

Folosirea de tehnici scalabile, peer-to-peer, pentru creșterea disponibilității serviciului oferit de rețelele SIP

Teză de doctorat

Voichița Almășan
voichita.iancu@cs.pub.ro

Îndrumător: **Prof. Dr. Ing. Iosif Ignat** Iosif.Ignat@cs.utcluj.ro

*Universitatea Tehnică Cluj-Napoca, Facultatea de Automatizări și Calculatoare, Catedra de
Calculatoare
Iulie 2011*

Cuvinte cheie: serviciu larg disponibil, toleranță la defecte,
scalabilitate, echilibrare a sarcinii, peer-to-peer, SIP.

Rezumat

1 Introducere

Teza de fata, intitulata "Folosirea de tehnici scalabile, peer-to-peer, pentru cresterea disponibilitatii serviciului oferit de retelele SIP", prezinta principalele probleme ale unui serviciu oferit in cadrul unui sistem distribuit, intre care am identificat ca principala problema disponibilitatea serviciului. In urma identificarii si analizarii aspectelor care au impact asupra disponibilitatii serviciilor, se propun solutii bazate pe o infrastruktura logica peer-to-peer, descentralizata, combinate pe alocuri si cu tehnici DNS. De-a lungul intregii lucrari, argumentam, inclusiv prin evaluarea performantelor contributiilor propuse, de ce consideram descentralizarea ca fiind cheia solutionarii disponibilitatii in sisteme distribuite. De asemenea, subliniem cum anume poate fi ea folosita, prin intermediul unui DHT de tip peer-to-peer, pentru a oferi un serviciu distribuit cu o larga disponibilitate. Asadar, in lucrarea de fata prezentam modul de obtinere pentru un serviciu de o larga disponibilitate prin utilizarea de tehnici de descentralizare. In acelasi timp, propunem si un mecanism de transparenta in ceea ce priveste descentralizarea, pentru clientii sistemului, prin oferirea catre acestia a unei imagini unice a sistemului (Single System Image). Tehnica generica de conferire a largii disponibilitati prin descentralizare este validata, in final, prin aplicarea si evaluarea acesteia in cazul unui server SIP.

Capitolul 2 realizeaza descrierea contributiilor deja existente in domeniului in care se plaseaza teza de fata. Printre altele, acest capitol isi propune sa defineasca, in viziunea noastra, notiunile cu care opereaza aceasta teza, asa cum sunt: (1) *toleranta la defecte*; (2) *scalabilitatea*; (3) *echilibrarea incarcarii*; (4) *securitatea sistemului distribuit*; (5) *larga disponibilitate a unui serviciu*, componenta a securitatii; (6) *elemente de consens distribuit*; (7) *elemente de teoria jocurilor*; (8) *bazele de date*; (9) *retelele peer-to-peer*, si in particular DHT-ul Chord; (10) *retelele SIP* (Session Initiation Protocol), protocolul de semnalizare utilizat in retelele VoIP, pentru care modul de functionare, componentele si rolurile acestora sunt prezentate in detaliu; si, nu in ultimul rand, (11) *stadiul actual al convergentei si compatibilitatii protocoalelor peer-to-peer si SIP*.

Urmatoarele 4 capitole prezinta in amanunt contributiile originale pe care le propunem in urma cercetarii desfasurate in cadrul studiilor doctorale. In capitolul 3 propunem o modalitate generica, originala, de distribuire si *descentralizare* a unui serviciu, spre a-l face pe acesta larg disponibil. In acest capitol explicam si integrarea in solutia finala a contributiilor prezentate detaliat in capitolele 4 si 5.

In capitolul 4 vom prezenta solutia noastra pentru obtinerea unei baze de date total descentralizate, care sa ofere o buna persistenta a datelor. Am insistat asupra acestei probleme in teza de fata, intrucat consideram ca stocarea datelor referitoare la clienti este o problema importanta, comuna tuturor sistemelor distribuite. Tot in capitolul 4, vom prezenta o solutie bazata pe votare si consens distribuit, care trateaza aspectul malitiozitatii unor noduri din cadrul unui sistem distribuit descentralizat. Problema consensului apare in orice sistem distribuit care nu beneficiaza de un server central de autentificare. Serverul central ar fi fost o entitate ce ar fi redus surplusul de disponibilitate a serviciului, anuland din beneficiile descentralizarii. Nu in ultimul rand, tot in capitolul 4, vom prezenta, la final, moduri de interogare de catre clienti a serviciilor distribuite, si in particular a bazelor de date care trebuie sa fie larg disponibile. Ne propunem ca acei clienti sa fie cat mai putin impactati de timpul de obicei mare, necesar transmiterii raspunsului. Ideea de baza a reducerii impactului de comunicare asupra clientului o reprezinta decuplarea.

În capitolul 5 vom prezenta modele de auto-extindere a rețelei peer-to-peer, auto-extindere ce are drept scop principal o mai bună toleranță la defecte. Primul asemenea model de auto-extindere a infrastructurii logice, peer-to-peer, pentru serviciul generic descentralizat pe care dorim să îl oferim, urmărește conferirea unui factor de replicare mai ridicat pentru anumite date. Datele în cauză nu sunt alese la întâmplare, ci sunt acele date în legătură cu care se decide, în mod descentralizat, ca sunt critice pentru existența sistemului. Cea de-a doua modalitate de auto-extindere prezentată are ca scop principal o mai bună echilibrare a încărcării nodurilor care alcătuiesc sistemul, având ca o consecință colaterală pozitivă o scalabilitate crescută a serviciului oferit.

În capitolul 6 prezentăm modul în care tehnica generică, prezentată în capitolul 3 se poate adapta pentru a face ca un serviciu particular, un server de prezenta SIP, să devină larg disponibil.

În capitolul 7 se face o evaluare a tuturor contribuțiilor prezentate în capitolele anterioare ale lucrării. Pornind de la acest capitol, vom putea argumenta, în capitolul 8, utilitatea contribuțiilor originale prezentate în scopul acestei teze, precum și performanțele lor în a îmbunătăți disponibilitatea serviciilor distribuite. Tot în capitolul de final 8, vom trage concluziile relative la avantajele și neajunsurile soluțiilor originale propuse. Bineînțeles, am indicat și principalele direcții de îmbunătățire și de dezvoltare ulterioară pe care am dori să le urmărim, pentru obținerea, cu impact de performanță cât mai mic, a unui serviciu care să ofere o disponibilitate cât mai mare clienților săi.

2 Necesitățile și problemele identificate în contextul disponibilității serviciilor

În acest capitol ne propunem să descriem stadiul actual al cercetării în domeniile aflate în strânsă relație cu subiectul de interes al acestei teze: obținerea unei largi disponibilități pentru un serviciu distribuit, cu aplicații pentru un server SIP. Se descrie un tip foarte răspândit de sisteme distribuite, și anume sistemele distribuite de stocare a datelor. Prezentăm acest tip de sisteme aici pentru că sunt vitale pentru funcționarea oricărui alt tip de sistem distribuit. În continuare se realizează o descriere a rețelelor peer-to-peer, un tip particular de sistem distribuit, care a fost proiectat inițial pentru stocare de date, evoluând mai apoi dincolo de această funcționalitate inițială. Motivul pentru care ne-am oprit atenția asupra lor este faptul că rețelele peer-to-peer au evoluat dincolo de scopul lor inițial tocmai datorită faptului că sunt mai flexibile și funcționează pe un model *descentralizat*. Descrierea se încheie cu o analiză a modalităților în care rețelele peer-to-peer ar putea rezolva problemele de disponibilitate enunțate în teză. Tot în capitolul de introducere dorim să ne dedicăm analizei bibliografice a protocolului SIP, care generează un tip specific de arhitectură de sistem distribuit, și a implementărilor care combină acest protocol cu conceptul de peer-to-peer. În finalul studiului bibliografic se prezintă câteva modalități și tehnici deja existente prin care unui serviciu distribuit i se poate spori disponibilitatea. Prezentăm aici elementele de la care am pornit în cercetarea noastră pentru găsirea unei metode generice și scalabile, care să sporească disponibilitatea unui serviciu, inclusiv a serviciului SIP.

În concluzie, bibliografia acestei teze se referă atât la sistemele distribuite și la proble-

mele lor, cât și la modurile în care rețelele peer-to-peer pot fi utilizate pentru a diminua sau chiar elimina dintre aceste probleme. S-a făcut, în contextul studiului bibliografic, o exemplificare concretă pentru protocolul SIP, a problemelor care apar în sistemele distribuite în general. De asemenea, am analizat câteva metode existente de a face larg disponibil un serviciu distribuit. Am observat, însă, că aceste metode nu sunt suficient de *generice*, de *flexibile* sau de *scalabile*, oferind, astfel, un motiv și o direcție pentru teza de față. Pornind de aici, în capitolele ce urmează vom defini elementele componente, precum și modul lor de asamblare, pentru o platformă generică de distribuire și descentralizare a unui serviciu, în scopul obținerii unei mai largi disponibilități pentru acesta. Ne bazăm în acest demers pe neajunsurile metodelor identificate și analizate în capitolul de bibliografie a tezei, dar și pe idei din același capitol, utilizate în conjuncție cu avantajele rețelelor peer-to-peer. Rezultatul final al muncii de cercetare relative la un serviciu larg disponibil generic va fi validat aplicându-l unui server SIP, particular.

3 Creșterea flexibilității, scalabilității și echilibrarea încărcării unui serviciu, prin descentralizare transparentă

În capitolul de față ne propunem să descriem o soluție generică, aplicabilă oricărui serviciu de rețea, care prin *distribuirea* și *descentralizarea* elementelor componente, oferă o *largă disponibilitate* ("high availability") clienților săi. În același timp, dorim să menținem imaginea de sistem sau *serviciu unitar* (SSI – Single System Image) față de clienții serviciului. De aceea, ne propunem ca descentralizarea să se facă astfel încât ea să fie *transparentă* pentru clienți, nepunându-le acestora probleme. De asemenea, am dori ca, pentru a îmbunătăți un serviciu generic spre a fi unul larg disponibil, modificările legate de interfața serviciului să fie minime. În acest scop, descriem modul de interacțiune între noile componente, *descentralizate*, ale serviciului distribuit devenit de-acum larg disponibil. În fond, capitolul de față prezintă modul original de asamblare, în vederea obținerii unui serviciu larg disponibil, a unor contribuții proprii, ce vor fi descrise în capitolele 4 și 5, și evaluate în capitolul 7. Prezentăm și câteva concluzii preliminare în ceea ce privește modelul propus pentru distribuirea și descentralizarea serviciului. Subliniem din nou faptul că ne-am propus această descentralizare în scopul unei mai prompte și totodată corecte deserviri a clienților serviciului.

Modelul de repartizare între noduri a muncii acestora. Vom propune soluția noastră în ceea ce privește pastrarea față de lumea exterioară a imaginii de serviciu unitar și totodată, în mod transparent, în spatele acestei imagini, să distribuim și descentralizăm cât mai mult arhitectura unui serviciu centralizat generic. Am încercat, în soluția de față, ca surplusul de comunicare introdus din cauza distribuirii să fie cât mai mic posibil. Elementele folosite de noi în acest scop au fost: (1) un element de *distribuire a încărcării*, ce va fi numit în continuare "*load balancer*", sau "*dispatcher*", care va încerca să distribuie în mod egal și transparent încărcarea între nodurile pe care dorim să distribuim serviciul; și (2) o *rețea peer-to-peer* bazată pe Chord, îmbunătățită, care va fi prezentată succint în acest capitol, urmând că avantajele sale să fie detaliate tot în această teză. Componenta originală numită *load balancer* va controla accesul la rețeaua Chord, așa cum rezultă din figura 1. Modelul general original de distribuire a unui serviciu, prin descentralizare, în vederea creșterii disponibilității sale, este

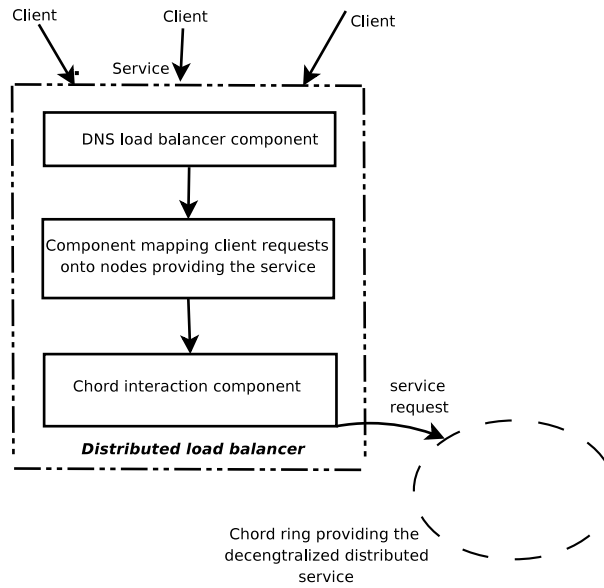


Figura 1: Arhitectura individuală a unui nod de echilibrare a încărcării.

prezentat schematic în figura 2, fiind descris pe larg în teza. Schema de față ar trebui să fie suficient de generică pentru a face larg disponibil orice serviciu, cu modificări minime.

Asadar, în acest capitol se prezintă, pe scurt, viziunea noastră originală asupra elementelor necesare descentralizării unui serviciu, precum și modul original ales pentru combinarea lor, toate în scopul obținerii unui serviciu larg disponibil. Ne-am bazat viziunea pe avantajele de flexibilitate și dinamicitate caracteristice rețelelor peer-to-peer, pe care le-am optimizat în maniera proprie. Am menționat, de asemenea, poziția rețelei peer-to-peer în cadrul arhitecturii îmbunătățite de serviciu distribuit, pentru a-i îmbunătăți disponibilitatea față de clienții săi. În capitolele care urmează, și anume capitolele 4 și 5, vom detalia modul de funcționare al acestor elemente originale, gândite ca niște componente independente. Aceste componente sunt, deci, reutilizabile și în alte contexte. Am descris, de altfel, în capi-

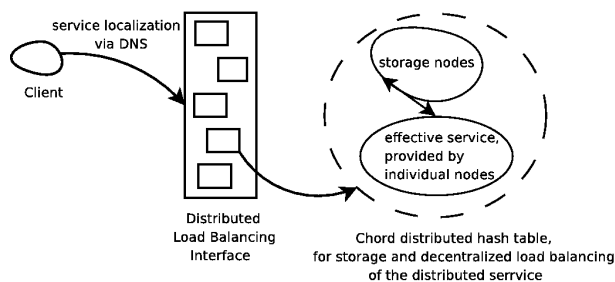


Figura 2: Metoda generală de distribuție și descentralizare a unui serviciu, pentru a-l face larg disponibil.

tolul 6, modul concret de imbinare si utilizare a acestor componente, conform sablonului prezentat in acest capitol, pentru a obtine un serviciu SIP larg disponibil.

4 Baza de date descentralizată pentru stocarea sigură și persistentă a datelor unui serviciu larg disponibil

In capitolul 3 am prezentat descrierea generica a modelului de descentralizare a unui serviciu, in scopul obtinerii unei disponibilitati cat mai largi pentru acesta. Cu aceasta ocazie am subliniat faptul ca, pentru a asigura disponibilitatea serviciului in orice moment, trebuie ca *in orice moment* sa avem la dispozitie datele aferente fiecarui client al aceluia serviciu. De aceea, este esentiala o *baza de date* care sa ofere o disponibilitate cat mai larga de acces la datele clientilor. Pentru aceasta, ne-am propus sa implementam o *baza de date distribuita*, tot *descentralizata*, care consideram ca reprezinta o *conditie necesara* pentru larga disponibilitate a oricarui serviciu.

In capitolul de fata prezentam solutia originala, propusa in cadrul acestei teze, pentru a obtine o baza de date distribuita descentralizata si de incredere, ce se bazeaza tot pe o retea Chord. Astfel, in sectiunea 4.1 a acestui capitol, prezentam arhitectura si modelul de interactiune propuse de noi, pentru a avea o baza de date cu un grad mare de disponibilitate. Baza de date propusa asigura o buna persistenta a datelor aferente clientilor serviciului distribuit, fara de care acesta nu ar putea fi larg disponibil. Pe urma prezentam o imbunatatire, elaborata tot in cadrul tezei de fata, pe care o propunem pentru baza de date larg disponibila. Imbunatatirea consta intr-o crestere a gradului de incredere in datele stocate, pe baza unor tehnici de consens in sisteme distribuite, adaptate pentru stocarea datelor bazata pe Chord. Nu in ultimul rand, prezentam o noua contributie originala, si anume viziunea noastra asupra modului ideal de proiectare a unui client pentru un serviciu larg disponibil. In particular, ne referim la clienti pentru o baza de date, ale caror cereri trebuie sa fie tratate in mod cat mai eficient si mai rapid de catre server-ul de baze de date. In general fie spus, pentru a optimiza timpul de raspuns perceput de client din partea server-ului, propunem decuplarea logica a elementului de interogare de restul clientului. O astfel de implementare de client a fost utilizata intern, in cadrul serviciului distribuit larg disponibil, pentru a optimiza comunicarea dintre nivelul care ofera propriu-zis serviciul si nivelul de stocare a datelor. Rolul si pozitia acestor 2 nivele au fost descrise in capitolul 3.

4.1 Construirea unei baze de date descentralizate, folosind o rețea Chord

Persistenta datelor, furnizata cel mai bine de catre o *baza de date distribuita*, este necesara in cazul oricarei aplicatii distribuite, si deci cu atat mai mult in cazul unui serviciu distribuit, in scopul de a nu pierde date referitoare la clientii acestuia. In plus, necesitatea unei baze de date distribuite si *descentralizate* se impune pentru a asigura cu succes stocarea, necesara aplicatiilor care proceseaza date. Din pacate, in prezent nu exista foarte multe solutii care sa garanteze o persistenta mare a datelor, intr-un mod transparent pentru aplicatie. Prin combinarea aspectelor de arhitectura pe care le-am descris si prin prezentarea de rezultate experimentale in capitolul 7, munca noastra in acest sens urmareste sa dovedeasca faptul

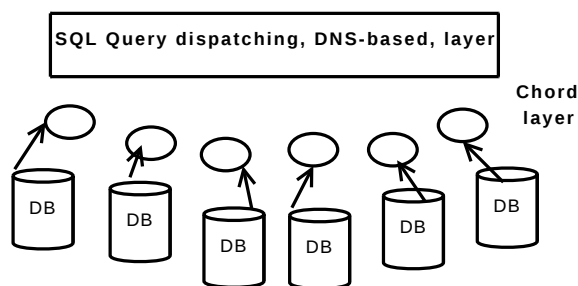


Figura 3: Arhitectura pe nivele a bazei de date distribuite.

ca DHT-ul Chord reprezintă o soluție viabilă pentru problema stocării datelor persistente, motiv pentru care considerăm că munca noastră are un impact în domeniul stocării persistente, pe care ne propunem să îl exploatăm și extindem și prin dezvoltări ulterioare. Modelul propus de noi poate fi văzut, reprezentat schematic, în figura 3.

Există mai multe *avantaje* ale soluției originale propuse de noi în teza, dintre care am dori să subliniem: (1) faptul că soluția propusă introduce un timp suplimentar de execuție relativ mic, în comparație cu bazele de date tradiționale, centralizate; (2) faptul că, în cazul în care avem un număr nu foarte mare de peer-i de stocare, conectați printr-o rețea de mare viteză, performanțele prototipului dezvoltat nu sunt influențate nici de către dimensiunea rețelei, nici de către volatilitatea nodurilor; (3) faptul că prin soluția dezvoltată de noi se ascunde locația exactă a datelor și a copiilor lor față de orice entitate exterioară bazei de date descentralizate, prin intermediul entității de distribuire a interogărilor SQL, pe care am definit-o în cadrul arhitecturii noastre; și (4) faptul că soluția noastră garantează un model de tranzacții atomice pentru toate interogările SQL suportate.

În ce privește performanțele soluției propuse, pe viitor am dori să *îmbunătățim* această bază de date distribuită și descentralizată, prin: (1) eliminarea necesității implicării câmpului de cheie primară în orice interogare; (2) introducerea posibilității de a actualiza datele deja existente, folosind interogări de tip UPDATE, și adaptarea protocolului de coerență în consecință; (3) investigarea altor soluții *ne-stactice* de echilibrare a sarcinii, pentru perfecționarea entității de distribuire a interogărilor SQL. Ne dorim, de asemenea, să testăm și performanțele prototipului propus prin compararea rezultatelor obținute în anumite condiții de stres cu cele obținute de alte baze de date distribuite, care asigură și ele o anumită persistență, așa cum este MySQL cluster, deoarece până în prezent am comparat performanțele soluției noastre doar cu versiuni nedistribuite ale bazelor de date.

4.2 Îmbunătățirea modalităților de stocare și a securității informației, pe baza teoriei jocurilor

Așa cum am afirmat deja, serviciile distribuite trebuie să își mențină niste date referitoare la fiecare client, dintre care unii clienți ar putea fi importanți, implicit datele lor fiind critice. Din acest motiv, ne dorim să îmbunătățim gradul de încredere al stocării, prin proiectarea unui mecanism de stocare sigură, care să mențină datele nealterate chiar și în prezenta unor

noduri malitioase, sau care nu functioneaza conform specificatiilor.

Propunem un algoritm original de vot distribuit, in spiritul consensului bizantin, algoritmul in cadrul caruia se schimba foarte putine mesaje suplimentare intre noduri, performantele sale fiind evaluate in cadrul capitolului 7. Acest algoritm propus reprezinta solutia noastra la problema mentinerii *gradului de incredere* si a functionalitatii corecte a unei baze de date distribuite si descentralizate, in ciuda unor posibile atacuri de securitate. Algoritmul nostru poate identifica si poate exclude din randul nodurilor de stocare orice nod care nu functioneaza conform cerintelor de functionare ale acestui sistem distribuit. Deoarece algoritmul se afla in legatura doar cu nivelul Chord, ceea ce inseamna ca nu are nevoie sa cunoasca alte detalii despre baza de date distribuita din care face parte si DHT-ul Chord, solutia propusa de noi ar putea fi usor utilizata pentru a oferi un grad sporit de incredere pentru orice sistem distribuit descentralizat, care in mod normal ar fi avut nevoie de o autoritate centrala de autentificare.

Pe viitor, am dori sa investigam si mai in detaliu aspecte legate de durata necesara pentru ca topologia distribuita sa converga la o topologie ce nu contine nici un nod malicios si, probabil, sa gasim chiar o formula matematica pentru a determina gradul de incredere, care sa sporeasca ritmul de convergenta al topologiei. Ne-am dori, de asemenea, sa adresam mai in detaliu si aspectul persistentei datelor care apartineau nodurilor malitioase, inainte ca acestea sa devina malitioase in urma vreunor atacuri, deoarece aceste noduri ar fi putut fi responsabile de date ce sunt critice pentru functionarea sistemului, si pentru care, in consecinta, ar fi de dorit sa luam masuri suplimentare ca sa nu le pierdem.

4.3 Propunere de tratare decuplată, la nivelul clientului, a interogărilor adresate unei baze de date

Chiar daca un serviciu distribuit beneficiaza de stocarea datelor clientilor sai intr-o baza de date rapida, totusi trebuie ca serviciul distribuit sa isi ia masurile necesare pentru cazul in care baza de date ar fi suprasolicitata de interogari ale clientilor pe care ii deservește, crescand astfel timpul de raspuns pentru serviciul distribuit si descentralizat, si, in consecinta performantele acestuia.

Punctul slab in marea parte a sistemelor distribuite este faptul ca nodurile care le compun sunt conectate prin intermediul unor legaturi care nu sunt sigure. Peste aceste legaturi ne-am dori sa putem comunica in mod sigur, deoarece schimbarea de mesaje si solicitarea unor anumite servicii reprezinta ideea centrala pentru functionarea oricarui tip de sistem distribuit. Din aceste considerente, putem afirma cu destula certitudine ca legaturile de retea dintre nodurile unui sistem distribuit reprezinta componentele cele mai sensibile ale acestora, ele avand si probabilitatea cea mai mare de esec. Un tip de esec impotriva caruia este greu de luptat il reprezinta *lipsa disponibilitatii*, care ar fi putut fi generata fie de catre un atac de refuz al acordarii serviciului (*Denial of Service – DoS*), fie din cauza unei proiectari necorespunzatoare a intregului sistem. Acest tip de esec s-ar putea datora fie unei gatuiri pe parcursul unei conexiuni de retea, fie unei supraincercari a unui nod din partea caruia solicitam un serviciu. Cum gaturile de retea pot fi impiedicate prin metode de imbunatatire a calitatii serviciilor (*Quality of Service – QoS*), supraincercarea nodurilor ramane problema cu care ne preocupam aici, incercand sa o rezolvam utilizand unele metode de decuplare a

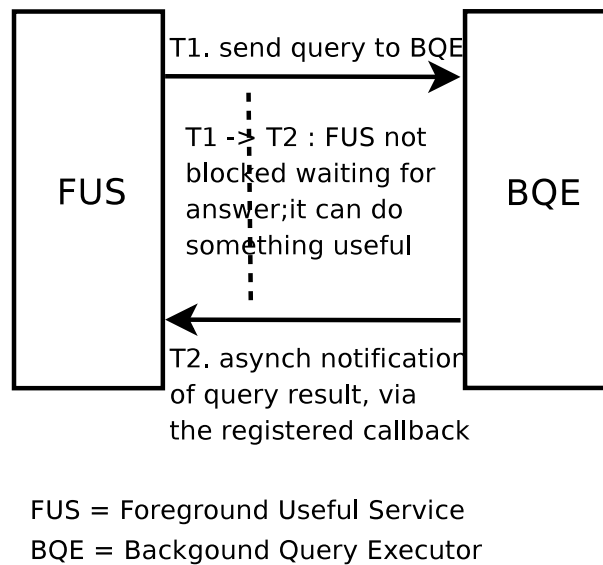


Figura 4: Descrierea schematica a mecanismului de decuplare.

entitatilor si, nu in ultima instanta, de echilibrare a sarcinii.

In teza descriem un nivel de decuplare din cadrul server-elor, care le separa pe acestea in serviciul util si in entitatea secundara de executie a interogarilor si care utilizeaza un mecanism de tip apel de procedura la distanta (*Remote Procedure Call* – *RPC*). Schematic, aceasta decuplare poate fi vazuta in figura 4. Cele doua nivele descrise interactioneaza utilizand o interfata simpla, ce ofera functionalitati similare cu cele oferite de catre DMA (*Direct Memory Access*) la nivelul hardware. Am evaluat performantele prototipului obtinut de noi pentru decuplarea executiei interogarilor in capitolul 7, si am estimat si timpul suplimentar introdus prin aceasta decuplare. Se poate observa faptul ca nivelul de decuplare are performante bune, timpul suplimentar introdus fiind nesemnificativ in comparatie cu durata executiei efective a interogarii. De asemenea, se poate observa faptul ca arhitectura propusa de noi minimizeaza si riscul anumitor atacuri de tip DoS.

In completarea solutiilor preexistente, care se concentreaza mai ales pe optimizari ce se pot face la nivelul executiei interogarilor pe server-ele vizate si pe optimizari la nivelul legaturilor de retea, solutia noastra a incercat sa adreseze cat mai generic problema decuplarii executiei interogarilor de catre un serviciu distribuit. Testarea s-a realizat pe un sistem distribuit VoIP, si anume un server SIP. Motivul utilizarii unui server SIP a fost doar faptul ca aceste tipuri de server-e au nevoie sa efectueze foarte des multe interogari, atat de baze de date, cat si ale sistemului DNS.

Pentru a imbunatati rezultatele obtinute de aceasta versiune a prototipului nostru, in urmatoarea sa versiune dorim sa optimizam anumite aspecte, cum ar fi modul efectiv de solicitare a executiei interogarilor, prin mai multa interactiune cu arhitectura de intrare-iesire de nivel jos, cand se transmite interogarea. Dorim sa realizam, de asemenea, mai multe masuratori ale prototipului nostru pe sisteme multiprocesor si multi-core, in medii de executie cat mai apropiate de cele reale. In plus, deoarece dorim sa adresam o gama multipla de inte-

rogari, ne propunem ca pe viitor să definim și un limbaj generic de formulare a interogărilor și a răspunsurilor la acestea.

5 Exploatarea dinamicității rețelelor peer-to-peer, pentru o mai bună stocare a datelor

În acest capitol ne propunem să exploatăm în detaliu dinamicitatea intrinsecă a rețelelor peer-to-peer de tip DHT, propunând soluții originale, ca extensii pentru DHT-ul Chord. Reamintim faptul că acest DHT este utilizat pentru a implementa o bază de date descentralizată, pentru a fi *larg disponibilă*. Baza de date în cauză este prezentată în amanunț în capitolul 4. Ceea ce aduce nou acest capitol sunt completări originale ale protocolului Chord, cu scopul de: (1) a spori siguranța de stocare prin sporirea *tolerantei la defecte* a sistemului, asigurându-ne astfel că nu se pierd date, așa cum vom vedea în secțiunea 5.1, precum și (2) cu scopul de a *echilibra*, în mod transparent și *scalabil*, încărcarea individuală a nodurilor din rețeaua de stocare, așa cum se prezintă în secțiunea 5.2. Ambele aspecte prezentate în acest capitol se bazează pe capacitatea de a se auto-extinde nelimitat și în mod transparent, pe care o au DHT-urile. Aceste sisteme distribuite reușesc să înglobeze orice nod care este “dispus” să îndeplinească rolul de peer.

Asadar, în capitolul de față se prezintă contribuții originale reprezentând extensii ale Chord, bazate pe capacitatea sa de auto-extindere. Ambele contribuții sporesc toleranța la defecte precum și echilibrarea încărcării, avantaje aplicabile și nodurilor din cadrul bazei de date prezentate în capitolul 4. Astfel, acest capitol adresează probleme de disponibilitate prin: (1) toleranța la defecte, (2) echilibrarea încărcării și (3) scalabilitate ale bazei de date aferente serviciului distribuit. Combinat cu capitolul 4, capitolul de față prezintă modificările care au trebuit aduse unui DHT Chord clasic, spre a-l adapta cât mai bine la funcția de stocare cu grad de disponibilitate și de încredere cât mai mare. În final, toate acestea ajută ca modelul de descentralizare a unui serviciu distribuit, prezentat în capitolul 3, să transforme acel serviciu într-unul larg disponibil, fapt ce reprezintă obiectivul central al acestei teze.

La fel ca în cazul elementelor prezentate în capitolul 4, fiecare contribuție originală prezentată în acest capitol poate fi utilizată împreună cu celelalte elemente nou-introduse de această teză. În particular, aceste elemente pot fi utilizate în conjuncție cu baza de date descentralizată, prezentată în capitolul 4. De asemenea, ele pot fi utilizate și individual, în alte tipuri de sisteme decât cele pentru care ele au fost inițial proiectate, în cadrul acestei teze. Din acest motiv, elementele în cauză au fost descrise independent unul de celălalt, în ciuda permanentei interacțiuni care are loc între ele, conform modelului descris în capitolul 3.

5.1 Îmbunătățirea siguranței de stocare a unei rețele Chord, prin auto-extindere

În ceea ce privește stocarea datelor referitoare la clienții mai importanți, la care ne-am referit deja ca fiind “date critice” sau “date sensibile”, discutăm în amanunț în teza că pentru acestea ar fi necesar un grad de replicare mai mare, pentru a spori certitudinea că aceste date nu se vor pierde în urma eșecurilor individuale ale nodurilor.

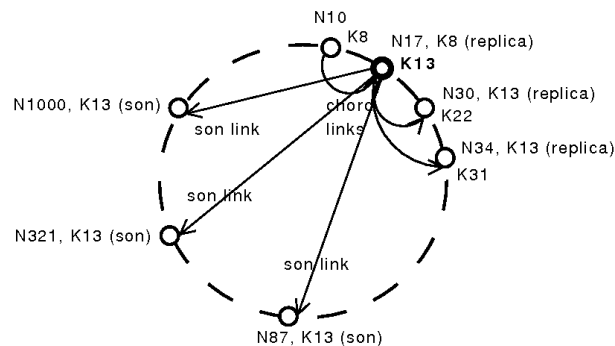


Figura 5: Stocarea utilizand un DHT de tip peer-to-peer, imbunatatit.

Problema stocării cu un grad de redundanță diferit, în funcție de datelor stocate, a fost adresată și în contextul aplicațiilor peer-to-peer care rulează peste grid. În acest context, ideea principală din spatele soluțiilor elaborate a fost ca aplicația de nivel înalt să solicite un grad de redundanță, care la rândul său se traduce, la nivelul overlay-ului peer-to-peer utilizat, printr-un număr de noduri fizice de stocare diferite. Dacă serviciul de partaj de date pe grid își da seama că nu are destui peer-i instanțiați pentru a asigura gradul de redundanță cerut, acesta va realiza integrarea unor noi peer-i în overlay-ul peer-to-peer, interogând și interacționând cu sistemul de gestiune a resurselor grid-ului. Rezervarea nodurilor suplimentare solicitate, necesare asigurării gradului dorit de redundanță, se face în maniera best-effort, deoarece nimeni nu ne poate asigura faptul că la un moment dat *grid-ul însuși* are disponibile toate resursele dorite. În urma rezervării nodurilor, serviciul peer-to-peer de stocare va fi extins astfel încât să le cuprindă și pe acestea, urmând ca mai apoi datele să poată fi stocate și pe noile noduri, în conformitate cu gradul de redundanță solicitat.

Avantajul unei soluții ca cea prezentată succint mai sus, utilizată pentru auto-extinderea la cerere a unui serviciu de partaj de date pe grid, îl constituie faptul că nu trebuie rezervate de la început toate nodurile de care serviciul distribuit ar putea avea nevoie din partea grid-ului, ceea ce duce la un consum de energie mai mic din partea aplicației luate în ansamblu, dar și la o mai eficientă partajare a resurselor grid-ului. Cele 2 dezavantaje pe care le vedem în ceea ce privește arhitectura distribuită de autoextindere la cerere sunt: (1) faptul că această soluție nu se poate adapta ușor spre a fi folosită pentru un set oarecare de noduri, care nu aparțin neapărat unui grid; și (2) faptul că această soluție nu este intim conectată cu dinamica topologiei peer-to-peer, lucru care ar putea fi de dorit, întrucât ar fi avantajos să putem reconfigura mai rapid și mai simplu întreaga topologie distribuită, în condițiile volatilității nodurilor. Având în vedere aceste dezavantaje identificate, și ținând cont de avantajele menționate deja pentru o asemenea soluție de auto-extindere la cerere a unui serviciu distribuit, propunem un model original, proiectat și implementat de către noi, care poate fi utilizat în scopul auto-extinderii unei infrastructuri de stocare distribuite și descentralizate, ce se bazează pe paradigma peer-to-peer. Mai mult chiar, nivelul de replicare pentru acest model a fost proiectat în special pentru a se afla într-o cat mai strânsă legătură cu DHT-ul utilizat, pentru ca soluția noastră să poată folosi ușor mecanismele intrinseci de stabilizare ale infras-

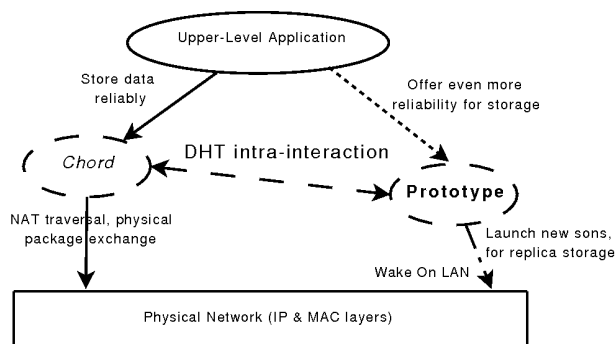


Figura 6: Interacțiuni între diversele entități implicate în stocarea persistentă a datelor.

structurii DHT, pentru cazurile de noduri volatile. Soluția noastră beneficiază, deci, atât de avantajele de “regenerare de la sine” ale infrastructurii Chord, cât și de mecanismul special proiectat de noi, care este capabil să extindă serviciul de stocare distribuit și descentralizat și pe alte noi noduri, așa cum se poate vedea schematic în figura 5. Rezultă un prototip care are o interacțiune strânsă cu DHT-ul, așa cum reiese din figura 6, în scopul de a reuși să stocheze într-o manieră mai eficientă date sensibile, cărora nu le este permis să dispară din cadrul infrastructurii de stocare coordonate de către DHT.

Considerăm ca principală contribuție obținută prin munca originală descrisă în detaliu în teza, care este mai apoi evaluată în capitolul 7, este faptul de a fi găsit o manieră de a rezolva problema pierderii datelor critice, prin stocarea acestora pe nodurile nou-integrate în cadrul infrastructurii peer-to-peer. Bineînțeles, infrastructura astfel obținută nu este nici ea 100% sigură, dar cu toate acestea ea oferă un grad de siguranță mult mai mare decât un DHT clasic. Surplusul de siguranță este obținut exploatând un alt aspect important conținut de munca noastră, și anume utilizarea unui grad de replicare variabil, în funcție de nevoile fiecărei date individuale, ceea ce permite sistemului să integreze noi noduri de stocare doar când acest lucru este necesar, fiind dictat de către dinamica internă a aplicației de nivel înalt cu care sistemul distribuit și descentralizat de stocare interacționează. Dealtfel, unul dintre rationamentele creării prototipului nostru a fost să dovedească, prin intermediul măsurătorilor pe care le-am efectuat asupra sa în capitolul 7, costurile implicate în raport cu performanțele oferite de o asemenea infrastructură generică de auto-extindere, care implică integrarea de noi noduri în infrastructura logică de tip peer-to-peer doar în momentul în care acestea sunt absolut necesare. Am văzut faptul că timpul suplimentar introdus de către această soluție, în cadrul unui set de noduri conectate printr-o rețea locală, este destul de mic, mai ales dacă luăm în considerare și beneficiile oferite de mecanismul de replicare îmbunătățit, ceea ce ne permite să sperăm că acest prototip reprezintă un bun candidat pentru a îndeplini necesitățile de stocare ale datelor sensibile din cadrul unui sistem descentralizat cu o largă distribuție geografică.

Pe viitor, ne-am dori să ne concentrăm atenția asupra adaptării prototipului nostru pentru a rula pe o infrastructură logică având o mai largă distribuție geografică, pentru a putea măsura performanțele sale și într-un asemenea context. De asemenea, ne-am dori să reali-

zăm niste experimente mai detaliate referitoare la influența pe care o are dimensiunea infrastructurii logice de tip peer-to-peer asupra probabilității de a pierde datele stocate în mod redundant. Înainte de a realiza asemenea experimente, am putea face și unele optimizări, pentru a maximiza avantajele utilizării modelului nostru de auto-extindere, cum ar fi: (1) investigarea unor posibilități de a reduce numărul de mesaje schimbate de către peer-i, cum ar fi cazul mesajelor schimbate între peer-ul fiu și părintele, sau cazul mesajelor schimbate între posesorul unei chei și succesorii acestuia, care păstrează și ei copii ale datelor aferente cheii; și (2) paralelizarea cererilor de lansare a noi noduri, ce vor fi integrate ca peer-i membri ai infrastructurii Chord, în cazul cererilor ce se referă la chei diferite, aflate în gestiunea nodului curent.

5.2 Auto-extinderea DHT-ului Chord, în scopul echilibrării încărcării

O problema care apare în cadrul oricărui sistem distribuit, și implicit în cadrul serviciilor distribuite și descentralizate, este aceea a găsirii unei modalități eficiente de utilizare a tuturor nodurilor din sistem, astfel încât să se evite supraincarea. În particular, ne preocupă rezolvarea problemei supra-încărcării și suprasolicitării unui nod, din ale cărui îndatoriri am dori să pasăm și altor noduri, în cazul în care nodul în discuție depășește un anumit prag de încărcare. În secțiunea 5.1 descriem un mecanism de *auto-extindere*, utilizat pentru a spori siguranța de stocare în special pentru anumite date sensibile, prin creșterea gradului de redundanță al acestora. Cunoscând proprietățile de dispersie ale DHT-ului, descrise în studiul bibliografic al acestei teze, vom descrie în amănunt modul în care auto-extinderea unei infrastructuri peer-to-peer ar putea fi benefică pentru echilibrarea încărcării nodurilor care o compun, fapt ce a fost și evaluat în capitolul 7.

Datorită faptului că *funcția de dispersie* este cea care este utilizată pentru a lua decizia referitoare la apartenența unor date la un peer, spre a fi responsabil de stocarea lor, putem privi această funcție ca reprezentând ea însăși un *load balancer distribuit*, poziționat în fața rețelei Chord, jucând un rol similar cu cel al înregistrărilor de tip DNS NAPTR și SRV poziționate în fața unui serviciu distribuit de rețea, așa cum am văzut în capitolul 3. Ne-am dorit, deci, să exploatăm beneficiile oferite de funcția de dispersie, spre a rezolva problema echilibrării încărcării nodurilor.

Astfel, am elaborat și mai apoi evaluat un model de echilibrare a sarcinii, prin îmbunătățirea unui DHT Chord, permitând acestui DHT să se adapteze mai bine în condițiile necesității stocării unor cantități crescânde de date. Am realizat această îmbunătățire pentru Chord prin adăugarea unei capacități suplimentare de stocare, în momentul în care ea a fost necesară, așa cum reiese din figura 7. Această capacitate de stocare este materializată fie de: (1) noi peer-i integrați în DHT prin intermediul mecanismului de auto-extindere; fie de (2) spațiu fizic de stocare, cum ar fi adăugarea unui hard disk suplimentar sau a unui hard disk mai mare pe nodul suprasolicitat. Ambele aceste metode de echilibrare fac ca DHT-ul să accepte mai ușor și mai repede să stocheze datele clienților săi, din punctul acestora de vedere.

Principala contribuție pe care o aduce soluția prezentată succint mai sus o reprezintă faptul de a *îmbunătăți* capacitatea naturală a infrastructurii de tip DHT de a echilibra încărcarea nodurilor membre, prin faptul de a permite acestei infrastructuri să se auto-extindă.

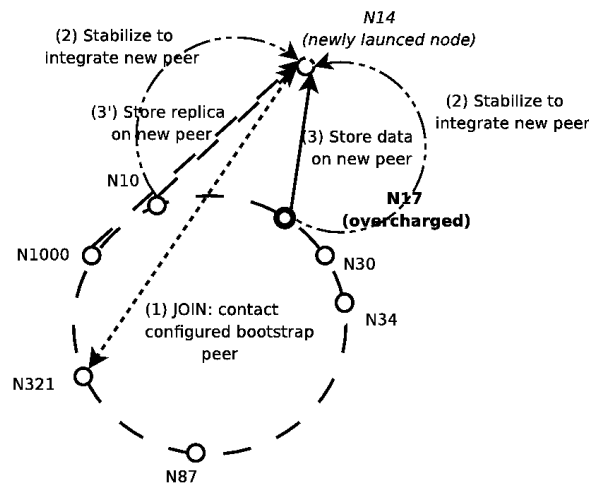


Figura 7: Adaugarea unui nou peer Chord, pentru a ajuta un peer supraincarcat.

Auto-extinderea se realizeaza prin adaugarea in principal de peer-i suplimentari, sau si de elemente fizice de stocare, ceea ce duce la o mai buna echilibrare a incarcarii intre peer-i si deci la evitarea aparitiei de peer-i supraincarcati sau suprasolicitati. Mecanismul proiectat de noi este complementar celui oferit implicit de DHT-urile Chord obisnuite, acesta fiind de asemenea transparent din punctul de vedere al clientilor DHT-ului, asa cum sunt si mecanismele interne ale acestuia. Originalitatea mecanismului nou-introdus consta in faptul ca reuseste sa extinda infrastructura in mod dinamic, modificand-o exact acolo unde este nevoie, prin utilizarea unui set maximal de peer-i disponibili, fapt ce face ca intregul sistem Chord imbunatatit sa devina mai scalabil si deci mai capabil decat inainte de a oferi o larga disponibilitate fata de clientii sai.

Pe viitor, ne-am dori sa ne concentram atentia asupra adaptarii, testarii si evaluarii modelului de echilibrare a sarcinii prezentat, pentru a rula pe o infrastructura logica distribuita geografic. De asemenea, ne-am dori sa gasim mijloace de a suprasolicita peer-ii cu cat mai multe cereri, spre a vedea posibile puncte slabe pe care le are solutia prezentata de echilibrare a incarcarii, si a vedea in ce momente sistemul ar putea deveni indisponibil sau partial indisponibil, pentru a fi constienti de modurile de a ataca sistemul prin atacuri de tip DoS, si in consecinta a putea preveni asemenea atacuri. Un alt aspect asupra caruia am dori sa ne concentram atentia este determinarea celei mai bune pozitii pentru nodul nou-lansat, astfel incat el sa preia cat mai multe din indatoririle nodurilor supraincarcate, aspecte pe care dorim sa le analizam in contextul a diverse situatii de distributie pentru cheile si pentru dimensiunile datelor.

6 Aplicarea modelului de descentralizare a unui serviciu pentru un server SIP

În capitolul 3 se prezintă modul generic de imbinare a elementelor prezentate în capitolele 4 și 5. Pe baza capitolelor anterioare, în capitolul de față dorim să validăm modelul generic de descentralizare a unui serviciu, spre a-l putea folosi pentru un server SIP de presență cu disponibilitate largă. Astfel, descriem modul original în care am reușit să facem ca serviciul de presență oferit într-o rețea SIP să fie larg disponibil. Pentru a obține această largă disponibilitate am utilizat, bineînțeles, *tehnicele peer-to-peer* definite pe parcursul acestei teze.

Scopul ce rezultă din chiar titlul tezei de față este acela de a realiza, și mai apoi valida, o distribuție și o descentralizare a unui server SIP, în vederea creșterii disponibilității acestuia. S-a vorbit deja, în capitolul 3, despre metodologia generică de a distribui un serviciu în scopul de a-l face larg disponibil clienților săi, fără a modifica ceva din modul de oferire al serviciului, în particular fără a schimba ceva în *protocolul de comunicare client-server* definit de serviciul respectiv. Analizând cu atenție modelul de distribuție generică a unui serviciu, se observă faptul că, în cazul unor protocoale orientate pe conexiune, va trebui ca, dacă mesajul inițial al clientului a trecut printr-un nod din load balancer și mai apoi a ajuns la nodul Chord care oferă și serviciul respectiv, toate mesajele din partea acelui client să treacă prin load balancer înainte de a ajunge la nodul Chord. Pe de altă parte, din cauza orientării pe conexiune a serviciului, load balancer-ul nu va putea realoca un client unui nou peer Chord la primirea unui mesaj oarecare, decât eventual în situația specială de eșec a primului nod Chord, situație ce va fi transparentă pentru client, dar ar putea implica eventuale retransmisii din partea acestuia. Într-o asemenea situație se realizează o încărcare inutilă a load balancer-ului, pentru cazul general, în care peer-ul responsabil de un client nu părăsește rețeaua, iar scopul nostru este tocmai să evităm încărcarea excesivă a unei componente a sistemului, care ar putea duce la gătirea acestuia. De aceea, deoarece componenta protocolului SIP pentru stabilirea de noi sesiuni se bazează tocmai pe existența unei conexiuni, sesiunea însăși, din cauza că *nu ne propunem să modificăm protocolul SIP în această lucrare*, vom exemplifica distribuția unui serviciu SIP în scopul unei largi disponibilități prin *descentralizarea și distribuția în scopul unei mult mai largi disponibilități a unui server SIP de presență*, serviciu ce se pretează foarte bine la descentralizarea folosind Chord și metoda de echilibrare a încărcării descrisă pe larg în capitolul 3. Motivul pentru care nu ne-am propus să implementăm și distribuția unui proxy de stabilire a sesiunii și de rutare SIP este: (1) pe de-o parte, faptul că, deși acest lucru este posibil folosind metoda generică de distribuție a unui serviciu pentru obținerea unei largi disponibilități, ar implica un număr mare de mesaje suplimentare între nodul Chord ce va servi un client SIP și acel client SIP; iar (2) pe de altă parte, faptul că dacă dorim să evităm acele mesaje suplimentare ar fi necesare schimbări în protocolul SIP, cum ar fi introducerea de noi câmpuri în mesajele de stabilire și de încheiere a sesiunii (INVITE, respectiv BYE sau CANCEL), modificări pe care nu dorim să le facem, pentru că am pierde compatibilitatea cu clienții SIP deja existenți. Asadar, în continuare ne vom ocupa de distribuția, în mod transparent pentru clienți, în scopul obținerii unei largi disponibilități a serviciului, a unui presență server, folosind modelul generic prezentat în capitolul 3. Preluând ideea din acest capitol și adaptând-o pentru SIP, cheia reprezentând ID-ul Chord al unui client SIP va fi obținută aplicând funcția de hash a DHT-ului Chord pe URI-ul (*Universal Resource Identifier*) clientului SIP în cauză.

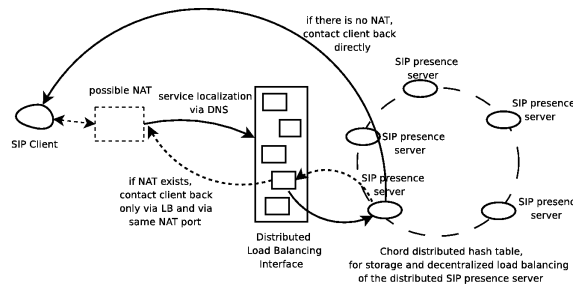


Figura 8: Distribuirea dinamica si echilibrata a clientilor intre mai multe server-e de presence SIP, aflate pe un inel Chord.

În cazul distribuirii dinamice a unui presence server, vom avea un load balancer de tipul celui descris în capitolul 3 în fața inelului Chord pe care rulează server-e de presence. Fiecare nod care este trecut în intrarea DNS aferentă serviciului de presence server, adică fiecare nod care îndeplinește funcționalitatea de *echilibrare a încărcării*, va funcționa după următorul algoritm simplu, care este ilustrat schematic și în figura 8: (1) va prelua mesaje de la clienții SIP de presence, prin intermediul unei componente minime de presence server SIP; (2) pentru fiecare client SIP de presence va parșa header-ul mesajului SIP, pentru a identifica URI-ul UAC-ului (User Agent Client) SIP curent și eventual al UAC-ului a cărui stare se dorește a se afla, cât și metoda SIP prezentă în mesaj: PUBLISH sau SUBSCRIBE; (3) dacă metoda este PUBLISH, URI-ul ce se va folosi în continuare este cel al clientului sursă, iar dacă metoda este SUBSCRIBE, URI-ul folosit este celălalt URI care apare în mesaj; (4) pe baza URI-ului SIP care trebuie utilizat, se va calcula cheia să în inelul Chord, folosind funcția de hashing a acestui DHT; (5) nodul dispatcher curent trimite, prin intermediul componentei sale Chord, întregul mesaj al clientului către peer-ul Chord responsabil de cheia acestui client, peer pe care rulează și presence server-ul ce deservește clientul; (6) presence server-ul de pe peer-ul care a primit mesajul de presence răspunde corespunzător clientului, astfel: (6.a) dacă clientul se află în spatele unui NAT (Network Address Translation), atunci server-ul de presence *retine* acest lucru relativ la clientul în cauză, în *baza de date distribuită* bazată pe Chord, în datele persistente ale acestuia, iar răspunsul va fi trimis trecând din nou prin nodul dispatcher; sau (6.b) altfel, răspunsul va fi trimis direct către client, pentru ca acesta să poată contacta de acum înainte direct nodul pe care rulează server-ul de presence responsabil pentru cererea sa. Se observă că NAT-ul pune probleme comunicării elegante de tip peer-to-peer, problema cu care, de altfel, se confruntă toate protocoalele peer-to-peer existente.

În urma recepționării unui mesaj de tip PUBLISH, dacă s-au primit și mesaje de tip SUBSCRIBE pentru clientul care a făcut PUBLISH, clienții care s-au înscris ca să afle starea de prezentă a acestuia vor fi notificați de către server-ul de presence printr-un mesaj de tip NOTIFY. Notificarea clienților se face fie direct, dacă clientul în cauză nu se află în spatele unui NAT, fie folosind dispatcher-ul, în caz contrar, și anume exact nodul din dispatcher prin care clientul în cauză a contactat serviciul de server presence, din cauza că, în majoritatea cazurilor, doar acelui nod din întreaga arhitectură a server-ului de presence NAT-ul îi permite să contacteze înapoi clientul.

Nu am mai insistat aici asupra modului de imbinare a componentelor individuale, prezentate de-a lungul acestei teze în capitolele 4 și 5, întrucât aceasta respecta ceea ce s-a definit în capitolul 3. Dorim să menționăm că și aplicarea modelului generic pentru a crea un serviciu SIP larg disponibil reprezintă o contribuție originală, dezvoltată în cadrul acestei teze.

7 Rezultatele evaluării unui serviciu descentralizat pentru a deveni larg-disponibil

În acest capitol vom prezenta, în ordine, rezultatele evaluării experimentale a elementelor pe care le-am descris deja în capitolul 4, și în capitolul 5, precum și în capitolele 3 și 6. Facem acest lucru pentru a putea ajunge la o concluzie legată de eficiența tehnicilor și metodelor utilizate pentru descentralizarea generică a unui serviciu, în scopul sporirii disponibilității sale. Elementele pe care am urmărit să le îmbunătățim, pentru a ajunge la un serviciu larg disponibil, și care vor fi evaluate în cele ce urmează, sunt: (1) persistența datelor pe care serviciul în cauză le stochează, punându-le la dispoziția clienților săi; (2) decuplarea componentei serviciului care realizează interogări, imposibil de evitat, cum ar fi interogări de DNS; (3) gradul de încredere și siguranța datelor păstrate de către serviciu, chiar și în condițiile unor amenințări de securitate sau a unor noduri care funcționează incorect; (4) toleranța la defecte și scalabilitatea, prin auto-extindere, pentru a putea asigura nevoile sporite de persistență, pe care le au datele critice sau sensibile; (5) echilibrarea sarcinii, statică sau dinamică prin auto-extindere, în scopul menținerii disponibilității largi a serviciului, chiar și în contextul creșterii solicitării acestuia de către clienți; și, nu în ultimul rând, (6) combinarea optimizărilor realizate asupra toleranței la defecte a serviciului, a scalabilității sale, a echilibrării încărcării sale, precum și a gradului de persistență a datelor cu care serviciul operează, în scopul obținerii unei soluții cât mai generice, de descentralizare, ce se poate aplica oricărui serviciu centralizat, spre a-i spori disponibilitatea.

Menționăm că toate aspectele evaluate în acest capitol sunt *ortogonale*, ceea ce are ca o consecință faptul că ideile originale introduse sunt mai ușor de urmărit dacă le discutăm și descriem separat. În acest capitol ne-a fost mai ușor chiar și să le evaluăm separat, fără ca în acest mod să negăm în vreun fel posibilitatea ca ele să fie utilizate împreună. Așa cum am menționat deja, împreună, aceste aspecte oferă toate premisele pentru ca serviciul are să utilizeze, distribuit în mod descentralizat, să devină cât mai larg disponibil pentru clienții săi.

7.1 Rezultatele evaluării bazei de date descentralizate

7.1.1 Evaluarea propriu-zisă a bazei de date descentralizate

În această secțiune vom evalua prototipul descris pe larg în capitolul 4.1. Ne interesează în special cât de potrivită este infrastructura logică Chord pentru a gestiona și controla o bază de date distribuită și descentralizată, principala noastră preocupare fiind timpul de răspuns suplimentar pe care aceasta soluție îl introduce, în comparație cu timpul de răspuns al unui server de baze de date tradițional, centralizat, care rulează pe un singur nod.

Server-ul descentralizat de baze de date pe care l-am proiectat și evaluat este compus dintr-un set de noduri pe care rulează server-e MySQL individuale, controlate fiecare de către un peer Chord, așa cum am descris în arhitectura prezentată în secțiunea 4.1. Am variat numărul de peer-i Chord și, în consecință, gradul de descentralizare al bazei de date distribuite, utilizând dimensiuni între 1 și 10 noduri, care toate aparțineau unei aceleasi rețele locale de mare viteză, pentru a nu trebui să ne preocupăm prea mult, cel puțin deocamdată, de timpul suplimentar necesar schimbului de mesaje implicat de comunicarea dintre peer-i. Având această arhitectură, precum și câte o componentă de distribuire a interogărilor care rulează pe fiecare peer Chord, conform arhitecturii prototipului nostru, am realizat următoarele experimente: (1) am măsurat timpul suplimentar introdus în cazul solicitării de către un client a unor date din baza de date distribuită și descentralizată; și (2) am măsurat timpul suplimentar introdus pentru a menține datele corect replicate, în condițiile în care se pot insera noi date în baza de date distribuită și descentralizată. Rezultatele acestor 2 experimente sunt ilustrate în figurile 9 și 10. Trebuie menționat faptul că valorile următoarelor parametri nu au fost considerate atunci când s-a făcut graficul din figurile 9 și 10: (1) latența rețelei noastre de test, care era în jurul valorii de 2ms, pentru cazul comunicării între oricare pereche de noduri; (2) intervalul de stabilizare predefinit, specific rețelei Chord, care reprezintă întârzierea maximă admisă până când un peer își da seama că unul dintre vecinii săi s-a schimbat, interval care a fost setat la 50ms; (3) timpul necesar realizării de către un nod MySQL server a unei interogări de tip SQL, timp egal în medie cu 10ms, pentru cazul celor mai simple interogări pe care le-am încercat, și egal în medie cu 500ms pentru interogările cele mai complicate pe care le-am încercat; (4) timpul necesar download-ului concret al copiei unor date de pe un nod MySQL server pe un altul, timp care este important să fie cunoscut în cazul în care există o volatilitate neneglijabilă a peer-ilor Chord, și pe care l-am considerat ca fiind aproximativ egal cu timpul mediu necesar pentru a trimite între peer-i un pachet IP de dimensiune medie, timp care am afirmat deja că în cazul rețelei noastre de test a fost egal cu aproximativ 2ms. Nu în ultimul rând, menționăm că dimensiunea datelor cu care lucrăm, pe care le stocăm și care pot fi obținute din baza de date, este suficient de mică pentru a putea încapă într-un singur pachet IP.

În urma efectuării experimentelor descrise mai sus, se poate observa, deci, că timpii suplimentari introdusi de către soluția noastră, atât pentru preluarea datelor de către clienți cât și pentru menținerea persistenței datelor, sunt independenți de numărul de peer-i Chord, în cazul considerat, în care numărul acestor peer-i este relativ mic, iar volatilitatea lor este moderată. Aceste limitări ale mediului de testare au fost făcute ținând cont și de modul în care ar fi alcatuit și s-ar comporta în realitate un server de baze de date distribuit și descentralizat, pentru că în mod sigur un asemenea server nu va avea nevoie să continue mai mult de câteva zeci de noduri, care pot fi alese pentru a se asigura o volatilitate moderată.

7.1.2 Evaluarea algoritmului de consens distribuit pentru menținerea unor date de încredere

În această secțiune ne propunem să evaluăm algoritmul de consens distribuit, descris pe larg în secțiunea 4.2. Ne vom canaliza atenția în special asupra timpului suplimentar introdus de acest algoritm distribuit, urmând să tragem concluzia dacă acest timp suplimentar este justificat de beneficiile functionale, aduse de către algoritmul în cauză.

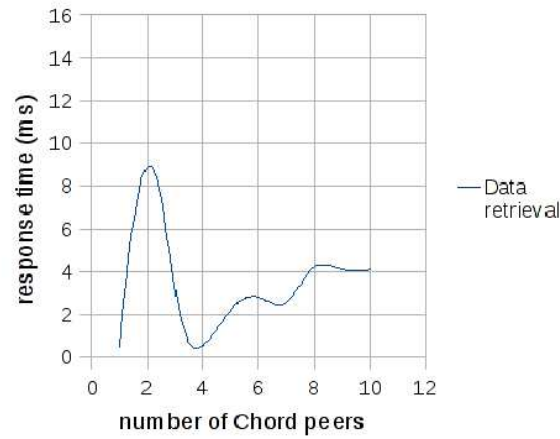


Figura 9: Evaluarea duratei preluării de date de la baza de date.

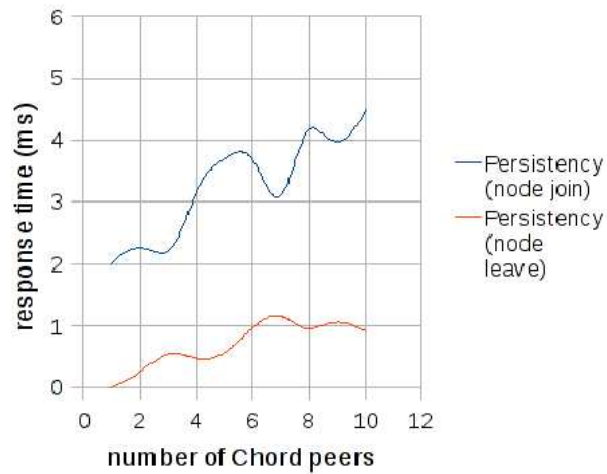


Figura 10: Evaluarea timpului suplimentar, necesar pentru mentinerea persistentei datelor.

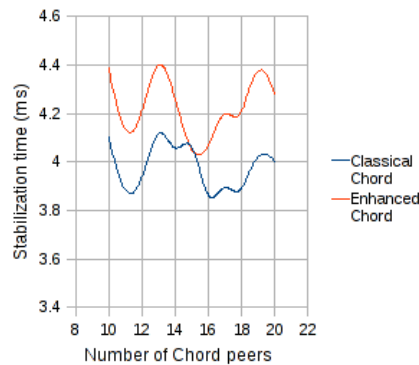


Figura 11: Evaluarea timpului suplimentar, necesar mentinerii topologiei, pe care il introduce algoritmul de consens distribuit.

Asa cum am anticipat deja in sectiunea 4.2, pentru a dovedi eficienta algoritmului nostru, am utilizat o baza de date distribuita, ce ruleaza peste infrastructura logica Chord, pe care am prezentat-o si in sectiunea 4.1, si am evaluat-o in sectiunea 7.1. Nodurile de stocare din cadrul bazei de date distribuite si descentralizate ruleaza server-e centralizate MySQL si am variat numarul de noduri de stocare, si, in consecinta, numarul de peer-i din cadrul topologiei Chord, intre 10 si 20 de noduri. Mediul nostru de test a constatat intr-o infrastructura fizica non-volatila, deoarece in evaluarea algoritmului de consens nu ne-a interesat sa vedem, deocamdata, impactul volatilitatii nodurilor. De asemenea, peer-i Chord utilizati de noi au rulat in cadrul unor masini virtuale, care au fost impartite in mod egal intre cele 5 noduri fizice diferite, pe care le-am avut la dispozitie, noduri conectate printr-o retea locala de mare viteza. Experimentul realizat de noi a inceput cu 10 peer-i Chord, dintre care unul era malitios, noi dorind sa vedem in cat timp acest nod malitios este eliminat din cadrul topologiei Chord. La fiecare pas, am reinitializat mediul de testare, cu un numar de noduri care crestea de fiecare data cu 1, pana cand am obtinut 10 seturi de rezultate experimentale. Din aceste date am dorit sa obtinem informatii despre: (1) cat de mult ii ia retelei Chord sa isi mentina topologia, prin rularea periodica a algoritmului de stabilizare, pe care noi l-am combinat si cu elemente ale algoritmului nostru de consens, prin tehnici de “piggy-backing” (in figura 11); si (2) cat de mult ii ia retelei Chord sa stocheze si sa returneze datele asociate unei chei, considerand neglijabila dimensiunea datelor, avand in vedere ca pentru fiecare mesaj receptionat de un peer Chord, peer-ul in cauza trebuie sa verifice, inainte de a procesa mesajul, daca are incredere in peer-ul sursa al mesajului (in figura 12). Nu am incercat sa masuram cat timp le ia peer-ilor Chord din cadrul unei topologii sa elimine un nod malitios, deoarece, asa cum am descris si in sectiunea 4.2, este evident ca aceasta durata de timp este influentata foarte mult si de cat de multe mesaje trimite nodul malitios.

Atat din figura 11, cat si din figura 12 se poate observa faptul ca, deoarece timpii suplimentari introdusi de algoritmul proiectat de noi sunt mici, dimensiunea topologiei Chord aproape ca nu influenteaza timpul suplimentar introdus de acest algoritm de consens distribuit optimizat. Consideram ca acest lucru reprezinta un rezultat foarte bun pentru algoritmul nostru, pentru ca stim ca in general, cand este vorba de algoritmi distribuiti de consens

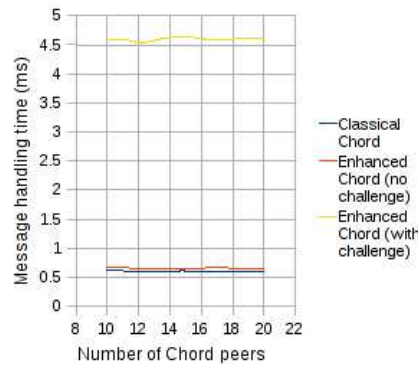


Figura 12: Evaluarea timpului suplimentar introdus de catre algoritmul de consens distribuit, prin comunicarea suplimentara.

prin vot, numarul de noduri din sistem influenteaza direct si semnificativ timpul necesar obtinerii unui verdict final al votului. Inseamna ca, daca consideram ca performantele de timp ale infrastructurii Chord clasice sunt bune, avand in vedere faptul ca algoritmul nostru aduce o imbunatatire semnificativa in ceea ce priveste siguranta infrastructurii Chord oferite, putem considera ca fiind bune si performantele infrastructurii noastre Chord, imbunatatite prin algoritmul de consens distribuit, optimizat.

7.1.3 Evaluarea, la nivelul clientului de baze de date, a timpului de raspuns pentru interogările transmise

In aceasta sectiune vom evalua prototipul descris pe larg in sectiunea 4.3. Ne intereseaza in special sa identificam niste corelatii, existente intre factorii externi prototipului elaborat si prototipul in sine, de-a lungul existentei acestuia. Am identificat doua asemenea corelatii, si anume: (1) *numarul de server-e diferite* care ar trebui contactate la un moment dat (in figura 13); si (2) *numarul de interogari identice* care ar trebui executate la un moment dat. Pentru simplitate, vom considera ca interogari efectuate in cadrul testelor noastre doar interogari de DNS, deoarece acestea sunt mai usor de conceput si de lansat si de asemenea este mai usor de masurat durata acestora (in figura 14). Pentru interogari de baze de date, durata depinde mult de gradul de dificultate al interogarii precum si de planificarea interna pentru executia interogarilor primite, de pe server-ul de baze de date, pe langa intarzierile introduse in retea si incarcarea server-ului de baze de date, aceste ultime doua aspecte impactand, dealtfel, si interogările DNS.

Scenariul de test considerat de noi consta in a avea mai multi clienti, care contin doar componenta de cereri, care a fost descrisa in sectiunea 4.3, precum si un singur worker, reprezentand entitatea secundara de executie a interogarilor. Pentru a demonstra faptul ca pana si pe o masina cu un singur procesor si un singur core prototipul deja prezentat de noi e capabil sa aduca imbunatatiri, am utilizat o asemenea masina ca fiind suportul hardware pentru testele noastre. In primul test, interogările clientilor au fost adresate mai

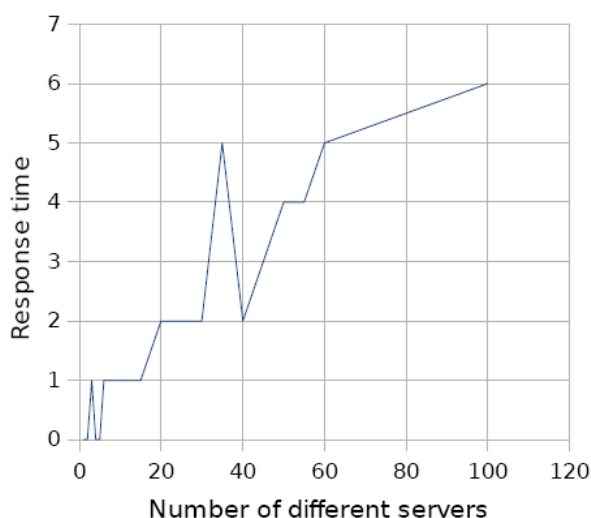


Figura 13: Influenta unui numar crescand de server-e diferite asupra performantelor entitatii secundare de executie a interogarilor.

multor server-e diferite, in timp ce in cel de-al doilea test interogarile tuturor clientilor au fost identice, fiind adresate unui aceluasi server. In contextul celui de-al doilea test, reamintim faptul ca entitatea secundara de executie a interogarilor a fost optimizata pentru a transmite in exterior un numar minim de interogari, grupand interogarile identice, pentru a minimiza numarul de mesaje transmise in retea, si deci numarul de operatii de intrare-iesire.

Combinand rezultatele obtinute in cazul unui numar variabil de server-e diferite catre care se transmit interogari cu rezultatele obtinute in cazul existentei unui numar variabil de interogari identice, obtinem o estimare suficient de buna a comportamentului sistemului intr-un mediu real, unde, la un moment dat, am putea avea cazul unui numar variabil de interogari, adresate catre mai multe server-e, dintre care unele interogari sunt identice. Este evident faptul ca, din punctul de vedere al server-elor distante, faptul de a li se transmite mai putine interogari, prin gruparea pe care am instituit-o la nivelul prototipului nostru, este si el benefic. Rezulta, din masuratorile descrise in aceasta sectiune, ca doar numarul de server-e diferite ce trebuie contactate influenteaza performantele nivelului de decuplare introdus de noi, pentru cazul in care aplicatia ruleaza pe o masina cu un singur procesor si un singur core. Ne putem imagina ca degradarea care s-ar vedea pe o masina multi-procesor multi-core va fi mult mai mica, poate chiar imperceptibila.

7.2 Rezultatele evaluării metodelor de exploatare a dinamicității peer-to-peer

7.2.1 Evaluarea funcționalității de autoextindere pentru persistență a datelor sensibile

In aceasta sectiune vom evalua prototipul descris pe larg in sectiunea 5.1. Ne intereseaza in special cat de mult ajuta prototipul elaborat de noi cresterea gradului de incredere pentru stocarea datelor sensibile, precum si persistenta acestora.

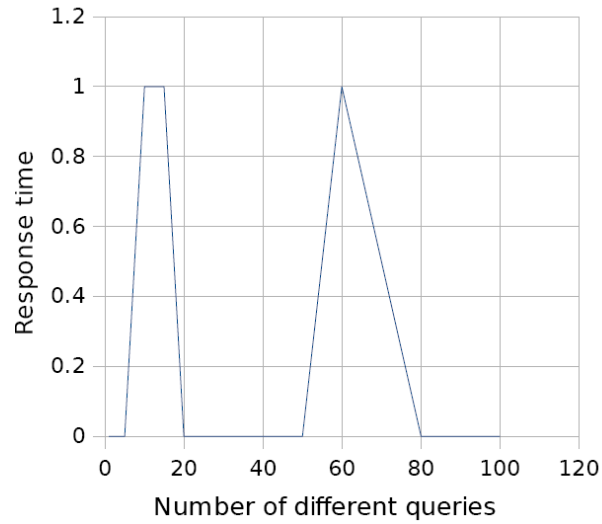


Figura 14: Influența unui număr crescând de interogări identice asupra performanțelor entității secundare de execuție a interogărilor.

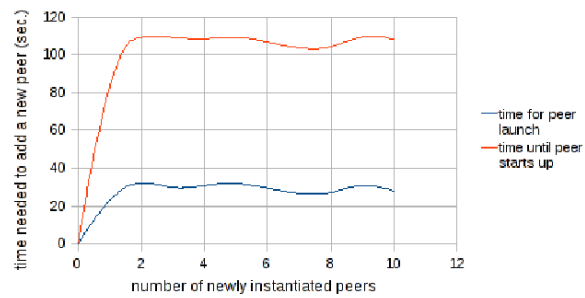


Figura 15: Instantiatierea unui nou nod-fiu la fiecare 5 minute.

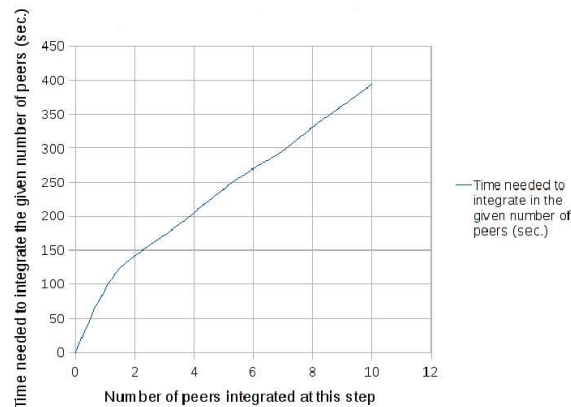


Figura 16: Instantiatieerea unui numar crescand de peer-i, la fiecare 5 minute.

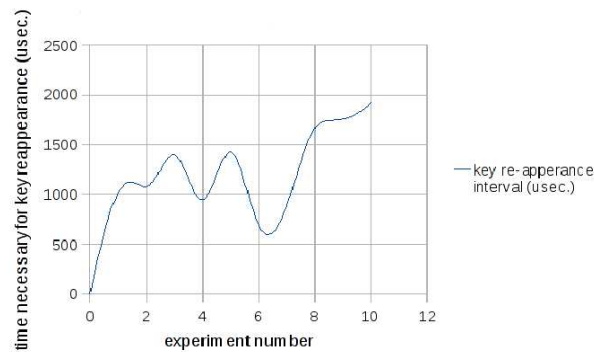


Figura 17: Durata replicarii datelor pe noul lor posesor.

Pentru evaluarea acestui prototip, am folosit un mediu de testare format din mai multe masini virtuale, care ruleaza peste cateva masini fizice, spre a putea avea un numar cat mai mare de peer-i, chiar daca nu am avut la dispozitie un numar atat de mare de noduri fizice. In aceasta sectiune ne vom ocupa in special de evaluarea mecanismului de instantiere de noi peer-i pe noduri suplimentare, adaugate la nevoie infrastructurii peer-to-peer. Astfel, dorim sa masuram: (1) cat de mult dureaza lansarea unui numar de noduri, situate fiecare la distante similare in reseaua fizica fata de nodul parinte, care le lanseaza (in figurile 15 si 16); si (2) cat de mult dureaza “restaurarea” datelor de pe un nod fiu pe noul nod parinte, reprezentand responsabilul si proprietarul de drept al datelor corespunzatoare cheii in cauza (in figura 17).

Se poate observa, din figura 17, ca timpul necesar schimbarii mesajelor, care readuc cheia si datele pe nodul care este de fapt responsabil de ele, este de aproximativ 2ms. Remarcam faptul ca acest timp nu include: (1) latenta de retea ce ar exista intr-un caz real, cand nodurile pot avea o larga distributie geografica; (2) intervalul de timp definit pe nodul fiu pentru a-si contacta periodic parintele prin mesaje de tip “ping”, interval care trebuie sa nu fie considerat nici prea mare, dar nici prea mic; si (3) timpul necesar download-ului efectiv al

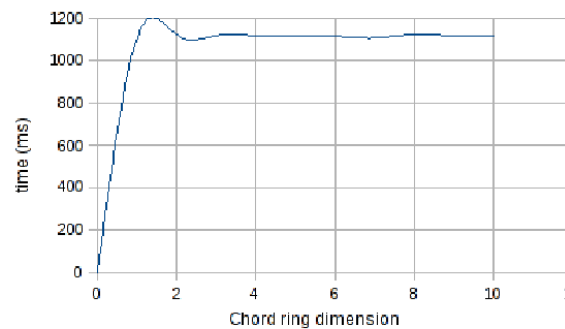


Figura 18: Adaugarea unui nou peer la infrastructura Chord, spre a prelua o parte din incarcarea unui peer supra-incarcat.

datelor aferente cheii pe noul parinte, timp care ar putea fi foarte mare, daca datele sunt mari. Aceste intervale de timp nu au fost considerate la realizarea figurii 17 pentru ca ele pot varia foarte mult, in functie de factori externi prototipului, si, de asemenea, pentru ca oricum insumate ar fi mult mai mari decat timpul necesar procesarii mesajelor schimbate intre peer-ul fiu si peer-ul parinte. Pe viitor ne dorim sa luam in considerare si o valoare medie pentru latentă in retea dintr-un sistem Chord mai larg distribuit geografic, astfel incat rezultatele obtinute de noi sa aiba o relevanta mai mare pentru cazul de utilizare reala. Cu toate acestea, chiar si avand in vedere timpii suplimentari introdusi prin aceste intarzieri, prin utilizarea prototipului elaborat de noi, impreuna cu mecanismele propuse de crestere a redundantei si deci a tolerantei la defecte, solutia noastra se comporta mult mai bine decat o retea Chord obisnuita, in a proteja anumite date foarte valoroase pentru aplicatia de nivel inalt care stocheaza date utilizand Chord.

7.2.2 Evaluarea funcționalității de echilibrare a încărcării, pentru o largă disponibilitate a sistemului

In aceasta sectiune vom evalua prototipul descris pe larg in sectiunea 5.2. Ne intereseaza in special sa masuram timpul suplimentar introdus pentru a spori scalabilitatea si echilibrarea incarcarii intre noduri, prin utilizarea de mecanisme de auto-extindere similare cu cele evaluate in sectiunea 7.2.1.

Intrucat metoda de auto-extindere este comuna cu solutia pentru sporirea tolerantei la defecte si a scalabilitatii, prezentata in sectiunea 5.1, si pentru evaluarea prototipului de echilibrare a incarcarii, am folosit un mediu de testare format din mai multe masini virtuale, care ruleaza peste cateva masini fizice, spre a putea avea un numar cat mai mare de peer-i, chiar daca nu am avut la dispozitie un numar atat de mare de noduri fizice. In aceasta sectiune ne vom ocupa in special de evaluarea mecanismului de instantiere de noi peer-i pe noduri suplimentare, adaugate la nevoie infrastructurii peer-to-peer. Astfel, dorim sa masuram: (1) timpul scurs intre momentul cand un peer isi da seama ca este supraincarcat

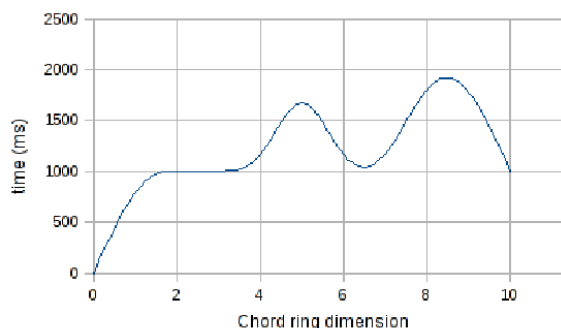


Figura 19: Un peer extrem de incarcat trebuie sa paraseasca inelul Chord, toate indatoririle fiindu-i preluate de catre succesorul sau.

si momentul in care supra-incarcarea sa este preluata de catre un peer nou-instantiat (in figura 18); si (2) timpul scurs intre momentul cand un peer isi da seama ca este extrem de incarcat, observatie in urma careia va decide sa paraseasca reseaua, pentru ca nu-si mai poate indeplini, oricum, corect indatoririle, si momentul in care toata incarcarea sa este preluata de catre nodul sau de rezerva, adica succesorul sau (figura 19).

Inainte de a descrie vreunul din experimentele realizate, am dori sa mentionam faptul ca, pe langa nodurile virtuale care ruleaza peste noduri fizice, in cadrul *infrastructurii fizice* pe care am utilizat-o, conexiunile de retea dintre noduri s-au realizat prin intermediul unei retele locale de mare viteza, avand o latentă de aproximativ 2ms. Pe de alta parte, pentru *infrastructura logica* utilizata, adica inelul Chord, am definit un interval de stabilizare al topologiei egal cu 1 secunda. Am considerat important sa mentionam aceste valori de timp, pentru ca ele prezinta un impact semnificativ asupra comportamentului sistemului de echilibrare a incarcarii, precum si asupra capacitatii acestuia de a face fata dinamicitatii.

Concluzia acestor doua tipuri de experimente, luate in ansamblu, realizate pentru a evalua capacitatea de auto-echilibrare si de scalabilitate a prototipului nostru, este faptul ca, intr-o maniera foarte simpla si transparenta, acestuia i se poate adauga capacitate suplimentara de stocare, materializata atat prin adaugarea de spatiu fizic suplimentar de stocare pentru peer-ii individuali existenti, cat si prin adaugarea de noi peer-i Chord, avand o capacitate proprie de stocare. In acest mod, unui operator uman sau unei aplicatii externe li se pot pune la dispozitie in mod eficient, din punct de vedere al consumului de energie si al nevoilor clientilor, toate tipurile de resurse de stocare avute la indemana.

7.3 Rezultatele evaluării unui serviciu generic, descentralizat pentru sporirea disponibilității sale

In aceasta sectiune dorim sa realizam o evaluare preliminară a modelului de distribuire a unui serviciu generic, prezentat in capitolul 3. Pentru a evalua acest prototip s-a utilizat un

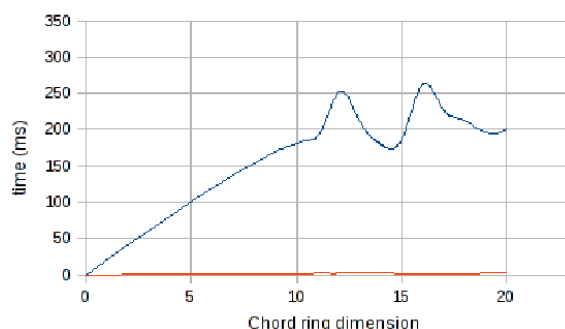


Figura 20: Timpul suplimentar introdus de către sistemul de descentralizare în contextul transmiterii cererilor și a răspunsurilor între client și serviciu.

load balancer, distribuit pe 3 noduri, pentru fiecare existând câte o pondere egală în cadrul înregistrărilor DNS de tip NAPTR, respectiv SRV. În spatele acestui load balancer există o rețea Chord, formată din 0 până la 20 de peer-i, fiecare peer aflându-se în strânsă legătură cu un server ce rulează pe același nod, toate server-ele care rulează pe peer-i fiind de același tip. Rețeaua care interconectează atât peer-ii între ei, cât și peer-ii de nodurile de load balancing, ce se află în fața serviciului propriu-zis, este o rețea LAN de mare viteză. Menționăm faptul că rețeaua Chord utilizată este o rețea Chord îmbunătățită conform contribuțiilor descrise în secțiunile 4.1, 4.2, 5.1, și 5.2, și care au fost evaluate și validate în secțiunile anterioare ale acestui capitol, pentru a oferi o mai mare scalabilitate, toleranță la defecte, persistență a datelor și echilibrare a încărcării peer-ilor.

La început, în primul experiment, rețeaua Chord este vidă, ceea ce înseamnă că nu există nici un nod care să ofere serviciul dorit. Apoi, succesiv, după ce a fost instantiat manual un prim peer, împreună cu serviciul aferent, ce urmează a fi distribuit, care a format o *rețea Chord elementară*, ce este capabilă să ofere serviciul generic considerat, s-au utilizat metode de auto-extindere a infrastructurii Chord, descrise în capitolul 5, ce au fost evaluate în secțiunile 7.2.1 și 7.2.2. Pentru fiecare dimensiune considerată a rețelei Chord s-au realizat măsurători ale: (1) timpului suplimentar, introdus pentru ca o *cerere* să ajungă de la client la serviciul pe care acesta îl contactează, de către infrastructura noastră de distribuție generică a unui serviciu, formată din load balancer-ul distribuit și din inelul Chord de peer-i pe care rulează și serviciul generic considerat; și (2) timpul suplimentar, introdus pentru ca *răspunsul* serviciului generic distribuit și descentralizat să ajungă de la peer-ul Chord pe care cererea a fost procesată până la client, scurt-circuitând, de această dată, nivelul de load balancing.

În figura 20 se pot vedea timpii suplimentari introdusi în recepționarea și procesarea cererii, respectiv în trimiterea răspunsului, conform modelului nostru de descentralizare și distribuție a unui serviciu. Menționăm că acești timpi s-au obținut scăzând timpii de procesare și de transmitere a mesajelor prin rețea, care se pierd și în mod obișnuit, când serviciul generic nu este distribuit conform modelului nostru din timpii totali, obținuți: (1) în cazul recepționării și procesării cererii, calculând diferența de timp dintre momentul când clientul a

trimis mesajul, determinat prin intermediul unei marci de timp prezente în mesaj, și momentul când server-ul a terminat de construit răspunsul către client; și (2) în cazul transmiterii răspunsului, calculând diferența de timp dintre momentul când server-ul a trimis răspunsul, determinat tot prin intermediul unei marci de timp prezente în mesaj, și momentul când clientul recepționează răspunsul.

Se observa că timpul suplimentar, introdus de către modelul de distribuire în scopul descentralizării unui serviciu, pentru recepționarea și procesarea cererii este în medie de 200ms. Acest fapt nu constituie o surpriză, întrucât, conform articolului de descriere a funcționalității și de evaluare a Chord, pentru un inel de aproximativ 10-30 de peer-i, timpul necesar realizării rout-ării unui mesaj, în condițiile chiar ale unei volatilități moderate a nodurilor, este de aproximativ 200ms. Pe de altă parte, timpul suplimentar introdus de către modelul de distribuire în scopul descentralizării unui serviciu, pentru transmiterea răspunsului către client la o cerere a acestuia, este aproape inexistent. Rezulta, deci, că ceilalți timpi suplimentari introdusi de către distribuirea descentralizată a serviciului, conform modelului introdus în această lucrare, sunt neglijabili.

Asadar, per ansamblu, timpul suplimentar introdus în general de către serviciul distribuit, pe care îl percepe clientul de la momentul trimiterii cererii și până la momentul primirii răspunsului pentru aceasta din partea server-ului descentralizat, este egal cu timpul necesar rout-ării unui mesaj Chord într-un inel de 10-20 de peer-i, ceea ce înseamnă, de fapt, că avem un cost pentru beneficiile pe care Chord le introduce, care nu este, însă, mai mare decât însuși costul rout-ării unui mesaj prin intermediul inelului Chord, care este de fapt componenta principală de descentralizare a serviciului, care îi permite acestuia să devină mai larg disponibil către clienții săi, prin metode enunțate deja în capitolul 3, cât și în secțiunile 4.1, 4.2, 5.1, și 5.2.

7.4 Rezultatele evaluării unui SIP presence server, descentralizat pentru sporirea disponibilității sale

În această secțiune dorim să realizăm o evaluare preliminară a modelului de distribuire a unui SIP presence server, prezentat în 6, și care se bazează pe modelul generic de distribuire și descentralizare a unui serviciu, care a fost prezentat în 3, fiind evaluat în secțiunea 7.3. Ipotezele privind load balancer-ul și rețeaua Chord sunt identice cu cele din secțiunea 7.3, și anume: (1) load balancer-ul este compus din 3 noduri, având o probabilitate și o pondere egală în cadrul înregistrărilor DNS de tip NAPTR și SRV; (2) dimensiunea rețelei Chord, îmbunătățite, identice cu cea folosită în secțiunea 7.3, variază între 0 și 20 de peer-i, dintre care primul peer este pornit manual, iar următorii sunt obținuți, pe baza unor configurări la initializarea fiecărui peer, conform metodelor de autoextindere descrise în capitolul 5, și evaluate în secțiunile 7.2.1 și 7.2.2; și (3) rețeaua care interconectează oricare două noduri din cadrul sistemului distribuit descentralizat este o rețea LAN rapidă. Menționăm, de asemenea, că server-ul de presence utilizat, care rulează pe fiecare nod din inelul Chord, este server-ul open-source OpenSIPS.

Pentru fiecare dimensiune considerată a rețelei Chord s-au realizat măsurători ale impactului de timp asupra unui client de SIP presence, datorat distribuirii server-ului de presence prin intermediul load balancer-ului și a rețelei Chord, însumând timpul scurs de la trimi-

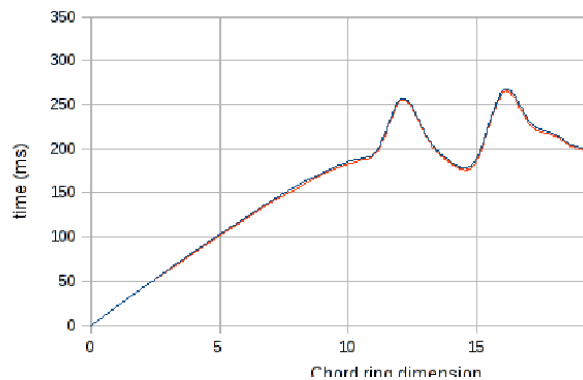


Figura 21: Impactul distribuției în scopul obținerii unei mai largi disponibilități, a unui server de presence SIP, asupra unui client al acestuia.

terea cererii de către client și până la primirea răspunsului din partea server-ului, în cazul în care: (1) clientul se afla în spatele unui NAT, și deci răspunsul trebuie să îi fie trimis de către server-ul care îl tratează utilizând și load balancer-ul, pentru care există un port deschis în cadrul NAT-ului; și (2) clientul nu se afla în spatele unui NAT, caz în care server-ul de presence îi poate trimite răspunsul direct, scurt-circuitând astfel load balancer-ul, spre a economisi timp de trimitere în rețea și spre a nu-l mai încărca suplimentar pe acesta și cu transmiterea răspunsurilor.

În figura 21 se poate vedea impactul de timp asupra clientului a descentralizării server-ului de presence, atât în cazul în care clientul se afla în spatele NAT-ului, cât și în cazul în care acesta nu se afla în spatele NAT-ului. Timpii de impact au fost obținuți, similar cu procedeul din secțiunea 7.3, prin scăderea, din durata totală de timp de când un client a trimis mesajul și până când a primit răspunsul, a duratei de timp necesare în mod obișnuit transmiterii mesajului și a răspunsului și procesării cererii primite de la client de către nodul server de presence SIP. Se poate remarca faptul că impactul descentralizării server-ului de presence, spre a-l face mai larg disponibil, este foarte apropiat pentru clienți aflați în spatele NAT-ului de cel pentru clienți care nu se afla în spatele NAT-ului, timpul suplimentar introdus fiind în medie de 200ms, adică exact durata routării unui mesaj prin inelul Chord de dimensiune maxim 20 de noduri. Comparând rezultatele obținute pentru distribuirea în scopul obținerii unei mai largi disponibilități a server-ului de presence din cadrul OpenSIPS cu rezultatele obținute pentru distribuirea generică a unui serviciu, din secțiunea 7.3, putem remarca faptul că ele sunt foarte apropiate, ceea ce înseamnă că server-ul de presence a fost distribuit corect, conform modelului generic, fapt ce realizează o validare preliminară a acestui model de distribuție și descentralizare, în scopul obținerii unei mai largi disponibilități.

7.5 Concluzii ale evaluării rezultatelor experimentale

În acest capitol am realizat evaluarea rezultatelor științifice originale introduse de către această teză. Subliniem și cu această ocazie faptul că noi considerăm ca esențial în creșterea

disponibilitatii serviciului este DHT-ul. In sine, orice DHT este un sistem scalabil, echilibrat si tolerant la defecte, permitand totodata si multe imbunatatiri si extensii, dintre care am prezentat cateva ca fiind contributii originale ale acestei teze, pe care le-am evaluat in acest capitol. Nu pretindem ca evaluarea ar fi fost una exhaustiva, asa cum mentionam, dealtfel, in fiecare sectiune a acestui capitol. In masuratorile efectuate am urmarit doua directii mari: (1) imbunatatirea timpului dintre 2 esecuri consecutive ale sistemului (*Mean Time Between Failures* – MTBF), in special in sectiunea 7.2, si (2) minimizarea impactului de timp al componentelor și algoritmilor pentru mentinerea unei infrastructuri disponibile, pe parcursul intregii evaluari a rezultatelor. Consideram ca evaluarea facuta aici a fost suficienta pentru a sublinia viabilitatea directiei urmate in cercetarea noastra, viabilitate sustinuta si prin cele 7 lucrările publicate in conferinte si reviste de specialitate, elaborate de-a lungul acestui doctorat.

8 Contribuții și concluzii

In capitolele anterioare am identificat si analizat problemele care pot sa impacteze disponibilitatea unui serviciu distribuit. Problemele identificate si analizate sunt: (1) persistenta datelor; (2) toleranta la defecte; (3) scalabilitatea; (4) echilibrarea incarcarii; si (5) securitatea sistemului. Pentru acestea, pe parcursul lucrarii de fata, oferim solutii proprii generice, bazate pe o arhitectura descentralizata de tip peer-to-peer. In capitolele 3 si 6 am propus o maniera proprie de a combina toate contributiile originale prezentate in capitolele 4 si 5, pentru a putea intr-adevar sa obtinem serviciul distribuit larg-disponibil, pe care il dorim. Platformei generice de descentralizare a unui serviciu i-am conferit o forma unitara si completa, care reuseste sa combine toate avantajele solutiilor prezentate in capitolele 4 si 5. De asemenea, in capitolul 7, am evaluat, una cate una, performantele pe care le aduc fiecare dintre contributiile noastre, fiecare in parte. Am masurat independent fiecare contributie, caci consideram ca solutiile noastre la problemele identificate in studiul bibliografic sunt generice si independente. Rezultatele obtinute sunt valabile si pot fi, deci, aplicate si utilizate in orice context ar fi nevoie de ele, nu doar in cel al descentralizarii unui serviciu, spre a-i creste disponibilitatea.

Vom prezenta, pe scurt, principalele contributii originale aduse prin intermediul tezei de fata:

1. In capitolul 4.3: o solutie pentru decuplarea executiei interogarilor, astfel incat un server care are nevoie sa interogheze un alt server sa nu astepte blocant raspunsul. El poate, deci, prelua si deservi alti clienti pana la primirea raspunsului de la server-ul distant. Aceasta contributie este utila si necesara intrucat orice serviciu am dori sa furnizam, server-ul proiectat va trebui sa transmita un minim de interogari catre alte server-e. De aceea, se poate ca performantele serviciului pe care dorim sa il oferim sa poata fi impactate, prin dependenta de un server distant, de entitati al caror comportament nu il putem controla.
2. In capitolul 4.1: o solutie pentru stocarea persistenta a datelor, utilizand o baza de date distribuita si descentralizata, bazata pe arhitectura peer-to-peer Chord. Peer-ii Chord controleaza noduri independente, pe care ruleaza server-e obisnuite de baze de date.

Acesta soluție este utilă și necesară deoarece orice serviciu, și cu atât mai mult unul care se dorește larg disponibil, are nevoie să stocheze pe o durată cât mai lungă datele clienților pe care îi deserveste.

3. În capitolul 4.2: o soluție pentru asigurarea unui grad de siguranță și încredere pentru stocare cât mai mare, chiar și în prezența unor noduri malicioase în cadrul sistemului distribuit. Aceasta soluție este bazată tot pe arhitectura Chord și definește un protocol de consens distribuit adaptat acestei arhitecturi. Acesta soluție este utilă și necesară, deoarece orice serviciu, și cu atât mai mult unul care se dorește larg-disponibil, are nevoie *sa stocheze cu un grad mare de încredere* date, pe o durată cât mai lungă, indiferent dacă în sistem există sau nu noduri malicioase.
4. În capitolul 5.1: o soluție pentru asigurarea toleranței la defectele individuale ale nodurilor din sistem, și într-o oarecare măsură și pentru sporirea scalabilității sistemului. Toleranța la defecte este crescută în special pentru datele critice pe care serviciul le stochează. Soluția prezentată este bazată și ea tot pe arhitectura DHT-ului Chord, prin definirea unui mecanism de autoextindere, ce este invocat pentru a spori gradul de redundanță cu care sunt stocate datele critice. Acesta soluție este utilă și necesară deoarece orice serviciu, și cu atât mai mult unul care se dorește larg-disponibil, are nevoie să asigure toleranța la defecte, spre a proteja datele clienților săi, pe care le stochează, și în special datele critice ale acestora.
5. În capitolul 5.2: o soluție pentru asigurarea echilibrării încărcării nodurilor, și, în consecință, și pentru sporirea scalabilității sistemului. Soluția noastră se bazează tot pe DHT-ul Chord, utilizând la maxim avantajele dispersiei uniforme oferite de funcția de hashing a DHT-ului. Datorită acestei funcții de dispersie, soluția noastră reușește să integreze, la nevoie, tot prin autoextindere, noi resurse, exact acolo unde ele sunt necesare, adică acolo unde încărcarea este prea mare. Aceasta soluție este utilă și necesară deoarece orice serviciu, și cu atât mai mult unul care se dorește larg-disponibil, are nevoie să asigure echilibrarea încărcării nodurilor și, deci, scalabilitate în ceea ce privește numărul de clienți pe care îi poate deservi.
6. În capitolul 3: o soluție pentru combinarea avantajelor de persistență a datelor, de încredere, de toleranță la defecte, de echilibrare a încărcării și de scalabilitate. Propunem, deci, la final, o metodă generică, bazată pe DNS și pe DHT-uri, de a avea un unic punct de acces la un serviciu care beneficiază de toate avantajele menționate mai sus. În acest mod, descentralizarea serviciului, utilizând arhitectura Chord, în scopul largirii disponibilității acestuia, devine transparentă pentru clienții serviciului. Aceasta soluție este utilă întrucât oferă o “rețetă” generică, ce se poate aplica oricărui serviciu distribuit, în scopul largirii disponibilității acestuia. Mai mult, această soluție este validată nu doar prin măsurători, ci și prin aplicarea ei pentru un server SIP de prezentă.
7. În capitolul 6: O soluție pentru a obține un serviciu SIP larg disponibil, prin aplicarea soluției generice de descentralizare a unui serviciu pentru un server de prezentă SIP. La rândul său, server-ul de prezentă descentralizat a fost evaluat, spre a confirma că i s-a sporit disponibilitatea față de clienți, în urma modificărilor aduse.

Pe baza rezultatelor obținute, evaluând experimental contribuțiile originale ale tezei, putem conchide că obiectivul tezei principal al tezei a fost îndeplinit. Asadar, am obținut o

reteta pentru realizarea unui serviciu larg disponibil cu impact de performanta rezonabil. Impactul de performanta in vederea disponibilitatii, observat in urma interpretarii datelor experimentale, este in general comparabil cu perioada de stabilizare a DHT-ului Chord. Cu alte cuvinte, daca mediul academic si mediul de productie considera Chord, si DHT-urile in general, ca fiind o solutie viabila pentru servicii utilizate *la scara larga* (large scale), atunci si solutia originala propusa de noi poate fi considerata viabila. Un argument in favoarea faptului ca utilizand Chord putem obtine solutii practice viabile este faptul ca acest DHT este utilizat in implementari legate de P2P SIP, pentru care in prezent se pregateste un RFC. Acesta, desi nu unic, este si un argument important in favoarea *actualitatii* subiectului abordat de teza in cauza.

Pe viitor ne propunem sa imbunatatim unele contributii prezentate, in directiile descrise in momentul prezentarii si evaluarii performantelor lor. De asemenea, ne propunem sa investigam si alte metode de a spori disponibilitatea unui serviciu, pe care mai apoi sa le aplicam unor servicii distribuite particulare, critice, asa cum este si SIP.