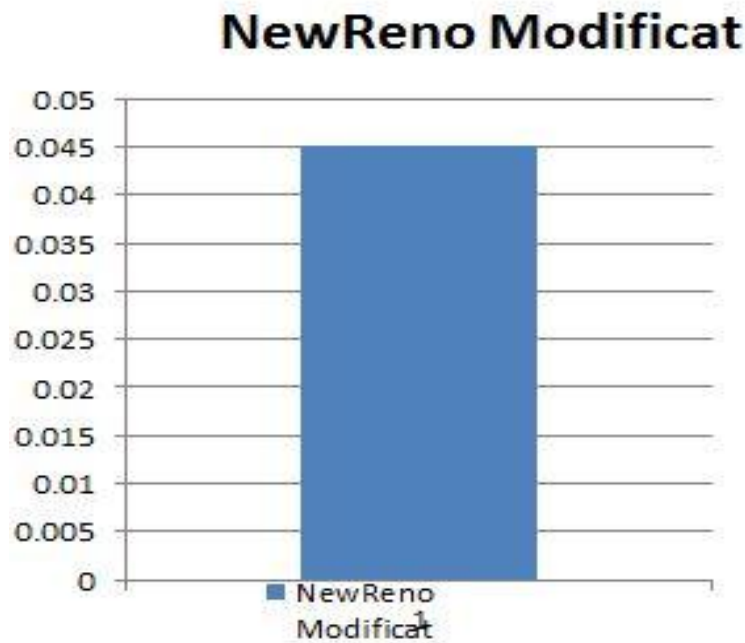


Figura 6.6 : Numărul pachetelor pierdute în procente pentru algoritmul New Reno modificat

Pierderea de pachete poate fi cauzată de o serie de factori, inclusiv degradarea semnalului peste medie a rețelei, pierderea de pachet se datorează canalului de congestie, pachetele sunt respinse în tranzit, fie din cauza unui hardware al rețelei defect, drivere de rețea defecte sau rutine de rutare corupte. În afară de aceasta, probabilitatea de pierdere a pachetelor este afectată și de raportul și distanța semnal-zgomot între emițător și receptor.

Figura 6.6 arată pierderile de pachete în funcție de timp pentru New Reno modificat. Este clar din figura 6.5 și figura 6.6 că cele patru protocoale au același comportament în timpul Slow Start și fazele de evitare a congestiei. Cu toate acestea, cu timpul, rata pierderilor de pachete de NewReno modificat este crescut. Acest lucru se datorează faptului că, dimensiunea ferestrei de congestie crește mai mult decât cea a Reno și NewReno, și mai multe pachete sunt transmise, astfel încât a crescut și probabilitatea pierderilor de asemenea. Acest lucru apare în figura 6.6, NewReno modificat cauzează pierderi mai mari decât NewReno. Acest comportament este observat, de asemenea, în NewReno TCP comparativ cu Reno.

End-to-End Delay : Întârzierea end-to-end se referă la timpul necesar pentru un pachet să fie transmise într-o rețea de la sursă la destinație.

Tip Algoritm	New Reno	Reno	Tahoe	NewrenoModificat
End to End Delay	0,279	0,07	0,15	0,23

Tabelul 6.2 Intarzierea End-to-End

Durata medie a unui pachet de date pentru a ajunge la destinație include, de asemenea, întârzierea cauzată de procesul de descoperire a traseului și coada în transmiterea de pachete de date. Numai pachetele de date livrate cu succes la destinații sunt introduse în această numărare. Se observa ca timpul cel mai mare de între sursa și destinație a unui pachet este înregistrat la algoritmul New Reno, la fel și în cazul New Reno modificat. Cel mai scurt timp de parcurgere a unui pachet în rețea este deținut de Reno.

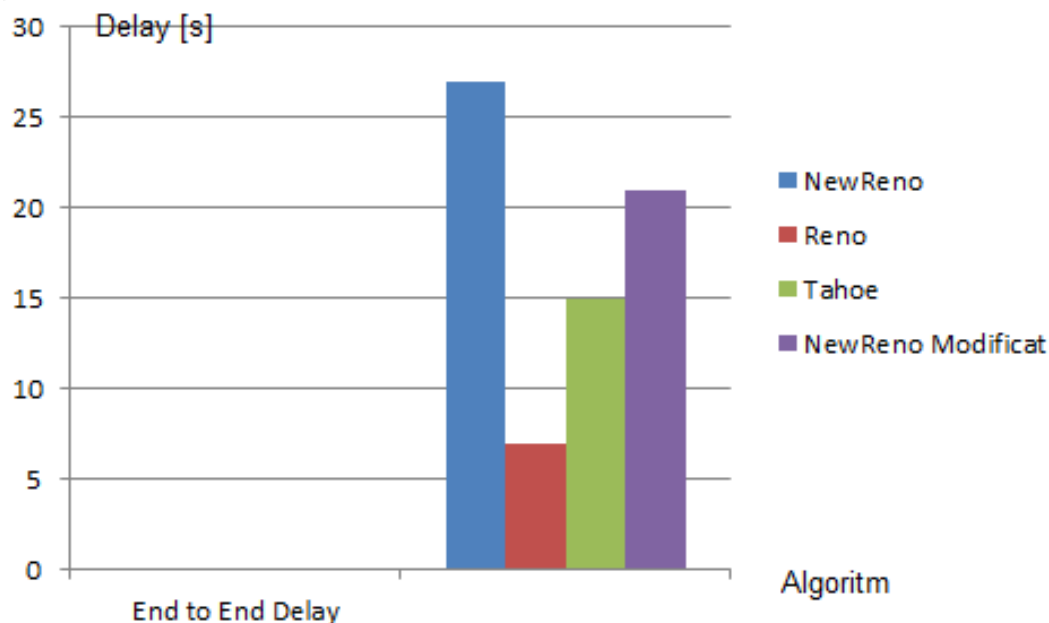


Figura 6.7 : End-to-End Delay

În cazul setării paramerului de slow start threshold în faza de fast recovery , fereastra de congestie se schimba astfel:

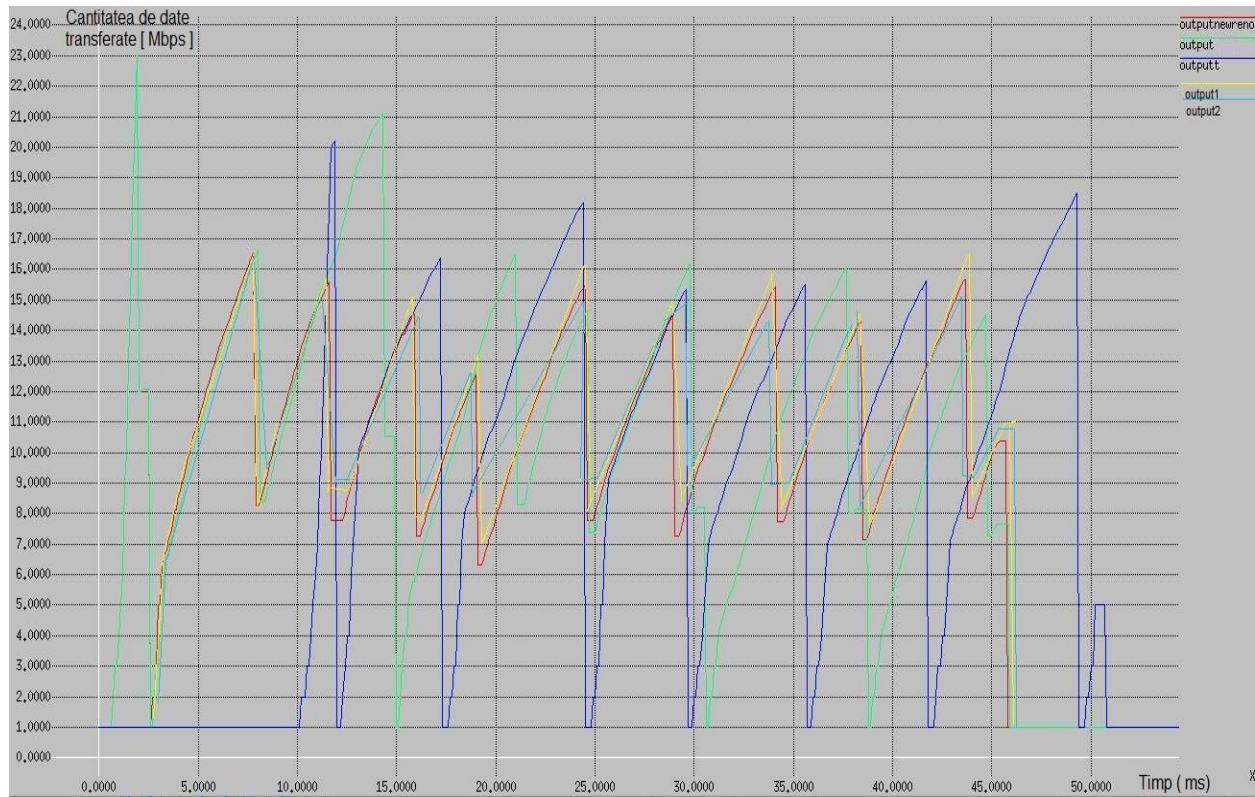
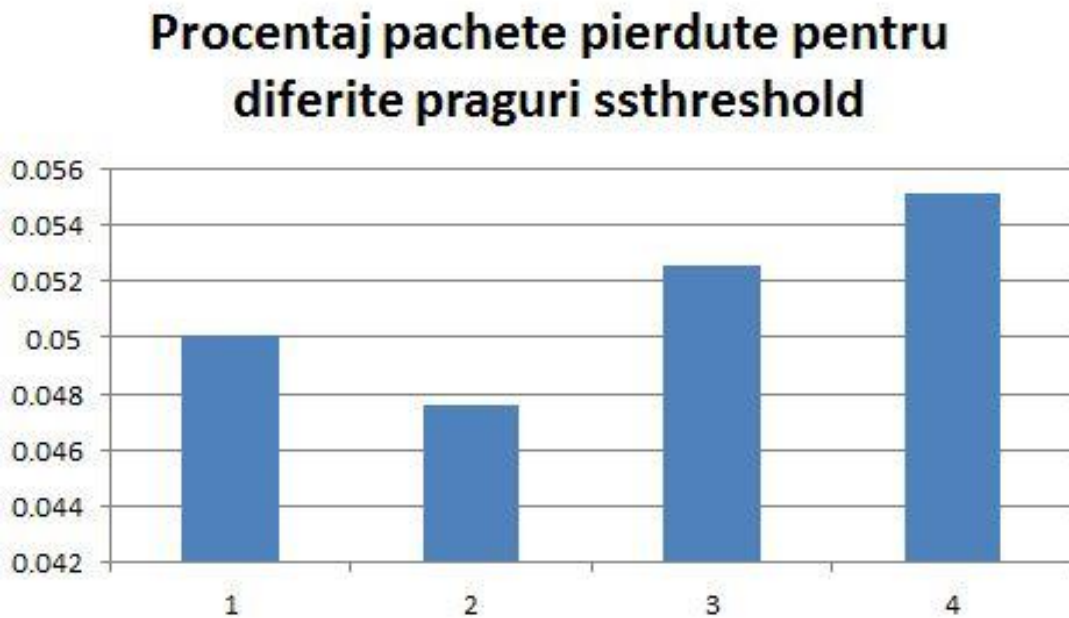


Figura 6.8 : Fereastra de congestive pentru diferite

Se observa o imbunatațire a ferestrei de congestive, deoarece aceasta nu mai prezinta oscilații atat de frecvente ca în cazul new Reno, fiind în jurul unui prag optim. Conform simulărilor realizate, am observant că pragul optim de setare a ssthresholt este egal cu patru sesimi din fereastra cwnd.

.

Dacă aleg pragul gasit ca optim, am ales să analizez și numărul de pachete pierdute



1 – $ssthreshold = (cwnd/2) + 15$

2 – $ssthreshold = (3 * cwnd)/5$

3 – $ssthreshold = (CWND * 2)/3$

4 – $ssthreshold = (CWND * 3)/4$

Pragul optim de $2/3$ din fereastra cwnd are un dezavantaj. Acela produs de faptul ca procentajul numatului de pachete pierdute in retea are o usoara crestere, fata de New Reno, dar are totuși performant mai bune decat Tahoe si Reno.

7. Concluzii

În această lucrare se propune un algoritm de recuperare rapidă modificat pentru creșterea performanțelor NewReno TCP. Mecanismul este dezvoltat prin adaptarea ferestrei de congestie a expeditorului TCP în funcție de nivelul de congestie în rețea. Acest nivel este determinat prin utilizarea Round Trip Time (RTT), care reprezintă un indicator pentru sarcinile de trafic pe rețea. Mecanismul propus este evaluat prin folosirea simulatorului NS2 și comparația cu NewReno TCP, Reno TCP și Tahoe TCP. Rezultatele simularilor arată că, în comparație cu algoritmul de recuperare rapidă modificat cu NewReno TCP, cel modificat îmbunătățește performanțele sale atât împotriva debitului și întârzierea pachetului din cauza de a transfera mai multe pachete la destinație. Deși modificările suplimentare ale algoritmului de recuperare rapidă a îmbunătățit performanța, mecanismul propus, NewRenoModificat, este inefficient în ceea ce privește pierderile de pachete, așa cum se arată în Figura . Îmbunătățiri suplimentare ar trebui să fie luate în considerare în activitatea viitoare pentru a îmbunătăți NewRenoModificat împotriva pierderilor de pachete.

New-Reno TCP este o variantă de Reno, cu o mică modificare în algoritmul de recuperare rapidă. Acest lucru a fost făcut în scopul de a rezolva problema timeout atunci când mai multe pachete sunt pierdute formă aceeași fereastră. Rețineți că performanțele mai ridicate au fost obținute ca urmare a micii modificări Reno TCP. Deși New Reno rezolvă problema timeout atunci când mai multe pachete sunt pierdute formă aceeași fereastră, se poate retransmite doar un singur pachet pe Round Trip Time.

Se observă din capitolul 6 că cele mai bune performanțe în ceea ce privește lățimea de bandă le are algoritmul New Reno modificat, urmat de New Reno. Lățimea de bandă nu oscilează ca în cazul Reno și Tahoe.

La capitolul număr de pachete pierdute, algoritmul New Reno modificat are performanțe mai mici decât algoritmul original New Reno.

Varianța modificată a algoritmului New Reno deși prezintă îmbunătățiri ale ferestrei de congestie și a lățimea de bandă, prezintă un deficit a performanțelor atunci când vine vorba de numărul de pachete pierdute și durata medie a unui pachet de a ajunge la destinație.

8 . Bibliografie și referințe

1. W. Richard Stevens – “TCP/IP Illustrated” ,Volume 1The Protocols “ Editia a doua , Addison-Wesley, 1994.
2. Aleksander Malinowski , Bogdan M. Wilamowski - “Transmission Control Protocol—TCP”
3. Xiang Chen, Hongqiang Zhai, Jianfeng Wang, and Yuguang Fang – “ TCP Performance over Networks”
4. W. Stevens - TCP Slow Start, Congestion Avoidance, Fast Retransmit, and Fast Recovery Algorithms” Network Working Group , NoAO , 1997
5. B. Braden, ed., "Requirements for Internet Hosts --nCommunication Layers," [RFC 1122](#), Oct. 1989.
6. V. Jacobson, "Modified TCP Congestion Avoidance Algorithm," April 30, 1990.
7. J. Padhye, and S. Floyd, “ On Inferring TCP Behavior”, Computer Communications Review ACM-SIGCOMM, Vol. 31, August 2001
8. M. Allman, V. Paxson, and W. Stevens, “TCP Congestion Control”, RFC 2581, IETF, April 1999
9. G. R. Wright, W. R. Stevens, "TCP/IP Illustrated, Volume 2: The Implementation", Addison-Wesley, 1995.
10. V. Jacobson, "Congestion Avoidance and Control," Computer Communication Review, vol. 18, no. 4, pp. 314-329, Aug. 1988.
11. M. Miyake, and H. Inamura, "TCP Enhancement Using Recovery of Lost Retransmissions for NewReno TCP," Transactions of Information Processing Society Journal, Vol. 46 (9), pp. 2185-2195, September 2005.
12. Flavius Copaciu , Gabriel Lazar , Virgil Dobrota “Practical Analysis of TCP Implementations:Tahoe, Reno, NewReno”
13. K. Fall, and S. Floyd, “Simulation–Based Comparison of Tahoe, Reno and SACK TCP”, Computer Communications Review ACM SIGCOM , July 1996
14. HARDIK V. MIYANI, VISHV B. KUKADIYA, MR. KAPIL S. RAVIYA, MR.DHRUMIL SHETH - “PERFORMANCE BASED COMPARISON OF TCP VARIANTS ‘TAHOE, RENO, NEWRENO, SACK’ IN NS2 USING LINUX PLATFORM”
15. Celio Albuquerque, Brett J Vickers “Preventing Congestion Collapse and promoting fairness in Internet”IEEE/ACM Transactions, Volume (12)-01, Feb 2004.
16. T.V.Lakshman, U. Madhow, and B. Suter, “Window-based error recovery and flow control with
17. a slow acknowledgment channel: A study of TCP/IP performance,” in Proc. IEEE Infocom 1997.
18. Saleem-ullah Lar, Xiaofeng Liao and Songtao Guo - “Modeling TCP NewReno Slow Start and Congestion- Avoidance using Simulation Approach” IJCSNS International Journal of Computer Science and Network Security, VOL.11 No.1, January 2011
19. Jitendra Padhye, Victor Firoiu, “Modeling TCP Reno Performance A simple model & its Empirical Validation” IEEE/ACM Transactions Volume(8)-02 April 2000.
20. Dongmin Kim, Beomjoon Kim, Jechan Han and Jaiyong Lee “Enhancement to the Fast Recovery Algorithm of TCP NewReno,ICOIN 2004, LNCS 3090, pp 332-341.”
21. Habibullah Jamal and Kiran Sultan, “Performance Analysis of TCP Congestion Control Algorithms”, International Journal of Computers and Communications, Volume (02)-01, 2008.
22. Jitendra Padhye, Victor Firoiu, Donald F. Towsley, Fellow, IEEE, and James F. Kurose, Fellow, IEEE
23. “Modeling TCP Reno Performance: A Simple Model and Its Empirical Validation” IEEE/ACM TRANSACTIONS ON NETWORKING, VOL. 8, NO. 2, APRIL 2000
24. J. Padhye, V. Firoiu, and D. Towsley, “A stochastic model of TCP Reno congestion avoidance and control,” Tech. Rep. UMASS-CS-TR-1999-02.
25. S.Floyd, “Congestion Control Principles”, ACIRI, September 2000 Robert Shorten, Fabian Wirth, Douglas Leith, “A positive systems model of TCP- like congestion control: Asymptotic results” , April 7, 2004

26. V. Jacobson, and M. J. Karels, "Congestion Avoidance and Control," Proceedings of ACM SIGCOMM, Vol.18 (4), pp. 314-329, August 1988.
27. Hanaa Torkey , Gamal Attiya , Ibrahim Z. Morsi - “ Modified Fast Recovery Algorithm for Performance Enhancement of TCP-NewReno”, International Journal of Computer Applications (0975 – 8887) Volume 40– No.12, February 2012
28. M. Niels, B. Chadi, A. Konstantin, and A. Eitan, "Inter-protocol fairness between TCP NewReno and TCP Westwood," The 3rd EuroNGI Conference on Next Generation Internet Networks, Vol.1, pp. 21-23, May 2007.
29. Jacobson, Van; Karels, MJ (1988). "[Congestion avoidance and control](#)". *ACM SIGCOMM Computer Communication Review* **18** (4): 314–329.
30. Corbet, Jonathan. "[Increasing the TCP initial congestion window](#)". LWN. Retrieved 10 October 2012.
31. Mark Allman, Vern Paxson, W. Richard Stevens (April 1999). "[Fast Retransmit/Fast Recovery](#)". *TCP Congestion Control*. *IETF*. sec. 3.2. RFC 2581. Retrieved 2010-05-01.
32. S. McCanne and S. Floyd, "ns Network Simulator",
33. <http://www.isi.edu/nsnam/ns>.
34. http://nsnam.isi.edu/nsnam/index.php/Installing_ns2.31_on_Ubuntu7.04
35. <http://www.nabble.com/Network-Simulator-ns-2-f15582.html>
36. <http://nile.wpi.edu/NS/>
37. www.acm.org
38. <http://www.isi.edu/nsnam/ns/tutorial/>
39. <http://www.opalsoft.net/qos/DS.htm>
40. www.ieee.org
41. http://www.netlab.ecnu.edu.cn/software/ns-class/tcp_8h.htm
42. http://www.unibuc.ro/prof/niculae_c_m/telecom/transm_ctrl_prot_tcp.htm
- 43.