

XİDMƏTLƏRİ İDARƏETMƏ SERVERİNDƏ TRANZAKSIYALARIN EMAL VAXTININ TƏYİN EDİLMƏSİ

Ramin Hüseynov, Əli Tağıyev
Azərbaycan Texniki Universiteti

Açar sözlər: idarəetmə sistemində tranzaksiyalar, birprosessorlu server, çoxprosessorlu tranzal serverlər, RAID texnologiyası, intellektual serverləri yaddaş lövhələri, sərf olunan ümumi vaxtın variasiya əmsali

Xidmətləri idarəetmə sistemində tranzaksiyaların emalı prosesinə təsir göstərən amillərdən biri onların həmin sistemlərin növbələrində yaranan gecikmə müddətləridir. Bu təsirləri azaltmaq üçün xidmətləri idarəetmə qovşağı çoxprosessorlu hesablama sistemi bazasında qurulur. Fərz edək ki, birprosessorlu server vahid saniyədə B_{PB} sayda tranzaksiya emal etmək imkanına malikdir. İdarəetmə sisteminin məhsuldarlığını artırmaq üçün vahid saniyədə B_p sayda tranzaksiya emal edən K_{BS} sayda təkpəsessorlu hesablama sistemində ekvivalent sayılan çoxprosessorlu tranzal serverdən istifadə edilir [2]. Bu halda prosessorların serverdəki tələb olunan sayı aşağıdakı ifadə ilə təyin edilir [4]:

$$K_{BS} = B_p / B_{PB} \quad (1)$$

Əgər bir tranzaksiyanın təkpəsessorlu hesablama sistemində emal müddətini q_{PB} ilə, onun çoxprosessorlu serverdə emal müddətini isə q_p ilə işarə etsək, həmin emal müddətləri aşağıdakı ifadələrlə təyin edilir:

$$q_{PB} = 1/B_{PB}, \quad q_p = 1/B_p \quad (2)$$

Onda bir tranzaksiyanın çoxprosessorlu serverdə emal müddəti bərabərdir:

$$q_p = q_{PB} / K_{BS} \quad (3)$$

haradaki K_{BS} və q_{PB} uyğun olaraq (2) və (3) ifadələri ilə hesablanırlar.

Fərz edək ki, intellektual x_i xidmətinin təqdim edilməsi prosesində xidmətləri idarəetmə serverinin yaddaş lövhələrinə statistik verilənləri əvvəlcədən məlum olan n_{P/x_i}

sayda yazı, n_{R_i} sayda isə oxu üçün müraciət olunur. Onda bir tranzaksiya üzrə n_w sayda yazı və n_R sayda isə oxu müraciətlərinin orta sayını uyğun olaraq aşağıdakı ifadələrlə təyin etmək olar [3]:

$$n_w = \sum_{i=1}^{M_x} n_{w_i} P_{x_i} / n_{TRC} \quad (4)$$

$$n_R = \sum_{i=1}^{M_x} n_{R_i} P_{x_i} / n_{TRC} \quad (5)$$

haradaki P_{x_i} və n_{TRC} uyğun olaraq (6) və (7) ifadələri ilə təyin edilir.

Sistemin yaddaşına müraciət proseslərinin məhsuldarlığını yüksəltmək üçün disk massivlərinin təşkili (Redundant Array of Independent Disks, RAID) texnologiyası tətbiq edilir. RAID texnologiyası «güzlüklü» yaddaş lövhələrindən paralel istifadə olunmasına imkan verir. Yaddaş yazı prosesində müxtəlif lövhələrdə saxlanması üçün verilənlər bu lövhələrin sayı n_{OL} qədər hissələrə bölünür. Həmin verilənlərin yaddaşdan oxunması zamanı bu lövhələrin hamısına eyni bir vaxtda müraciət edilir, nəticədə oxu prosesinə sərf olunan müddət demək olar ki, n_{OL} dəfə azalır. Bu zaman bir tranzaksiya üzrə yaddaş lövhəsinə yazı və oxu müraciətlərinə sərf olunan q_L orta vaxtı bərabər olacaqdır [4]:

$$q_L = q_M (n_w + n_R / n_{OL}), \quad (6)$$

burada q_M - lövhə yaddaşına bir müraciət üçün sərf olunan vaxtdır (onun qiyməti xidmətləri idarəetmə serverinin texniki xarakteristikaları siyahısında göstərilir); n_w və n_R uyğun olaraq (4) və (5) ifadəsi ilə təyin edilir.

Onda, xidmətləri idarəetmə serverində hər bir tranzaksiyaya uyğun gələn verilənlərin emalına lövhə yaddaşına müraciətlər üçün sərf olunan q_L vaxt ilə tranzaksiyanın çoxprosessorlu hesablama sistemində q_p emalı müddətinin cəmi qədər vaxt sərf edilir (şək.1). Diaqramdan da göründüyü kimi idarəetmə serverində bir tranzaksiyaya xidmət vaxtı lövhə yaddaşına müraciətlərə və onların emalına sərf olunan vaxtların cəminə bərabərdir:

$$t_{PL} = q_p + q_L, \quad (7)$$

haradaki q_p və q_L uyğun olaraq (2) və (3) ifadələri ilə təyin edilir.

Cari xidmətin təqdim olunmasında xidmətləri idarəetmə sistemində API interfeysi ilə qoşulmuş intellektual serverlər iştirak etdiyi halda, onların fəallaşdırılması üçün törəmə tranzaksiyaların yerinə yetirilməsi müddəti intellektual xidmətə müraciəti reallaşdıran əsas tranzaksiyaların gecikməsinə səbəb olur. Bu serverlər ilə xidmətləri idarəetmə sistemi arasında informasiya mübadiləsi sistemi sonsuz növbəyə malik olan birkanallı kütləvi xidmət sistemi kimi təsvir edilir.

Fərz edək ki, intellektual serverləri fəallaşdıran tranzaksiyalar puasson seli yaradır, onlara xidmət etmə müddəti də eksponensial paylanma qanununa tabedir. Bu halda kütlə-

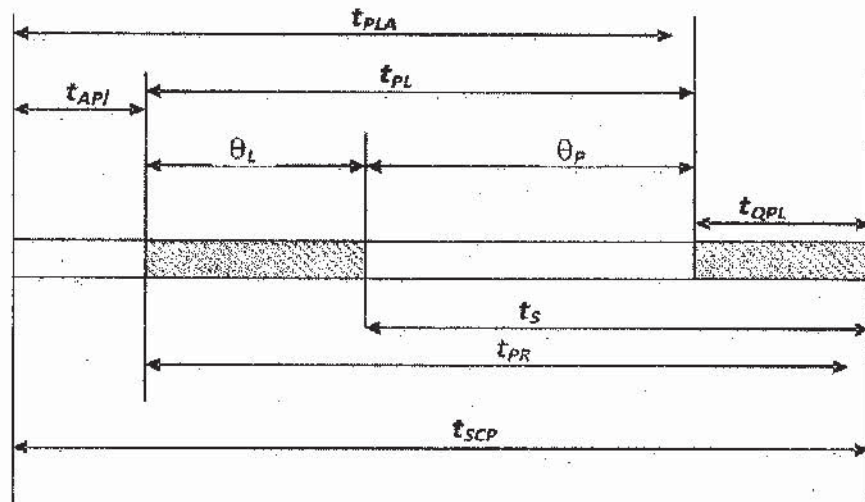
vi xidmət sisteminin $M/M/1/\infty$ modelini alırıq. Xidmətləri idarəetmə sistemində qoşulmuş hər bir i -ci intellektual serverin xidmətə intensivliyini μ_i ilə işarə etsək, cari xidmətin təqdim olunmasında iştirak edən intellektual serverləri fəallaşdıran bir tranzaksiyanın yerinə yetirilməsinə sərf olunan vaxt aşağıdakı kimi təyin edilir:

$$t_{API} = \sum_{i=1}^{M_S} \Lambda_i (\mu_i (\mu_i - \Lambda_i)) , \quad (7)$$

burada M_S - cari xidmətin təqdim olunmasında iştirak edən intellektual serverlərin sayı, L_{TR} - (2) ifadəsi ilə hesablanır.

Beləliklə, xidmətləri idarəetmə sistemində bir tranzaksiyaya xidmət vaxtı bərabərdir: $t_{PLA} = t_{PL} + t_{API}$ (8)

burada t_{PL} və t_{API} uyğun olaraq (7) və (8) ifadələri ilə hesablanırlar.



Şəh. 1. Tranzaksiya emalının zaman diaqramı

Bir tranzaksiya üzrə yaddaş lövhəsinin yüklənmə əmsalı aşağıdakı ifadə ilə təyin edilir:

$$\rho_L = L_{TR} \cdot q_L , \quad (9)$$

burada L_{TR} və q_L uyğun olaraq (2) və (6) ifadələri ilə hesablanırlar.

Hər bir tranzaksiya ilə xidmətləri idarəetmə serverinin prosessorunun yüklənmə əmsalı bərabərdir:

$$\rho_P = L_{TR} \cdot q_P \quad (10)$$

burada L_{TR} və q_P uyğun olaraq (2) və (3) ifadələri ilə hesablanırlar.

Cari xidmətin təqdim olunmasında xidmətləri idarəetmə sistemində qoşulmuş heç olmazsa bir intellektual server iştirak edirsə ($M_S > 0$), bu serverləri fəallaşdıran törəmə tranzaksiyalar ilə xidmətləri idarəetmə sisteminin yüklənmə əmsalı bərabərdir:

$$R_{API} = L_{TR} \cdot t_{API} / M_S , \quad (11)$$

burada M_s - intellektual xidmətin təqdim olunmasında iştirak edən serverlərin sayı; L_{TR} və t_{API} uyğun olaraq (2) və (6) ifadələri ilə təyin edirlər.

Onda xidmətləri idarəetmə sisteminin hər bir tranzaksiya ilə ümumi yüklənmə əmsali bərabərdir:

$$R_{PLA} = \rho_p + \rho_L + R_{API} = L_{TR} \cdot \tau_L + L_{TR} \cdot \tau_p + L_{TR} \cdot t_{API} / M_s = \\ = L_{TR} \cdot (\tau_L + \tau_p) + t_{API} / M_s = L_{TR} (t_{PL} + t_{API} / M_s), \quad (12)$$

burada t_{PL} - (11) ifadəsi ilə təyin edilir.

Əgər (11) - (12) ifadələrini sifarişin prioritetli növbədə gözləmə müddətini təyin edən (2) ifadəsində nəzərə alsaq, onda hər bir tranzaksiya üzrə məlumatın xidmətləri idarəetmə serverində emal növbəsində orta gözləmə vaxtını alırıq:

$$t_{QPL} = R_{PLA} \cdot t_{PLA} (1 + G_{PLA}^2) / (2(1 - R_{PLA})), \quad (13)$$

burada, G_{PLA} - lövhə yaddaşına müraciətlərə, tranzaksiyaların emalına və serverlərin fəallaşdırılmasına sərf olunan ümumi vaxtın variasiya əmsalıdır (puasson seli üçün $G_{PLA} = 1$ qəbul edilmişdir); t_{PLA} və R_{PLA} uyğun olaraq (11) və (12) ifadələri ilə təyin edirlər.

Beləliklə, xidmətləri idarəetmə sistemində müraciəti reallaşdıran bir tranzaksiyanın orta gecikmə vaxtı bərabərdir:

$$t_{SCP} = t_{PLA} + t_{QPL} \quad (14)$$

haradaki t_{PLA} və t_{QPL} uyğun olaraq (11) və (13) ifadələri ilə təyin edirlər.

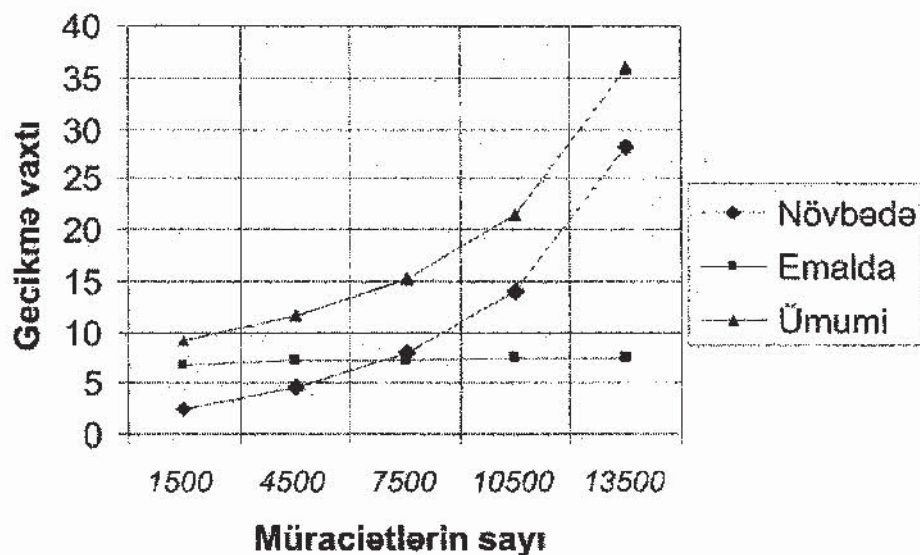
Hesabi eksperimentin nəticələrinə görə tranzaksiyanın xidmətləri idarəetmə sistemində orta gecikmə vaxtının rabitə seansları üzrə müraciətlərin sayından asılılıq qrafiki qurulmuşdur (şəkl.2.).

Bu asılılığın təhlili aşağıdakı nəticələrə gəlməyə imkən verir:

- müraciətlərin sayı artıqca, idarəetmə sistemində tranzaksiyanın gecikmə vaxtı növbədə və ümumi halda artır. Bu zaman tranzaksiyanın serverdəki emal müddəti isə dəyişməz qalır, tranzaksiyanın gecikmə vaxtının maksimum qiyməti ümumi halda, minimum qiyməti isə emal prosesində olur;
- xidmətləri idarəetmə serverində tranzaksiyanın emal üçün prosessorlardan birinin azad olmasını növbədə gözləmə vaxtı həmin sistemdə gecikmə vaxtının 25%-ni təşkil edir.
- müraciətlərin sayı 9 dəfə artıqca, tranzaksiyanın emal üçün prosessorların azad olmasını növbədə gözləmə vaxtı 12 dəfə, xidmətləri idarəetmə sistemində isə ümumi gecikmə vaxtı isə 3.9 dəfə artır. Bu zaman tranzaksiyanın növbədə gözləmə vaxtı 53% artaraq, onun sistemdə ümumi gecikmə vaxtının 78%-ni təşkil edir.

İntellektual üstqurum şəbəkələrinin xidmətləri idarəetmə sistemində istifadə olunan serverlərin prosessorlarının miqdarını və məhsuldarlığını, yaddaş müraciət vaxtını, paralel istifadə olunan yaddaş lövhələrinin sayını artırmaqla intellektual xidmət müraciətlərini reallaşdıran tranzaksiyaların xidmətləri idarəetmə sistemində gecikmə vaxtının

artımının qarşısı alına bilər. Bunlar isə xidmətləri idarəetmə sisteminin texniki xarakteristikaları ilə tənzimlənir.



Şəh. 2.. SCP serverində tranzaksiyanın gecikmə vaxtının müraciətlərin sayından asılılıq qrafiki

Xülasə

Məqalədə telekommunikasiya şəbəkələrində rabitə seanslarının təşkili zamanı gecikmələrin hesablanması üçün riyazi modellərin, hesablama metodlarının, alqoritmlərin və proqram təminatının işlənilib-hazırlanması öz əksini tapmışdır.

İSTİFADƏ OLUNAN ƏDƏBİYYAT

1. Абилов А. В. Сети связи и системы коммутации. Ижевск: ИЖГТУ, 2002, 315 с..
2. Абросимов Л. И. Методология анализа вероятностно-временных характеристик вычислительных сетей на основе аналитического моделирования. Москва: МЭИ, 1996, 114 с.
3. Абросимов Л.И. Анализ и проектирование вычислительных сетей. М.: Изд-во МЭИ. 2000, 52 с
4. Азарнова Т.В., Каширина И.Л., Чернышова Г.Д. Методы оптимизации. – Воронеж: ВГУ, 2003, 86 с.
5. Барсков А.А. Интеллектуальные сети связи Tellin. Сети и системы связи, 2004, №2, с. 34-38.
6. Брусиловский С. А., Копылов Д. А., Конвергенция интеллектуальной сети и сети следующего поколения. IKS-ONLINE, 2004, №2, с.31-37.
7. Варакин Л.Е., Кучерявый А.Е., Соколов Н.А., Филюшин Ю.И. Интеллектуальная сеть: концепция и архитектура // Электросвязь, 1992, – №1, с. 24-31.
8. Вексельман М.И. Построение распределенной системы управления дополнительными услугами интеллектуальной сети связи. Информайонные процессы, 2005, Том 5, №3, с. 23-31.
9. Волков А.Н., Кузин А.В. Сети и телекоммуникации. М: «Academia» · 2006, 352 с.