

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



Modelarea sistemelor RAID și măsurarea performanțelor

Lucrare de licență

Prezentată ca cerință parțială pentru obținerea titlului de Inginer
În domeniul Electronică și Telecomunicații
Programul de studii de licență Calculatoare și Tehnologia Informației



2014

Conducător științific: **Conf. Dr. Ing. Ștefan Stăncescu**
Absolvent: **Dragoș – Alexandru Bălăucă**

APROBAT Director de Departament
Prof. Dr. Ing. Sever PAȘCA

TEMA PROIECTULUI DE LICENȚĂ A STUDENTULUI

BĂLĂUCĂ I. DRAGOȘ - ALEXANDRU

1. Titlul temei:

Modelarea sistemelor RAID și măsurarea performanțelor

2. Date inițiale pentru proiectare:

Simulator DiskSim 4.0

Mediul de implementare Linux

Limbaj Visual Basic

3. Conținutul memoriului:

Modele de arhitecturi RAID și prezentarea acestora

Notiuni introductive rebuilding

Moduri de funcționare pentru matricile RAID cu disc de rezervă

Experimente pe mediul de simulare DiskSim 4.0 cu parametri și workload variabile

Simulare, comparații, performanțe și concluzii

4. Material grafic obligatoriu:

Rezultate și grafice comparative

5. Realizare practică:

Implementare software

6. Locul de desfășurare al lucrării practice

Catedra de Calculatoare și Tehnologia Informației

7. Lucrarea servește la:

Autodotare catedra

8. Mijloacele materiale puse la dispoziție de:

Catedra de Calculatoare și Tehnologia Informației, UPB

9. Realizarea practică rămâne în proprietatea:

Facultatea de Electronică Aplicată și Tehnologia Informației, UPB

10. Data eliberării temei:

2014

CONDUCĂTOR PROIECT:
Conf. Dr. Ing. Ștefan Stăncescu

Absolvent:
Bălăucă I. Dragoș - Alexandru

Declarație de onestitate academică

Prin prezența declar că lucrarea cu titlul “ *Modelarea sistemelor RAID și măsurarea performanțelor*”, prezentată în cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației a Universității “Politehnica” din București ca cerință parțială pentru obținerea titlului de *Inginer* în domeniul *Calculatoare și Tehnologia Informației*, este scrisă de mine și nu a mai fost prezentată niciodată la o facultate sau instituție de învățământ superior din țară sau străinătate.

Declar că toate sursele utilizate, inclusiv cele de pe Internet, sunt indicate în lucrare, ca referințe bibliografice. Fragmentele de text din alte surse, reproduse exact, chiar și în traducere proprie din altă limbă, sunt scrise între ghilimele și fac referință la sursă. Reformularea în cuvinte proprii a textelor scrise de către alți autori face referință la sursă. Înțeleg că plagiatul constituie infracțiune și se sancționează conform legilor în vigoare.

Declar că toate rezultatele simulărilor, experimentelor și măsurărilor pe care le prezint ca fiind făcute de mine, precum și metodele prin care au fost obținute, sunt reale și provin din respectivele simulări, experimente și măsurători. Înțeleg că falsificarea datelor și rezultatelor constituie fraudă și se sancționează conform regulamentelor în vigoare.

București,

Absolvent Bălăucă Dragoș - Alexandru

(semnătura în original)

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

Subsemnatul BĂLĂUCĂ I. DRAGOȘ ALEXANDRU ²⁾

, posesor al
B.I./C.I./pașaport seria RD, nr. 619555, având CNP

1	9	1	0	5	2	3	4	4	0	0	1	0
---	---	---	---	---	---	---	---	---	---	---	---	---

,

am întocmit proiectul de absolvire/diplomă/lucrarea de disertație cu
tema:³⁾ MODELAREA SISTEMELOR RAIS ȘI MĂSURAREA
PERFORMANTELOR

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

Luând în considerare conținutul art. 143 din Legea Educației Naționale nr. 1/2011,
al. (4) „Îndrumătorii proiectelor de diplomă/lucrărilor de disertație răspund în solidar cu
autorii acestora de asigurarea originalității conținutului acestora” și al. (5) „Este
interzisă comercializarea de lucrări științifice în vederea facilitării falsificării de către
cumpărător a calității de autor al unui/unei proiect de diplomă/lucrări de disertație”,
declar pe proprie răspundere, că acest/această proiect/lucrare este original(ă), fiind
rezultatul propriei activități intelectuale, nu conține porțiuni plagiare, 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: 25.06.2014

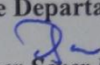
Semnătura:


(în original)

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

Universitatea "Politehnica" din București
Facultatea de Electronică, Telecomunicații și Tehnologia Informației
Departamentul : Electronică Aplicată și Ingineria Informației

Aprobat Director de Departament :


Prof. Dr. Inginer Sever PAȘCA

TEMA PROIECTULUI DE DIPLOMĂ
a studentului Bălăucă I. Dragoș Alexandru , grupa 443A

1. Titlul temei: **Modelarea Sistemelor RAID și măsurarea performanțelor**

2. Contribuția practică, originală a studentului va consta în (în afara părții de documentare):
Se vor prezenta structuri de RAID cunoscute și se va realiza în limbaj C, un simulator de
matrici RAID; se vor simula arhitecturi RAID diferite și se vor compara rezultatele cu datele
teoretice; se vor măsura parametri de performanță în diferite scenarii de funcționare normală
sau cu defecte; se vor simula proceduri de rebuild-ing și se vor compara timpii de execuție.

3. Proiectul se bazează pe cunoștințe dobândite în principal la următoarele discipline: Sisteme
de Operare, Rețele de Calculatoare, Inginerie Software, Arhitecturi de Rețele și Internet,
Arhitectura Sistemelor de Calcul

4. Realizarea practică/ proiectul rămân în proprietatea: *Universitatea Politehnica Bucuresti*

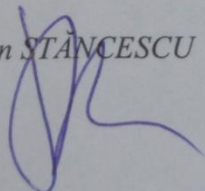
5. Proprietatea intelectuală asupra proiectului aparține: *Universitatea Politehnica Bucuresti*

6. Locul de desfășurare a activității: *Universitatea Politehnica Bucuresti*

7. Data eliberării temei: 12.12.2013

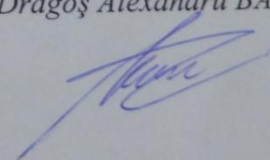
CONDUCĂTOR LUCRARE:

Conf. Dr. Ing. Ștefan STĂNCESCU



STUDENT:

Dragoș Alexandru BĂLĂUCĂ



Cuprins

NOTĂ INTRODUCATIVĂ.....	15
INTRODUCERE	17
CAPITOLUL 1	19
Specificații RAID.....	19
1.1 Noțiuni generale.....	19
1.2 Tehnologia RAID.....	20
1.3 Implementarea tehnologiei RAID	20
1.3.1 Implementare software	21
1.3.2 Implementare hardware	21
CAPITOLUL 2	23
Niveluri RAID Standard.....	23
2.1 RAID 0 – Disk Stripping.....	23
2.2 RAID 1 – Disk Mirroring	24
2.3 RAID 2 – Parallel Array with Error Corection Code.....	25
2.4 RAID 3 – Parallel Array With Parity.....	25
2.5 RAID 4 – Striped Array with Parity.....	26
2.6 RAID 5 – Striped Array with Distributed Parity	26
2.7 RAID 6 – Striped Array with Dual Distributed Parity.....	27
Niveluri RAID Hibrid	28
2.8 RAID 01 (0+1).....	28
2.9 RAID 10 (1+0).....	29
2.10 RAID 100 (10+0).....	29
2.11 RAID 51 (5+1).....	30
2.12 RAID 50 (5+0).....	30
Avantaje și dezavantaje	31
CAPITOLUL 3	33
Rebuilding.....	33
3.1 Noțiuni Generale	33
3.2 Functionarea matricilor RAID 5.....	35
3.3 Proceduri de rebuilding	37
CAPITOLUL 4	39
Simulatorul DiskSim	39

4.1	Prezentarea simulatorului DiskSim 4.0	39
4.2	Utilizarea mediului de simulare DiskSim 4.0	40
4.3	Fișierele de parametri	41
4.4	Planificarea accesului la datele de pe disc	42
4.5	Specificațiile componentelor subsistemului I/O	45
4.6	Parametrii pentru organizarea datelor în matrici de discuri	47
CAPITOLUL 5		49
	Simulatorul de arhitecturi RAID - DiskSim. Rezultate experimentale	49
5.1	Simulatorul pentru arhitecturi RAID	49
5.1.1	Interfața principală	49
5.1.2	Configurarea simulatorului	52
5.2	Rezultate experimentale.....	57
5.3	Concluzii.....	63
CAPITOLUL 6		65
	Analizatorul de arhitecturi RAID - Iometer. Rezultate experimentale.....	65
6.1	Prezentarea mediului de dezvoltare Iometer	65
6.2	Interfața programului.....	65
6.3	Rezultate experimentale.....	66
CAPITOLUL 7		73
	Considerente economice	73
ANEXE		77
	Anexa 1 – Managerul de text.....	77

Lista figurilor

Figura 2.1 – Repartiția blocurilor de date în RAID 0.....	18
Figura 2.2 – Repartiția blocurilor de date în RAID 1.....	19
Figura 2.3 – Repartiția blocurilor de date în RAID 2.....	20
Figura 2.4 – Repartiția blocurilor de date în RAID 4.....	21
Figura 2.5a – Repartiția blocurilor de date în RAID 5.....	22
Figura 2.5b – Repartiția blocurilor de date în RAID 6.....	22
Figura 2.6 – Repartiția blocurilor de date în RAID 0+1.....	23
Figura 2.7a – Repartiția blocurilor de date în RAID 1+0.....	24
Figura 2.7b – Repartiția blocurilor de date în RAID 10+0.....	24
Figura 2.8 – Repartiția blocurilor de date în RAID 5+0.....	25
Figura 3.2 – Modelul PCM de procesare a cererilor.....	29
Figura 3.3 – Modelul VSM de procesare a cererilor.....	30
Figura 3.4 – Rotația parității într-o matrice RAID 5.....	31
Figura 4.5 – Geometria fizică a harddisk-ului.....	32
Figura 5.1 – Interfața grafică a aplicației.....	43
Figura 5.2 – Generatorul de Workload-uri.....	44
Figura 5.3 – Formatul fișierului workload.....	45
Figura 5.5 – Crearea fișierului workload – cod sursă.....	47
Figura 5.6a – genNorm2 – cod sursă.....	48
Figura 5.6b – genUnif – cod sursă.....	48
Figura 5.7a – Legarea butoanelor din interfața grafică.....	49
Figura 5.7b – Legarea butoanelor din interfața grafică.....	49
Figura 5.8a – Mainframe.java – cod sursă.....	50
Figura 5.8b – Mainframe.java – cod sursă.....	50
Figura 5.10 – Latența rotațională medie – FCFS.....	52
Figura 5.11a – Timpul mediu de căutare pe disc - FCFS.....	53
Figura 5.11b – Timpul mediu de răspuns - FCFS.....	53
Figura 5.12 – Timpul mediu de acces la disc - FCFS.....	54
Figura 5.13a – Latența rotațională medie – SSTF LBN.....	55
Figura 5.13b – Timpul mediu de căutare pe disc – SSTF LBN.....	55
Figura 5.14a – Timpul mediu de răspuns – SSTF LBN.....	56
Figura 5.14b – Timpul mediu de acces la disc – SSTF LBN.....	56
Figura 6.2 – IOMETER – interfața grafică.....	60
Figura 6.3a – Analiză RAID5 – 1worker – All in One.....	61
Figura 6.3b – Analiză RAID5 – 1worker – All in One – Recovery Mode.....	61
Figura 6.4 – Analiză RAID5 – 1worker – All in One – Recovery Mode.....	62
Figura 6.5a – RAID5 în timpul procesului de rebuilding.....	63
Figura 6.5b – Analiză RAID5 – 1worker – All in One – Rebuilding.....	63
Figura 6.6 – Analiză RAID5 – 4workeri – All in One – Rebuilding.....	64
Figura 6.7 – Analiză RAID5 – 1worker – All in One – Rebuilding Low Prio.....	65
Figura 6.8 – Analiză RAID5 – 1worker – All in One – Rebuilding High Prio.....	66

NOTĂ INTRODUCȚIVĂ

Pentru lucrarea de licență am ales tema “Modelarea sistemelor RAID și măsurarea performanțelor” deoarece în contextul evoluției tehnologiei soluțiile de stocare reprezintă un domeniu de interes foarte important și intens dezbătut.

În zilele noastre, când accesul la informație este absolut necesar la orice oră a zilei, trebuie acordată o atenție deosebită sistemelor de stocare a datelor.

Metodologia lucrării constă în prezentarea teoretică a principalelor tipuri de sisteme profesionale de stocare, caracteristicile de bază ale acestora, avantajele și dezavantajele lor, precum și modul de funcționare al acestora. În continuare se vor alege principalele 4 arhitecturi RAID folosite în viața de zi cu zi și se vor analiza folosind mediul de simulare DiskSim.

În urma rezultatelor obținute în mediu de simulare DiskSim se vor alcătui rapoarte de performanță și grafice comparative între aceste arhitecturi.

Pentru a obține și o privire asupra unei configurații reale aflată în folosință la momentul actual se vor realiza o serie de teste asupra unei matrice RAID 5 și rezultatele vor fi prezentate ca o completare la setul de rezultate obținut pentru simulările realizate cu DiskSim.

Contribuția practică a studentului constă în realizarea mediului de simulare prin recompilarea simulatorului DiskSim și înglobarea acestuia într-o aplicație de sine stătătoare, realizarea unui software de extragere și interpretare a rezultatelor precum și în manipularea cu succes a unor software-uri profesionale de analiză mediilor de stocare.

INTRODUCERE

RAID (Redundant Arrays of Inexpensive Disks sau Redundant Arrays of Independent Disks) este o tehnologie dezvoltată pentru utilizarea simultană a două sau mai multe unități HDD într-o configurație (matrice) în scopul obținerii de performanțe crescute alături de o creștere a nivelului de siguranță a datelor.

În ultimii 25 de ani, standardul RAID s-a schimbat pentru a deveni o caracteristică indispensabilă pentru metodele de stocare ale claselor enterprise, dar în același timp și o variantă de stocare orientată către consumator. RAID imprastie datele pe mai multe discuri dure pentru a crește toleranța la erori și pentru a îmbunătăți performanța disc-ului.

RAID trebuie privit ca un instrument ce furnizează diferite nivele de stocare și care crește performanța și disponibilitatea sistemelor clasice. Unitățile de disc continuă să devină mai fiabile și cu capacități mai mari. Cu toate acestea, din ce în ce mai multe date sunt stocate, expunându-le la riscuri atunci când unul dintre discuri se defectează. În plus, față de aspectul de disponibilitate a RAID și tehnologiile de protecție a datelor, inclusiv oglindire locală și la distanță sau de replicare pentru alte sisteme de stocare, RAID oferă, de asemenea, îmbunătățiri ale performanței pentru a ajuta la restabilirea balanței între capacitate și rezultate excepționale. ^[1]

RAID reprezintă un acronim pentru *Redundant Array of Independent Disks* - matrice redundantă de discuri independente și este realizat prin combinarea mai multor discuri într-o unitate logică, unde datele sunt distribuite în unități în diferite moduri numite "nivele RAID".

Tehnologia RAID se împarte în mai multe scheme de stocare, care pot diviza și reproduce datele între mai multe discuri fizice. Discurile fizice funcționează ca un singur disc într-o matrice RAID ce este accesată de sistemul de operare. Diferitele scheme sau arhitecturi sunt numerotate cu o cifră după cuvântul RAID (spre exemplu: RAID 0, RAID 1 etc.). Fiecare schemă furnizează un raport diferit între două obiective principale: creșterea fiabilității datelor și creșterea performanței intrare/ieșire. ^[2]

CAPITOLUL 1

Specificații RAID

1.1 Noțiuni generale

Norman Ken Ouchi reușește să obțină în anul 1978 o licență ce se intitulează „Sistem pentru recuperarea datelor stocate de pe unități de memorie”. Conținutul licenței sale descria ceea ce mai târziu urma să fie cunoscut sub termenul de “RAID 5 cu scrieri în fâșii pline”. Această licență, din 1978, avea să inoveze prin oglindirea și/sau duplicitatea discurilor (mai târziu denumită RAID1) precum și protecția datelor prin paritate dedicată (ulterior denumită RAID 4).

Cu toate acestea, ideea de RAID a fost introdusă ca atare pentru prima dată în 1987, de către David A. Petterson, Garth A. Gibson și Randy Katz, o echipă de cercetători de la Universitatea Berkley din California. Tehnologia RAID a fost prezentată la acea vreme ca fiind o metodă ieftină și eficientă de backup. Aceștia au studiat posibilitatea utilizării a două sau mai multe unități ca una singură pentru sistemul gazdă și au publicat o lucrare intitulată: „Un caz de matrice redundantă de discuri ieftine (RAID)”. Lucrarea lor specifica un număr de niveluri RAID, fiecare având avantaje și dezavantaje teoretice.

De-a lungul anilor au apărut diverse implementări ale conceptului RAID. Singurul aspect care a rămas la fel este numerotarea. Majoritatea diferă de nivelul original RAID. Acest lucru poate crea confuzie, deoarece, de exemplu, una dintre implementările RAID 5 poate fi complet diferită de cealaltă. RAID 3 și RAID 4 sunt adesea confundate și chiar folosite interschimbat.

În lucrarea „*Un caz de matrice redundantă de discuri ieftine (RAID)*” autorii defineau formal nivelele RAID de la 1 la 5 în secțiunile 7 - 11:

- „**Primul nivel RAID:** Discuri Oglindite”
- „**Nivelul doi RAID:** Coduri Hamming pentru Corectarea Erorilor”
- „**Nivelul trei RAID:** Un Singur Disc De Verificare Per Grup”
- „**Nivelul patru RAID:** Citiri și Scrieri Independente”
- „**Nivelul cinci RAID:** Date împărțite/paritate pentru toate discurile (nu este un disc unic de redundanță)”
- „**Nivelul șase RAID:** Redundanță P+Q”. Aici matricea RAID necesită accesarea a șase discuri datorită necesității de a reînnoi ambele informații: P și Q
- „**Nivelul zece RAID:** striped mirrors (oglinzi întrețesute)”. Termenul este acum folosit pentru a exprima combinația dintre RAID 0 (întrețesut) și RAID 1 (oglindit).^[3]

1.2 Tehnologia RAID

O singură unitate logică formată din mai multe hard discuri fizice ce folosește o componentă hardware sau o aplicație software poartă numele de tehnologie RAID.

Principale tipuri în RAID sunt:

- **mirroring** (*oglindirea*) – se bazează pe copierea aceluiasi set de date pe mai multe discuri. Oglindirea poate crește viteza de citire a datelor deoarece sistemul poate accesa date diferite de pe cele două discuri. În aceeași timp însă, scrierea va fi mai lentă dacă sistemul insistă ca ambele discuri să confirme corectitudinea datelor scrise.
- **striping** (*întreșute*) – constă în împărțirea datelor pe mai multe discuri; Formatul întreșut este folosit în mod special pentru mărirea performanței, deoarece citirea secvențelor de date se face de pe mai multe discuri în mod simultan.
- **error correction** (*cu corectarea erorilor*) - unde discurile redundante de verificare stochează datele pentru a fi detectate și corectate eventualele erori. Verificarea erorilor în mod obișnuit va încetini sistemul deoarece datele vor fi citite din mai multe locații și apoi comparate. ^[4]

Diferitele configurații afectează în mod diferit stabilitatea și performanța (viteza de acces). Folosirea mai multor discuri în paralel crește probabilitatea ca unul dintre ele să se defecteze; dar utilizând funcții automate de detectare și corectare a erorilor, sistemul poate deveni mai stabil și poate repara în mod automat („din mers”) anumite erori. Gamele de discuri moderne oferă posibilitatea de a alege și a schimba configurația RAID după dorință.

Sistemele RAID au fost proiectate să ruleze în continuare chiar și în caz de defectare completă a unui disc— discurile pot fi schimbate „la cald” iar datele pot fi recuperate în mod automat, în timp ce sistemul rulează în continuare (eventual un pic mai lent, până la terminarea recuperării datelor). Prin comparație, sistemele de discuri normale trebuie oprite până când datele sunt recuperate.

RAID este adesea folosit în sistemele ce necesită o rată de acces cât mai ridicată și unde este important ca sistemul să ruleze cât mai multă vreme cu puțință. În general, RAID este folosit la servere, dar poate fi folosit și în cazul stațiilor de lucru (workstation). ^[3]

1.3 Implementarea tehnologiei RAID

RAID combină hard discuri fizice într-o singură unitate logică folosind o component hardware sau o aplicație software. Soluțiile hardware prezintă sistemului RAID atașat ca un singur hard disc, fără ca sistemul de operare să cunoască arhitectura fizică. Soluțiile software sunt implementate în sistemul de operare, dar aplicațiile utilizează arhitectura RAID ca o singură unitate.

Distribuția de date între mai multe unități poate fi realizată atât hardware cât și software. Pe lângă aceste două modalități există și un mod hibrid, format atât dintr-o parte hardware, cât și dintr-o parte software.

1.3.1 Implementare software

Majoritatea sistemelor de operare furnizează implementări software. Acest layer se situează deasupra driverelor discurilor și oferă un strat de abstractizare între discurile logice și discurile fizice. În mod normal RAID-ul software se limitează la RAID 0, 1 și 5. În sistemele de operare multi-thread (Linux, Mac OS X, Novell, Windows NT, 2000..) sistemul de operare poate executa mai multe procese I/O (permite scrieri și citiri multiple). Această caracteristică permite, în plus, realizarea RAID 0/1 în implementare software cu toate că nimeni nu folosește această practică. Acest lucru se realizează deoarece calcularea parității (datele redundante) necesită o putere de procesare ridicată. Implementările software necesită timpi reduși de prelucrare și consumă foarte puține resurse. În anul 2007 calculatoarele personale, obișnuite, ofereau puterea de calcul necesară crescând folosirea procesorului cu doar un procent în plus. Ele erau eficiente ca și cost până într-un punct neoferind însă același grad de performanță ca și RAID-ul bazat pe hardware. RAID-ul software folosește resursele unui server legat la spațiul de stocare, CPU-ul acestuia fiind doar partial ocupat cu procesarea datelor. În cazul unei defecțiuni a serverului, datele aflate în spațiul de stocare nu pot fi accesate atâta timp cât defecțiunea persistă. În hardware matricea de HDD-uri poate fi atașată ușor altei gazde astfel timpul de recuperare se reduce simțitor. Totuși acest sistem este la fel de eficient ca și cel software în cazul în care nu există un sistem de backup al hostului pe care să se facă trecerea rapid.

Implementarea software permite și crearea de matrici RAID din partiții ale unui HDD.

1.3.2 Implementare hardware

RAID-ul hardware solicită minim un controller RAID. Într-un calculator personal acest controller poate consta într-un card PCI sau poate fi inclus în chipset-ul plăcii de bază. În aplicațiile industriale controllerul lucrează ca și unitate separată. HDD-urile pot fi SATA, IDE/ATA, SSA, SCSI, canal optic sau combinații ale acestora. Sistemul ce utilizează acest sistem de stocare poate fi conectat în mod direct la controller sau printr-o rețea SAN (StorageArea Network). SAN-ul este o rețea ce utilizează un protocol de comunicare asemănător cu SCSI. Acest protocol se poate suprapune peste o rețea Ethernet. Controllerul se ocupa de managementul discurilor și calculează datele redundante caracteristice nivelului RAID ales. Majoritatea acestor sisteme au implementată o memorie cache nevolatilă, ce îmbunătățește performanțele. Implementarea hardware prezintă performanțe garantate, fără utilizarea procesorului gazdei, controllerul prezentând sistemului de operare un simplu disc logic. De asemenea, sistemele industriale prezintă suport pentru "hot swapping" (schimbarea HDD-urilor defecte fără a fi necesară oprirea sistemului).

1.3.3 Implementare hibridă

Sistemele hibride RAID si-au facut apariția o dată cu introducerea controllerelor ieftine RAID implementate în controlere HDD ca și extensii de BIOS reprezentate de drivere de sistem. Din păcate aceste controllere efectuează toate calculele necesare în regim software. Ele adună majoritatea dezavantajelor implementării software și hardware. Ca și implementările hardware ele sunt proprietate unui fabricant de controllere RAID și nu se pot face combinații între mai multe controllere. Avantajele constau în abilitatea de a boot-a de pe discul logic și integrarea mai strânsă cu driverul device-ului ce oferă o manipulare mai bună a erorilor.

CAPITOLUL 2

Niveluri RAID Standard

Inițial au fost concepute 5 niveluri RAID, dar pe măsură ce tehnica a evoluat au apărut mai multe variații ale acestora, cum ar fi nivelurile hibrid și mai multe niveluri non-standard. Acestea niveluri RAID (inclusiv formatele de date asociate lor) sunt standardizate de către SNIA în standardul DDF – *Common RAID Disk Drive Format*.

Diferitele niveluri RAID folosesc unul sau mai multe dintre tipurile enumerate la punctul 1.2 în funcție de cerințele sistemului. Principalele obiective ale arhitecturilor RAID sunt: mărirea siguranței datelor și creșterea vitezei de acces. Stabilitatea și performanța sunt afectate în mod diferit de către diferitele configurații ce pot fi folosite.

2.1 RAID 0 – Disk Stripping

RAID 0 nu este o arhitectură RAID în adevăratul sens al cuvântului deoarece nu asigură nicio redundanță a datelor. RAID 0 este folosit strict pentru maximizarea performanțelor în lucrul cu harddisk-urile. După cum spune și denumirea, datele sunt întrețesute (stripping) în mod secvențial pe mai multe discuri, ce sunt tratate ca un singur disc (sau volum) virtual. De obicei 4 discuri formează un volum. Implementările RAID 0 divizează volumele în felii și scriu datele în felii consecutive, localizate fizic pe fiecare disc din cadrul ariei de discuri.

Operațiile de scriere și citire au loc în paralel fiind executate pe toate discurile aflate în volum. Mărimea feliilor este definită de către utilizator și poate fi de 512 sectoare, spre exemplu, dimensiunea uzuală a sectoarelor putând fi de 512 octeți. Utilizând această tehnică de întrețesere se obține o rată de transfer a datelor foarte ridicată, îmbunătățindu-se în felul acesta și performanța efectivă a sistemului.

În eventualitatea defectării chiar și a unui singur disc al ariei, RAID 0 nu oferă redundanță datelor. În lipsa posibilității de regenerare a discului, datele ce se găsesc pe acesta se vor pierde.

Din acest motiv, RAID 0 nu este folosit în aplicațiile de înaltă disponibilitate. ^{[5][6]}

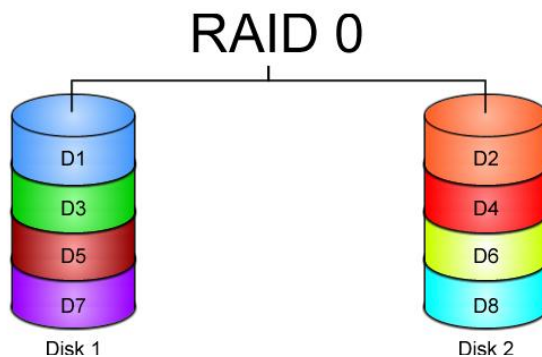


Fig. 2.1

Repartiția blocurilor de date în RAID 0

2.2 RAID 1 – Disk Mirroring

RAID 1 asigură redundanța datelor prin oglindirea acestora (mirroring). Metoda aceasta creează două sau mai multe copii ale aceluiași volum de date și le plasează pe discuri separate. Fiecare copie este o imagine oglindită a celeilalte.

RAID 1 poate scădea performanța generală a sistemului deoarece trebuie să efectueze tot atâtea scrieri câte copii. Performanța citirilor este însă mult îmbunătățită deoarece cererile sunt trimise simultan controlerelor de discuri ale tuturor copiilor. Discul care răspunde primul cererii de citire va returna datele sistemului.

Această metodă este mai puțin sigură comparativ cu scrierea secvențială, dar realizează o creștere semnificativă a performanțelor la scriere. RAID 1 este cea mai scumpă implementare deoarece asigură redundanța datelor în proporție de 100% și totodată asigură cea mai mare disponibilitate a datelor prin utilizarea celui mai mic număr de discuri din configurație.

Anumite sisteme de operare de tip UNIX ale unor sisteme de calcul tolerante la defectări oferă prin intermediul unei componente software numită “*Logical Volume Manager*” un alt tip de RAID 1. Această componentă crește performanța sistemului oferind posibilitatea scrierii în paralel pe toate copiile, fără a se mai aștepta ca scrierea pe o copie să se efectueze cu succes și apoi să se altereze conținutul altei copii. Comparativ cu scrierea secvențială, această metodă este mai puțin sigură, dar oferă o creștere semnificativă a vitezei de scriere. Întrucât RAID 1 asigură redundanța datelor în proporție de 100%, ea este cea mai scumpă implementare dintre toate tipurile de arii de discuri, dar asigură și cea mai mare disponibilitate a datelor prin utilizarea celui mai mic număr posibil de discuri în configurație. ^{[5][6]}

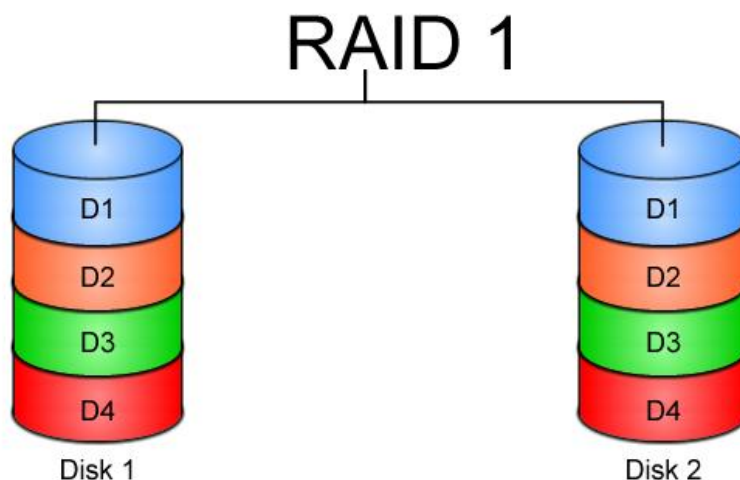


Fig. 2.2
Repartiția blocurilor de date în RAID 1

2.3 RAID 2 – Parallel Array with Error Corection Code

Acest model presupune că discurile în paralel distribuie datele pe mai multe discuri efectuând o întrețesere la nivel de bit, iar în cadrul ariei de discuri sunt utilizate și discuri de verificare. Citirea și scrierea datelor se face utilizând tehnici de codificare bazate pe coduri corectoare de erori de tip Hamming pentru a asigura detecția și corecția erorilor.

Din păcate RAID 2 vine cu mai multe dezavantaje decât avantaje. Printre acestea trebuie amintit faptul că utilizarea tehnicilor de codare bazate pe coduri corectoare de erori de tip Hamming necesită utilizarea unor grupuri mari de discuri pentru a asigura consistența (ex: 4 discuri de verificare pentru 10 discuri de date, 5 discuri de verificare pentru 25 de discuri de date). Din cauză că această tehnică de codificare este foarte complexă și scump de realizat, RAID 2 nu prezintă un real interes pentru mediul comercial.

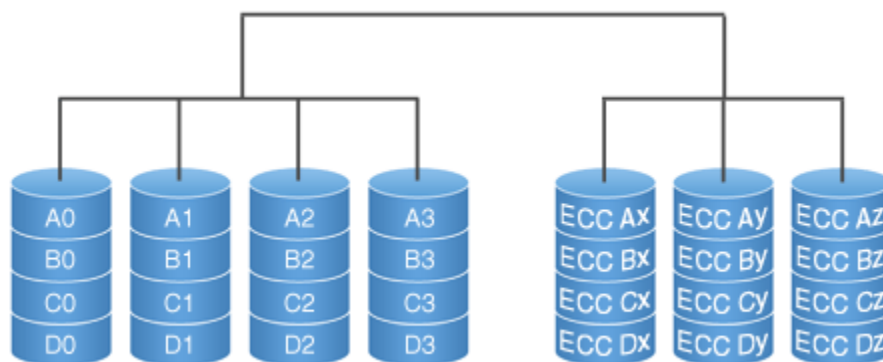


Fig. 2.3

Repartiția blocurilor de date în RAID 2

2.4 RAID 3 – Parallel Array With Parity

RAID 3 folosește doar un singur disk dedicat pentru memorarea informațiilor de paritate, celelalte discuri din cadrul ariei de discuri fiind folosite pentru memorarea informației utile ce este împrăștiată pe toate discurile ariei. (exact ca și în cazul RAID 2).

Acesta efectuează un SAU EXCLUSIV asupra datelor și nu un cod corector de erori.

Cantitatea minimă de date care este scrisă sau citită într-o operație de intrare sau de ieșire este egală cu numărul de discuri din arie înmulțit cu numărul de octeți ai sectorului discului. Aceasta este cunoscută ca unitate de transfer. La un moment de timp dat nu poate fi activă decât o singură cerere de intrare sau de ieșire deoarece capetele de acces se mișcă în paralel în aria de discuri. Acest lucru ne asigură o rată bună de transfer a datelor atunci când datele accesate sunt dispuse secvențial pe disc, dar face ca RAID 3 să nu fie suficient de practic pentru aplicații cu date dispuse aleator în aria de discuri. O rată de transfer foarte bună se obține atunci când unitatea de transfer are aceeași mărime cu dimensiunea blocurilor de date care se scriu sau se citesc. Cu cât blocurile transferate sunt de dimensiuni mai mici, cu atât rata de transfer scade.

RAID 3 oferă și posibilitatea înlocuirii la cald a unui disc și recuperarea datelor. ^{[7][8][9]}

2.5 RAID 4 – Striped Array with Parity

În cazul RAID 4 datele sunt distribuite la fel ca și în cazul RAID 3 pe mai multe discuri de date și pe un disc de paritate. Diferența constă în faptul că un bloc de date se află pe un singur disc, astfel încât citirea se face accesând un singur disc și nu pe ambele. Astfel aria de discuri poate prelucra mai multe cereri deodată întrucât discurile nu se mai utilizează în paralel.

Discul de paritate devine un element de scădere a performanței deoarece este implicat în toate operațiile de scriere. Acest lucru se întâmplă deoarece paritatea, pentru fiecare operație de scriere, se calculează având în vedere și informația de pe celelalte discuri, în sectoarele de date asociate. Această procedură face ca RAID 4 să fie considerat nepractic.^[9]

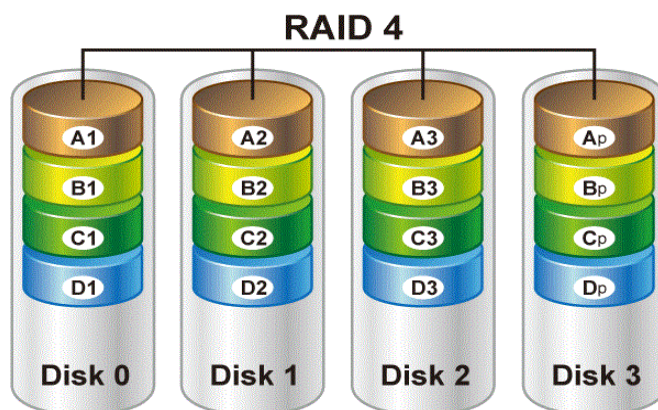


Fig. 2.4

Repartiția blocurilor de date și de paritate în RAID 4

2.6 RAID 5 – Striped Array with Distributed Parity

În cazul RAID 5 datele sunt distribuite pe mai multe discuri de date și pe un disc de paritate. RAID 5 nu are însă un disc dedicat special numai pentru informația de paritate, ci distribuie informațiile de tip date și informațiile de tip paritate pe toate discurile ariei.

RAID 5 permite mai multe accese concurente la dispozitivele din cadrul ariei de discuri, fiind satisfăcute astfel multiple cereri concurente de intrare sau ieșire. Astfel, performanța transferurilor pentru aceste arii de discuri crește semnificativ, acestea devenind cele mai potrivite și utile în cazul operării cu blocuri de dimensiune mică, sau cu datele dispuse în mod aleator pe discuri.

Ciclul de scriere se realizează în 3 pași:

- Pasul 1 – Este citită paritatea existentă și data ce urmează a fi schimbată
- Pasul 2 – Vechea dată din valoarea parității existente este scăzută și apoi se memorează paritatea corectă într-un buffer
- Pasul 3 – Este calculată noua paritate pe baza valorii din buffer și a valorii noii date de memorat și se scriu pe discurile corespunzătoare noua dată și paritatea actualizată.

Implementările software ale RAID 5 sunt suportate de către multe sisteme moderne de operare (Linux, Windows Server etc.). Procesorul efectuează toate calculele pentru paritate, iar în cazul unei defecțiuni, operațiunile de intrare sau de ieșire sunt de până la de 3 ori mai consumatoare de resurse. Având un cost relativ scăzut per GB utili și o performanță medie foarte bună, RAID 5 este utilizat cu succes pentru toate tipurile de transfer (servele pentru baze de date raționale, servele de fișiere) ^{[1][3][9][10]}

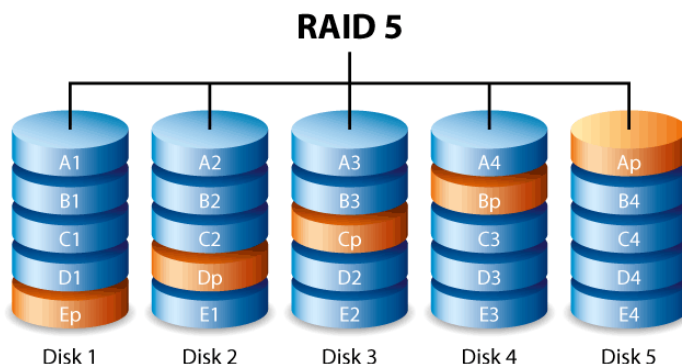


Fig. 2.5a

Repartiția blocurilor de date și de paritate în RAID 5

2.7 RAID 6 – Striped Array with Dual Distributed Parity

Ceea ce diferențiază nivelul RAID 6 de nivelul RAID 5 este faptul că se calculează și se stochează paritate dublă pentru date. În consecință sunt necesare minim 4 discuri, performanța scăzând datorită calculului dublu de paritate. Redundanța crește însă, aria suportând două discuri defecte. Reconstrucția datelor poate fi un mare consumator de resurse, performanța fiind puternic afectată, dar redundanța dublă compensează prin posibilitatea de amânare a procesului de rebuilding pentru un plus de viteză.

Pentru a crea nivele etajate de RAID, oricare nivel poate fi combinat cu alt nivel, dar cel mai des sunt folosite nivelele 0 și 1. ^{[11][12]}

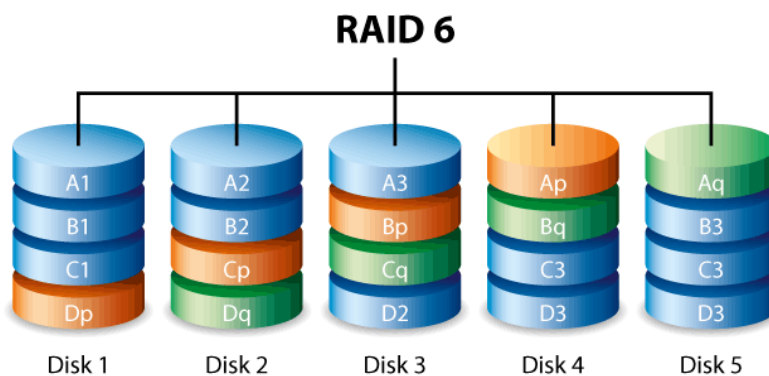


Fig. 2.5b

Repartiția blocurilor de date și de paritate în RAID 6

Niveluri RAID Hibrid

Nivelurile RAID hibride combină două sau mai multe niveluri standard de RAID pentru a obține performanță și redundanță suplimentară. Atunci când combinăm nivelurile RAID, un tip de RAID care oferă redundanță este, de obicei, unit cu un tip de RAID 0 ce mărește performanța. Cu aceste configurații, este de preferat să existe RAID 0 pe partea de sus și matrici redundante în partea de jos, deoarece mai puține discuri vor trebui să fie regenerate atunci când un disc se defectează. Prin urmare, RAID 1+0 este preferabil lui RAID 0+1, dar avantajele administrative ale "oglinzii împărțite" de RAID 1 ar fi pierdute. Trebuie remarcat, cu toate acestea, faptul că aspectul blocurilor pe disc pentru configurațiile RAID 1+0 și RAID 0+1 sunt oarecum identice, deci aceste limitări apar doar în software. [13]

2.8 RAID 01 (0+1)

Acest nivel dispune minim 4 discuri grupate în fâșii (oglindite), înlătură neajunsurile și îmbunătățește performanțele, însă crește complexitatea. RAID 0+1 creează un al doilea set de fâșii oglindite cu un set primar de fâșii de discuri. Matricea continuă să funcționeze cu unul sau mai multe unități de stocare afectate în aceeași oglindă stabilită, dar dacă driverul cedează de pe ambele părți ale oglinzii, datele de pe sistemul RAID sunt pierdute. [2][3]

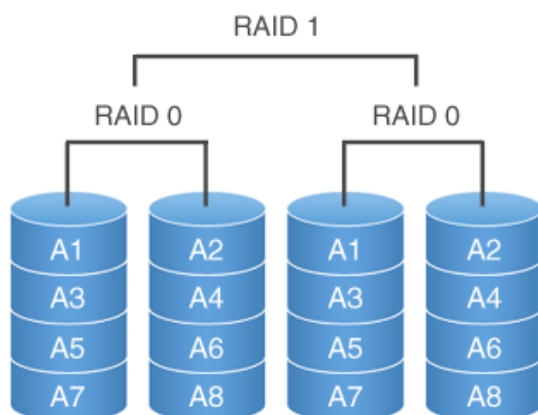


Fig. 2.6
Repartiția blocurilor de date: RAID 0+1

2.9 RAID 10 (1+0)

Acest nivel dispune minim 4 discuri grupate în fâșii (oglindite), înlătură neajunsurile și îmbunătățește performanțele, însă crește complexitatea. Diferența față de RAID 0+1 este că RAID 1+0 creează discuri cu date întrepesute de la o serie de unități de stocare în oglindă. Într-o situație de eșuare, RAID 1+0 funcționează mai bine, deoarece toate celelalte discuri continuă să fie utilizate. Matricea poate susține mai multe discuri pierdute, atât timp cât oglinda nu pierde toate unitățile sale. ^{[2][3]}

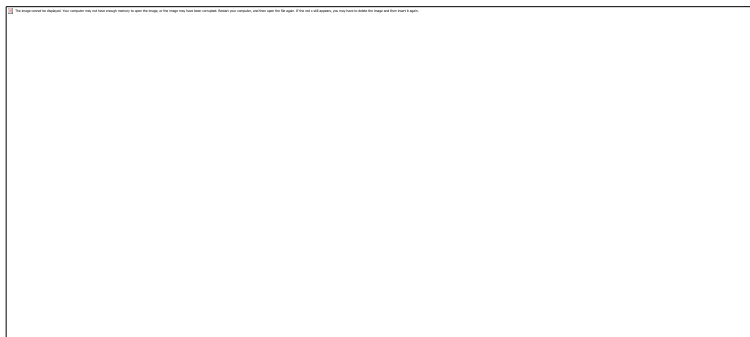


Fig. 2.7a
Repartiția blocurilor de date: RAID 1+0

2.10 RAID 100 (10+0)

Un raid 100, numit uneori și RAID 10+0, este o felie de tipuri RAID 10. Din punct de vedere logic el este de fapt o matrice RAID 10 implementată cu ajutorul software-ului RAID 0 peste hardware-ul de RAID 10.

Principalul avantaj al RAID 100 față de un singur nivel RAID este răspândirea încărcării pe mai multe controllere, obținându-se astfel performanță mai bună la citiri aleatoare și atenuarea riscurilor pe matrice. Tocmai din acest motiv, RAID 100 reprezintă cea mai bună alegere când vine vorba de baze de date foarte mari, unde controllerele hardware de RAID limitează numărul de discuri fizice permise în fiecare matrice standard. ^{[2][3]}

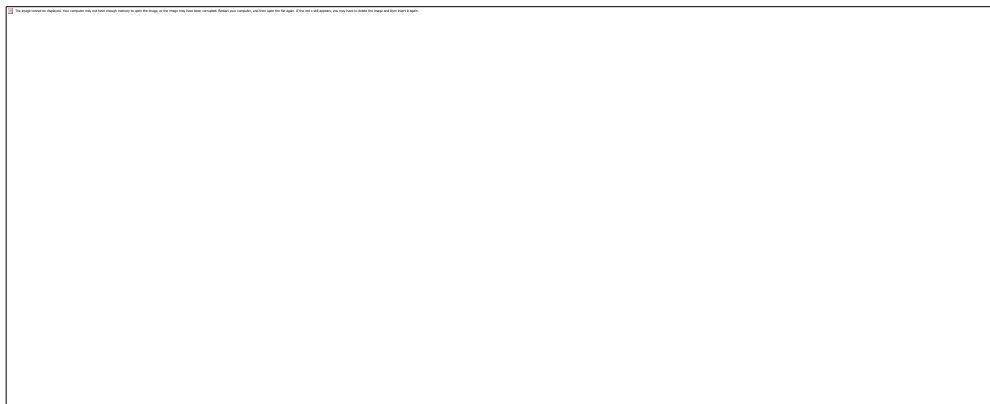


Fig. 2.7b
Repartiția blocurilor de date: RAID 10+0

2.11 RAID 51 (5+1)

Acest nivel hibrid este o matrice formată din două matrice RAID 5, care sunt oglindite una față de cealaltă. În general această configurație este folosită astfel încât fiecare set de RAID 5 se află pe un controller separat. Astfel, scrierile și citirile sunt echilibrate în ambele matrice RAID 5. Unele controllere suportă RAID 51 pe mai multe canale pentru a ține sincronizate părțile diferite. În această configurație, cele două seturi de RAID 5 nu știu că reprezintă oglindirea celuiilalt, iar RAID 1 nu știe că discurile sale de bază sunt RAID 5.

Această configurație poate susține eșecul oricărui disc din orice matrice, plus încă un disc suplimentar din cealaltă matrice înainte să existe pierderi de date. Spațiul maxim ce se poate găsi pe un RAID 51 este egal cu dimensiunea unui set individual RAID 5. ^{[13][2][3]}

2.12 RAID 50 (5+0)

RAID50 combină feliile de blocuri de la nivelul RAID0 cu paritatea distribuită de RAID5. Aceasta este de fapt o matrice RAID 0 întrețesută de-a lungul elementelor RAID 5. Pentru a realiza un nivel hibrid RAID 50 avem nevoie de cel puțin 6 discuri.

Avantajul acestei configurații este că ca pot eșua câte un drive din fiecare set, fără să se piardă date. Cu toate acestea, unitățile rămase în acel set devin un punct slab pentru întreaga matrice până când discul defect este înlocuit. Este suficient să se mai defecteze un singur disc (pe lângă cel inițial) și întreg sistemul cedează, datele stocate în întreaga matrice fiind pierdute. Timpul petrecut în procesul de rebuilding reprezintă o perioadă de vulnerabilitate a setului RAID. ^{[13][2][3]}

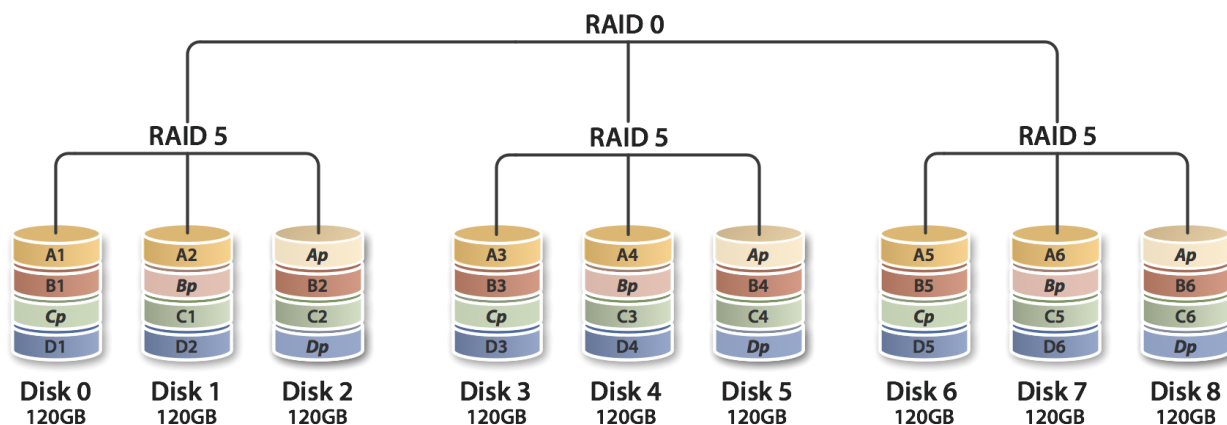


Fig. 2.8
Repartiția blocurilor de date și de paritate: RAID 5+0

Avantaje și dezavantaje

RAID	Avantaje	Dezavantaje
0	Ușor de proiectat și implementat Nu folosește calculi de paritate	Nu tolerează greșeli. Defectarea unui singur disc duce la pierderea tuturor datelor dintr-un șir.
1	Viteza de citire este de 2 ori mai mare Redundanță 100% Poate susține simultan mai multe defecțiuni În caz de defectare datele nu trebuie reconstruite ci pur și simplu copiate	Capacitatea efectivă de memorie este de fapt jumătate din capacitatea maximă totală Posibil să nu suporte schimbarea discului în timpul funcționării
2	Rate de transfer foarte ridicate	Rata de transfer este cel mult egală cu cea a unui singur disc
3	Viteză mare de citire și scriere pentru pachetele mari de date Disk-failure nu încetinește ratele de transfer	Tehnologia este complexă și mare consumatoare de resurse, deci realizarea software a acesteia nu se rentează
4	Transfer de date foarte mare la citire Eficiență ridicată datorită discurilor de paritate aflate în proporție mai mică decât a celor de date	Cea mai slabă rată de transfer la scriere În cazul defectării unui disc, recuperarea datelor este extreme de dificilă
5	Cea mai ridicată rată de transfer la citire Rată de transfer bună la scriere	Datele sunt dificil de recuperat în comparație cu nivelul RAID 1
6	Toleranță foarte mare la erori Poate susține mai multe defectări simultane	Schema de paritate este bidimensională deci trebuie implementate N+2 discuri
01	Înlătură neajunsurile și îmbunătățește performanțele	Crește complexitatea Dacă driverul cedează pe ambele părți ale oglinzii, datele sunt pierdute
10	Toleranță la erori. Poate susține mai multe defectări simultane, în anumite cazuri	Tehnologie foarte scumpă Performanța scade întrucât discurile se mișcă în paralel la adresa cerută
50	Un drive de la fiecare din seturile RAID5 poate eșua fără pierderi de date	Sistemul rămâne vulnerabil până la schimbarea discului căzut
51	Poate susține eșecul fiecărui disc din matrice plus încă un disc suplimentar Performanță bună dar costisitoare în același timp	Costul foarte mare al implementării din cauza discurilor identice ce trebuie folosite și a controllerelor de ultimă generație
100	Răspândirea încărcării de-a lungul mai multor controller Performanță mai bună la citiri aleatoare	Complexitate foarte ridicată Prețul foarte ridicat al implementării

CAPITOLUL 3

Rebuilding

3.1 Noțiuni Generale

Noțiunea de rebuilding se referă la reconstrucția sistematică a datelor de pe discul defect, pe un disc de rezervă. Aceasta procedură începe imediat ce discul de rezervă devine disponibil.

Matricile RAID folosesc discuri de rezervă pentru a minimiza pierderea datelor. Astfel după ce un disc se defectează, datele de pe acesta pot fi reconstruite imediat pe discul de rezervă (rezervă dedicată – “dedicated sparing”). Rezerva distribuită – “distributed sparing” și rezerva de paritate – “parity sparing” reprezintă două alternative de rezervă. În cazul rezervei de paritate sistemul pornește fără rezervă, dar dacă un disc pică, atunci grupurile de paritate aflate în matrici diferite sunt combinate pentru a forma matrici și grupuri de paritate mai mari.

O matrice RAID poate funcționa în unul din următoarele 3 moduri:

- **Modul Normal** – Discurile principale sunt operaționale, iar discul de rezervă poate fi atât funcțional cât și defect sau în reparație.
- **Modul Deteriorat** – Presupune deteriorarea unuia din cele $N+1$ discuri primare. În continuare, pentru a recrea blocul de date pierdut, cererea pentru citirea datelor de pe discul defect presupune accesul la blocurile corespunzătoare de pe discurile funcționale. O cerere de scriere pe discul eşuat rezultă în citirea celor N blocuri de date corespunzătoare pentru a calcula blocul de paritate (folosind noua versiune a blocului de date), urmată de o simplă scriere a blocului de paritate. Astfel, rata de cereri pentru discurile funcționale aproape se dublează.
- **Modul Rebuild** – Operația de rebuilding pornește după ce unul din discuri iese din funcțiune, nemai fiind capabil să reintre în modul normal de lucru al sistemului. În modul de rebuild sistemul citește trasee consecutive de date de la toate discurile funcționale în paralel, calculează conținuturile traseelor discului defect (prin aplicarea operației XOR asupra celor N trasee corespunzătoare din discurile funcționale), și scrie traseele reconstruite pe discul de rezervă. (În scrierea pe discul de rezervă nu luăm în considerare opțiunea de a verifica, conform căreia blocurile de date abia scrise sunt recitite pentru a verifica scrierea. Această opțiune aproape dublează încărcarea pe discul de rezervă, astfel încât timpul de rebuild este de așteptat să fie dominat de timpul de scriere pe acest disc când această opțiune este activă). Sistemul este predispus la pierderea datelor în cazul în care unul din discurile funcționale cedează înainte ca reconstrucția să fie completă.

Cea mai mică unitate de date ce poate fi reconstruită se numește Unitate de Rebuild (UR) și este de obicei o unitate întrețesută sau o parte a acesteia.

T_{rebuild} reprezintă timpul necesar reconstruirii discului defect, iar $T_{\text{răspuns}}$ reprezintă timpul de răspuns la cererile utilizatorului. Aceste două mărimi sunt principalele repere pentru măsurarea performanței sistemului.

Rebuildingul se împarte în 2 categorii: orientat pe fâșii și orientat pe disc. În cazul rebuildingului orientat pe disc se dedică un proces pentru fiecare disc care citește UR în mod asincron de pe discurile ce au supraviețuit. În rebuildingul orientat pe fâșii, fiecare UR este reconstruită printr-un proces dedicat care citește UR de pe discurile supraviețuitoare, execută între ele operația logică XOR și scrie UR rezultată pe discul de rezervă.

Rebuildingul orientat pe disc necesită un buffer mai mare față de rebuildingul orientat pe fâșii atunci când se consideră un număr mic de procese, însă rebuildingul orientat pe disc este mai performant.

Din punct de vedere al spațiilor tampon se poate afirma că există 2 categorii: spații tampon temporare și spații tampon dedicate scrierii.

Spațiile tampon temporare sunt dedicate citirilor de pe fiecare disc, în timp ce spațiile tampon dedicate scrierii sunt dedicate scrierii pe discul de rezervă.

Spațiile tampon sunt interschimbabile și au ca unitate de măsură tot o unitate de reconstrucție (UR). Imediat ce o UR este citită în spațial tampon, i se aplică o operație logică XOR cu unitățile UR corespunzătoare iar rezultatul este stocat în zona tampon de pe discul de rezervă. Spațiile tampon nu se epuizează niciodată deoarece operația XOR este foarte rapidă. Datorită acestui fapt, citirea unităților de UR de pe discuri, în timpul unui rebuild orientat pe disc, nu este oprită din cauza limitărilor de memorie. Totuși, spațiul tampon pentru discul de rezervă, are un efect asupra performanței de reconstrucție. Toate trimerile viitoare la spațiul tampon implică memoria tampon a discului de rezervă.

Cererile de rebuild pot fi procesate în două feluri: cu aceeași prioritate ca și cererile utilizatorului și la o prioritate mai mică decât cererile utilizatorului. **Modelul Clientului Permanent (PCM)** presupune o prioritate egală a cererilor de rebuild cu cererile de utilizator. În acest model o singură cerere este procesată la un moment dat într-un sistem de tip coadă de așteptare. Imediat ce o cerere a fost servită, îi vine rândul următoarei, iar la sfârșitul cozii de așteptare vine o alta nouă.

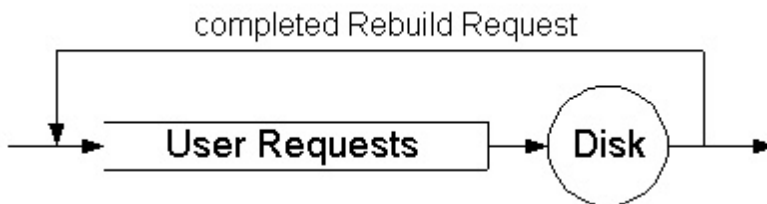


Fig. 3.2
Modelul PCM de procesare a cererilor

Modelul Serverului de Eliberare (VSM) reprezintă o alternativă mai bună pentru PCM, oferind timpi de răspuns și de reconstrucție mai mici. Acesta efectuează reconstrucția atunci când nu se află nicio cerere de utilizator în așteptare (adică imediat ce discul devine inactiv). Imediat ce sosește o cerere de utilizator reconstrucția este oprită.

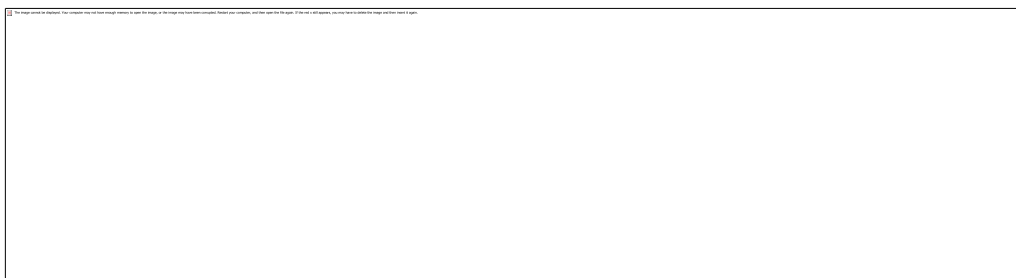


Fig. 3.3
Modelul VSM de procesare a cererilor

3.2 Funcționarea matricilor RAID 5

Un sistem RAID 5 este format din $N+1$ discuri primare care conțin blocuri de date, blocuri de paritate și un disc de rezervă dedicat. Un șir de biți asociați pe blocurile de date folosind operația logică XOR formează un bloc de paritate. Blocurile de paritate sunt distribuite uniform în discurile primare pentru a preveni o posibilă strângulare.

Pentru a scrie un bloc de date, matricile RAID 5 necesită 4 accesări ale discului pentru că paritatea trebuie alterată de fiecare dată când blocurile de date sunt modificate.

Cele 4 accesări ale discului sunt:

1. Citire date vechi
2. Citire paritate veche
3. Scrierea noilor date
4. Scrierea noii paritati

Înainte de a începe recuperarea datelor de pe o matrice RAID 5 trebuie să determinăm configurația și parametrii. Configurația unei matrice RAID 5 constă în cunoașterea următorilor parametrii:

- Cate discuri sunt in configurația RAID 5
- Care este secvența discurilor
- Ce valoare s-a folosit pentru mărimea blocurilor
- Care este pattern-ul de paritate folosit

Toți acești parametri pot fi determinați manual sau automat. Pentru a determina poziția parității trebuie să știm faptul că rotația acestora decurge după cum este ilustrat în figura următoare:

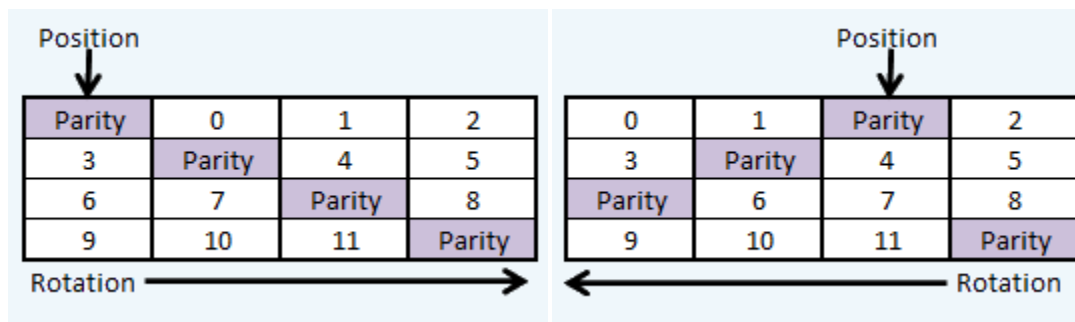


Fig. 3.4
Rotația parității

Pentru a determina poziția trebuie să privim blocurile de memorie din matrice. Blocul care arată cel mai puțin ca un bloc de date, este blocul de paritate.

Determinarea ordinii discurilor se realizează prin urmărirea fișierelor text foarte lungi. Se preferă acele fișiere care au și timestamps. Pentru a găsi astfel de fișiere pe discuri se folosesc tool-uri special, cum ar fi WinHex. Când s-a găsit un fragment dintr-un astfel de fișier pe unul din discuri, este nevoie să găsim discul ce conține următorul fragment și așa mai departe. În acest fel se poate determina ordinea discurilor, deși este imposibil de afirmat care disc a fost primul.

Pentru a determina care disc a fost primul, trebuie să folosim din nou un tool de vizualizare a discurilor și să căutam următoarele 2 lucruri, în funcție de tipul de RAID pe care îl avem:

- MBR-ul (Master Boot Record) – în cazul unui RAID realizat hardware
- Sectorul de boot – în cazul unui RAID realizat prin mijloace software

În cazul unui RAID realizat hardware, discul care conține MBR-ul este primul disc din acel RAID. *Master Boot Record* este un tip special de sector de boot aflat la începutul partițiilor unui element masiv de stocare (*mass storage device*). MBR-ul reține informații cu privire la organizarea partițiilor logice pe acel mediu de stocare. Partițiile logice conțin fișierele de sistem.

În cazul unui RAID realizat software, discul ce conține Sectorul de Boot la început este primul disc. ^[14]

3.3 Proceduri de rebuilding

Un hard poate întâmpina 2 tipuri de probleme, și anume:

- probleme logice: problemele apar la nivelul partițiilor și pot fi caracterizate prin: coruperea sistemului de fișiere din cauza atacului cu viruși, ștergerea accidentală a unui folder sau chiar a întregii partiții, formatarea accidentală și blocarea sistemului de operare prin intermediul sectoarelor bad. Recuperarea datelor în acest caz nu ridică mari probleme și presupune utilizarea unor soft-uri concepute special pentru a putea accesa datele corupte. Mai trebuie menționat faptul că recuperarea datelor cu ajutorul unui soft este valabilă atâta timp cât hard disk-ul este funcțional.
- probleme fizice: se produc atunci când cel puțin una dintre componentele hardului se defectează.

Problemele fizice includ:

- probleme de natură firmware: nu țin de partea electrică și pot fi definite ca fiind drept imposibilitatea de a accesa o informație aflată pe suprafața hardului utilă în citirea datelor existente pe acesta.
Recuperarea datelor se face cu echipament specializat dar și cu o vastă experiență în domeniu.
- probleme mecanice: apar datorită utilizării îndelungate a produsului, fluctuațiilor de curent electric sau a diferitelor probleme legate de fabricația acestuia. Când apare un defect mecanic următoarele componente pot fi afectate: stricarea totală sau parțială a capetelor de citire sau scriere, deteriorarea motorului, zgârierea suprafețelor.
- probleme electrice: apar mai ales datorită fluctuațiilor de curent. Numai cu ajutorul procesului de reconstrucție a zonelor afectate se poate ajunge la recuperarea datelor aparent pierdute. ^[15]

Sistemele RAID pot fi concepute să ruleze mai departe chiar și în caz de defectare completă a unui disc dur din RAID – discurile pot fi înlocuite „la cald” și datele recuperate automat, în timp ce sistemul rulează în continuare (eventual ceva mai lent, până la terminarea recuperării datelor).

În procesul de rebuild există 3 proceduri de copiere:

1. **Rebuild folosind procedura de bază** - Actualizările de blocuri de date sau paritate vizate pentru discul defect sunt utilizate pentru a actualiza blocurile discului de rezervă dacă blocul

este deja reconstituit. Pentru a reconstrui un bloc de date, cererile de citire adresate discurilor defecte conduc la citirea blocurilor corespunzătoare din toate discurile încă funcționale, chiar dacă blocul de date a fost deja recreat în prealabil pe discul de rezervă.

2. **Rebuild cu citire redirectionată** – În cazul procedurii de bază, blocurile de date sunt recreate. În această a doua procedură, blocurile de date reconstruite pe discul de rezervă sunt citite direct de pe el. Acest lucru ajută la scurtarea timpului de rebuild și la reducerea utilizării discurilor funcționale. Chiar dacă citirile redirectionate interferează cu scrierile de rebuild pe discul de rezervă, atâta timp cât discul de rezervă nu constituie o ștrangulare, încheierea procesului de rebuilding depinde în timp de citirea discurilor funcționale.
3. **Rebuild Piggy-backing pe un workload normal** – “Ideea aici este de a "captura" un bloc de date de pe discul defect, care este reconstruit datorită unei cereri de citire, ce a fost emisă ca parte din workload-ul normal. O implementare relativ simplă este posibilă presupunând doar o ușoară modificare a unui spațiu tampon convențional. Atunci când o citire de pe discul defect este satisfăcută de reconstrucția blocului, conținuturile blocului vor fi stocate într-un spațiu tampon. În acest moment copia tampon poate fi marcată ca fiind "modificată", adresa discului asociată cu pagina de tampon poate fi schimbată pentru a referi discul în așteptare, iar harta de bit pentru reconstrucție poate fi schimbată pentru a afișa această pagină ca fiind copiată. Înlocuirea normală a spațiului tampon va copia în cele din urmă conținutul la blocul corespunzător pe discul în așteptare. (La sfârșitul reconstrucției orice bloc rămas în cache poate fi aruncat la discul în așteptare). Piggy-backing și citirea redirectionată pot ajuta la reducerea încărcării pe discurile funcționale într-o matrice care are un disc defect. Măsura în care acest lucru poate aduce un beneficiu depinde de mai mulți factori. Pentru aceeași rată de reconstrucție și aceeași rată minimă de acces pentru workload-ul normal, atât citirea redirectionată cât și piggy-backing pot fi utilizate pentru a reduce sarcina pe discurile funcționale. Așadar, capacitatea de acces I/O fiind disponibilă pe discurile funcționale, poate fi folosită pentru a susține fie un workload mai mare în timpul reconstrucției sau un rebuild mai rapid.” [16]

“Comparația timpului de rebuild pentru diferite astfel de scheme se bazează pe utilizarea valorilor medii în estimarea timpului necesar de dispozitivul de ștrangulare (discurile funcționale sau discul de rezervă) pentru a finaliza procesul de rebuild. Un sistem RAID cu clustere (grupuri) are mărimea grupului $G \leq N + 1$ și o rată de declustering $\alpha = (G - 1) / N \leq 1$. Performanța sistemului în modul deteriorat este mai bună pentru valori mai mici ale lui α , din moment ce un număr mai mic de discuri va fi implicat în prelucrarea cererilor fork-join. Una din concluziile studiului care tratează acest aspect este că unele dintre optimizările propuse au ca rezultat o creștere, și nu o scădere a timpului de rebuild. Acest lucru se datorează faptului că studiul de simulare ia în considerare caracteristicile detaliate de disc, cum ar fi necesitatea de a invoca o urmărire înainte de unele accesări de disc.” [17]

CAPITOLUL 4

Simulatorul DiskSim

4.1 Prezentarea simulatorului DiskSim 4.0

Programul DiskSim este un software de simulare eficient, exact și extrem de configurabil al sistemelor cu discuri. Acest program a fost conceput pentru a fi de folos celor care lucrează în cercetare oferind detalii cu privire la numeroase aspecte din arhitectura subsistemelor de stocare. DiskSim (ajuns la versiunea 4.0) a fost creat de către John S. Bucy, Jiri Schindler, Steven W. Schlosser, Gregoiy R. Ganger, Bruce R. Worthington și Yale N. Patt la "University of Michigan". Această ultimă versiune (4.0) vine cu îmbunătățiri aduse variantei precedente (3.0), unele deficiențe fiind rezolvate, precum și unele librării fiind restructurate pentru o mai mare eficiență și pentru a face simulatorul mult mai ușor de utilizat. În particular, modulul de discuri poate simula în mare detaliu configurații moderne de discuri fiind recunoscut pentru multitudinea de detalii oferite precum și pentru fidelitatea acestora.

La nivel structural, simulatorul este compus din mai multe module ce permit configurarea diverselor componente ale unui sistem de stocare. Astfel, au fost făcute configurabile în simulator, discurile, controller-ul, buss-urile, cache-urile precum și componentele software care dirijează activitatea acestora: driverele, algoritmi de planificare a transferului, organizarea structurilor RAID etc. Simulatorul include de asemenea și un modul micro-electro-mecanic (*MEMS – based module*).

Simulatorul primește ca parametru de intrare un fișier cu înregistrări (*trace-uri*) ale activității pe magistrala I/O. Aceste *trace-uri* provin de la un sistem real sau pot fi create de un generator intern de sarcini de lucru (workload). Modulul de generare a trace-urilor (generatorul de workload) este foarte flexibil, permițând modelarea cu acuratețe a diverselor rate de sosire a cererilor. DiskSim 4.0 simulează și raportează numai aspectele legate de performanța sistemelor de stocare și nu modelează comportamentul restului de componente ale unui sistem de calcul sau interacțiunea dintre acestea și sistemul de stocare pe discuri. Cu toate că simulează aspectele legate de stocarea datelor pe discuri, DiskSim nu salvează sau reface efectiv datele provenite de la o cerere. Partea de cod a simulatorului este scrisă în limbajul C și necesită pentru compilare compilatorul GCC al sistemului de operare Linux.

Pentru alte aspecte legate de compilarea programului în Linux se va citi fișierul README situat în directorul cu fișiere de cod DiskSim de pe CD-ul anexat prezentei lucrări. Libparam unifică parametrii de intrare ai DiskSim-ului. Acest lucru face mai ușoară simularea utilizând fișierele de input ale modelului de disc dorit fără a mai copia un cod de intrare pentru a putea rula simulatorul. ^{[18][19]}

4.2 Utilizarea mediului de simulare DiskSim 4.0

Mediul de simulare DiskSim poate fi încărcat și compilat pe orice sistem ce rulează Linux, distribuția acestuia nefiind importantă. Se remarcă o ușurință mai mare la instalarea acestuia pe o distribuție bazată pe Debian.

Pentru a putea rula, DiskSim solicită în linia de comandă 5 parametri. Fișierul curent trebuie stabilit ca fiind “valid”.

DiskSim ***<parfile>*** ***<outfile>*** ***<tracetype>*** ***<tracefile>*** ***<synthgen>*** ***[par_override]***

unde:

- **diskSim** – este numele fișierului executabil
- **parfile** – reprezintă numele fișierului cu parametrii
- **outfile** – este numele fișierului de la ieșire. Dacă în locul acestui parametru se aplică ***“stdout”*** rezultatele vor fi afișate pe monitor
- **tracetype** – specifică tipul de trace care se află în fișierul de trace-uri (ascii, validate etc.)
- **tracefile** – este numele fișierului de trace-uri ce va fi folosit la intrare. Pentru ca în loc să folosim un fișier cu trace-uri, putem introduce datele de la tastatură folosind comanda ***“stdin”***
- **synthgen** – setează generatorul de workload-uri în stare activă sau pasivă (valoarea 0 dezactivează generatorul, în timp ce orice altă valoare îl activează).
- **par_override** – permite ca o serie de parametrii implicați ai simulatorului sau o parte din parametrii transmiși prin fișierul de parametri să fie înlocuiți cu valori transmise prin linia de comandă.

Sintaxa exactă este precizată mai jos:

<component> ***<parameter>*** ***<newvalue>***

unde:

- **component** - este numele unei componente ai cărei parametri se doresc a fi modificați;

- **parameter** - este un șir de caractere ce identifică parametrul care va fi înlocuit; Dacă numele conține spații libere el trebuie pus între ghilimele, astfel ca shell-ul să-l considere un singur argument. Pentru a referenția un parametru al unei subcomponente, cum ar fi de exemplu planificatorul unui disc, se folosește scrierea Scheduler parameter:
- **newvalue** - este noua valoare a modulului instanțiat.

4.3 Fișierele de parametri

Pentru a insera fișierul de parametrii, DiskSim folosește *libparam*. Libparam unifică parametrii de intrare ai DiskSim-ului. Într-un astfel de fișier cu parametrii găsim:

- blocuri – delimitate de accolade {}, ce constă într-un număr de asignări de tipul “name = value”
- instanțieri
- specificații topologice

În exemplu de mai jos este ilustrat un exemplu concret de linie de comandă:

```
Disksim parm.1B stdout ascii Jan6 0 "disc0" "Segment size fin blks0" 64 "disc*"
"Scheduler.'Scheduling policy" 2
```

Această comandă realizează următoarele acțiuni:

- citește parametrii inițiali din fișierul parms.1B. Specificațiile sunt stocate în fișiere separate și incluse apoi în fișierul principal de parametri, fiind de obicei foarte detaliate.
- exportă rezultatele pe monitor (din cauza comenzii stdout)
- citește din fișierul t.Jan6 fluxul de intrare ASCII
- nu generează activitate sintetică
- se alege planificatorul de disc (în acest caz 2 = Elevator SCAN) pentru discul instanțiat cu numele ‘disc0’ și pentru driverul acestuia (‘driver0’)

În fișierul parms.1B există câteva caracteristici ale discului cum ar fi: statistica blocurilor, specificațiile driverelor, magistralelor și ale controllerelor precum și instanțierea acestora sub diferite denumiri. Tot în cadrul acestui fișier, prin comanda “source” este ales fișierul .discspecs și este descrisă topologia sistemului. În fișierul .discspecs este inclus fișierul .model tot prin comanda „source” apoi sunt enumerați parametrii ce se referă la organizarea internă a harddisk-ului (planificator de disc, condiții de scriere / citire etc.). În fișierul .model este descrisă partea cea mai de jos a hard disc-ului

(suprafețe, cilindri, numărul de blocuri, împărțirea pe zone, piste, sectoare). În ultima parte a acestui fișier se poate alege un fișier .seek din care simulatorul își calculează curba de distanțe de căutare. Fișierul .seek ajută la modelarea matematică a harddisk-ului. Specificațiile sunt de obicei foarte detaliate și de aceea ele sunt stocate în fișiere separate și incluse apoi în fișierul principal de parametri.

Fișierul principal de parametri pentru un disc simplu conține:

- **Max queue length** - numărul maxim de cereri ce pot fi găsite în coadă la un moment dat;
- **Scheduler** - o coadă de I/O;
- **Bus transaction latency** - întârzierea pe care o implică transferul fiecărui mesaj de către disc;
- **Acces time** - sinonim pentru "Constant acces time";
- **Bulk sector transfer time** - timpul necesar discului să transfere pe magistrală un bloc de 512 biți;
- **Print stats** - specifică dacă statisticile discului sunt sau nu scoase în fișierul de ieșire;
- **Block count** - capacitatea discului exprimată în blocuri;
- **Command overhead** - specifică o întârziere de procesare care are loc imediat ce o cerere a sosit către disc;
- **Constant acces time** - specifică valoarea fixată pentru timpul de acces;
- **Never disconnect** - specifică dacă discul ține întreg controlul discului pe toată perioada îndeplinirii cererii sau nu (0/1);

4.4 Planificarea accesului la datele de pe disc

Un calculator utilizează cel mai frecvent dispozitivele de I/O, adică discurile. Acest subcapitol prezintă câteva moduri de planificare a accesului la discuri. Un HDD este compus din mai multe platan. Fiecare platan utilizează două capete pentru scrierea și citirea datelor. Unul dintre capete este folosit pentru partea de sus a platanului, iar celălalt pentru partea de jos. Ambele părți ale platanului sunt alcătuite din mai multe piste. Fiecare pistă este împărțită la rândul ei în mai multe sectoare.

Un ansamblu de brațe cuprinde capetele ce accesează platanele. Un braț se poate deplasa doar în două direcții: înspre ax sau direcția opusă axului. Toate capetele de scriere / citire se mișcă înainte și înapoi împreună, astfel încât fiecare cap este întotdeauna situat la aceeași pistă. Nu se poate să avem un cap la pista 0 și altul la pista 1000. În afară de deplasarea capetelor, axul permite rotirea astfel încât capetele să poată accesa sectorul de pe pista pe care se situează. Din cauza acestui fapt, se întâmplă ca de multe ori locația pistei să nu fie specificată ca un număr de pistă, ci mai degrabă ca un număr de cilindru. Un cilindru formează de fapt o colecție a tuturor pistelor. Referirea la sectoare individuale ale discului se face în mod tradițional prin trimitere la cilindri, capete și sectoare (CHS). ^[20]

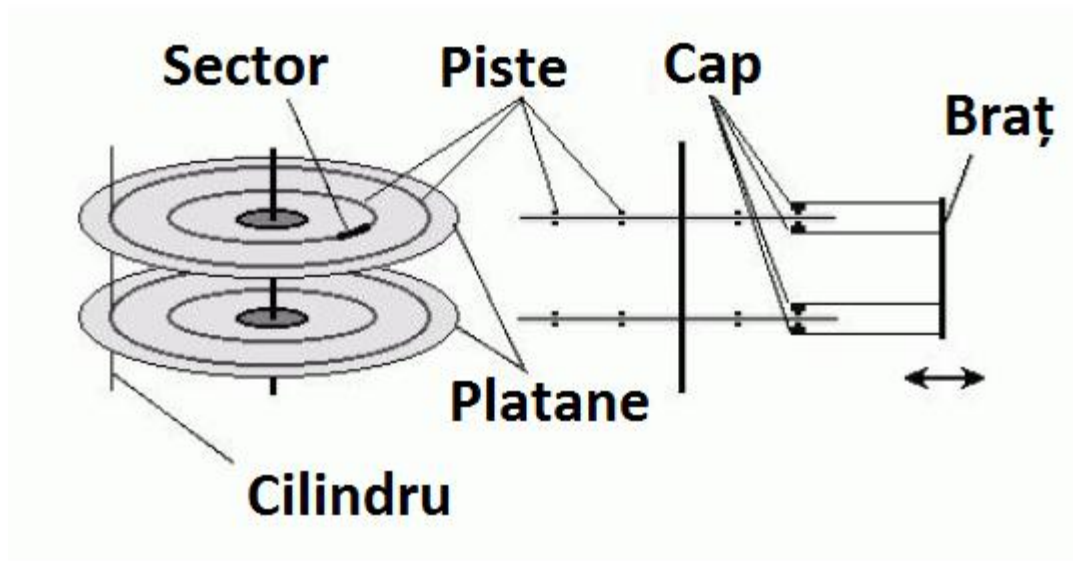


Fig. 4.5
Geometria fizică a hard disk-ului

Atunci când HDD-ul este folosit, discul se învâрте la o viteză constantă. În acest fel, o dată ce capul s-a poziționat deasupra pistei dorite, nu mai trebuie decât să aștepte până când sectorul dorit se află sub el. Timpul de căutare reprezintă timpul necesar mutării capului, iar timpul necesar pentru ca sectorul dorit să devină disponibil la nivelul capului este cunoscut sub numele de latență rotațională. Operațiile de citire sau scriere pot începe o dată ce capul este în poziția corectă, fapt ce reprezintă de transferul de date. De asemenea, dacă discul nu este disponibil, o cerere de I/O trebuie să aștepte într-o coadă de cereri până când discul devine disponibil. ^[20]

O parte din algoritmi de planificare a accesului la disc suportați de mediul de simulare DiskSim sunt:

1. **First Come First Served (FCFS)** – este cel mai simplu algoritm de planificare a accesului la disc. În acest algoritm cererile sunt prelucrate în aceeași ordine în care sunt primite. Această metodă are avantajul de a fi corectă în sensul formal, dar se comportă într-o manieră similară cu planificarea aleatoare deoarece cererile I/O vin într-un mod aleator și cel mai probabil nu vor accesa piste de pe suprafața discului în ordinea potrivită. Algoritmul FCFS mai este cunoscut și sub numele de “First In First Out” (FIFO), “Run To Completion” sau “Run Until Done”.
2. **SCAN (Elevator)** – FCFS este singurul algoritm de planificare a accesului la date de pe disc ce nu lasă cereri neprelucrate până la golirea cozii. Cu alte cuvinte, apare o problemă legată de prioritate: pot sosi mereu noi cereri care vor fi alese înainte de o cerere existentă. În cazul algoritmului SCAN, brațul se deplasează în ambele direcții, satisfăcând toate

cererile, până când ajunge la ultima pistă într-o direcție sau până când nu mai există cereri în acea direcție. După ce ajunge la un capăt pornește în sens opus. Este simplu de observat faptul că algoritmul SCAN favorizează pisteles din centru deoarece brațul trece prin regiunea de mijloc mai des decât trece prin extremitățile discului.

3. **Circular SCAN (C-SCAN)** – restrânge scanarea într-o singură direcție fiind similar cu algoritmul SCAN. Diferența este că atunci când ultima pistă a fost parcursă într-o anumită direcție, brațul se întoarce la capătul opus al discului și scanarea se reia de la început. Acest algoritm reduce întârzierea maximă cauzată de cererile noi.,
4. **Shortest Seek Time First (SSTF)** - selectează cererile I/O care necesită ca brațul să se depalseze cât mai puțin de la poziția actuală. Acest algoritm dovedește o performanță mult mai ridicată decât FCFS deși o serie de decizii optime, fiecare la câte un pas nu garantează o soluție optimă per ansamblu.
5. **VSCAN** - reprezintă o combinație între algoritmul SSTF (Shortest Seek Time) și algoritmul SCAN (Elevator). VSCAN deservește cererile folosind algoritmul SSTF atunci când nu este necesară schimbarea direcției de deplasare a capului de citire/scriere. Furnizează un echilibru bun între timpul mediu de răspuns și rezistența la “înfometare” (fenomenul de întârziere a unor cereri).
6. **Shortest Positioning Time First (SPTF)** - Ca și la SSTF și SPTF apare problema fenomenului de “înfometare”. Cererile care necesită un timp mai mare pentru poziționare sunt amânate pe o perioadă lungă de timp. Pentru a reduce variația răspunsului în timp, cererile din coadă care sunt în așteptare de o perioadă prea lungă de timp primesc prioritate. Prioritatea crește o dată cu timpul petrecut de cerere în coadă, sau se poate stabili o limită de timp, după care cererilor în așteptare să li se dea prioritatea maximă. O variantă a acestui algoritm este SATF (Shortest Access Time First). Acesta dorește reducerea timpului de poziționare împreună cu reducerea timpului de transfer.
7. **LOOK** – similar lui SCAN, capul circulă de-a lungul suprafeței de disc satisfăcând cereri în direcții alternantive. Capul de citire / scriere își schimbă direcția de parcurgere dacă nu mai există cereri în acea direcție, rezultând o reducere a numărului de cilindri parcurși.
8. **Circular LOOK (C-LOOK)** - Algoritmul C-LOOK reprezintă o combinație între algoritmii C-SCAN și LOOK. Capul de citire / scriere se mișcă la fel ca la algoritmul SCAN cu deosebirea că aici se mișcă doar până la ultima cerere în fiecare sens, după care inversează direcția de deplasare imediat, iar pe calea de întoarcere , capul de citire al discului se oprește. ^{[20][22][23][24][25]}

În continuare voi prezenta o listă cu algoritmii puși la dispoziție de mediu de simulare DiskSim 4.0

- FCFS
- ELEVATOR LBN

- CYCLE LBN
- SSTF LBN
- ELEVATOR
- CYCLE CYL
- SSTF CYL
- SPTF OPT
- SPCTF OPT
- SATF OPT
- WPTF OPT
- WPCTF OPT
- WATF OPT
- ASPTF-OPT

4.5 Specificațiile componentelor subsistemului I/O

- **Scheduling policy** – indică algoritmul de planificare folosit pentru a selecta următoarea cerere ce va fi îndeplinită
- **Cylinder mapping strategy** – precizează nivelul de cunoaștere a dispunerii datelor pe disc
 - 0** → înseamnă că singura informație disponibilă este cea cu privire la numărul blocurilor logice ale fiecărei cereri
 - 1** → indică faptul că planificatorul mai cunoaște și frontierele zonelor, numărul de sectoare / zonă și numărul de sectoare fizice / pistă pentru fiecare zonă
 - 2** → planificatorul are de asemenea acces la informațiile despre dispunerea sectoarelor și pistelor de rezervă din fiecare zonă
 - 3** → planificatorul are acces și la lista cu sectoarele și pistele "adormite";
 - 4** → planificatorul are acces la lista sectoarelor și pistelor remapate, asigurând astfel o mapare LBN corectă
 - 5** → planificatorul folosește numărul de cilindru dat împreună cu cererea, admitând experimente cu diverse mapări
 - 6** → planificatorul are acces numai la numărul mediu de sectoare/cilindru
- **Write initiation delay** – reprezintă o aproximare a întârzierilor de procesare ce au loc la scrierea cererilor, înaintea oricăror întârzieri de poziționare. Algoritmii de planificare folosesc această valoare care selectează cererile în funcție de întârzierea de poziționare
- **Read initiation delay** - reprezintă o aproximare a întârzierilor de procesare la citirea cererilor, ce au loc înaintea oricăror întârzieri de poziționare. Această valoare este folosită de algoritmii de planificare care selectează cererile în funcție de întârzierea de poziționare

- **Sequential stream scheme** – returnează o valoare întreagă, interpretată ca un câmp de date boolean. Aceasta furnizează informațiile despre concatenarea cererilor:

- bitul 0 → specifică dacă cererile de citire secvențiale sunt sau nu concatenate de planificator
- bitul 1 → specifică dacă cererile de scriere secvențiale sunt sau nu concatenate de planificator
- bitul 2 → specifică dacă cererile de citire sunt sau nu planificate împreună
- bitul 3 → specifică dacă cererile de scriere sunt sau nu planificate împreună
- bitul 4 → specifică dacă cererile de citire / scriere sunt planificate într-o ordine aleatoare, împreună, pentru a fi îndeplinite de disc

- **Maximum concat size** - lungimea maximă a unei cereri, rezultată din concatenarea a mai multe cereri secvențiale;

- **Overlapping request scheme** - specifică politica de planificare folosită pentru tratarea cererilor suprapuse (suprascrise):

1 → indică faptul că cererile care sunt complet suprascrise de alte cereri care au sosit după ele sunt incluse în acestea;

2 → adaugă la cele de mai sus și posibilitatea ca cererile sosite după suprascriere să fie adăugate, dacă sunt disponibile datele necesare;

- **Sequential stream diff maximum** - specifică distanța maximă dintre două cereri secvențiale dintr-un șir secvențial;

- **Scheduling timeout scheme** - acest parametru specifică schema de întârzieri a unei cozi multiple:

0 → indică faptul că cererile nu sunt mutate din coada de bază într-o coadă cu prioritate mai mare datorită întârzierilor excesive din cauza parcurgerii cozilor;

1 → cererile din coada de bază, care au depășit un anumit timp de așteptare în coadă sunt mutate în una din cele două cozi de mare prioritate (coada pentru cererile foarte vechi sau cea pentru cererile cu prioritate ridicată);

2 → cererile din coada de bază care au așteptat cel puțin jumătate din timpul limită sunt mutate în coada pentru cereri vechi, câștigând o prioritate puțin mai mare.

4.6 Parametrii pentru organizarea datelor în matrici de discuri

Parametrii utilizați pentru organizarea datelor în matrici de discuri sunt configurați în blocurile „logorg” și sunt următorii:

- **Addressing mode** – specifică cum este adresată organizarea logică a datelor. Array indică faptul că există un singur dispozitiv logic pentru întreaga organizarea logică. Parts indică faptul că dispozitivele de stocare back-end sunt tratate ca și cum nu au nici o organizație logică, și cererile sunt re-mapate în mod corespunzător.
- **Distribution scheme** – aceasta arată schema de distribuție de date (care este necesară la sistemul de redundanță). Asis indică faptul că nu apare nicio re-mapare. Striped indică faptul că datele sunt împărțite la dispozitivele din matrice. Random indică faptul că un disc aleator este selectat pentru fiecare cerere. Ideal indică faptul că o distribuție ideală de date ar trebui să fie simulată prin alocarea cererilor de disc folosind algoritmul round robin.
- **Redundancy scheme** - aceasta precizează sistemul de redundanță. Noredun (RAID0) indică faptul că nu există nici o redundanță. Shadowed (RAID1) indică faptul că una sau mai multe copii ale fiecărui disc de date sunt menținute. Parity_disk indică faptul că un singur disc de paritate este menținut pentru a proteja datele din celelalte discuri ale matricii. Parity_rotated (RAID5) indică faptul că blocurile de paritate sunt împărțite de-a lungul tuturor discurilor din matrice.
- **Devices** – lista cu numele dispozitivelor care vor face parte din organizație.
- **Stripe unit** – specifică dimensiunea unității de întrețesere.
- **Number of copies** – specifică numărul de exemplare ale fiecărui disc de date dacă organizarea logică prezintă redundanța Shadowed. Altfel, parametrul este ignorat.
- **Copy choice on read** – specifică politica folosită pentru selectarea unui disc dintr-un set de discuri Shadowed care se va ocupa de cererile de citire sosite din moment ce oricare dintre ele poate face acest lucru. Acest parametru este ignorat dacă redundanța Shadowed nu este selectată.
- **RMW vs. reconstruct** - precizează punctul de întrerupere în selectarea actualizărilor de paritate ca raport al discurilor de date care sunt actualizate. În cazul în care numărul de discuri actualizate de cererile de scriere este mai mic decât punctul de întrerupere, atunci blocul "vechi" de date, blocul "vechi" de paritate, și blocul "nou" de date sunt folosite pentru a calcula paritatea nouă.
- **Parity stripe unit** – specifică dimensiunea unității de întrețesere pentru sistemul de redundanță Parity_Rotated. Nu este obligatoriu să fie egală cu „stripe unit”, dar una trebuie să fie multiplul celeilalte.
- **Parity rotation type** – specifică cum este rotită paritatea de-a lungul discurilor din matrice:
 - 1 → simetric la stânga,
 - 2 → asimetric la stânga,
 - 3 → asimetric la dreapta,
 - 4 → simetric la dreapta. ^{[19][18]}

CAPITOLUL 5

Simulatorul de arhitecturi RAID - DiskSim. Rezultate experimentale

5.1 Simulatorul pentru arhitecturi RAID

Pentru partea practică a acestei lucrări a fost realizată o aplicație JAVA ce încorporează simulatorul DiskSim 4.0 și un generator de workload-uri. Aplicația poate fi rulată atât pe sistemele de operare Windows cât și pe platformele Linux.

Aplicația realizează comunicarea între utilizator, generatorul de workload-uri și simulatorul DiskSim prin intermediul unei interfețe grafice ce va fi prezentată în subcapitolul 5.1.1. Generatorul de workload-uri crează fișiere speciale ce sunt necesare simulatorului DiskSim pentru generarea rezultatelor.

5.1.1 Interfața principală

Interfața principală a aplicației este împărțită în 2 tab-uri. Primul tab este dedicat generatorului de workload-uri și tipului de disc folosit. Cele două tipuri de disk ce pot fi folosite sunt: Seagate Cheetah 146GB 15000RPM și Quantum Tornado 9GB 10000RPM.

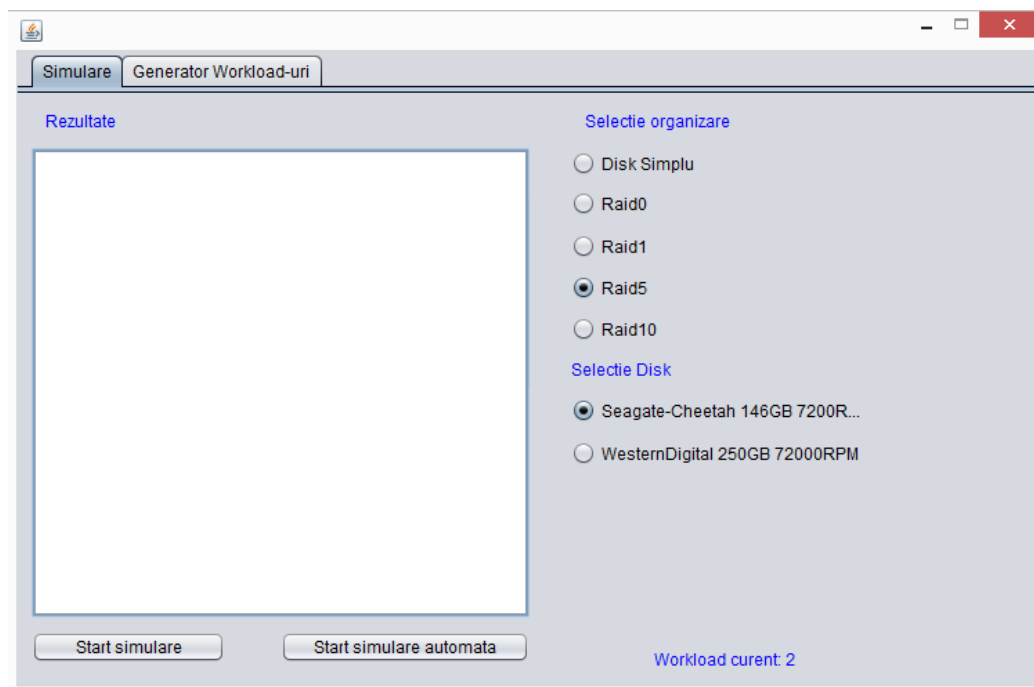


Figura 5.1
Interfața grafică a aplicației

Generatorul de workload-uri generează workload-uri în folder-ul “simulationfiles/workloads” sub numele “trace.....”. Programul nu poate selecta în vederea realizării unei simulări un workload din cele anterioare, ci îl folosește doar pe cel în curs. O dată ce s-a generat un workload nou, cel vechi nu mai poate fi folosit. În partea de jos a ecranului este afișat workload-ul curent.

Generatorul de Workload-uri permite configurarea anumitor parametrii ai workload-ului cum ar fi:

- Numărul de cereri
- Timpul între cereri
- Tipul de distribuție a cererilor

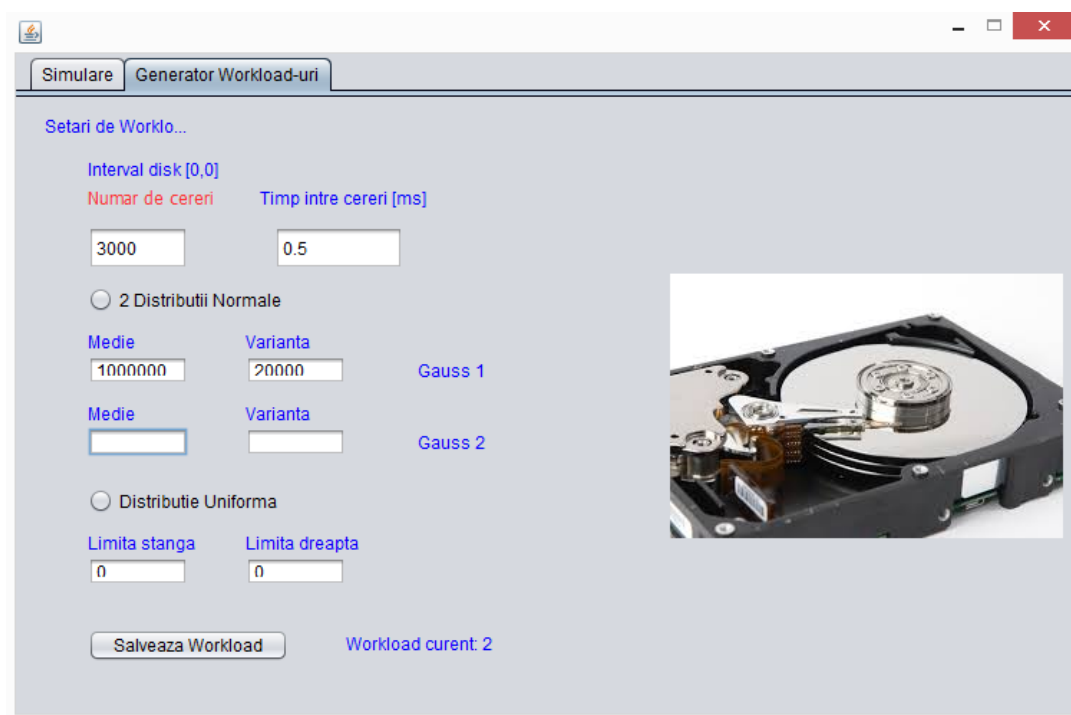


Fig. 5.2
Interfața aplicației – Generatorul de Workload-uri

Parametrii configurabili ai distribuției gaussiene sunt media și varianța acesteia. Atunci când se selectează și o a doua gaussiană, se utilizează 2 limite, una superioară și una inferioară pentru poziționarea distribuției pe disk. Cea de-a doua gaussiană va avea aceeași varianță ca și prima.

Cele două gaussiene alternează în funcție de valoarea unei variabile aleatoare cu distribuție uniformă ce ia valori între [0, 1]. Dacă variabila este mai mică decât 0.5 atunci se alege prima gaussiană, altfel se alege cea de-a doua. Parametrii de control ai distribuției uniforme sunt capetele intervalului în care se vor poziționa cererile pe disk.

În partea superioară a căsuței există un câmp dedicat, scris cu font roșu, ce specifică intervalul în care cererile pe disk pot fi poziționate. Unitatea de măsură este blocul de date și are 512 bytes.

Pentru a genera workload-ul final este nevoie ca butonul “Generează Workload” din josul căsuței să fie apăsat. Fișierul cu workload-uri va conține cereri generate sub un format standard.

0.0	0	31647817	1	1
0.2	0	33053507	1	1
0.4	0	30280558	1	1
0.6	0	27439933	1	1
0.8	0	30723408	1	1
1.0	0	97262470	1	1
1.2	0	30623897	1	1
1.4	0	98859013	1	1
1.6	0	32150787	1	1
1.8	1	30157532		1
2.0	0	31348364	1	1
2.2	0	32708195	1	1
2.4	0	33375934	1	1
2.6	1	98744585		1
2.8	0	102166903	1	1
3.0	0	29738271	1	1
3.2	0	102355882	1	1
3.4	0	104354412	1	1
3.6	1	98560747		1
3.8	0	98355988	1	1
4.0	0	102848224	1	1
4.2	3	95017627		1
4.4	0	98648855	1	1
4.6	0	29836755	1	1
4.8	0	33356153	1	1
5.0	0	103526302	1	1

Fig. 5.3
Formatul fișierului de workload

- Prima coloană reprezintă timpul la care sosește cererea față de momentul 0.
- A 2-a coloană reprezintă numărul organizării de disk-uri care se simulează
- A 3-a coloană indică numărul blocului accesat
- A 4-a coloană indică numărul de blocuri accesate în acea etapă
- Ultima coloană reprezintă tipul cererii. 1 pentru citire și 0 pentru scriere

Cadrul din mijloc este destinat selecției arhitecturii de discuri ce dorim a fi folosită și rulării efective a simulării. Putem alege între un Disk Simplu sau o matrice RAID 0, RAID 1, RAID 5 sau RAID 10. Dacă vom alege o matrice RAID 0 aceasta vine cu un număr presetat de 4 discuri. RAID 1 este presetat pe 2 discuri, RAID 5 are 5 discuri, iar RAID 10 are 4 discuri.

A doua setare ce poate fi făcută este legată de algoritmul de planificare a accesului la disc. Pornind de la cei 27 de algoritmi implementați în DiskSim se poate alege dintr-o listă simplificată ce include algoritmii FCFS, ELEVATOR LBN, CYCLE LBN, SSTF LBN sau SPTF OPT.

După ce am ales parametrii simulării putem face simularea apăsând butonul Start Simulare. În acest fel, folosind workload-ul generat se va genera o simulare, iar rezultatele vor fi afișate în consolă. La ieșire simulatorul va returna o serie de parametri. Aceștia sunt: “IOdriver Response time average” , “IOdriver Requests per second” , “IOdriver Number of reads” , “IOdriver Number of writes” ,

“IOdriver Request size average”. Aceste date vor fi salvate în fisierul “Rezultate.out” din folder-ul “simulationfiles”.

Cel de-a treia funcție a aplicației poate fi folosită independent de celelalte două și realizează o simulare automată. Simularea va fi realizată, pe rând, pentru fiecare tip de arhitectură și pentru fiecare tip de algoritm de planificare a accesului la disc. Simularea va fi efectuată doar asupra tipului de disc selectat.

Simularea automată va crea o serie de fișiere workload și apoi va apela simulatorul pentru fiecare tip de arhitectură și pentru fiecare tip de algoritm disponibil pe fiecare workload. Workload-urile reprezintă o înșiruire de cereri de citire sau scriere aplicate prin intermediul a două gaussiene. Una dintre ele este fixă, iar cealaltă se deplasează de la un capăt la celălalt al intervalului disk-ului folosit. Aceasta se deplasează cu un pas setat de parametrul “Rezoluție parcurgere”. La apăsarea butonului “Start Simulare” se vor crea toate workload-urile și apoi se va începe seria de simulări pe aceste workload-uri. Rezultatele sunt afișate în consolă sub forma unor perechi de numere (poziția blocului mediei celei de-a doua gaussiene și timpul mediu de răspuns la cereri măsurat în milisecunde). Aceste date se salvează și în fișierele “RezultateAuto...” din folder-ul “simulationfiles”.

5.1.2 Configurarea simulatorului

Integrarea simulatorului în programul JAVA, precum și realizarea legăturilor dintre simulator și interfața JAVA s-au realizat prin folosirea comenzii de consolă a simulatorului cu înlocuirea după caz a parametrilor:

Exemplu:

```
disksim raid1.1.parv trace1sim1.out ascii trace1.trace 0 "disc*" "scheduler: scheduling policy" 2
```

Pentru execuția simularilor, DiskSim folosește o serie de parametri specifici cum ar fi:

- Fișierul de parametri – ce are extensia .parv
- Fișierul cu specificațiile disk-ului – ce are extensia .diskspecs
- Fișierul de workload – ce are extensia .trace
- Alte fișiere de descriere a informațiilor (opționale sau nu)

Pentru că DiskSim nu permite crearea de discuri noi în interiorul său s-a folosit utilitarul Dixtrac, special creat pentru a extrage parametrii de pe un disk existent. Utilitarul extrage toate trăsăturile unui disk real și le transformă în fișiere necesare simulatorului DiskSim. Acest program nu poate extrage date decât de pe un disk cu magistrală SCSI, pe o platformă Linux. După extragerea caracteristicilor discului, toate datele memorate pe acesta vor fi pierdute.

Comanda pentru apelarea utilitarului Dixtrac este “./new_disk.pl” și trebuie dată având setat ca director curent, folder-ul “/dixtrac”.

Pentru mai multe detalii privind implementarea simulatorului în program, se poate analiza codul sursă al simulatorului prezentat în cele ce urmează:

Aplicația este realizată din MainFrame.java , GenerateProcess.java , GenerateWorkload.java , ReadWithScanner.java și Rezultate.java.

S-au scris rând pe rând funcțiile pentru crearea workload-ului, setarea parametrilor workload-ului, scrierea datelor din fișirul workload etc.

```
public void genNorm(int nr, float mean, float var, float dist, int opt) {

    PrintWriter dst;
    try {
        Main.tracefile = "trace" + this.number + ".trace";
        String dstName = "simulationfiles/workloads/" + Main.tracefile;
        dst = new PrintWriter(new FileWriter(dstName));
    } catch (IOException e) {
        System.out.println("Cannot opent trace file!!!");
        return;
    }

    int i;
    String s;

    double val = 0;
    for (i = 0; i < nr; i++) {
        val = r.nextGaussian();
        if ((val >= -3) & (val <= 3)) {
            s = (float) (i * dist) + "\t0\t" + (int) (var * val + mean)
                + "\t" + 1 + "\t" + this.rw(opt);
            dst.println(s);
        } else {
            i = i - 1;
        }
    }
    dst.close();
}
```

Fig 5.5
Crearea fișierului de workload

Programul crează un fișier cu extensia .trace ce trebuie să aibă un anumit format, în 5 coloane, prezentat în capitolul 5. O mostră dintr-un astfel de fișier se regăsește în figura 5.3

Generarea acestui fișier a fost realizată prin generarea unei valori random gaussiene. Prima funcție se ocupă de prima distribuție normală, fiind denumită intuitiv genNorm, în timp ce cea de-a 2a funcție denumită intuitiv **genNorm2** și reprezentată în figura 5.6a se ocupă de cea de-a 2a distribuție normală.

A 3-a funcție denumită **genUnif** și prezentată în figura 5.6b se ocupă de modelarea distribuției uniforme.

Aceste funcții sunt folosite pentru crearea workload-ului în funcție de dorințele utilizatorului. Secțiunile de cod au fost legate cu interfața grafică folosindu-se mediu de dezvoltare NetBeans.

În figura 5.7 sunt prezentate layar-ele interfeței grafice.

Este bine de știut faptul că aplicația în sine nu face simulările propriu zise ci apelează simulatorul de discuri DiskSim 4.0 prin intermediul interfeței grafice. Aplicația generează însă workload-urile.

```

60 public void genNorm2(int nr, float mediei, float varianta1, float mean2,
61 float var2, float dist, int opt) {
62     float mean, var;
63     PrintWriter dst;
64     try {
65         Main.tracefile = "trace" + this.number + ".trace";
66         String dstName = "simulationfiles/workloads/" + Main.tracefile;
67         dst = new PrintWriter(new FileWriter(dstName));
68     } catch (IOException e) {
69         System.out.println("Cannot opent trace file!!!");
70         return;
71     }
72     int i;
73     String s;
74
75     double val = 0;
76     for (i = 0; i < nr; i++) {
77         if (r.nextInt(2) == 0) {
78             mean = mediei;
79             var = varianta1;
80         } else {
81             mean = mean2;
82             var = var2;
83         }
84
85         val = r.nextGaussian();
86         if ((val >= -3) & (val <= 3)) {
87             s = (float) (i * dist) + "\t0\t" + (int) (var * val + mean)
88                 + "\t" + 1 + "\t" + this.rw(opt);
89             dst.println(s);
90         } else {
91             i = i - 1;
92         }
93     }
94     dst.close();
95 }

```

Fig. 5.6a
genNorm2 – cod sursă

```

95 public void genUnif(int nr, int a, int b, float dist, int opt) {
96
97     PrintWriter dst, dst2;
98     try {
99         String dstName = "simulationfiles/workloads/trace" + this.number + ".trace";
100         String dstName2 = "simulationfiles/workloads/trace" + this.number + Main.soft + ".trace";
101         dst = new PrintWriter(new FileWriter(dstName));
102         dst2 = new PrintWriter(new FileWriter(dstName2));
103     } catch (IOException e) {
104         System.out.println("Cannot opent destination file!!!");
105         return;
106     }
107     int i, j;
108     String s, s2;
109
110     double val = 0;
111     for (i = 0; i < nr; i++) {
112         val = r.nextFloat() * (b - a) + a;
113         s = (float) (i * dist) + "\t0\t" + (int) val + "\t" + 1 + "\t" + this.rw(opt);
114         dst.println(s);
115         s2 = (float) (i * dist) + "\t0\t" + (int) val + "\t" + 1 + "\t" + this.rw(opt);
116         dst2.println(s2);
117         s2 = (float) (i * dist) + "\t1\t" + (int) val + "\t" + 1 + "\t" + this.rw(opt);
118         dst2.println(s2);
119         s2 = (float) (i * dist) + "\t2\t" + (int) val + "\t" + 1 + "\t" + this.rw(opt);
120         dst2.println(s2);
121         s2 = (float) (i * dist) + "\t3\t" + (int) val + "\t" + 1 + "\t" + this.rw(opt);
122         dst2.println(s2);
123         s2 = (float) (i * dist) + "\t4\t" + (int) val + "\t" + 1 + "\t" + this.rw(opt);
124         dst2.println(s2);
125     }
126     dst.close();
127     dst2.close();
128 }
129
130 }

```

Fig. 5.6b
genUnif – cod sursă

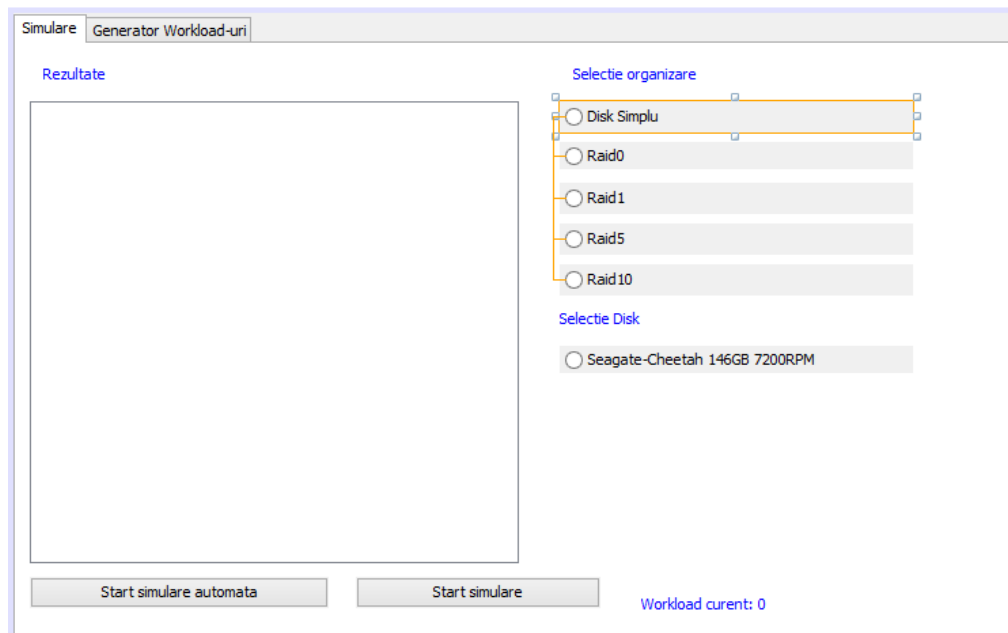


Fig. 5.7a
Legarea butoanelor din interfața grafică

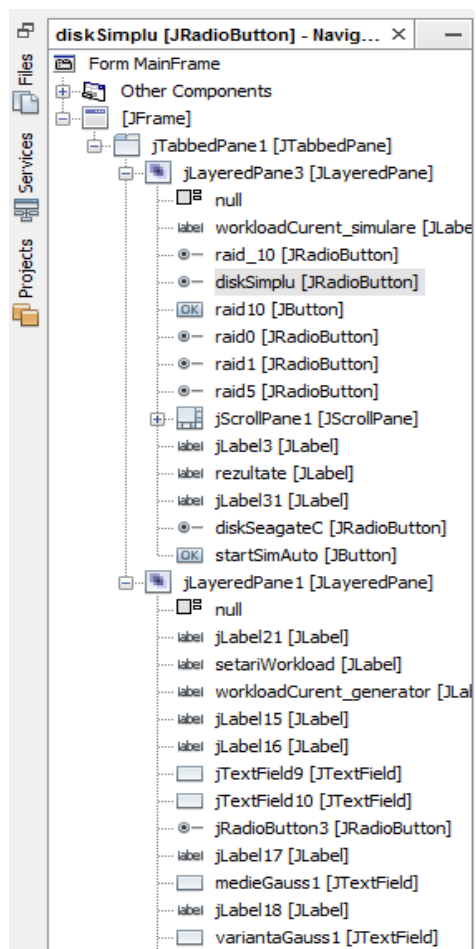


Fig. 5.7b

MainFrame.java începe cu blocul de control al workload-ului. Acesta citește parametrii setați și realizează simularea cerută.

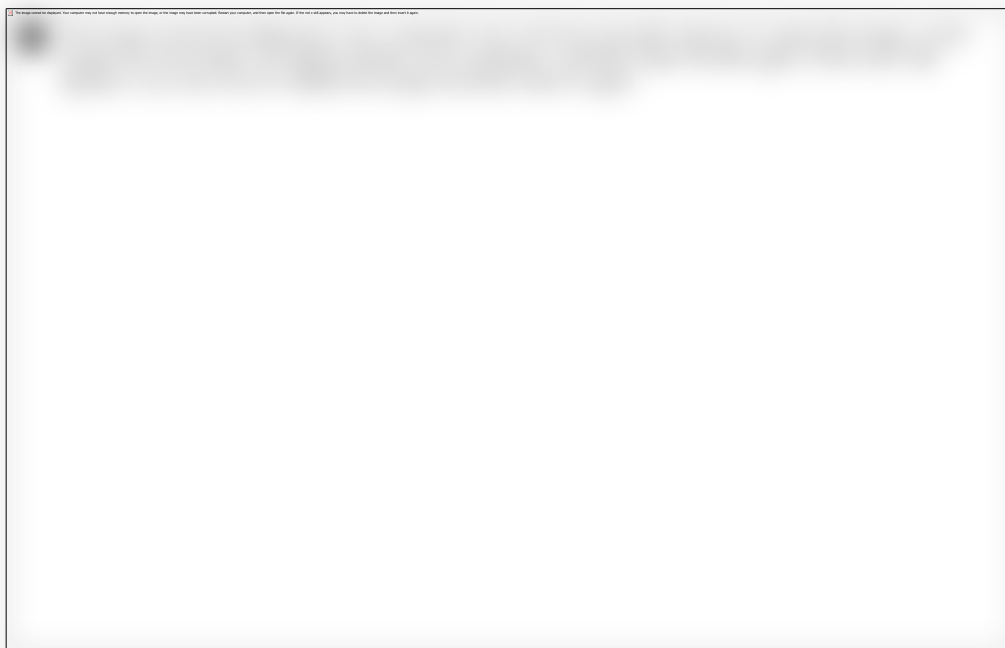


Fig. 5.8a

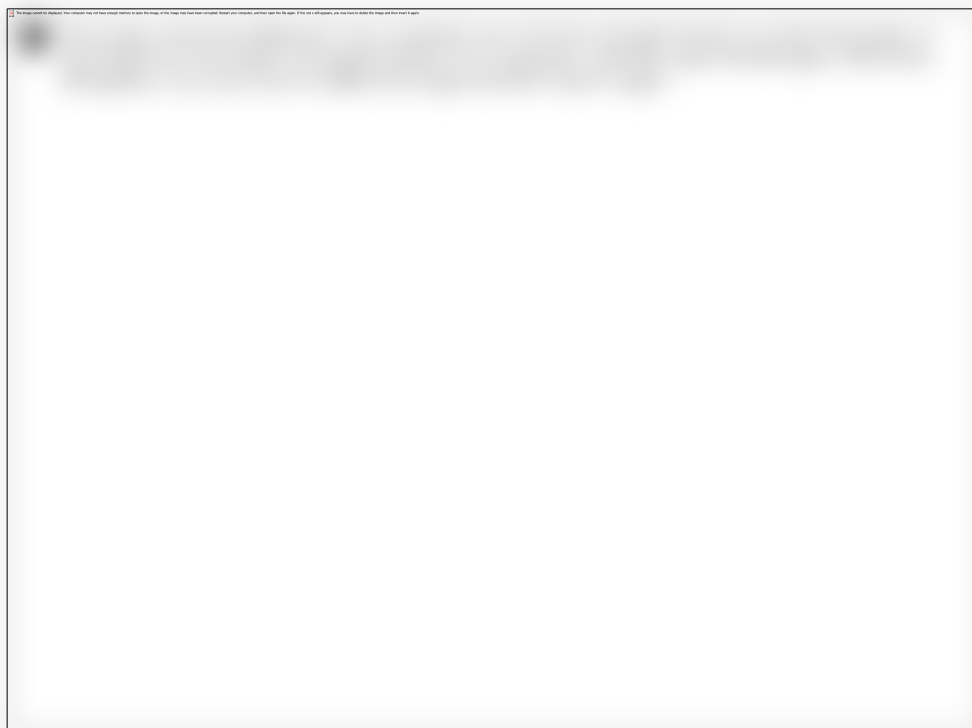


Fig. 5.8b

5.2 Rezultate experimentale

Pentru realizarea simulărilor s-a folosit o configurație ce specifică următorii parametrii:

- Număr de cereri – 3.000
- Timp între cereri – 0.5
- Medie – 100×10^6
- Varianță – 2×10^6
- Rezoluție parcurgere – 5×10^6

Folosind generatorul de workload-uri au fost realizate 50 de fișiere workload. Au fost generate atât sarcini doar de scriere sau citire cât și sarcini mixte, ce din punct de vedere matematic simulează evoluția unui proces de rebuilding.

Workload-urile ce simulează doar scrieri au coloana responsabilă formată strict din valoarea 0, iar cele responsabile doar cu citiri au coloana respectivă formată doar din valoarea 1.

Am constatat experimental că din punct de vedere matematic pentru a simula o procedură de rebuilding este nevoie ca (în cazul unei matrice RAID 5, spre exemplu, cu 5 discuri – în care cade un disc) după o secvență de 4 citiri consecutive (care simulează citirea de pe discurile rămase “în picioare”) să avem o scriere, care să simuleze reconstrucția datelor abia culese pe discul ce este reconstruit. Desigur aceste procese pot fi intercalate cu citiri și / sau scrieri aleatoare ce simulează faptul că celelalte discuri sunt folosite în timp ce se reconstruiește discul căzut.

Pentru a simula cât mai multe situații s-a folosit funcția de simulare automată a simulatorului. Au fost generate simulări care să includă toate combinațiile posibile între algoritmi de planificare a accesului la disc, tipul de matrice RAID folosită și cele 50 de workload-uri.

Rezultatele se găsesc pe CD-ul atașat lucrării, precum și în anexă. Folosind aceste date s-au construit graficele de performanță ale structurilor RAID. Aceste grafice au fost realizate folosind mediul de dezvoltare Microsoft Excel.

În continuare vor fi prezentate o serie de grafice care ilustrează rezultatele experimentelor realizate asupra matricelor RAID. Acestea demonstrează anumite caracteristici prezentate în partea teoretică a lucrării. Aceste grafice sunt grupate în funcție de:

- **Latenta rotațională** – acest parametru reprezintă timpul necesar suprafeței (unde se face operația de citire sau scriere) să se rotească până la poziția capului de citire
- **Timpul de acces** – acest parametru reprezintă intervalul de timp dintre momentul în care discul primește o cerere și momentul soluționării cererii respective împreună cu returnarea rezultatului.
- **Timpul de căutare** – acest parametru reprezintă unul din cele 3 tipuri de întârzieri asociate cu citirea sau scrierea datelor pe un disc. Pentru a putea scrie sau citi (de) pe un disc, capul de citire trebuie mutat fizic la poziția cu pricina. Acest procedeu se numește “căutare”. Distanța variază în funcție de cât de departe de destinație este capul de citire și de timpul necesar fiecărei instrucțiuni de citire sau scriere. Se va considera timpul mediu.
- **Timpul de transfer** – acest parametru reprezintă timpul necesar unui disc pentru a transfera un singur bloc de 512bytes pe magistrală.

Rezultatele generate de simulator au fost prelucrate cu ajutorul unui manager de text, realizat în limbaj JAVA. Acesta analizează fișierele .outv și extrage parametrii necesari aranjându-i într-un tabel pentru a se putea construi graficul de performanță.

Pentru algoritmul de acces FCFS s-au obținut următoarele rezultate:

Din punct de vedere al Latenței Rotaționale:

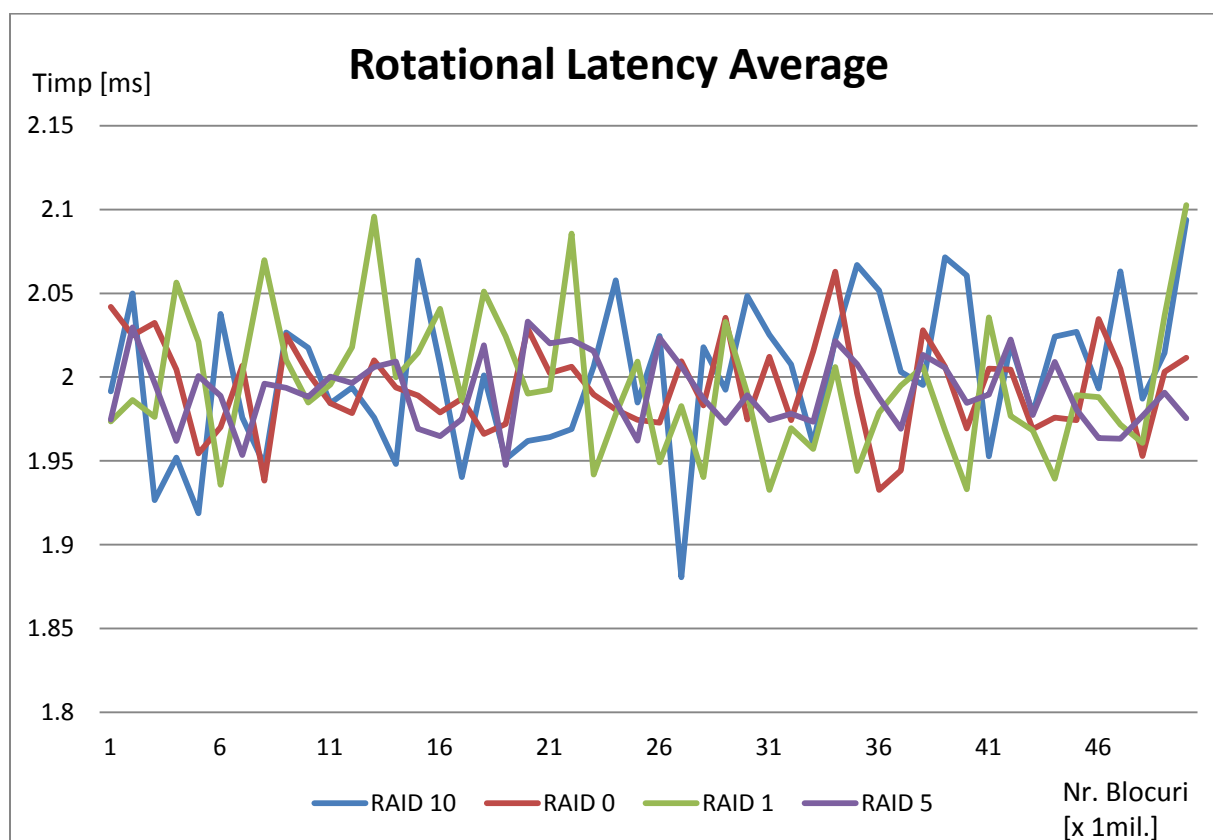


Fig. 5.10

Rotational Latency Average

Observând graficul de mai sus, unde sunt suprapuse cele 4 rezultate grafice aferente sistemelor RAID 0, RAID 1, RAID 5 și RAID 10 putem observa faptul că latențele rotaționale ale acestora nu variază în limite foarte mari, dar sunt totuși diferite. RAID 10 înregistrează minime ale latenței rotaționale ce se situează undeva în jurul valorii de 1.88ms, în timp ce valorile cele mai mari sunt atinse de către RAID 1 (2ms).

Din punct de vedere al timpului de căutare:

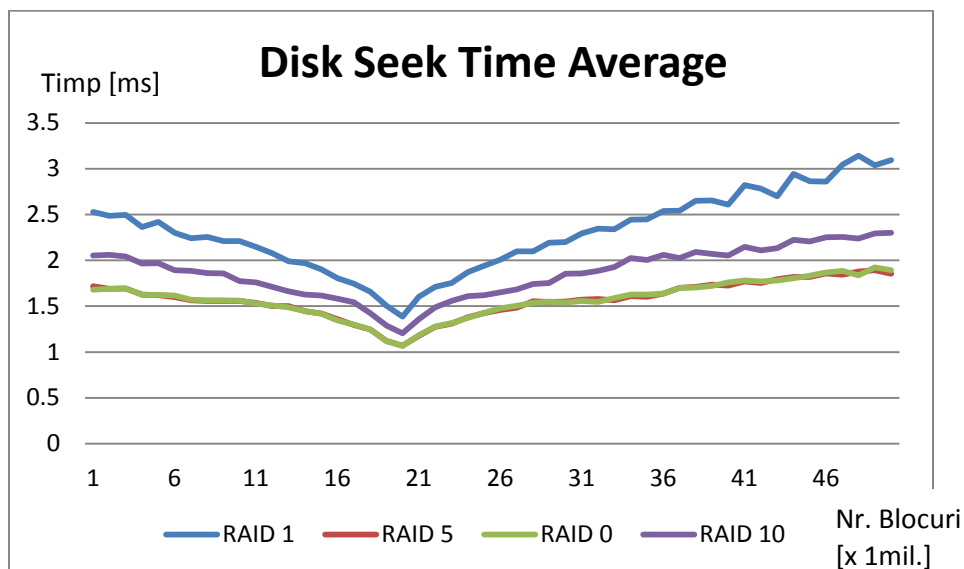


Fig. 5.11a

Disk Seek Time Average

Timpul de căutare pe disk variază în funcție de sistemul RAID folosit conform prezentărilor teoretice realizate în capitolele anterioare. Se poate observa că timpul cel mai scurt de acces la disk este deținut de RAID 0 și RAID 5, în timp ce timpul cel mai lent îi aparține lui RAID 1.

Din punct de vedere al Timpului Mediu de Răspuns

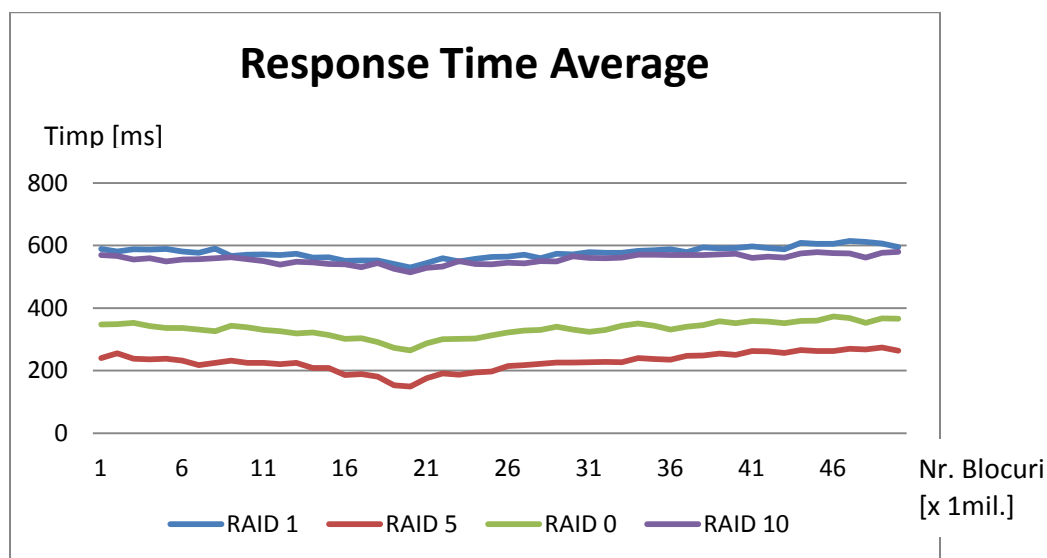


Fig. 5.11b

Response Time Average

Din punct de vedere al timpului mediu de acces la disk

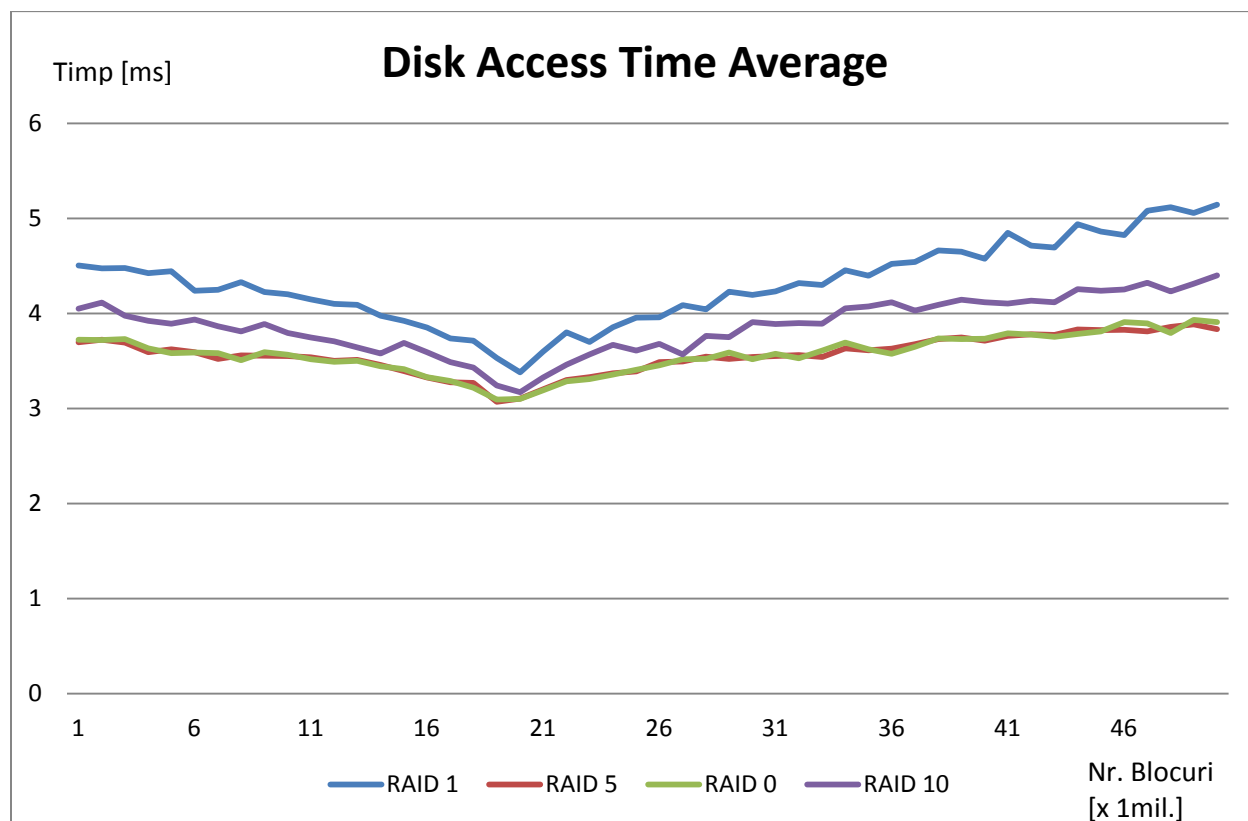


Fig. 5.12

Disk Access Time Average

Timpul de acces la disk variază proporțional între cele 4 tipuri de RAID analizate. Conform datelor prezentate teoretic, se observă că RAID 0 și RAID 5 împart locul întâi la capitolul timp de accesare al discului.

Pe mijlocul graficului, undeva aproape de mijlocul numărului de blocuri de date se poate observa un minim al timpului măsurat. Acesta se produce (în toate graficele) datorită faptului că seek-ul este făcut aleator, fiind favorizate pistele de pe mijlocul discului. Căutarea pe diverse piste pleacă din orice pistă și ajunge pe orice pistă.

Pentru a evidenția rezultatele diferite obținute în cazul folosirii unui alt algoritm de acces la disk în continuare sunt prezentate aceleași analize asupra discului, dar folosind algoritmul SSTF LBN.

Pentru algoritmul de acces la disc SSTF LBN s-au obținut rezultatele:

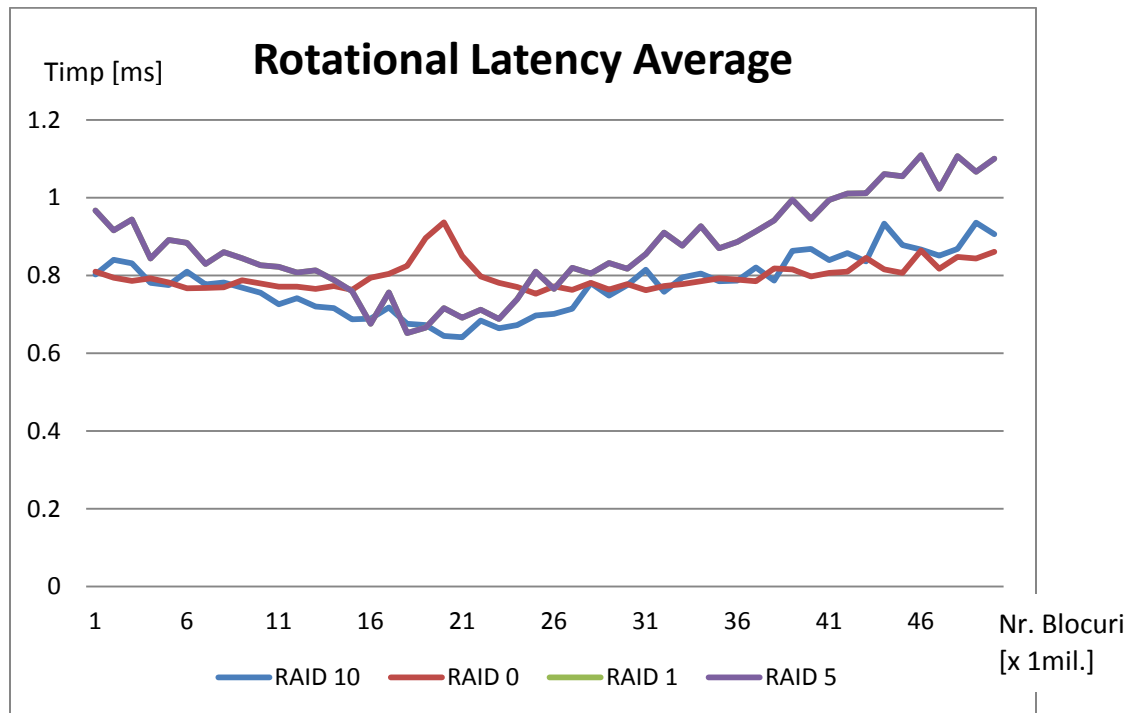


Fig. 5.13a
Rotational Latency Average

Pentru algoritmul de acces la disc SSTF LBN s-au obținut rezultatele:

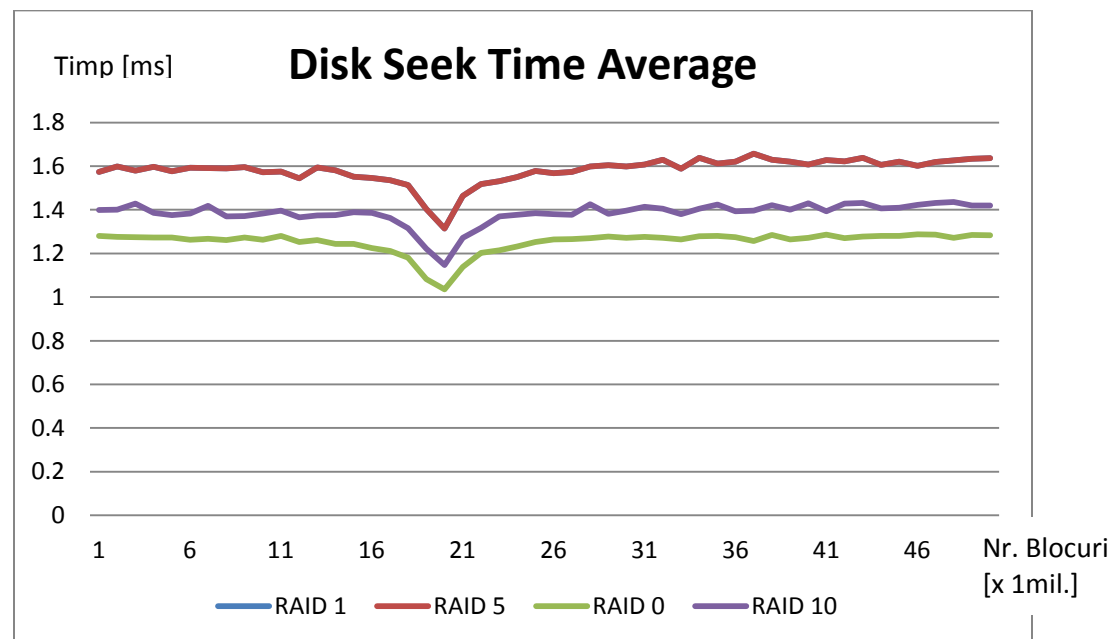


Fig. 5.13b
Disk Seek Time Average

Pentru algoritmul de acces la disc SSTF LBN s-au obținut rezultatele:

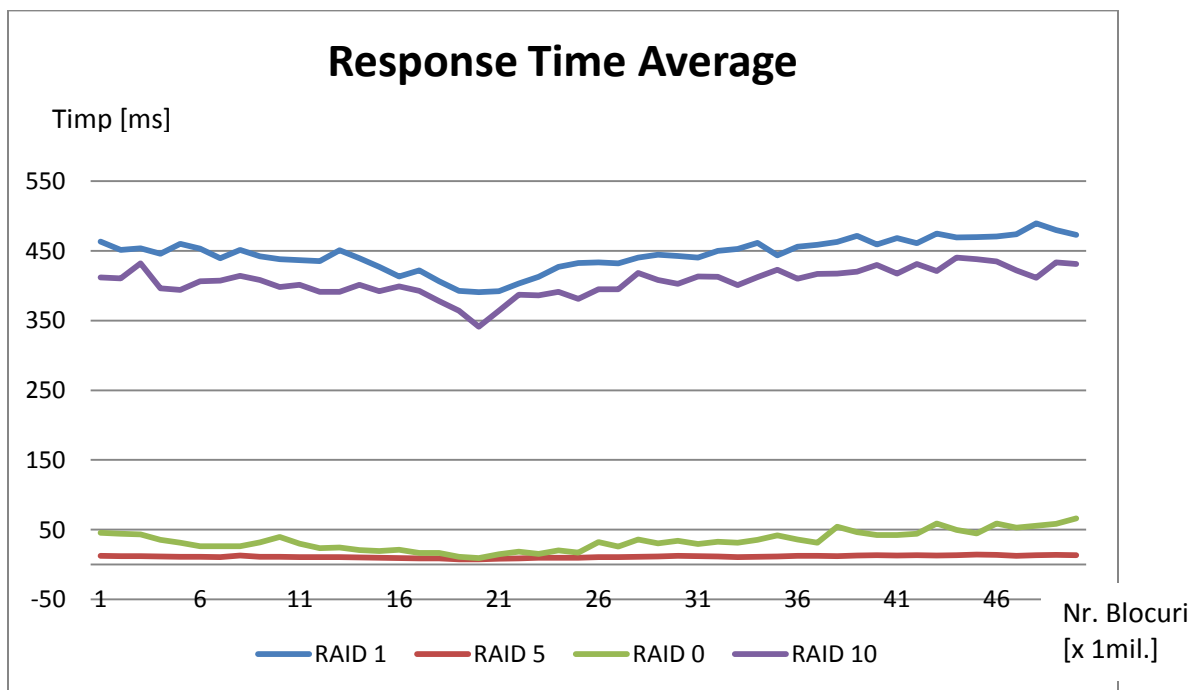


Fig. 5.14a
Response Time Average

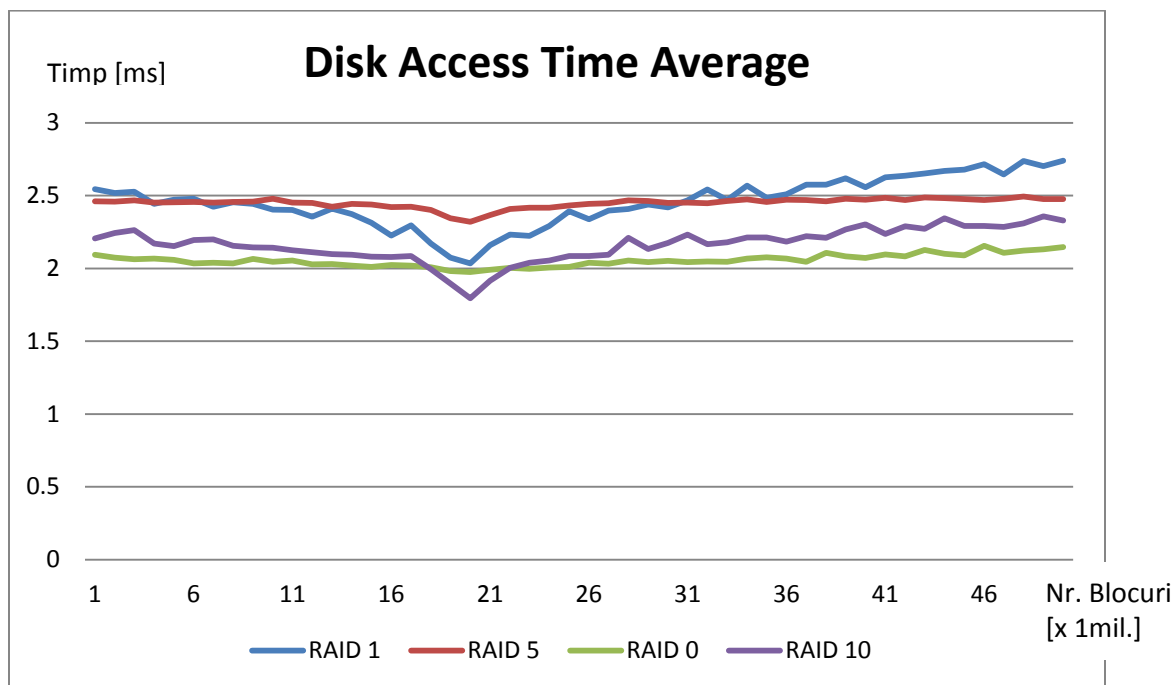


Fig. 5.14b
Disk Acces Time Average

5.3 Concluzii

Rularea simulărilor prezentate anterior ne-au permis conturarea următoarelor concluzii:

În primul rând, aruncând o privire de ansamblu asupra graficelor observăm faptul că RAID 5 oferă performanțe mai bune decât RAID 0, RAID 1 și RAID 1+0 la aproape toate capitolele. În cazul Latenței Rotazionale se observă rezultate ce diferă nesemnificativ, iar în cazul Timpului de Căutare se observă performanțe similare între RAID 5 și RAID 0.

Performanțele unui RAID privind timpii de căutare nu pot fi îmbunătățite prin legarea HDD-urilor într-un RAID anume. Singura metodă viabilă de îmbunătățire a timpului de căutare rămâne procurarea unui harddisk ce oferă nativ un timp de căutare mai bun, cum ar fi spre exemplu SSD-urile. (*Solid State Disk*) Nici o matrice mecanică nu va reuși să atingă performanțele obținute de un SSD. Performanța citirilor în RAID 5 crește cu numărul de discuri, deoarece cu cât sunt mai multe discuri, cu atât sunt mai multe capete de citire/scriere iar matricile RAID 5 au capacitatea de a citi simultan de pe toate discurile.

Este adevărat totuși că un RAID poate îmbunătăți ușor timpul de căutare atunci când are un număr mare de “*queue depths*” (QD). Un utilizator de rând poate atinge un QD de 3-10 cereri, în timp ce un server are peste 20. Totodată RAID mărește de asemenea și latența întrucât procesarea cererilor necesită un *overhead* mai mare. Pentru home useri, în esență, timpul de căutare va scădea în urma folosirii unui sistem RAID față de un singur disk mai rapid.

O alternativă ar putea consta în folosirea unei curse scurte (short – stroking). În acest fel se forțează harddiskul să folosească doar partea exterioară a platanului, unde viteza liniară este mai mare.

Din punct de vedere al timpului mediu de acces la disk putem afirma că se evidențiază din nou RAID 5 și RAID 0. Aceste două tipuri de RAID prezintă timpi de acces similari, cu o constantă fixă adăugată în funcție de latența dorită a rotației platanului.

Timpul mediu de acces reprezintă suma dintre timpul mediu de căutare (variabil) și latența rotațională medie (stabilită de viteza platanului).

Timpul de căutare este bazat pe durata medie necesară deplasării între piste aleatoare, iar latența rotațională reprezintă timpul necesar platanului pentru a efectua o jumătate de rotație. O jumătate de rotație pentru că acesta este timpul mediu necesar pentru a aduce bit-ul căutat sub capul de citire.

RAID 0 stochează sau citește o parte din informații de pe un harddrive și cealaltă parte de pe celălalt harddrive beneficiind de viteză sporită tocmai datorită faptului că are nevoie să acceseze doar jumătate din date per harddrive fizic. Performanțele vor crește o dată cu numărul de harddisk-uri implicate, întrucât secțiunile de date sunt distribuite egal pe ambele HDD-uri ale matricei, oferind fiecărui HDD mai puțin de făcut pentru fiecare bucată de dată în parte.

Un alt factor ce influențează destul de mult performanțele discurilor este algoritmul de acces la disc ales. Primul set de rezultate este bazat pe algoritmul FCFS care din cauza simplității sale furnizează rezultatele cele mai slabe. Al doilea set de rezultate este obținut folosind algoritmul SSTF LBN observându-se rezultate mai bune.^[26]

CAPITOLUL 6

Analizatorul de arhitecturi RAID - Iometer. Rezultate experimentale

6.1 Prezentarea mediului de dezvoltare Iometer

Pentru a putea observa comportarea sistemelor RAID într-un mediu real de funcționare am efectuat o testare a unei matrice RAID 5 cu discuri conectate printr-un controller Smart Array 5300, folosind mediul de simulare Iometer.

Iometer este un tool de măsurare și caracterizare a subsistemelor I/O ce fac parte din sisteme *single* sau *cluster* ce poate măsura performanțele sub un load controlat. Acesta este în același timp un generator de workload-uri și un tool de măsurare, putând genera operații de I/O asupra sistemului și să măsoare răspunsul acestuia. Iometer poate fi folosit pentru a emula load-ul de pe un disk sau de pe o rețea sau poate fi folosit pentru a genera load-uri I/O complet sintetice.

Programul poate măsura:

- Performanța disk-urilor și a controllerelor de rețea
- Lățimea de bandă și capacitățile de latență ale magistralelor
- Performanțele magistralelor share-uite
- Performanțele hard-drive-urilor la nivel de sistem

6.2 Interfața programului

Iometer este compus din două programe separate: Iometer și Dynamo.

Iometer este programul de control. Folosind interfața grafică a Iometer-ului se poate configura workload-ul, se pot seta parametrii de operare și se pot începe sau se pot opri testele. Iometer îi transmite programului Dynamo ce să facă, colectează datele și apoi exportează într-un fișier de ieșire rezultatele obținute.

Dynamo este generatorul de workload-uri. Acesta nu are o interfață grafică. La comanda programului Iometer acesta efectuează operații I/O și înregistrează informațiile cu privire la performanțe, apoi le returnează programului Iometer. Generatorul de workload-uri Dynamo este de tip multithread și fiecare copie a acestuia, ce poate fi rulată în paralel, poate simula workload-ul mai multor programe. Fiecare copie a programului Dynamo care lucrează poartă numele de *manager*, iar fiecare thread al fiecărei copii Dynamo este numită generic *worker*.

Interfața grafică a programului Iometer este prezentată în figura următoare



Fig. 6.2
Iometer - Interfața grafică

Opțiunile de configurare ale acestui program constă în stabilirea numărului de sectoare, stabilirea sectorului de unde va începe testarea, stabilirea modului de scriere sau citire asupra discului și împărțirea acestor sarcini pe thread-uri.

Rezultatele obținute sunt afișate parțial în interfața grafică și exportate în totalitate într-un fișier excel.

6.3 Rezultate experimentale

În urma testelor efectuate s-a obținut un set de rezultate ce sunt anexate lucrării, pe CD, în formatul original, excel.

S-au efectuat mai multe teste asupra matricei RAID 5, începând cu generarea de workload-uri pe un singur thread asupra configurației aflate în regim normal de funcționare și asupra configurației ce ulterior a fost introdusă în Recovery Mode. S-a continuat prin testarea configurației folosind 4 thread-uri și respectiv 1 thread atunci când aceasta efectuează Rebuilding în condiții de prio Normal, Low și High (pentru Rebuilding).

Primul test efectuat s-a realizat folosind un singur worker. Ca și specificații de acces la disk s-a folosit opțiunea globală “All-in-One” care înglobează un mix realizat automat între toate opțiunile posibile oferite de Iometer. Al 2-lea test efectuat s-a realizat introducând diskul în Recovery Mode.

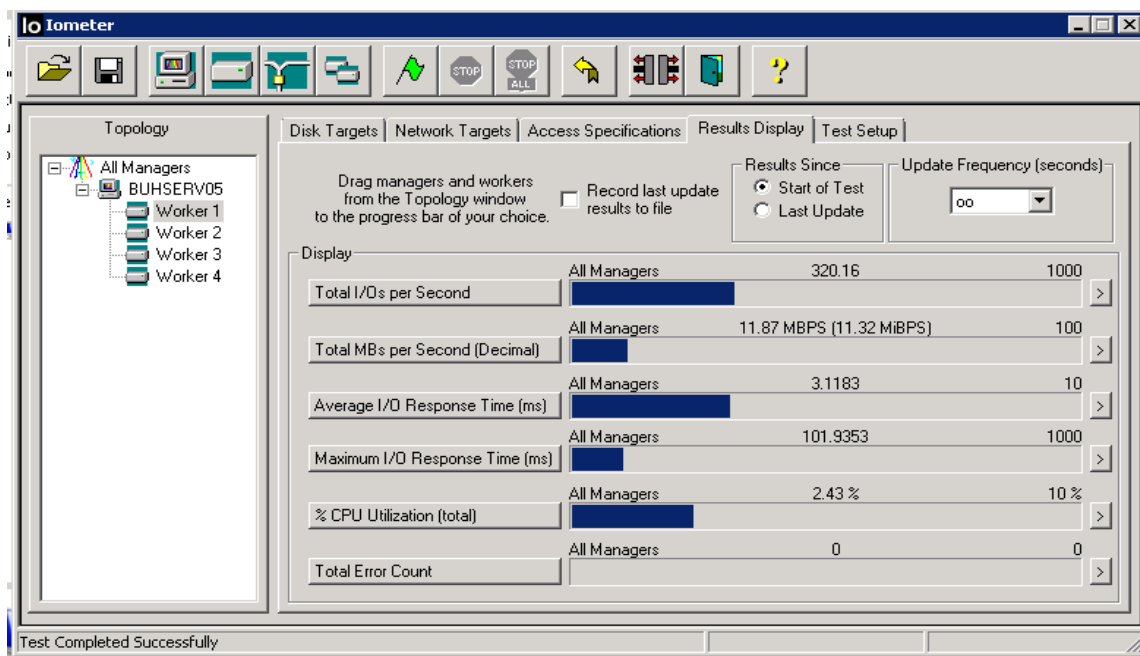


Fig. 6.3a

Analiza 1 - RAID 5 – 1 worker – Algorithm All-in-One

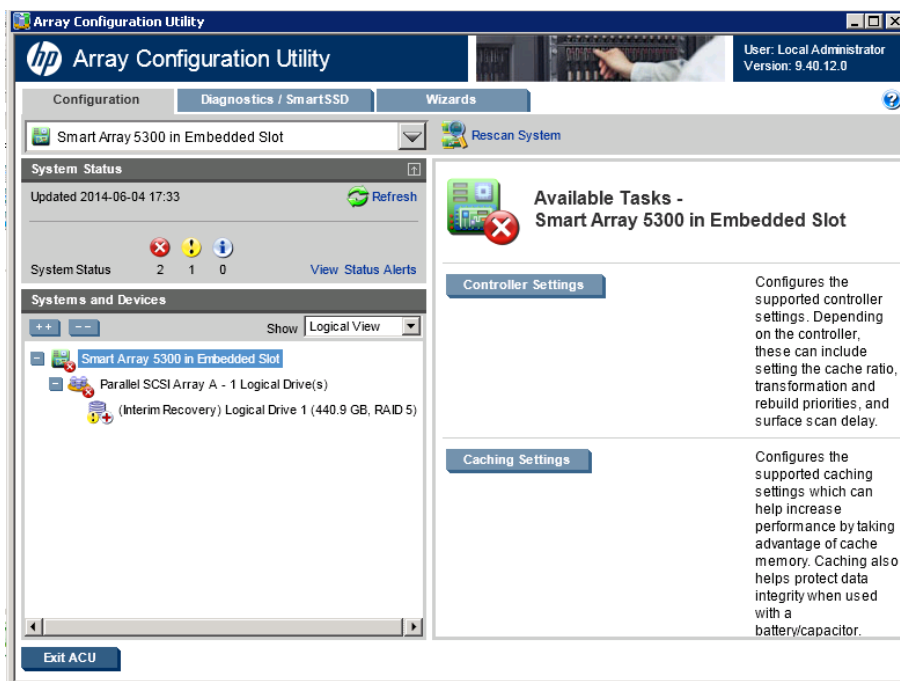


Fig. 6.3b

Analiza 2 – RAID 5 – 1 worker – Recovery Mode

Pentru a nu distruge un disk am simulat producerea unui eveniment nefericit asupra acestuia prin scoaterea manuală, reconectarea la un alt sistem și ștergerea manuală a datelor de pe acesta. Ulterior pentru a o induce disk-ul în starea de Rebuilding, acesta a fost introdus la loc în sistemul inițial.

Rezultatele obținute pot fi observate în figura de mai jos.

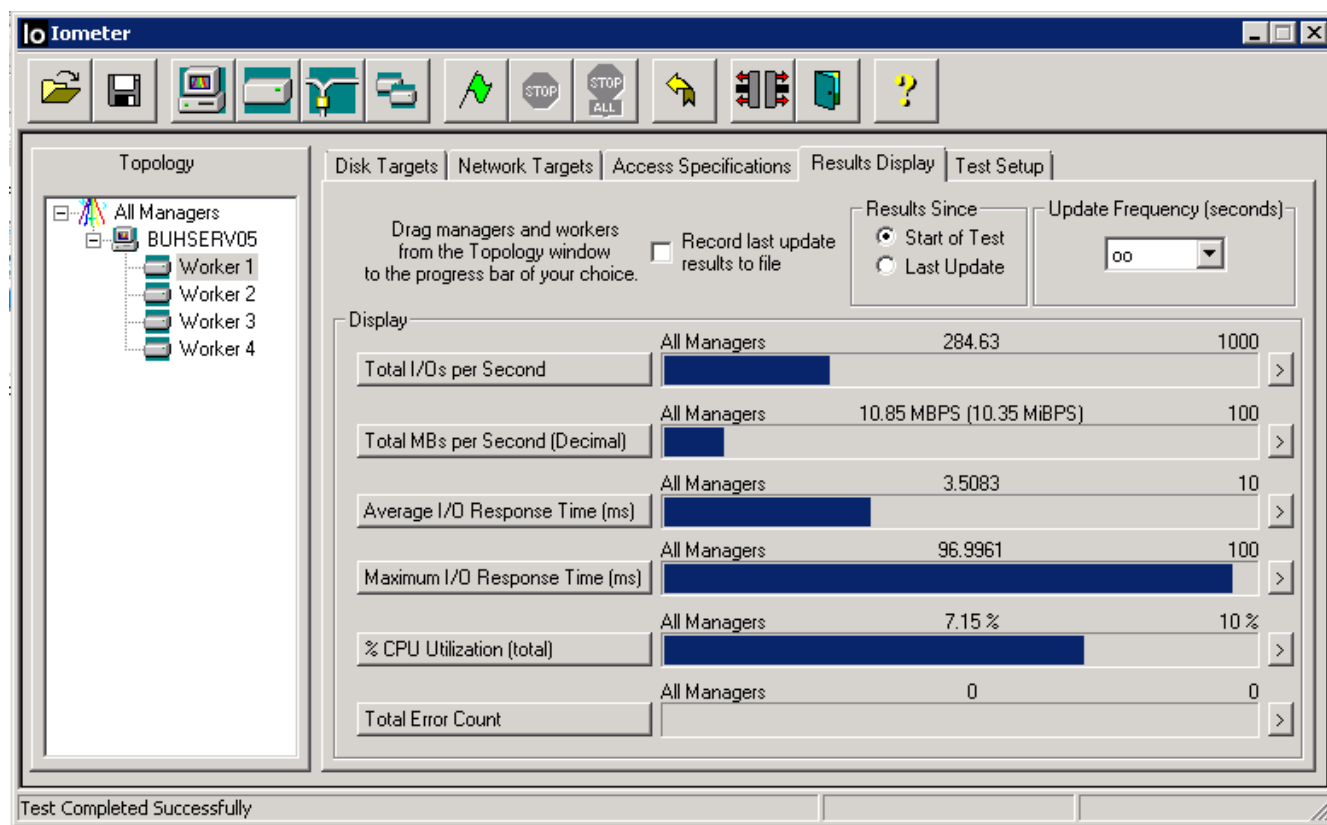


Fig. 6.4

Analiza 2 – RAID 5 – 1 worker – Recovery Mode

Se observă o încărcare mult mai puternică a procesorului (7.15% față de 2.43%) precum și o scădere a vitezei medii totale. (10.85MBPS atunci când diskul este în Recovery Mode față de 11.87MBPS atunci când diskul funcționează în parametrii normali). Numărul de operații de I/O pe secundă este și el în scădere, de la 320.16 în condiții normale de funcționare până la 284.63 atunci când diskul se află în Recovery Mode. Timpul maxim de răspuns la operațiile de I/O scade de la 101.9353 ms în cazul funcționării în parametrii normali la 96.9961ms în cazul funcționării în regim de Recovery Mode. De asemenea atunci când disk-ul intră în Recovery Mode, timpul de răspuns mediu la operațiile I/O crește de la 3.1183ms la 3.5083ms.

Testarea performanțelor în condiții de rebuilding a produs următoarele rezultate:

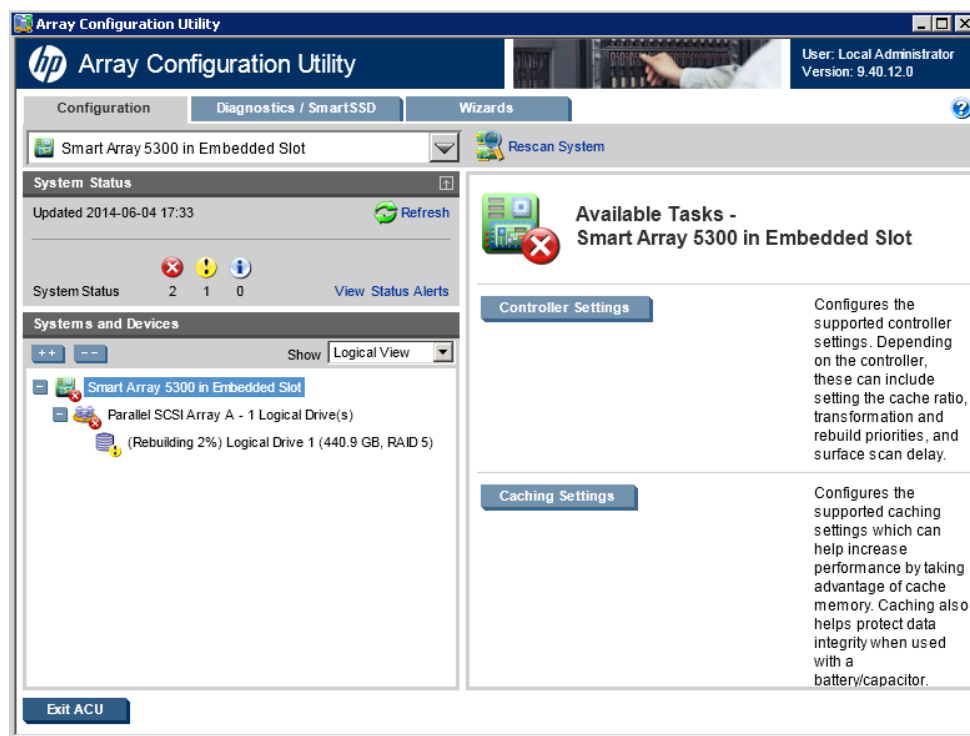


Fig. 6.5a
RAID 5 în procesul de rebuild

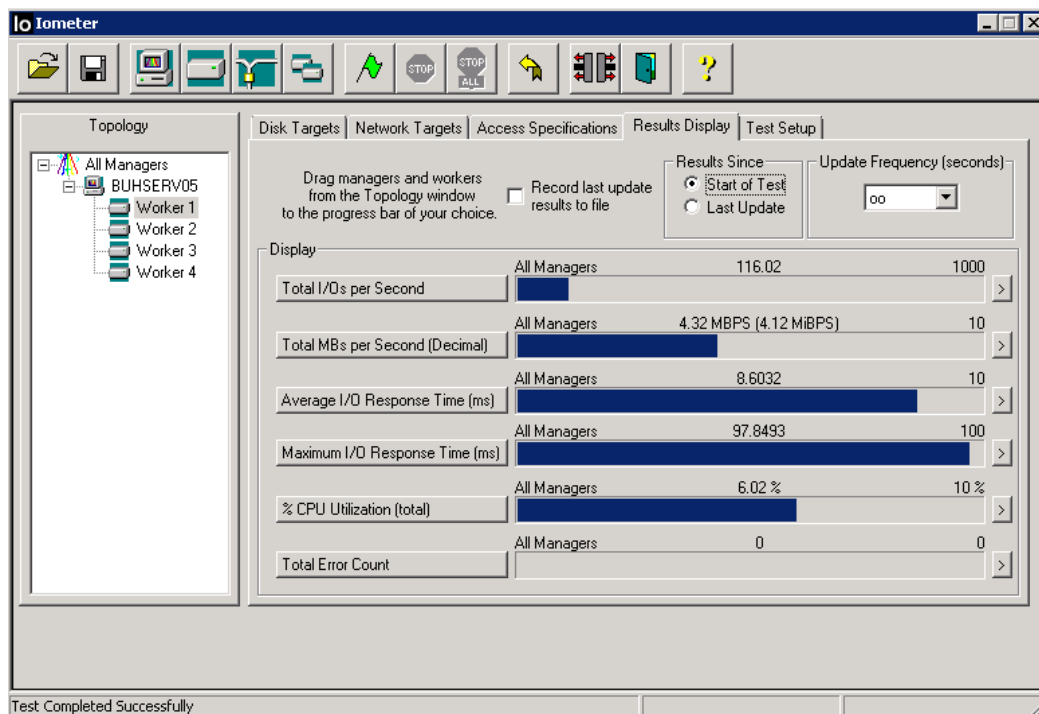


Fig. 6.5b
Rezultate experimentale – RAID 5 – Rebuilding – 1 worker – Normal Prio

Față de rezultatele obținute în urma rulării în parametrii normali, se observă din nou o creștere semnificativă la nivelul de încărcare al procesorului. (6.02% în timpul rebuilding-ului față de 2.43% în condiții normale de rulare)

Întrucât în timpul efectuării operației de rebuild sistemul este mai încărcat se observă o creștere considerabilă a timpului mediu de răspuns la operațiile de I/O. Creșterea are loc de la 3.1183ms pentru sistemul folosit în parametrii normali de funcționare până la 8.6032ms pentru folosirea acestuia în timpul operației de rebuild.

Totodată se observă că traficul făcut scade de la 11.87 MBPS în condiții normale de utilizare până la 4.32 MBPS atunci când are loc procesul de rebuild.

De asemenea numărul total de IOps scade de la 320.16 în condiții normale de utilizare la 116.02 în timpul realizării rebuilding-ului.

Paralelizând cererile furnizate de către generatorul de workload-uri dar păstrând prioritatea acordată rebuilding-ului observăm câteva schimbări semnificative. Aceste schimbări sunt extrem de importante pentru performanța procesului de rebuilding, evident existând anumite soluții de rebuilding mai bune decât altele.

Paralelizarea workloadurilor aduce un plus la nivelul tuturor parametrilor măsurați. Numărul total de IOps crește de la 116.02 (pentru 1 worker) la 346.38 (pentru 4workeri), viteza de transfer crește de la 4.32 MBPS la 12.83 MBPS, timpul mediu de răspuns crește și el din nefericire de la 8.6032ms la 11.4990 însă nu este o creștere la fel de semnificativă ca cea a tuturor parametrilor “pozitivi”. Toate aceste creșteri semnificative ale performanțelor vin cu un cost ce se observă la nivelul procesorului. Încărcarea acestuia crește considerent de la 6.02% până la 16.05%.

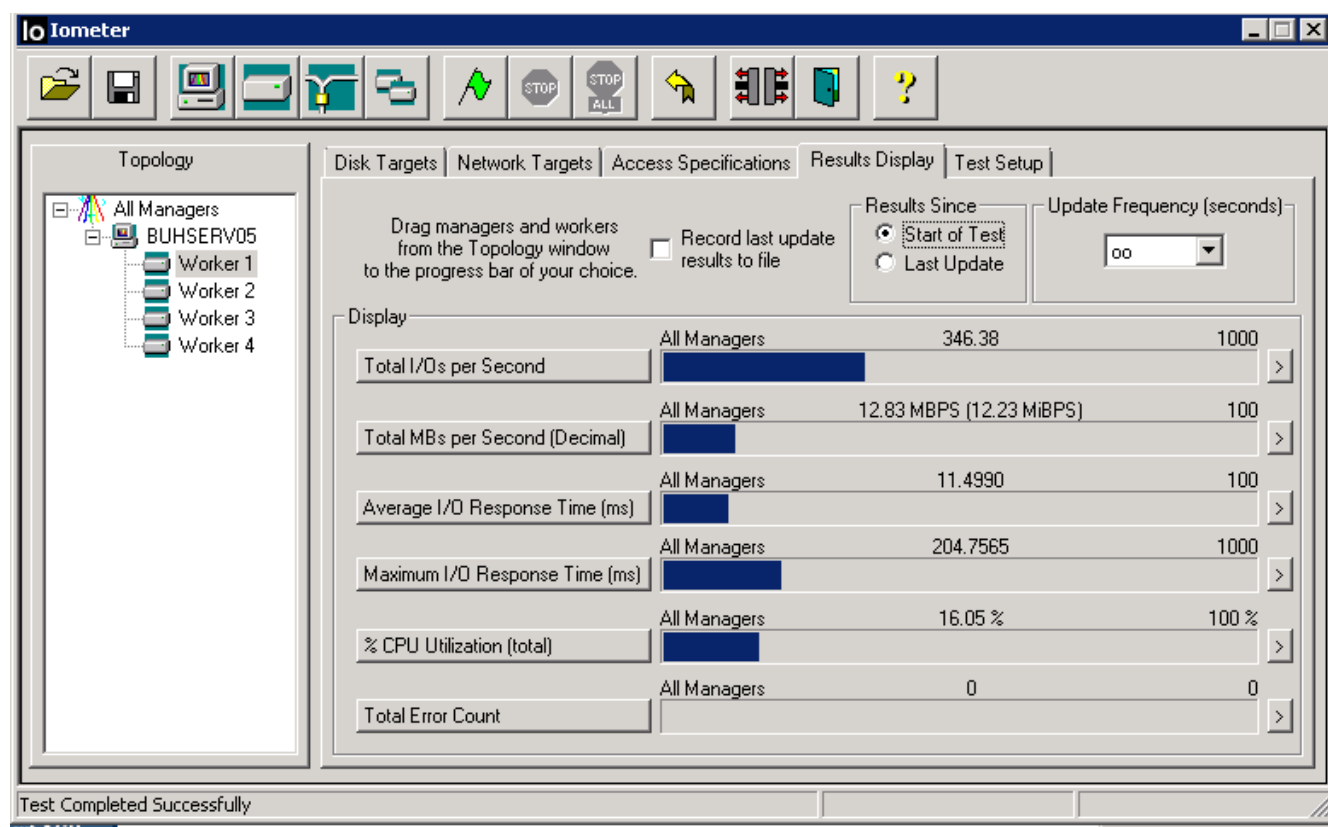


Fig. 6.6

Rezultate experimentale – RAID 5 – Rebuilding – 4 workeri

Următorul test a fost făcut folosind aceiași parametrii (1 worker) și în timpul aceluiași proces de rebuilding dar prioritatea procesului de reconstrucție a fost trecută pe „Low”.

Se observă cum load-ul procesorului scade de la 16.05% la 3.98% dar o dată cu acesta și restul indicilor de performanță raportează valori mai scăzute. Acest rezultat se datorează faptului că oferind o prioritate scăzută procesului de rebuilding, controllerul se poate concentra pe operațiile efectuate asupra discurilor rămase în stare de funcționare.

Această metodă este preferată dacă se urmărește folosirea discurilor rămase, dar nu este indicată atunci când scopul principal este de a realiza un rebuilding cât mai rapid.

Numărul total de I/Ops scade de la 346.38 (în cazul priorității normale și 4 workeri) până la 277.56 (în cazul priorității scăzute cu 1 worker). Rezultatul obținut este totuși mai bun decât cel rezultat pentru prioritatea normală cu 1 worker și anume 116.02 I/Ops. Totodată viteza de transfer scade de la 12.83 MBPS (în cazul priorității normale și 4 workeri) până la 10.65 MBPS. Scorul obținut este oricum mai bun decât cel pentru prioritatea normală cu 1 worker și anume 4.32 MBPS.

La nivelul timpilor de răspuns medii respectiv maximi la operațiile I/O se observă o coborâre a valorilor de la 11.4990ms la 3.5992ms și respectiv de la 204.7565ms la 169.8599ms. Aceste rezultate sunt bune dacă ne dorim lucrul cu discurile rămase, dar nu se apropie de calitatea celor rezultate în cazul procesului de rebuilding realizat cu prioritate normală și anume 8.6032ms pentru timpul mediu de răspuns și respectiv 97.8493ms pentru timpul maxim de răspuns.

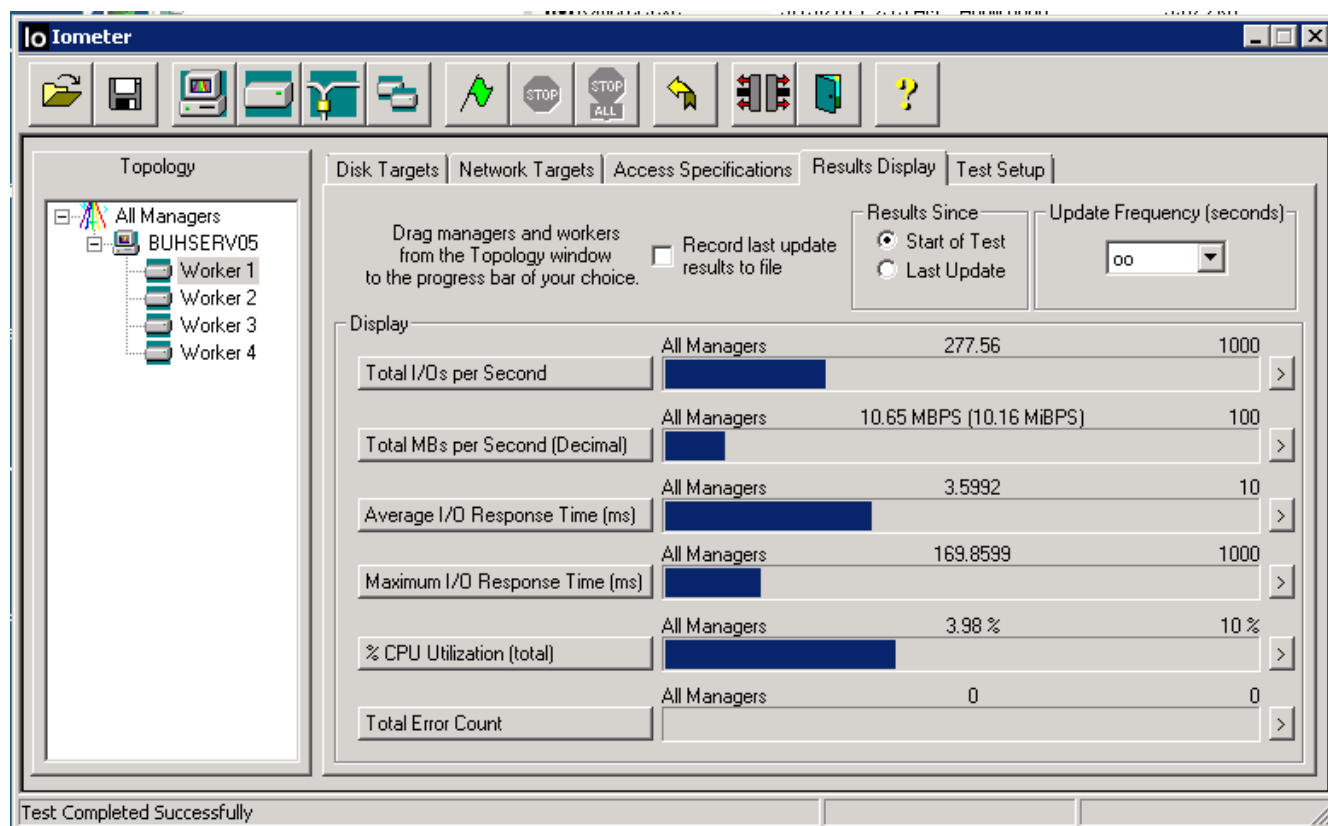


Fig.6.7

Rezultate experimentale – RAID 5 – Rebuilding – 1 worker – Low Prio

Ultimul experiment realizat folosește o aceeași configurație asupra căreia s-a realizat rebuilding cu prioritate ridicată (High). Workloadul a fost realizat de către un singur worker.

Se observă cum întrucât toate resursele sunt canalizate pentru realizarea cât mai rapidă a rebuildingului, numărul de I/Os scade dramatic de la 277.56 (pentru Low Priority) și 346.38 (pentru Normal Priority) până la 81.53.

De asemenea transferul de date se face la viteze mult mai mici. Vechile performanțe atinse (10.65 MBPS – pentru Low Prio și respectiv 4.32 MBPS – pentru Normal Prio ; ambele cu 1 worker) depășesc cei 3.10 MBPS posibili în cazul High Prio cu 1 worker.

Încărcarea procesorului crește până la 5.30%.

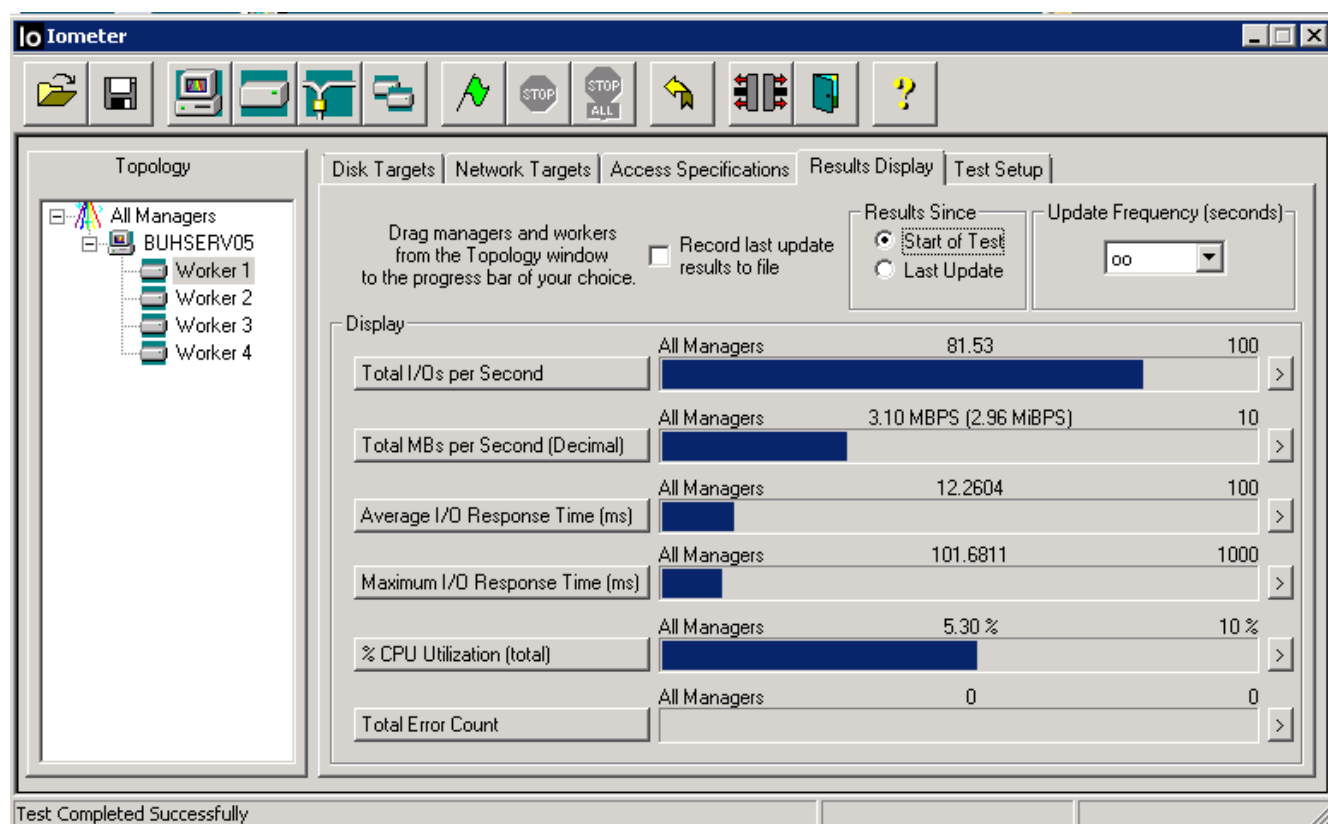


Fig. 6.8
Rezultate experimentale – RAID 5 – Rebuilding – 1 worker – High Prio

CAPITOLUL 7

Considerente economice

Utilizarea tehnologiilor RAID aduce o serie indiscutabilă de avantaje economice. Dezvoltarea sistemelor RAID nu a survenit doar ca un pas natural al evoluției tehnologiei ci și ca un răspuns la nevoile utilizatorilor.

În primul rând, sistemele RAID oferă siguranță și viteză. Cele mai multe RAID-uri se găsesc în mediu buisness și enterprise. Este evident că marile companii multinaționale își doresc viteză și redundanță din considerente economice. Chiar dacă folosindu-se discuri simple în configurații standard ar crește capacitățile de stocare, pierderile aduse în urma defectărilor discurilor ar depăși cu mult economiile făcute prin folosirea unor metode clasice de stocare și gestionare a datelor.

Un alt avantaj major adus de RAID este puterea de acces. Serverele ce folosesc sisteme RAID permit conectarea utilizatorilor din mai multe locații și surse, precum și stocarea datelor diferite ce nu au legătură una cu alta. În acest fel se economisește mult spațiu, pentru ca sistemele profesionale de stocare ocupă foarte mult spațiu și totodată din punct de vedere logistic, toate operațiile asupra echipamentelor se vor face în același loc. În acest fel se vor diminua costurile privind deplasarea tehnicienilor în locații îndepărtate una față de alta.

Un alt avantaj important îl reprezintă unificarea tehnologiilor. Folosind sisteme RAID cu același mod de funcționare, depanarea și servisarea acestor echipamente este mai ușor de făcut. Un lucru ce s-a evitat din totdeauna din raționamente economice a fost folosirea mai multor tehnologii în cadrul aceluiași proiect. În primul rând pentru că din cauza unei defecțiuni acest lucru ar putea conduce la costuri suplimentare privind mentenanța sistemelor și în al doilea rând pentru ca pregătirea personalului calificat și instruirea utilizatorilor ar fi mult prea dificil de realizat.

Creșterea productivității este un alt mare aport economic adus de folosirea tehnologiilor RAID. Pierderea datelor ar reprezenta o catastrofă economică pentru firme. Prin folosirea sistemelor RAID ce au redundanță protecția datelor este asigurată. În caz că unul din discurile matricei cedează, singurele costuri ce vor apărea vor fi legate de înlocuirea sau repararea componentei hardware defecte. Recuperarea datelor se va realiza în mod automat folosind funcțiile RAID nemaifiind nevoie de intervenția altor persoane. Astfel se economisesc atât timp cât și bani.

În concluzie, utilizarea tehnologiilor RAID este o necesitate datorită avantajelor economice pe care le aduce. Folosirea acestora este recomandată atât pentru beneficiile economice aduse cât și pentru cele tehnologice.

Bibliografie

1. [9] – “A Performance Evaluation of RAID Architectures” – Shenze Chen, *Member, IEEE*, and Don Towsley, *Fellow, IEEE*
2. [10] – „Data Allocation In Disk Arrays with Multiple RAID Levels” – Jun Xu, Dissertation Paper, New Jersey Institute of Technology, 2008
3. [11] – „A Hybrid Approach to Failed Disk Recovery Using RAID-6 Codes: Algorithms and Performance Evaluation” – Liping Xiang and Yinlong Xu, University of Science and Technology of China – 2011
4. [12] – „SCAN: An Efficient Decoding Algorithm for RAID-6 Codes” – Jianqiang Luo, Lihao Xu – Wayne State University Detroit
5. [24] – “Rețele de Calculatoare, Ediția a 4-a” – Andrew S. Tanenbaum – Universitatea Vrije, Amsterdam
6. [25] – „Operating Systems: Design and Implementation, 3rd edition” – Andrew S. Tanenbaum - Universitatea Vrije, Amsterdam
7. [16] - "Performance analysis of disk arrays under failure" - R. R. Muntz and J. C. S. Lui.
8. [17] - "Performance analysis of RAIDS disk arrays with a vacationing server model" - A. Thomasian and J. Mcnon.
9. [18] – „The DiskSim Simulation Environment Version 4.0 Reference Manual” – John S. Bucy, Jiri Schindler, Steven W. Schlosser, Gregory R. Ganger , 2008
10. [19] - „The DiskSim Simulation Environment Version 3.0 Reference Manual” – John S. Bucy, Gregory R. Ganger , 2003
11. [20] - “I/O Management and Disk Scheduling” CSc33200: Operating Systems, CS-CCNY, Fall 2003 Jinzhong Niu December 8, 2003
12. [22] – Structuri de Date și Algoritmi – Note de Curs - Sorin Zoican
13. [23]- Sisteme de Operare – Note de Curs – Conf. Dr. Ing. Ștefan Stăncescu
14. [1] – <http://searchstorage.techtarget.com/feature/Learn-all-about-RAID-storage> - Greg Schulz, 2008 – accesat la data 02.07.2014
15. [2] - <http://en.wikipedia.org/wiki/RAID> - 09.03.2014
16. [3] - <http://ro.wikipedia.org/wiki/RAID> - accesat la data 08.03.2014
17. [4] - <http://www.revistait.ro/tehnologie/totul-despre-raid/> - accesat la data 01.04.2014
18. [5] - <http://www.maxpedia.ro/educatie2/schema-de-alocare-a-fisierelor-pe-medii-de-stocare.php> - accesat la data 11.04.2014
19. [6] - <http://blogu.lu/kassak/schema-de-alocare-a-fisierelor-pe-medii-de-stocare> - accesat la data 11.04.2014
20. [7] - <http://www.acnc.com/raidedu/3> - accesat la data 22.06.2014
21. [8] - http://en.wikipedia.org/wiki/Standard_RAID_levels - accesat la data 18.05.2014
22. [13] - http://en.wikipedia.org/wiki/Nested_RAID_levels - accesat la data de 15.04.2014
23. [14] - <http://www.freeraidrecovery.com/library/raid-rebuild.aspx> - accesat la data 12.06.2014
24. [15] - <http://techdepo.ro/recuperarea-datelor-de-pe-un-hard-disk-cu-probleme/#> - accesat la data 14.06.2014
25. [21] - <http://wiki.lug.ro/RAID> - accesat la data de 8.05.2014
26. [26] - <http://www.tomshardware.co.uk/forum/225521-32-should-raid-faster-raid-reads> - accesat la data 22.06.2014

ANEXE

Anexa 1 – Managerul de text

```
using System;
using System.Collections.Generic;
using System.ComponentModel;
using System.Data;
using System.Drawing;
using System.Globalization;
using System.IO;
using System.Linq;
using System.Text;
using System.Text.RegularExpressions;
using System.Threading.Tasks;
using System.Windows.Forms;
using Excel = Microsoft.Office.Interop.Excel;

namespace Balauca
{
    /// <summary>
    /// Structura prin care se memoreaza valorile din fiecare fisier.
    /// </summary>
    public struct Items
    {
        public decimal? item1;
        public decimal? item2;
        public decimal? item3;
        public decimal? item4;
        public decimal? item5;
    }

    public partial class Form1 : Form
    {
        /// <summary>
        /// Cheia 1.
        /// </summary>
        public string item1 = "Rotational latency average";
        /// <summary>
        /// Cheia 2.
        /// </summary>
        public string item2 = "Transfer time average";
        /// <summary>
        /// Cheia 3.
        /// </summary>
        public string item3 = "Disk seek time average";
        /// <summary>
```

```

/// Cheia 4.
/// </summary>
public string item4 = "Response time average";
/// <summary>
/// Cheia 5.
/// </summary>
public string item5 = "Disk acces time average";

public Form1()
{
    InitializeComponent();
    writeToControls();
}

private void button1_Click(object sender, EventArgs e)
{
    //Citeste valorile de cautat.
    readFromControls();

    OpenFileDialog openFileDialog = new OpenFileDialog();
    openFileDialog.Multiselect = true;

    List<Items> items = new List<Items>();
    if (openFileDialog.ShowDialog() == System.Windows.Forms.DialogResult.OK)
    {
        //Pentru feicare fisier deschis se cauta Cheile definite.
        foreach (string fileName in openFileDialog.FileNames)
        {
            String[] fileLines = File.ReadAllLines(fileName);
            Items item = new Items();

            //Se parcurg toate liniile fisierului.
            foreach (string line in fileLines)
            {
                if (!item.item1.HasValue && line.Contains(item1 + ": \t"))
                {
                    int x = line.LastIndexOf("\t");
                    item.item1 = decimal.Parse(line.Remove(0, x + 1), CultureInfo.InvariantCulture);
                }

                if (!item.item2.HasValue && line.Contains(item2 + ": \t"))
                {
                    int x = line.LastIndexOf("\t");
                    item.item2 = decimal.Parse(line.Remove(0, x + 1), CultureInfo.InvariantCulture);
                }

                if (!item.item3.HasValue && line.Contains(item3 + ": \t"))
                {
                    int x = line.LastIndexOf("\t");
                    item.item3 = decimal.Parse(line.Remove(0, x + 1), CultureInfo.InvariantCulture);
                }
            }
        }
    }
}

```

```

    }

    if (!item.item4.HasValue && line.Contains(item4 + ": \t"))
    {
        int x = line.LastIndexOf("\t");
        item.item4 = decimal.Parse(line.Remove(0, x + 1), CultureInfo.InvariantCulture);
    }

    if (!item.item5.HasValue && line.Contains(item5 + ": \t"))
    {
        int x = line.LastIndexOf("\t");
        item.item5 = decimal.Parse(line.Remove(0, x + 1), CultureInfo.InvariantCulture);
    }
}

items.Add(item);
}
}

```

```

SaveFileDialog saveFile = new SaveFileDialog();
saveFile.Filter = "Excel Worksheets|*.xls";
if (saveFile.ShowDialog() == System.Windows.Forms.DialogResult.OK)
{
    //Se realizeaza scrierea in fisierul Excel.
    WriteExcel(saveFile.FileName, items);
}
}

```

```

/// <summary>
/// Creeaza un fisier Ecel cu datele obtinute.
/// </summary>
/// <param name="path">Calea si numele fisierului.</param>
/// <param name="items">Lista cu parametrii de scris.</param>
public void WriteExcel(string path, List<Items> items)
{
    Excel.Application xlApp;
    Excel.Workbook xlWorkBook;
    Excel.Worksheet xlWorkSheet;
    object misValue = System.Reflection.Missing.Value;

    xlApp = new Excel.Application();
    xlWorkBook = xlApp.Workbooks.Add(misValue);

    xlWorkSheet = (Excel.Worksheet)xlWorkBook.Worksheets.get_Item(1);
    int count = 1;
    foreach (Items item in items)
    {
        xlWorkSheet.Cells[count, 1] = item.item1;
        xlWorkSheet.Cells[count, 2] = item.item2;
    }
}

```

```

xlWorkSheet.Cells[count, 3] = item.item3;
xlWorkSheet.Cells[count, 4] = item.item4;
xlWorkSheet.Cells[count, 5] = item.item5;
count++;
}

xlWorkBook.SaveAs(path, Excel.XlFileFormat.xlWorkbookNormal, misValue, misValue, misValue, misValue,
Excel.XlSaveAsAccessMode.xlExclusive, misValue, misValue, misValue, misValue, misValue);
xlWorkBook.Close(true, misValue, misValue);
xlApp.Quit();

System.Runtime.InteropServices.Marshal.FinalReleaseComObject(xlWorkBook);
System.Runtime.InteropServices.Marshal.FinalReleaseComObject(xlWorkSheet);
System.Runtime.InteropServices.Marshal.FinalReleaseComObject(xlApp);
GC.Collect();
MessageBox.Show("Excel file created!");
}

/// <summary>
/// Citeste valorile de cautat in controale.
/// </summary>
private void readFromControls()
{
    item1 = textBox1.Text;
    item2 = textBox2.Text;
    item3 = textBox3.Text;
    item4 = textBox4.Text;
    item5 = textBox5.Text;
}

/// <summary>
/// Scrie valorile in controale.
/// </summary>
private void writeToControls()
{
    textBox1.Text = item1;
    textBox2.Text = item2;
    textBox3.Text = item3;
    textBox4.Text = item4;
    textBox5.Text = item5;
}

}
}

```