

Improving Data Center Resource Utilization Using Clustering, Fuzzy Logic, and Evolutionary Algorithms

M. Jahanbani ¹, S. E. Dashti ², S. Hamzeloo ³

^{1,3} Department of Computer Engineering, Pasargad Institute of Higher Education, Shiraz, Iran

² Department of Electrical and Computer Engineering, Ja.C., Islamic Azad University, Jahrom, Iran

ABSTRACT

RESEARCH PAPER

Received: 2025-12-4

Accepted: 2026-4-11

KEYWORDS:

Cloud Computing,
Scheduling,
Virtual Machine Migration,
Evolutionary Algorithms,
Data Centers,

Cloud computing is now considered a model beyond traditional distributed computing systems such as Grid and Cluster, which is capable of responding to dynamic requests and diverse user requirements. With the increase in users, there is a need to establish appropriate mechanisms for load balancing and task scheduling. Load balancing is essential for evenly distributing the workload on physical servers, preventing resource congestion, and improving system performance. Due to the user requests for service and the requirement for providing correct and timely responses from service providers, as well as due to the limited resources available in this environment, there is a need for scheduling. In this paper, an approach to optimizing virtual machine migration is proposed by combining genetic algorithms and ant colony optimization for resource scheduling operations, and it uses K Means clustering and fuzzy logic to quantify the dependency between virtual machines and physical machines for migration in order to maintain load balance. In this research, the proposed model is compared with three load balancing algorithms in the CloudSim simulation environment. In this evaluation, our proposed model achieved a 4.5% reduction in task execution time, a 4.9% increase in deadline success ratio, a 3.9% improvement in task diversity. The computational complexity was reduced by 8.3%, the virtual machine migration efficiency was improved by 2.5%, and the decision latency was significantly reduced by 9.5%, and it also achieved energy savings of 30-35%.

¹Corresponding author:

 seyedebrahim.dashti@iau.ac.ir

Copyright © Author(s).



نشریه تخصصی آرمان پردازش، دوره ۷، شماره ۱، بهار ۱۴۰۵



فصلنامه تخصصی آرمان پردازش (APJ)

Homepage: www.armanprocessjournal.ir

بهبود استفاده از منابع مراکز داده با استفاده از خوشه بندی، منطق فازی و الگوریتم های تکاملی

مژده جهان بانی^۱، سید ابراهیم دشتی^۲، سام حمزه لو^۳^۱ گروه مهندسی کامپیوتر، موسسه آموزش عالی پاسارگاد، شیراز، ایران^۲ گروه مهندسی برق و کامپیوتر، واحد جهرم، دانشگاه آزاد اسلامی، جهرم، ایران

چکیده

محاسبات ابری امروزه به عنوان مدلی فراتر از سیستم های سنتی محاسبات توزیع شده مانند Grid و Cluster مطرح است که توانایی پاسخگویی به درخواست های پویا و نیازمندی های متنوع کاربران را دارد. با افزایش کاربران، نیاز به استقرار مکانیزم های مناسبی جهت توازن بار و زمانبندی کار احساس می شود. متعادل سازی بار برای توزیع یکنواخت حجم کار در سرورهای فیزیکی، جلوگیری از ازدحام منابع و بهبود عملکرد سیستم ضروری است. با توجه به درخواست کاربران جهت دریافت سرویس و لازمه ی ارائه پاسخ درست و به موقع از سوی سرویس دهندگان و همچنین به دلیل محدود بودن منابع موجود در این محیط، نیاز به زمانبندی احساس می شود. در این مقاله رویکردی برای بهینه سازی مهاجرت ماشین مجازی، با ترکیب الگوریتم های ژنتیک و بهینه سازی کلونی مورچه ها برای عملیات زمان بندی منابع پیشنهاد شده است و از خوشه بندی K Means و منطق فازی برای تعیین کمیت وابستگی بین ماشین های مجازی و ماشین های فیزیکی جهت مهاجرت به منظور حفظ تعادل بار استفاده می کند. در این تحقیق مدل پیشنهادی با سه الگوریتم متعادل سازی بار در محیط شبیه سازی کلودسیم مقایسه شده است. در این ارزیابی مدل پیشنهادی ما به کاهش ۴،۵ درصدی در زمان انجام کار، افزایش ۴،۹ درصدی در نسبت موفقیت در مهلت مقرر، بهبود ۳،۹ درصدی در تنوع وظایف دست یافت، پیچیدگی محاسباتی ۸،۳ درصد کاهش یافت، راندمان مهاجرت ماشین مجازی ۲،۵ درصد بهبود یافت و تأخیر تصمیم گیری به طور قابل توجهی ۹،۵ درصد کاهش یافت و همچنین به صرفه جویی در مصرف انرژی به میزان ۳۰-۳۵٪ دست یافته است.

مقاله پژوهشی

واژگان کلیدی:

محاسبات ابری،

زمانبندی،

مهاجرت ماشین مجازی،

الگوریتم های تکاملی،

مراکز داده،

مقدمه

که در آن (N) تعداد ماشین‌های مجازی (VM)، (M) تعداد ماشین‌های فیزیکی (PM)، (R_i) منبع درخواستی (i) VM، (C_j) ظرفیت (j) PM، (X_{ij}) متغیر تصمیم دودویی (اختصاص VM (i) به PM (j) است. همچنین MS زمان تکمیل پردازش، DD تأخیر تصمیم‌گیری، VCE کارایی محاسباتی VM، DHR نسبت موفقیت در مهلت مقرر، TD تنوع وظایف و EE کارایی اجرا می‌باشند که در بخش ارزیابی تعریف شده‌اند.

پیشینه تحقیق

در حوزه این پژوهش، مقالات و روشهایی ارائه شده است که در ادامه به بررسی برخی از آن تحقیقات پرداخته میشود. در مقاله ای نویسندگان الگوریتم زمانبندی ارایه کرده اند که مبتنی بر امنیت و با استفاده از تکنیک بهینه سازی ازدحام ذرات و یادگیری انطباقی چندگانه میباشد. الگوریتم پیشنهادی با استفاده از یادگیری انطباقی باعث تنوع در جمعیت میشود و همچنین منجر به تعادلی بین عملیات اکتشاف و بهره برداری میشود. این الگوریتم پارامترهای امنیت، زمان بازگشت، بار، مصرف انرژی و هزینه را در حین توزیع کارها مد نظر قرار میدهد. در نتیجه به توزیع بار و کاهش مصرف انرژی منجر میشود. این الگوریتم در محیط رایانش ابری بخصوص در بار کاری بالا بسیار اثربخش میباشد. با این حال، این روش در بارهای کاری با ناهمگنی بالا عملکرد مناسبی ندارد و از قابلیت خوشه‌بندی برای تشخیص وابستگی بین VMها استفاده نمی‌کند. همچنین به دلیل ماهیت PSO، در فضاهای جستجوی بزرگ با خطر همگرایی زود هنگام مواجه است [۲].

نویسندگان روشی برای زمانبندی وظایف در رایانش ابری پیشنهاد کرده اند که بر مبنای الگوریتم رقابت استعماری و الگوریتم کرم شب تاب میباشد. در این روش، از الگوریتم هوش جمعی برای پردازش درخواستهای کاربران و زمانبندی وظایف برای توازن بار استفاده شده است. هدف این کار، بهبود زمانبندی و توازن بار در محیط ابری با استفاده از الگوریتمهای جستجوی محلی و سراسری میباشد. از الگوریتم رقابت استعماری، برای بررسی فضای راه حل و الگوریتم کرم شبتاب، به منظور جستجوی محلی استفاده میشود. در شبیه سازی انجام شده معیارهای پایداری، زمان پردازش، توازن بار و درصد راندمان در نظر گرفته شده است و الگوریتم ارائه شده قادر به انجام تعداد وظایف بیشتری هست که به دلیل ترکیب مناسب دو الگوریتم میباشد. اگرچه این ترکیب جذاب است، اما پیچیدگی محاسباتی بالایی دارد (دو الگوریتم با جمعیت‌های جداگانه) و برای محیط‌های

محاسبات ابری به دلیل عملکرد و کارایی بالا و انعطاف پذیری و سودآوری، امروزه مورد توجه سازمانها و شرکت های بزرگ قرار گرفته اند. از آنجایی که سازمان ها به طور فزاینده ای به زیرساخت های ابری متکی هستند، بهینه سازی مهاجرت ماشین های مجازی امری مهم تلقی می شود. در حالی که مهاجرت ماشین های مجازی مزایای بسیاری دارد، این فرآیند مملو از چالش هایی مانند دستیابی به توزیع بار بهینه در طول مهاجرت است. یکی از چالش های مهمی که ارائه دهندگان سرویس های ابری با آن مواجه هستند مدیریت موثر منابع فیزیکی میزبانی شده توسط زیرساخت های فیزیکی می باشد. مصرف کنندگان ابر درخواست سرویس را از هر جایی از جهان به ابر می فرستند. ارائه دهندگان سرویس های ابری باید به کاربران اطمینان دهند که نیازهای آنها به طور کامل برآورده خواهد شد. نکته مهم این است که هر یک از کاربران به دنبال دریافت پاسخی سریع و صحیح از سمت سرویس دهندگان ابر می باشند. با توجه به اینکه حجم درخواست ها بالا میباشد، باید از روشهای مختلف زمانبندی جهت سرویس دهی به مشتریان استفاده کرد.

بنابراین ما با استفاده از تکنیکها و الگوریتمهای فرا ابتکاری در مراکز داده ابری به دنبال زمانبندی منابع و کاهش تعداد سرورهای فیزیکی در مراکز داده ابری هستیم بدین صورت که از طریق مهاجرت زنده ماشین مجازی بتوانیم بدون توقف سرویس، به یک تعادل نسبی بارکاری سرورها برسیم، سپس بتوانیم از منابع سخت افزاری سرورها حداکثر استفاده را ببریم. این امر منجر به کاهش قابل قبول تعداد سرورها و همچنین کاهش چشمگیر مصرف انرژی و هزینه می‌شود. با توجه به افزایش نگرانی های مرتبط با این موضوع، با استفاده از تکنیک های مختلفی نظیر مجازی سازی، زمانبندی و مدیریت منابع میتوان یک فرایند زیرساختی باهدف متعادل سازی بار در مراکز داده ی ابری ایجاد کرد [۱]. زمانبندی به معنای انتخاب بهترین منبع برای یک وظیفه، یا اختصاص ماشین های مجازی به وظایف طوری که زمان اجرا حداقل شود. توابع هدف به شرح زیر می‌باشند:

$$\begin{aligned} \min(\text{MS, DD, Energy}) \\ \max(\text{VCE, DHR, TD, EE}) \end{aligned}$$

قیود مسئله:

$$\begin{aligned} \sum_{j=1}^M x_{ij} &= 1 \forall i \in \{1, \dots, N\} \\ \sum_{i=1}^N x_{ij} \cdot R_i &\leq C_j \forall j \in \{1, \dots, M\} \\ x_{ij} &\in \{0, 1\} \forall i, j \end{aligned}$$

زمان بندی یک جنبه حیاتی هستند که به طور قابل توجهی بر عملکرد محیط های محاسبات ابری تأثیر می گذارند.

یکی از مدل های برجسته زمان بندی در مقالات ، فرآیند مدل زمان بندی مبتنی بر بهینه سازی ازدحام ذرات (PSO) است [۸]. الگوریتم بهینه سازی ازدحام ذرات یک الگوریتم الهام گرفته از طبیعت است که رفتار اجتماعی پرندگان و ماهی ها را شبیه سازی می کند. اگرچه مدل های مبتنی بر بهینه سازی ازدحام ذرات اثربخشی خود را در بهینه سازی تخصیص وظایف و کاهش زمان انجام کار نشان داده اند ، اما اغلب در مدیریت ماهیت پویای حجم کار ابری، عملکرد ضعیفی دارند. علاوه بر این، این مدل ها ممکن است به طور کافی تعادل بین توزیع بار و سطوح مصرف انرژی را نشان ندهند [۹]. یکی دیگر از مدل های زمان بندی قابل توجه، مدل زمان بندی مبتنی بر الگوریتم بازپخت شبیه سازی شده (SA) است. که یک تکنیک بهینه سازی احتمالاتی است که از فرآیند تبرید در عملیات متالورژی الهام گرفته شده است [۱۰]. در حالی که مدل های مبتنی بر الگوریتم بازپخت شبیه سازی شده بهبودهای بالقوه ای در تخصیص وظایف و تعادل بار ارائه می دهند، پیچیدگی محاسباتی آنها می تواند بالا باشد و منجر به عملکرد نامطلوب در محیط های ابری در مقیاس بزرگ شود [۱۱ و ۱۲]. در مقابل الگوریتم های ژنتیک و بهینه سازی کلونی مورچگان الگوریتم هایی هستند که برای غلبه بر محدودیت های مدل های مبتنی بر بهینه سازی ازدحام ذرات و الگوریتم بازپخت شبیه سازی شده استفاده شده اند. و ضمن تطبیق با بارهای کاری پویا، مسائل بهینه سازی پیچیده را به طور مؤثر حل می کنند و در نتیجه به توزیع بار بهینه و ویژگی های بهره وری انرژی دست می یابند [۱۳-۱۵]. تحقیقات اخیر نیز نشان داده اند که ترکیب الگوریتم های تکاملی با یادگیری ماشین می تواند نتایج بهتری در مسئله مهاجرت ماشین مجازی ارائه دهند [۱۶-۱۹].

طریقه کار مدل و راه حل پیشنهادی

مدل پیشنهادی از ترکیب الگوریتم های ژنتیک (GA) و بهینه سازی کلونی مورچه ها (ACO) برای زمان بندی منابع استفاده می کند و رویکردی برای توزیع بار ارائه می دهد. همچنین کمی کردن وابستگی ماشین مجازی به ماشین فیزیکی با استفاده از خوشه بندی K Means و منطق فازی، تعادل بار بهینه را حتی در مواجهه با عدم قطعیت های مهاجرت تضمین می کند. در جدول زیر کلیه مراحل کار و روش پیشنهادی آورده شده است.

ابری در مقیاس بزرگ مقیاس پذیر نیست. همچنین از مهاجرت زنده VM پشتیبانی نمی کند [۳].

در مقاله ای تحت عنوان "ارائه الگوریتم زمان بندی مهاجرت ماشین های مجازی جهت بهینه سازی همزمان مصرف انرژی و تولید آلاینده ها در شبکه محاسباتی ابر" از الگوریتم ژنتیک برای پیش بینی یا تطبیق الگو استفاده میشود. در این کار از روش رگرسیون استفاده میکنند. در اینجا موتور الگوریتم ژنتیک یک جمعیت اولیه ای از فرمول را ایجاد میکند و افراد با مجموعه ای از داده ها مورد آزمایش قرار گرفته و مناسبترین آنها انتخاب میشوند. میبینیم که باگذشت از بین تعداد زیادی از نسلها الگوریتم ژنتیک به سمت فرمولهایی میرود که از دقت بیشتری برخوردارند.

در پژوهشی مشابه، از الگوریتم ژنتیک برای بهینه سازی همزمان مصرف انرژی و تولید آلاینده ها در مهاجرت ماشین های مجازی استفاده شده است. این پژوهش نشان داد که الگوریتم ژنتیک با گذشت نسل های متعدد به سمت فرمول هایی با دقت بیشتر همگرا می شود [۴]. در مقاله ای الگوریتمی پیشنهاد داده اند که با کمک الگوریتم ژنتیک و همچنین در نظر گرفتن تعادل بار سیستم، زمان اجرا و هزینه ی اجرا کاهش پیدا کند. الگوریتم سعی دارد با ایجاد تغییر در الگوریتم ژنتیک و کاهش تعداد ایجاد جمعیت، عملکرد بهتری نشان دهد. هدف اصلی الگوریتم، تخصیص کارها به منابع با در نظر گرفتن طول کارها و تعداد دستورات قابل اجرای بوده است. الگوریتم، کارها را با در نظر گرفتن تعداد کار و توانایی های منابع، به منابع تخصیص میدهد و نتایج اینگونه است که الگوریتم به طور خاص زمان اجراء هزینه ی اجرا و میانگین درجه عدم تعادل را بهبود میدهد. الگوریتم پیشنهادی کارها را با در نظر گرفتن طول کار و توانایی منبع، به منابع اختصاص داده است. بنابراین کل زمان اجرای هر ماشین مجازی کاهش پیدا کرده است. کاستی اصلی این روش عدم توجه به وابستگی بین VM ها و PM ها و همچنین نادیده گرفتن عدم قطعیت در فرآیند مهاجرت است که روش پیشنهادی ما با استفاده از خوشه بندی فازی به آن می پردازد [۵]. در شبکه ابری، پیکربندی ماشین مجازی یک رویه ضروری است که نیازمند به حداقل رساندن هزینه های مربوط به پیکربندی است. با استفاده از تکنیک های زمان بندی مؤثر که ماشین های مجازی پیکربندی شده پویا را روی میزبان های کمتری یکپارچه می کند، اجرا می شود [۶]. مهاجرت ماشین مجازی زنده، استفاده کارآمد از منابع، تحمل خطا، تعادل بار، محاسبات تلفن همراه، مدیریت انرژی و اشتراک منابع را افزایش می دهد. با کاهش فرکانس فیزیکی یا مجازی پردازنده با ایجاد تعادل بین زمان بندی پردازنده مجازی برای چارچوب پس از کپی با سرعت پردازش و انتقال در ماشین منبع اجرا می شود [۷].

جدول ۱. معماری کلی روش پیشنهادی

شماره	نام بلوک	ورودی	خروجی
۱	ورودی پارامترها	ویژگی‌های VM/PM، پارامترهای الگوریتم	داده‌های نرمالایز شده
۲	خوشه‌بندی K-Means و فاز‌سازی	ویژگی‌های VM/PM	درجه عضویت فاز و امتیاز وابستگی $A(x)$
۳	محاسبه VCM و TRM	UT, BW, RAM, MIPS, $A(x)$	VCM و TRM برای هر VM-task
۴	الگوریتم ACO (تولید مسیرهای اولیه)	TRM, VCM, پارامترهای ACO	جمعیت اولیه راه‌حل‌ها (بهترین مسیرهای فرومون)
۵	الگوریتم GA (بهبود راه‌حل‌ها)	جمعیت اولیه از ACO	راه‌حل‌های بهینه‌شده
۶	تابع تناسب و آستانه Fth	راه‌حل‌ها، TRM, VCM	انتخاب بهترین راه‌حل
۷	تصمیم مهاجرت و خروجی	بهترین راه‌حل	تخصیص نهایی VM به PM

الگوریتم بهینه سازی کلونی مورچه ها

الگوریتم بهینه سازی کلونی مورچه ها از رفتار مورچه‌ها در یافتن کوتاه‌ترین مسیر به منبع غذا الهام گرفته است. الگوریتم کلونی مورچه‌گان از دو بخش تشکیل شده است. در بخش اول، مقادیر پارامترهای مسأله تنظیم و متغیرهای جمعیت اولیه مقداردهی می‌شوند. بخش دوم شامل یک حلقه تکرار است که سه مرحله الگوریتم کلونی مورچه‌گان را اجرا می‌کند. این سه مرحله به این صورت می‌باشد که :

- مرحله اول یا مرحله تولید جواب‌های کاندید را شروع کن.
- مرحله دوم یا مرحله جستجوی محلی جواب‌ها میباشد. در این مرحله، از جواب‌های بهینه محلی استفاده می‌شود تا مشخص شود کدام یک از فرومون‌ها باید به روز رسانی شوند. این مرحله اختیاری است.
- مرحله سوم هم مرحله به روز رسانی فرومون میباشد.

در هر حلقه از الگوریتم کلونی مورچه‌گان، جواب‌های کاندید توسط تمامی مورچه‌های مصنوعی ساخته می‌شوند. جواب‌های تولید شده، از طریق یک الگوریتم جستجوی محلی بهبود می‌یابند. در مرحله آخر، فرومون‌ها به روز رسانی می‌شوند. در زمینه مهاجرت ماشین مجازی، مسیرهای کوتاه‌تر به معنای بهینه‌سازی مواردی مثل زمان مهاجرت کمتر، کاهش مصرف انرژی یا بهره‌وری بالاتر از منابع است.

جدول ۱ معماری کلی روش پیشنهادی را نشان می‌دهد. ابتدا ویژگی‌های VMها و PMها وارد سیستم می‌شود. بلوک خوشه‌بندی فاز، وابستگی بین VMها و PMها را تعیین می‌کند. سپس معیارهای VCM و TRM محاسبه شده و به الگوریتم ACO داده می‌شوند تا راه‌حل‌های اولیه ساخته شوند. در ادامه، الگوریتم GA این راه‌حل‌ها را بهبود می‌بخشد. در نهایت، تابع تناسب بهترین تخصیص را انتخاب کرده و تصمیم مهاجرت صادر می‌شود. بر اساس بررسی مدل‌های موجود مورد استفاده برای افزایش کارایی ماشین‌های مجازی متعادل‌کننده بار، می‌توان مشاهده کرد که اکثر این مدل‌ها در هنگام اعمال بر روی استقرارهای ابری در مقیاس بزرگ، پیچیدگی بالایی دارند. برای غلبه بر این مشکلات، این بخش به طراحی یک عملیات زمان‌بندی منابع می‌پردازد که مهاجرت ماشین‌های مجازی را با استفاده از الگوریتم‌های ترکیبی فراابتکاری بهبود می‌بخشد و می‌توان مشاهده کرد که مؤلفه اصلی، عملیات زمان‌بندی منابع با استفاده از ترکیب الگوریتم‌های ژنتیک و الگوریتم بهینه‌سازی کلونی مورچه‌ها میباشد. این الگوریتم‌ها به این دلیل انتخاب شده‌اند زیرا سابقه شبیه‌سازی موفقیت‌آمیز برای حل وظایف بهینه‌سازی پیچیده را دارند. با توجه به بررسی‌ها می‌توان گفت بهترین الگوریتم‌ها از هر خانواده، الگوریتم ژنتیک از خانواده تکاملی و الگوریتم مورچه‌گان از خانواده هوش جمعی می‌باشند، که از نظر توازن بار، زمان اتمام کار و میانگین زمان پاسخ گویی مناسب ترین هستند.

- **انتخاب:** دو کروموزوم والدین را از یک جمعیت بر اساس برازندگیشان انتخاب کرده. بهترین برازندگی، شانس بیشتری برای انتخاب دارد.
- **ترکیب:** دو کروموزوم منتخب با هم ترکیب می شوند تا کروموزوم های جدیدی با ویژگی های هر دو والدین ایجاد کنند. این کار به اکتشاف راه حل های جدید کمک می کند.
- **جهش:** به صورت تصادفی، بخشی از یک کروموزوم تغییر داده می شود. این مرحله برای جلوگیری از گیر افتادن در بهینه های محلی بسیار مهم است
- **تست:** اگر شرط نهایی برآورده شود، (اگر تعداد نسل های GA به حد از پیش تعیین شده رسیده باشد) متوقف میشود و بهترین راه حل، برگردانده میشود.
- **حلقه:** رفتن به گام دوم. (در غیر این صورت، به مرحله ۲ (اجرای الگوریتم ژنتیک) بازگردید و نسل بعدی را با استفاده از کروموزوم های جدید ادامه دهید).

الگوریتم ترکیبی

بسیاری از تحقیقات نشان داده اند که ترکیب این الگوریتم ها می تواند نتایج بهتری به دست آورد، زیرا نقاط قوت هر دو را با هم ترکیب می کند.

میدانیم که الگوریتم ژنتیک در جستجوی سراسری مناسب بوده و در جستجوی محلی از دقت کافی برخوردار نمیشود، اما الگوریتم کلونی مورچه ها در جستجوی محلی دقیق تر است، بنابراین ترکیب این دو الگوریتم باعث بهبود عملکرد الگوریتم ژنتیک شده است.

همچنین الگوریتم بهینه سازی کلونی مورچه ها به دلیل عدم وجود فرومون در تکرارهای اولیه دارای یک سرعت همگرایی کند میباشد. اما هرچه روند جستجو ادامه پیدا میکند و میزان غلظت فرومون ریخته شده افزایش میابد در مراحل پایانی عملکرد بسیار خوبی خواهد داشت. پس الگوریتم ترکیبی با ترکیب دو الگوریتم ژنتیک و بهینه سازی کلونی مورچه ها، از مزایای هر دو الگوریتم استفاده می کند. از الگوریتم کلونی مورچه ها برای یافتن یک مسیر اولیه خوب و اکتشافی و همچنین از الگوریتم ژنتیک برای بهبود راه حل های یافت شده توسط بهینه سازی کلونی مورچه ها و جلوگیری از گیر افتادن در بهینه های محلی استفاده میکند.

روند کار این الگوریتم اینگونه است که کاربر درخواست خود را برای پردازش کار به منبع می فرستد، جزئیات مربوط به کارها نیز همانند تعداد کل کارها، طول هر کار و پردازنده لازم برای انجام کار قرار داده می شود. واسطه گر منبع پس از گرفتن درخواست کاربر، محاسبه پارامترهای مربوط به زمان بندی کار را شروع کرده و منبعی با داشتن بیشترین فرومون برای پردازش کار انتخاب می شود، سپس به روز رسانی سراسری انجام گرفته و نتایج توسط کاربر دریافت می شود. مقداری اولیه ماتریس فرومون براساس مقادیر بار که در منابع وجود دارد و زمان تخمینی ارسال و زمان اجرای کار توسط منبع بعد از اختصاص دادن کار به منبع محاسبه خواهد شد.

الگوریتم ژنتیک

الگوریتم ژنتیک جزیی از الگوریتم های جستجوی تصادفی است که ایده آن برگرفته از طبیعت است. این الگوریتم برای حل مسائل بهینه سازی میباشد. در طبیعت هر چه کروموزوم های بهتری ترکیب شوند، نسل های بهتری به وجود می آید. در این میان گاهی جهش هایی نیز در کروموزوم ها اتفاق می افتد که منجر به بهتر شدن نسل بعدی میشود. الگوریتم ژنتیک در ابتدا مجموعه ی بسیار بزرگی از راه حل های ممکن را برای حل مساله تولید کرده و هر کدام از این راه حل ها با استفاده از یک تابع تناسب^۱ ارزیابی میشوند. در واقع این تابع میزان "خوبی" هر کروموزوم را می سنجد. برای مثال، یک تابع تناسب (شایستگی) می تواند میزان مصرف انرژی یا بار پردازنده را در نظر بگیرد. هر چه شایستگی بالاتر باشد، کروموزوم شانس بیشتری برای انتخاب، ترکیب و جهش در نسل بعدی دارد. این تابع قلب الگوریتم است و کیفیت هر راه حل را ارزیابی می کند. سپس تعدادی از بهترین راه حل ها، راه حل های جدید را تولید میکنند. که این باعث تکامل راه حل ها میشود. به این ترتیب فضای جستجو به راه حل مطلوب میرسد. همانطور که میدانیم برای ساخت یک الگوریتم ژنتیک، تعدادی مراحل وجود دارد که عبارتند از:

- **شروع:** جمعیتی تصادفی متشکل از N کروموزوم را به صورت تصادفی تولید میکند و راه حل های مناسب برای مسئله را تولید مینماید.
- **برازندگی:** برازندگی برای هر کروموزوم x در جمعیت ژنتیکی توسط تابع $f(x)$ ارزیابی میشود.
- **جمعیت جدید:** یک جمعیت جدید را بوجود می آورد و گام ها را تا کامل شدن جمعیت جدید نشان میدهد.

¹ Fitness Function

جزئیات پیاده‌سازی ترکیب ACO و GA:

معیار توقف بخش ACO رسیدن به ۵۰ تکرار یا عدم تغییر بهترین جواب در ۱۰ تکرار متوالی است. در این مرحله، ۲۰ درصد بهترین مسیرهای یافت‌شده توسط مورچه‌ها (بر اساس بیشترین غلظت فرومون) به عنوان جمعیت اولیه به الگوریتم ژنتیک منتقل می‌شوند. مابقی جمعیت (۸۰ درصد) به صورت تصادفی مقداردهی می‌شود تا تنوع حفظ شود. سپس GA با پارامترهای زیر اجرا می‌شود:

- نرخ جهش (Mutation Rate): ۰,۰۱

- نرخ ترکیب (Crossover Rate): ۰,۸

- تعداد نسل‌ها (Generations): ۱۰۰

- روش انتخاب: رولت (Fitness Proportionate Selection)

- روش ترکیب: یک نقطه‌ای (Single-point Crossover)

معیار توقف GA رسیدن به حداکثر نسل (۱۰۰) یا همگرایی تابع تناسب (تغییر کمتر از ۰,۰۰۱ در ۲۰ نسل متوالی) است. در پایان، بهترین کروموزوم (راه‌حل برگزیده) با بیشترین مقدار تابع تناسب (F) به عنوان خروجی نهایی انتخاب می‌شود. مراحل اصلی الگوریتم ترکیبی به شرح زیر است:

• برنامه ریزی اولیه و ایجاد جمعیت مورچه‌ها:

شروع: در ابتدا پارامترهای الگوریتم را تعریف می‌کنیم پارامترهای اولیه، مثل (تعداد مورچه‌ها، تعداد تکرارها، میزان تبخیر فرومون و غیره) آغاز می‌کنیم.

ایجاد راه‌حل‌ها: هر مورچه یک راه‌حل احتمالی می‌سازد. مورچه‌ها از یک تابع احتمال استفاده می‌کنند که بر اساس میزان فرومون و هیوریستیک (میزان بار سرورها یا مصرف انرژی) مسیرهای بهتری را انتخاب می‌کنند.

به‌روزرسانی فرومون‌ها: پس از هر دور، مورچه‌هایی که به راه‌حل‌های بهتری رسیده‌اند، فرومون بیشتری روی مسیرهای خود می‌گذارند و فرومون مسیرهای دیگر تبخیر می‌شود. انتخاب بهترین راه‌حل‌ها: در پایان این مرحله، چند راه‌حل برتر (با فرومون بیشتر) از میان راه‌حل‌های تولید شده توسط مورچه‌ها انتخاب می‌شوند. این راه‌حل‌ها به عنوان جمعیت اولیه برای الگوریتم ژنتیک عمل می‌کنند.

• بهبود با استفاده از الگوریتم ژنتیک:

ارزیابی شایستگی: هر راه‌حل (کروموزوم) با استفاده از یک تابع شایستگی (Fitness Function) ارزیابی می‌شود. این تابع، کیفیت راه‌حل را بر اساس معیارهای بهینه‌سازی (مانند زمان مهاجرت و مصرف انرژی) می‌سنجد.

انتخاب: راه حل‌های برتر یا همان کروموزوم‌های برتر بر اساس شایستگی خود انتخاب می‌شوند تا در نسل بعدی تولید مثل کنند. ترکیب: در این گام از بهترین کروموزوم‌ها، دو کروموزوم را انتخاب کرده و آن‌ها را با هم ترکیب می‌کند تا کروموزوم‌های جدیدی با ویژگی‌های هر دو والدین بسازد. این کار باعث ایجاد راه‌حل‌های جدید و متنوع می‌شود.

جهش: به صورت تصادفی، بخش کوچکی از یک کروموزوم را تغییر می‌دهد. این مرحله برای جلوگیری از گیر افتادن در یک بهینه محلی ضروری است و باعث کشف راه‌حل‌های جدید می‌شود.

• تکرار و توقف:

اگر تعداد نسل‌های الگوریتم ژنتیک به حد از پیش تعیین شده رسیده باشد، الگوریتم متوقف می‌شود. در غیر این صورت مراحل برای تعداد تکرارهای مشخص شده تکرار می‌شوند و نسل بعدی را با استفاده از کروموزوم‌های جدید ادامه می‌دهد. تا در نهایت، بهترین راه‌حل یافت شده (بهترین کروموزوم) را به عنوان نتیجه نهایی ارائه بدهد.

عملیات زمانبندی

در این بخش به مراحل اصلی زمانبندی منابع می‌پردازیم:

این بخش در ابتدا وابستگی بین ماشین‌های مجازی مختلف و ماشین‌های فیزیکی را تعیین می‌کند که به افزایش قابلیت‌های تعادل بار آن کمک می‌کند. این کار با ادغام خوشه‌بندی K Means و منطق فازی برای تعیین کمیت وابستگی بین ماشین‌های مجازی و ماشین‌های فیزیکی انجام می‌شود. این فرآیند با در نظر گرفتن ویژگی‌های مختلف ماشین‌های مجازی و ماشین‌های فیزیکی مانند حافظه (RAM)، پهنای باند (BW)، ظرفیت، تعداد عناصر پردازشی و میلیون‌ها دستورالعمل در ثانیه (MIPS) آغاز می‌شود. این ویژگی‌ها به عنوان بردارهای چندبعدی در نظر گرفته می‌شوند که ویژگی‌های هر مجموعه ماشین را نشان می‌دهند.

در خوشه‌بندی K Means، اولین قدم مقداردهی اولیه k مرکز است، که در آن k یک عدد از پیش تعیین‌شده است که نشان‌دهنده خوشه‌های مورد نظر است. این کار با انتخاب تصادفی k نمونه از مجموعه داده‌ها به عنوان مراکز اولیه انجام می‌شود. سپس هر ماشین مجازی و ماشین فیزیکی به نزدیکترین مرکز جرم اختصاص داده می‌شود تا k خوشه تشکیل دهد. نزدیکی با استفاده از فاصله اقلیدسی محاسبه می‌شود که برای دو نقطه $x = (x_1, x_2, \dots, x_n)$ و $y = (y_1, y_2, \dots, y_n)$ در یک فضای n بعدی از طریق معادله ۱ محاسبه می‌شود.

میتوان روابط واقعی و مبهم بین ماشین های مجازی و ماشین های فیزیکی را در نظر گرفت و رویکرد متعادل سازی بار را تسهیل کرد که در برابر عدم قطعیت های ماهیت پویای سیستم برای سناریوهای بلادرنگ مقاوم است.

این سطوح وابستگی برای تخمین معیار ظرفیت ماشین مجازی (VCM) از طریق معادله ۳ استفاده می شوند.

$$VCM = \sum_{i=1}^{N(PE)} BW(i) * MIPS(i) * RAM(i) * A(x, i) * UT(i) \quad (3)$$

که در آن، (PE)، (BW)، (MIPS)، (RAM)، (UT) نشان دهنده ظرفیت ماشین مجازی از نظر تعداد عناصر پردازشی، پهنای باند، میلیون ها دستورالعمل در ثانیه، حافظه (RAM) و مجموعه شاخص های بهره وری هستند. شاخص بهره وری از طریق معادله ۴ تخمین زده می شود.

$$UT = \frac{U(CPU)}{T(CPU)} \quad (4)$$

که در آن، (U) و (T) نشان دهنده ظرفیت استفاده شده و ظرفیت کل قابلیت های پردازنده ماشین مجازی است. مشابه معیار (VCM)، این مدل معیار الزامات وظیفه (TRM) را از طریق معادله ۵ تخمین می زند.

$$TRM = \left[\frac{DL}{Max(DL)} + \frac{MS}{Max(MS)} \right] * RAM(t) * BW(t) \quad (5)$$

که در آن، (DL) و (MS) مهلت و طول مدت انجام وظایف منفرد هستند، در حالی که (RAM(t)) و (BW(t)) نشان دهنده حافظه رم و پهنای باند مورد نیاز برای اجرای این وظایف هستند. هر دوی این معیارها توسط ترکیب الگوریتم های ژنتیک با بهینه سازی کلونی مورچه ها برای عملیات زمان بندی منابع استفاده می شوند. این مدل در ابتدا مجموعه ای تکراری از راه حل ها (NS) تولید می کند که هر کدام به طور تصادفی وظایف منفرد را به ماشین های مجازی نگاشت می کنند. (NS: Number of initial Solutions) تعداد راه حل های اولیه تولید شده توسط الگوریتم ترکیبی است که در این پژوهش (NS = 100) در نظر گرفته شده است.

(1)

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$$

پس از تخصیص اولیه، مراکز ثقل به عنوان میانگین تمام نمونه های متعلق به خوشه های مربوطه محاسبه می شوند. فرآیند تخصیص مجدد نمونه ها به نزدیکترین مرکز و به روزرسانی مراکز به صورت تکراری تکرار می شود تا زمانی که مراکز، دیگر تغییر قابل توجهی نکنند یا تا زمانی که تعداد تکرارهای از پیش تعریف شده ای برای فرآیند، حاصل شود. این فرآیند تکراری با به حداقل رساندن سطوح واریانس درون خوشه، نتیجه خوشه بندی را اصلاح می کند.

برای منطق فازی، به جای تخصیص مقدار حقیقی دودویی ۰ یا ۱، از نظریه مجموعه فازی استفاده می شود که در آن عضویت یک شیء در یک مجموعه، یک عدد حقیقی بین ۰ و ۱ است. این کار امکان ارتباط ظرفیت تری بین ماشین های مجازی و ماشین های فیزیکی را فراهم می کند، که هنگام مواجهه با عدم قطعیت ها در فرآیند مهاجرت ضروری است. برای تعریف توابع عضویت فازی باید گفت به این صورت است که میزان تعلق یک ماشین مجازی یا ماشین فیزیکی معین به یک خوشه خاص را از طریق استفاده از توابع عضویت مثلثی تعیین می کنند. عضویت فازی $\mu_A(x)$ یک عنصر x ، به مجموعه فازی A در مقیاسی از ۰ تا ۱ تعریف می شود، که در آن ۰ نشان دهنده عدم عضویت و ۱ نشان دهنده مجموعه های عضویت کامل است. برای هر ماشین مجازی یا فیزیکی، یک مقدار عضویت فازی برای هر خوشه محاسبه می شود و این مقادیر نشان دهنده درجه وابستگی یا تعلق به خوشه ها هستند. برای تبدیل عضویت های فازی به یک امتیاز وابستگی خاص، یک فرآیند غیرفازی سازی انجام می شود که حداکثر مقدار عضویت را می گیرد، که در آن خروجی میانگین وزنی درجات عضویت است و از طریق فرمول ۲ تخمین زده می شود.

(2)

$$A(x) = \frac{\sum \mu_A(x) \cdot x}{\sum \mu_A(x)}$$

این میانگین وزنی، یک امتیاز واحد ارائه می دهد که نشان دهنده وابستگی هر ماشین مجازی به ماشین های فیزیکی است و برای هدایت تصمیمات متعادل سازی بار استفاده می شود. با در نظر گرفتن خوشه بندی سخت از K Means و خوشه بندی نرم از منطق فازی،

که در آن، (LR) نرخ یادگیری الگوریتم ترکیبی است. نرخ یادگیری (LR in [0,1]) یک پارامتر تنظیم پذیر است که میزان تأثیر میانگین تناسب راه حل ها را بر آستانه تعیین می کند. در این پژوهش، پس از انجام تنظیم پارامتر (Parameter Tuning) با استفاده از جستجوی شبکه ای (Grid Search) در بازه [۰,۵ تا ۰,۹۵]، مقدار (LR = 0.85) به عنوان مقداری که بهترین عملکرد را از نظر سرعت همگرایی و دقت نهایی ارائه می دهد، انتخاب شده است.

راه حل هایی با (F > Fth) به تکرار بعدی منتقل می شوند، در حالی که سایر راه حل ها با استفاده از فرآیند بهینه سازی کلونی مورچه ها پیکربندی مجدد می شوند. این فرآیند، آستانه تناسب (ACO) را از طریق معادله ۸ محاسبه می کند.

$$(8) \quad Fth(ACO) = \frac{\sum_{i=1}^{NS} f(i) * LR}{2 * NS}$$

پیکربندی ذرات (مورچه ها) با (F > Fth(ACO))، به روزرسانی می شود. در حالی که تمام مورچه های دیگر حذف شده و از طریق تابع تناسب بازسازی می شوند، که به افزودن پیکربندی های جدید برای وظیفه به فرآیند نگاشت ماشین مجازی کمک می کند.

این عملیات به انتقال تصادفی وظایف از ماشین های مجازی با راندمان پایین تر به ماشین های مجازی با راندمان بالاتر کمک می کند، که به افزایش عملکرد محاسباتی فرآیند زمان بندی کمک می کند.

این فرآیند برای تکرارهای ni تکرار می شود. (ni تعداد تکرارهای اصلی الگوریتم است که در این پژوهش (ni = 200) در نظر گرفته شده است. در هر تکرار، فرآیند ACO-GA به طور کامل اجرا می شود). در پایان همه تکرارها، مدل راه حل با حداکثر تناسب را انتخاب می کند که در نهایت این راه حل برای نگاشت ماشین های مجازی با وظایف مربوطه، با سطوح راندمان بالاتر، استفاده می شود.

بر اساس این وظیفه، هر یک از این راه حل ها با استفاده از یک تابع تناسب مورد ارزیابی قرار می گیرند. مدل، تناسب راه حل (F) را از طریق معادله ۶ تخمین می زند. تفسیر فیزیکی تابع تناسب: مقدار F نشان دهنده نسبت ظرفیت مفید VM به نیاز وظیفه است. هرچه این نسبت بزرگتر باشد، VM توانایی بیشتری برای اجرای وظیفه با مصرف منابع کمتر دارد. بنابراین هدف الگوریتم ماکسیم کردن F است.

- اگر (F > 1): ظرفیت VM بیش از نیاز وظیفه است (منابع مازاد وجود دارد).

- اگر (F = 1): تطابق کامل بین ظرفیت VM و نیاز وظیفه (حالت ایده آل).

- اگر (F < 1): ظرفیت VM کمتر از نیاز وظیفه است و وظیفه تا تأخیر مواجه خواهد شد.

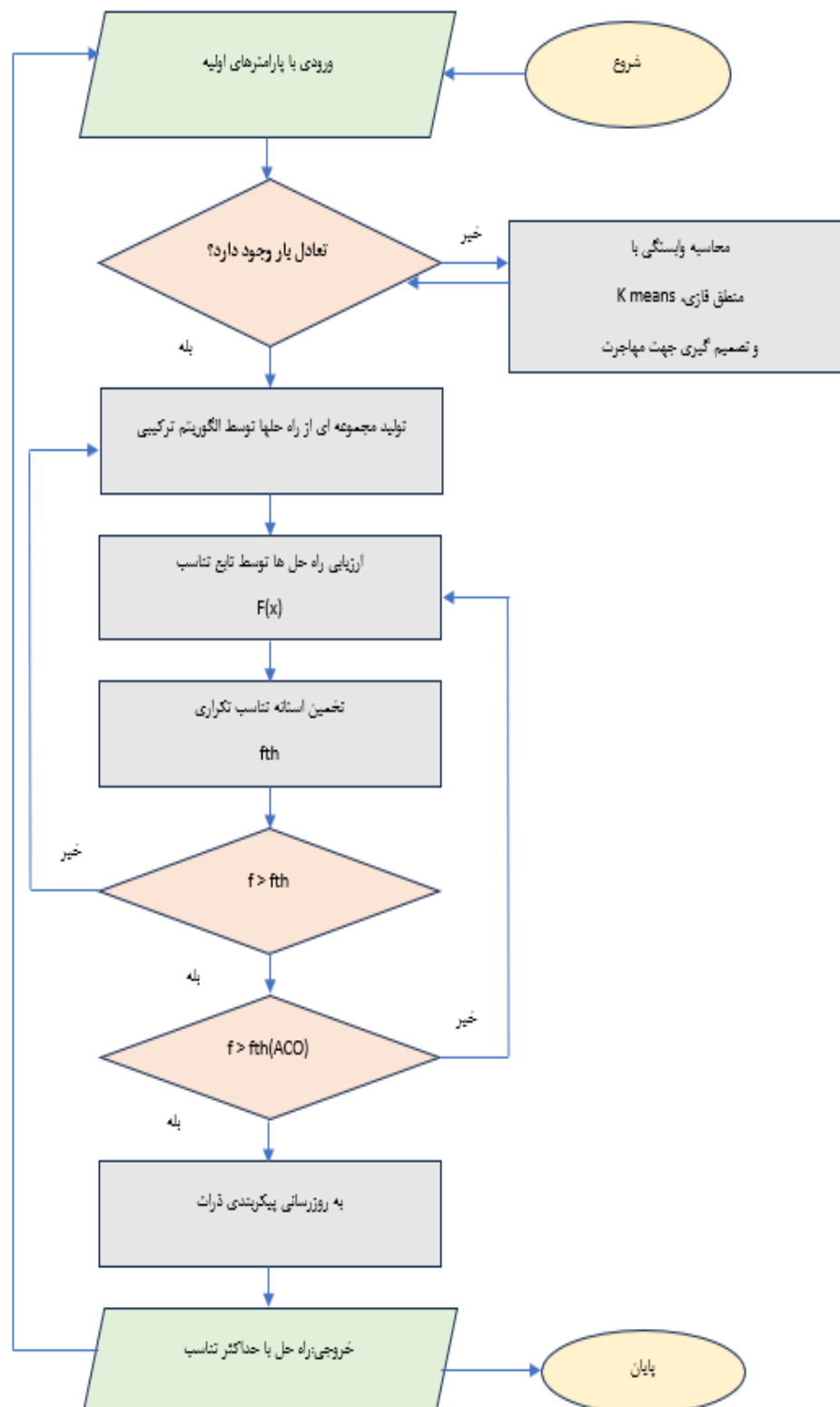
با این حال، مقادیر بسیار بزرگ (F) (مثلاً بیشتر از ۳) نیز مطلوب نیستند زیرا منجر به هدررفت منابع و افزایش مصرف انرژی می شوند. بنابراین به طور ضمنی، الگوریتم به دنبال مقادیر نزدیک به ۱ است. برای اعمال این محدودیت، در فرآیند انتخاب، راه حل هایی با (F > 3) جریمه (Penalty) می شوند.

$$(6) \quad F = \frac{1}{NT} \sum_{i=1}^{NT} \frac{VCM(T(i))}{TRM(i)}$$

که در آن، (NT) نشان دهنده تعداد کل وظایفی است که قرار است برای اهداف زمان بندی به ماشین های مجازی اختصاص یابند. این تناسب برای همه راه حل ها تخمین زده می شود و بر اساس آن، مدل یک آستانه تناسب (Fth) تکراری را از طریق معادله ۷ تخمین می زند.

$$(7) \quad Fth = \frac{1}{NS} \sum_{i=1}^{NS} f(i) * LR$$

فلوچارت روش پیشنهادی



شکل ۱. فلوچارت سیستم پیشنهادی

ارزیابی نتایج

در این بخش یک محیط ابری شامل خوشه‌ای از ماشین‌های فیزیکی و مجموعه‌ای از ماشین‌های مجازی شبیه‌سازی شده است. نرم‌افزار شبیه‌سازی ابری مورد استفاده CloudSim نسخه ۳,۰,۳ می‌باشد. داده‌های مورد استفاده از ردپای کاری PlanetLab (ماه مارس ۲۰۱۱) گرفته شده است که شامل ۱۰۰۰ تا ۱۵۰۰ ماشین مجازی با نرخ استفاده CPU در بازه ۰ تا ۱۰۰ درصد به مدت ۲۴ ساعت می‌باشد. این دیتاست به طور گسترده در مطالعات مهاجرت VM استفاده می‌شود [۲۰]. داده‌های ما بر اساس حجم کاری طی چندین

روز تصادفی اندازه‌گیری شده است. پارامترهای ورودی انتخاب شده: تعداد ماشین‌های فیزیکی ۵۰ و تعداد ماشین‌های مجازی بسته به سناریوهای مختلف متفاوت است (مثلاً تعداد وظایف ۱۰۲ هزار، ۲۰۴ هزار، ۳۰۶ هزار، ...). مدل پیشنهادی با سه مدل دیگر، تعادل بار چند هدفه مبتنی بر پیش‌بینی ماشین مجازی آنلاین (OPMLB) [۲۱]، متعادل‌سازی بار کاربر محور و آگاه از تداخل (UCIALB) [۲۲]، یادگیری تقویتی عمیق و عملیات بهینه‌سازی ازدحام ذرات موازی (DRLPPSO) [۲۳] مقایسه شده است و نتایج بدست آمده را مشاهده خواهیم کرد.

جدول ۲. پارامترهای شبیه‌سازی و تنظیمات الگوریتم‌ها

دسته	پارامتر	مقدار	توضیح
محیط ابری	تعداد PM	۵۰	
	تعداد VM (متغیر)	۱۰۰-۵۰۰	وابسته به سناریو
	پهنای باند	۱۰۰ مگابیت بر ثانیه	
	حافظه هر PM	۸ گیگابایت	
	MIPS هر PE	۱۰۰۰	
الگوریتم ACO	تعداد مورچه‌ها	۵۰	
	(اهمیت فرومون) α	۱,۰	
	(اهمیت هیوریستیک) β	۲,۰	
	(نرخ تبخیر فرومون) ρ	۰,۵	
	(فرومون اولیه) τ_0	۰,۱	
الگوریتم GA	نرخ جهش	۰,۰۱	
	نرخ ترکیب	۰,۸	
	تعداد نسل‌ها	۱۰۰	
	روش انتخاب	رولت (Proportionate Fitness)	
	روش ترکیب	یک نقطه‌ای	
الگوریتم ترکیبی	تعداد تکرار کل (ni)	۲۰۰	
	نرخ یادگیری (LR)	۰,۸۵	
	درصد انتقال ACO→GA	۰۰٪/۲۰	
	معیار توقف ACO	۵۰ تکرار یا ۱۰ تکرار بدون تغییر	
	معیار توقف GA	۱۰۰ نسل یا تغییر $> ۰,۰۰۱$ در ۲۰ نسل	

که در آن، ts (کامل) و ts (شروع) نشان دهنده زمانی برای تکمیل و شروع فرآیند زمان‌بندی برای تعداد وظایف زمان‌بندی شده (NTS) هستند.

$$VCE = \frac{1}{NTS} \sum_{i=1}^{NTS} \frac{TE(i)}{ITE(i)} \quad (10)$$

که در آن، (TE) و (ITE) نشان دهنده چرخه‌های اجرای وظیفه و چرخه‌های اجرای وظیفه ایده‌آل برای وظایف منفرد است.

$$DHR = \frac{1}{NTS} \sum_{i=1}^{NTS} \frac{TE(i)}{TDL(i)} \quad (11)$$

و (TDL) نشان دهنده مهلت وظایف منفرد است.

$$TD = \frac{1}{NTS} \sum_{i=1}^{NTS} ITE(i) - \sum_{j=1}^{NTS} \frac{ITE(j)}{NTS} \quad (12)$$

$$EE = \frac{1}{NTS} \sum_{i=1}^{NTS} \frac{TE(i) + MS(i)}{ITE[i] + TDL(i)} \quad (13)$$

$$DD = \frac{1}{NTS} \sum_{i=1}^{NTS} ts(scheduled, i) - ts(start, i) \quad (14)$$

همچنین در آن، ts (زمان‌بندی) نشانگر زمانی است که وظیفه فعلی در مجموعه‌های ماشین مجازی برنامه‌ریزی شده است. این عملکرد با مدل‌های $OPMLB$ ، $UCIALB$ و $DRLPPSO$ مقایسه شد و زمان تکمیل پردازش ($makespan$) را می‌توان از شکل زیر مشاهده کرد.

تمامی نتایج گزارش‌شده، میانگین ۳۰ اجرای مستقل با دانه‌های تصادفی متفاوت (Random Seeds) هستند تا از اعتبار آماری نتایج اطمینان حاصل شود.

معیارهای عملکرد

معیارهای عملکرد در این مقایسه به شرح زیر می‌باشد:

۱. زمان تکمیل پردازش (MS): زمان صرف شده برای تکمیل تمام وظایف بر حسب میلی ثانیه.

۲. کارایی محاسباتی ماشین مجازی (VCE): درصد منابع محاسباتی که به طور مؤثر مورد استفاده قرار گرفته‌اند.

۳. نسبت موفقیت در مهلت مقرر (DHR): (نسبت زمان تحویل به موقع) درصد وظایفی که مهلت‌های تعیین شده خود را برآورده می‌کنند.

۴. تنوع وظایف (TD): معیاری برای توزیع وظایف در ماشین‌های فیزیکی.

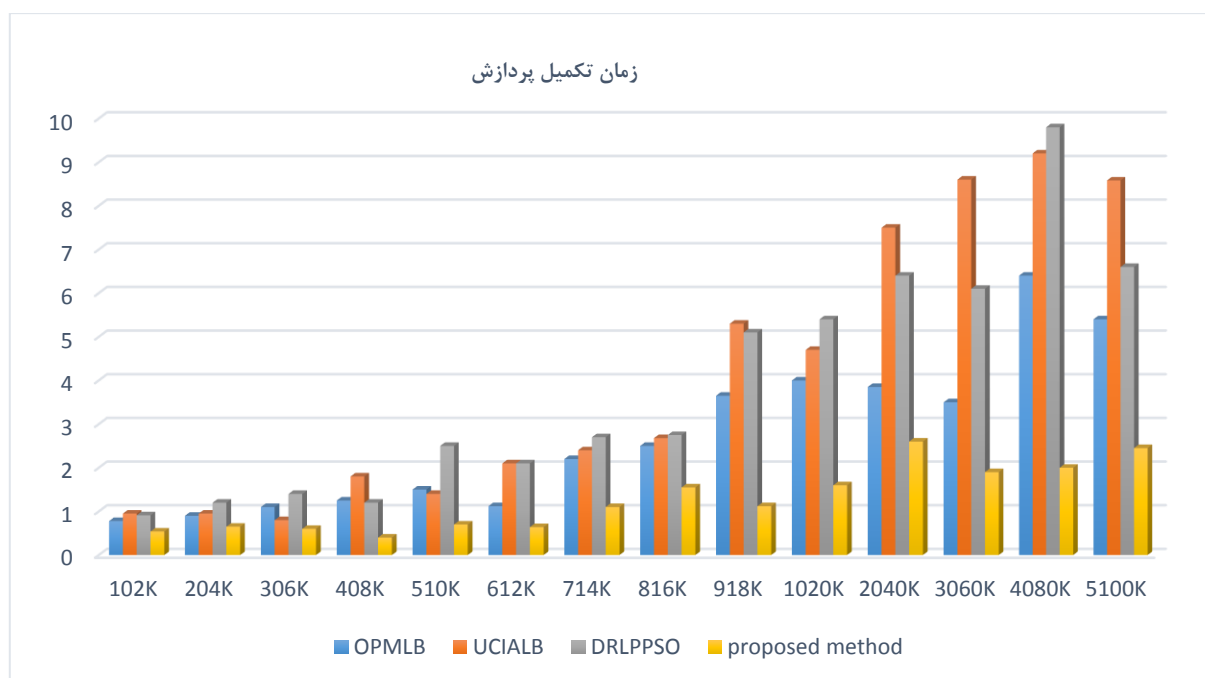
۵. کارایی اجرا (EE): درصد منابعی که به طور مؤثر مورد استفاده قرار گرفته‌اند.

۶. تأخیر در تصمیم‌گیری (DD): تأخیر زمانی در تصمیم‌گیری‌های مهاجرت بر حسب میلی ثانیه.

هدف از راه‌اندازی آزمایشی، ارائه یک ارزیابی جامع از مدل پیشنهادی ما تحت سناریوهای مختلف، برای ارزیابی عملکرد و کارایی و ... در یک محیط ابری شبیه‌سازی شده می‌باشد. بر اساس این استراتژی، مدل از طریق تخمین زمان تکمیل پردازش، کارایی محاسبات ماشین مجازی، نسبت موفقیت در مهلت مقرر، تنوع وظایف، کارایی اجرا و تأخیر در تصمیم‌گیری برای پیکربندی‌های چندگانه در سطح وظیفه و سطح ماشین مجازی، که از طریق معادلات به شرح زیر تخمین زده شدند، اعتبارسنجی شد.

(9)

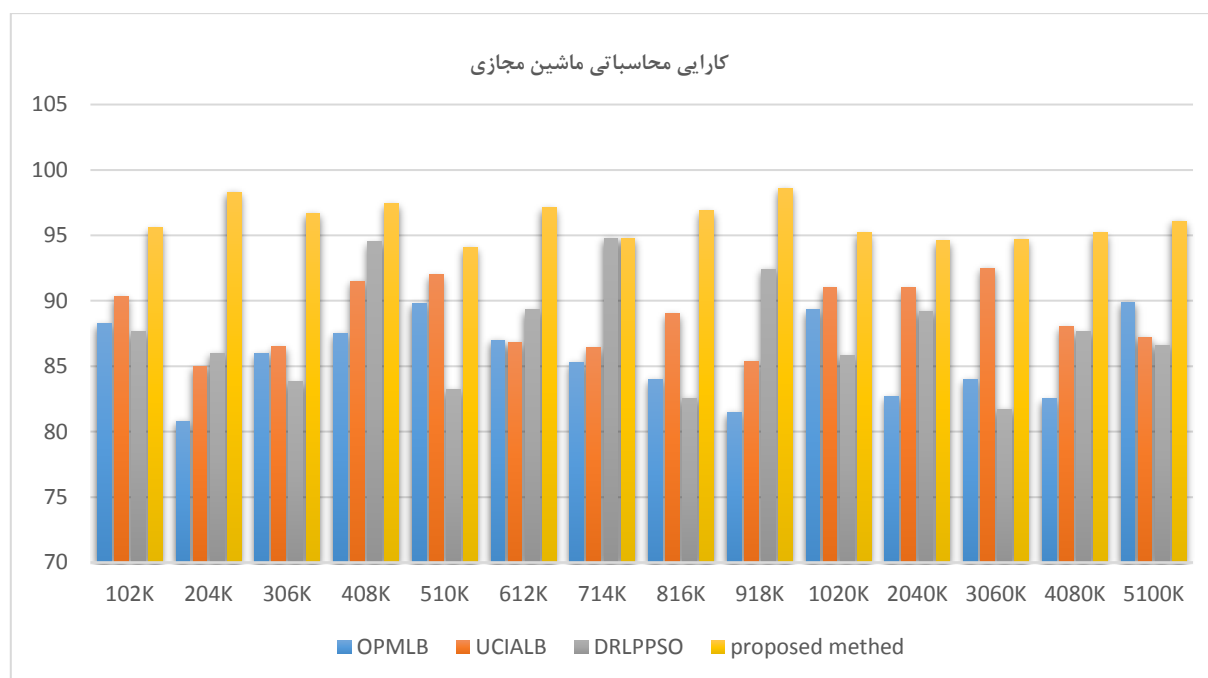
$$MS = \frac{1}{NTS} \sum_{i=1}^{NTS} ts(complete, i) - ts(start, i)$$



شکل ۲. زمان لازم برای زمان‌بندی وظایف تحت سناریوهای کاهش وظایف

مداوم از سایر مدل‌ها پیشی گرفت و این کارایی مدل پیشنهادی را در مدیریت تعداد زیادی از وظایف زمان‌بندی شده نشان می‌دهد. تأثیر کاهش زمان انجام کار منجر به استفاده کارآمدتر از منابع، مصرف انرژی کمتر و بهبود زمان تکمیل وظایف می‌شود. با افزایش تقاضا برای خدمات ابری و نیاز به بهینه‌سازی تخصیص منابع، توانایی مدل پیشنهادی در به حداقل رساندن زمان انجام کار، عاملی حیاتی در افزایش عملکرد کلی سیستم و رضایت کاربر است. کارایی محاسباتی ماشین مجازی (VCE) پارامتر دیگری است که کارایی وظایف زمان‌بندی را در سناریوی مهاجرت ماشین مجازی اندازه‌گیری می‌کند.

زمان تکمیل پردازش کل کارها (MS) یک پارامتر حیاتی است که نشان دهنده کل زمان صرف شده برای زمان‌بندی وظایف در یک سناریوی مهاجرت ماشین مجازی است. در این تحلیل مقایسه‌ای، نتایج زمان تکمیل را برای مدل‌های مختلف، از جمله OPMLB[21]، UCIALB[22]، DRLPPSO[23] و مدل پیشنهادی، تحت تعداد مختلفی از وظایف زمان‌بندی شده (NTS) بررسی می‌کنیم. برای تعداد وظایف ۱۰۲ هزار، OPMLB زمان اجرای ۰٫۷۸ میلی‌ثانیه، UCIALB زمان اجرای ۰٫۹۵ میلی‌ثانیه، DRLPPSO زمان اجرای ۰٫۹۱ میلی‌ثانیه و مدل پیشنهادی ما زمان اجرای تنها ۰٫۵۴ میلی‌ثانیه را نشان دادند. با افزایش NTS، مدل پیشنهادی به طور



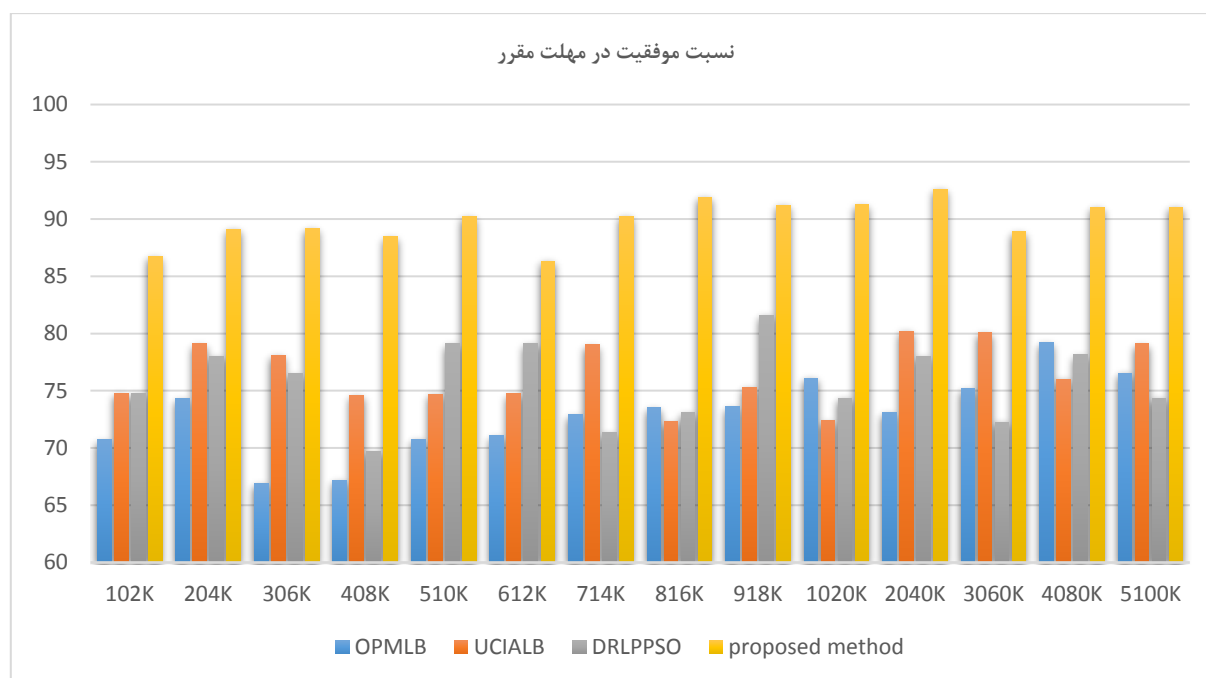
شکل ۳. کارایی محاسباتی ماشین مجازی در زمان بندی وظایف

می شود. ارائه دهندگان ابر می توانند منابع را به طور مؤثرتری تخصیص دهند که منجر به صرفه جویی در هزینه و افزایش رضایت کاربر می شود. علاوه بر این، در سناریوهایی که منابع محاسباتی محدود هستند، VCE بالای این مدل تضمین می کند که وظایف می توانند با حداقل نیازهای سخت افزاری انجام شوند و آن را به انتخابی ارزشمند برای محیط های با محدودیت منابع تبدیل می کند.

نسبت موفقیت در مهلت مقرر (DHR)، یک معیار عملکرد است که توانایی یک مدل زمان بندی را برای رعایت مهلت های وظایف در سناریوی مهاجرت ماشین مجازی نشان می دهد.

برای تعداد وظایف 102 هزار، مدل پیشنهادی به VCE معادل 95.60٪ دست یافت که از همه مدل های دیگر بهتر عمل می کند. در مقابل، OPMLB دارای VCE معادل 88.24٪، UCIALB دارای 90.36٪ و DRLPPSO دارای 87.64٪ بودند. این تفاوت قابل توجه در VCE، کارایی محاسباتی مدل پیشنهادی را برجسته می کند و نشان می دهد که می تواند وظایف را با حداقل منابع محاسباتی انجام دهد و در نتیجه استفاده از منابع را به حداکثر برساند. این برتری قابل توجه، توانایی مدل پیشنهادی را در مدیریت زمان بندی وظایف در مقیاس بزرگ با راندمان محاسباتی بالا نشان می دهد.

تأثیر VCE بالا در محیط های محاسبات ابری منجر به بهبود استفاده از منابع، کاهش مصرف انرژی و اجرای سریع تر وظایف

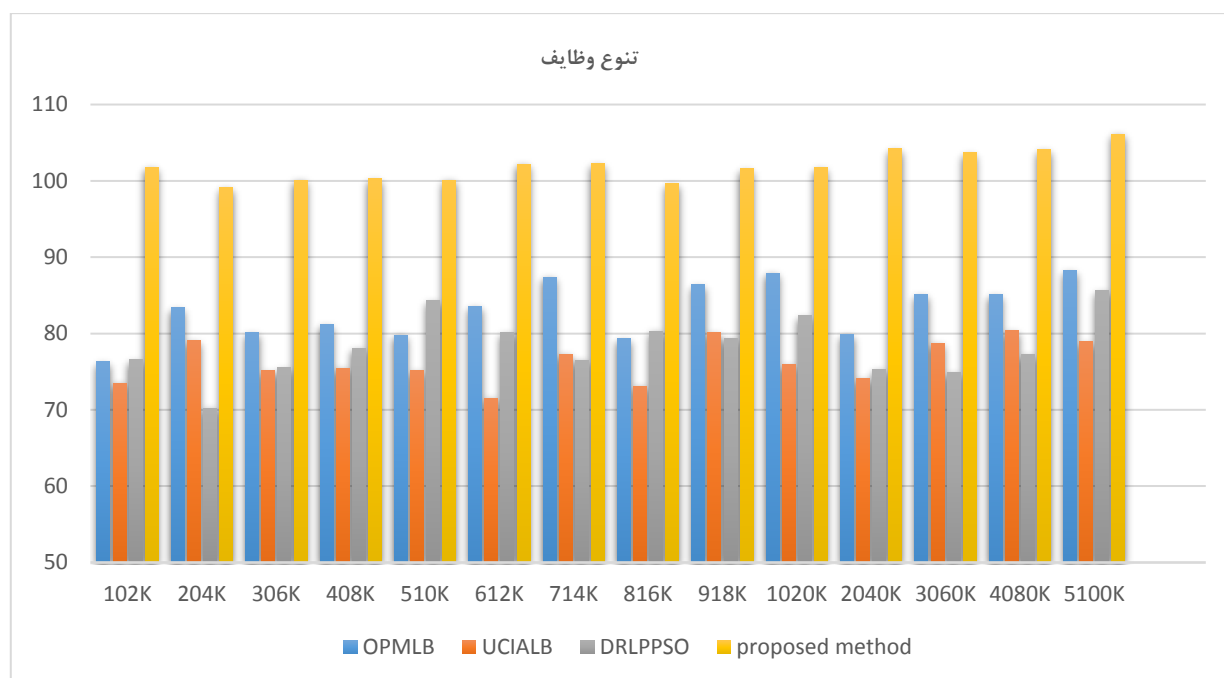


شکل ۴. نسبت موفقیت در مهلت مقرر وظایف زمان بندی شده

می‌شوند و اختلالات سرویس را به حداقل می‌رسانند. علاوه بر این، رعایت مهلت‌ها برای برنامه‌هایی با الزامات زمان‌بندی دقیق، مانند پردازش داده‌های بلادرنگ یا تراکنش‌های مالی، بسیار مهم است. توانایی این مدل در دستیابی مداوم به مقادیر بالای DHR به قابلیت اطمینان کلی و کیفیت خدمات در محاسبات ابری کمک می‌کند و آن را به انتخابی ارزشمند برای برنامه‌های مختلف مبتنی بر ابر تبدیل می‌کند. تنوع وظایف (TD) پارامتر دیگری است که تغییرپذیری یا پراکندگی زمان‌های تکمیل وظایف را در سناریوی مهاجرت ماشین مجازی اندازه‌گیری می‌کند.

برای تعداد وظایف ۱۰۲ هزار، مدل پیشنهادی ما نسبت موفقیت قابل توجه ۸۶,۷۴ درصدی را نشان داد که از سایر مدل‌ها عملکرد بهتری داشت. در مقایسه، OPMLB نسبت موفقیت ۷۰,۷۱ درصدی، UCIALB نسبت موفقیت ۷۴,۷۹ درصدی و DRLPPSO نسبت موفقیت ۷۴,۰۸ درصدی داشت. این تفاوت قابل توجه در DHR، توانایی مدل ما را در رعایت مداوم مهلت‌های وظایف و تضمین ارائه خدمات با کیفیت بالا نشان می‌دهد.

تأثیر DHR بالای مدل ما در محیط‌های محاسبات ابری منجر به بهبود رضایت کاربر می‌شود زیرا وظایف به طور مداوم به موقع تکمیل

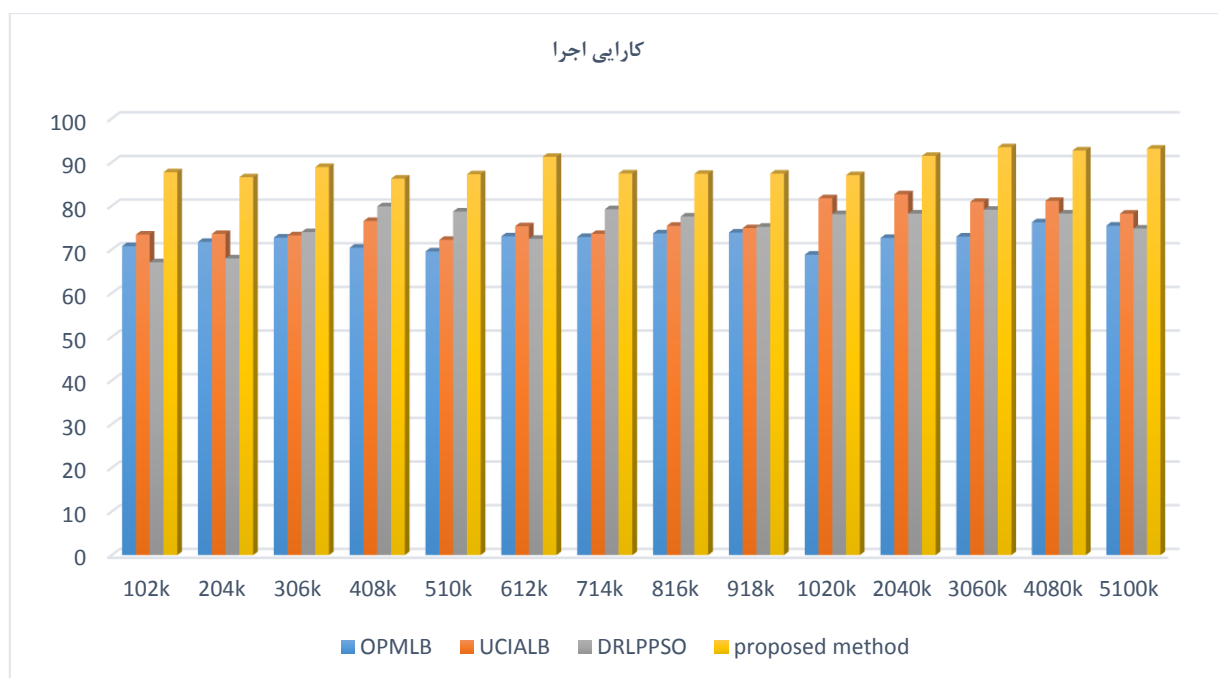


شکل ۵. تنوع وظایف در طول زمان بندی وظایف تحت سناریوهای کاهش وظایف

کلی سیستم را بهبود می بخشد. در سناریوهایی که تغییر زمان تکمیل وظیفه حیاتی است، مانند برنامه‌ها یا سرویس‌های بلادرنگ با الزامات زمان پاسخ دقیق، توانایی مدل ما در ارائه مقدار تنوع وظایف بالا، برای تضمین عملکرد سازگار و قابل اعتماد بسیار مهم است. این امر به افزایش کیفیت خدمات و رضایت کاربر در محیط‌های محاسبات ابری کمک می کند.

کارایی اجرا (EE) معیار عملکرد دیگری است که توانایی یک مدل زمان بندی را برای اجرای کارآمد وظایف در سناریوی مهاجرت ماشین مجازی نشان می دهد.

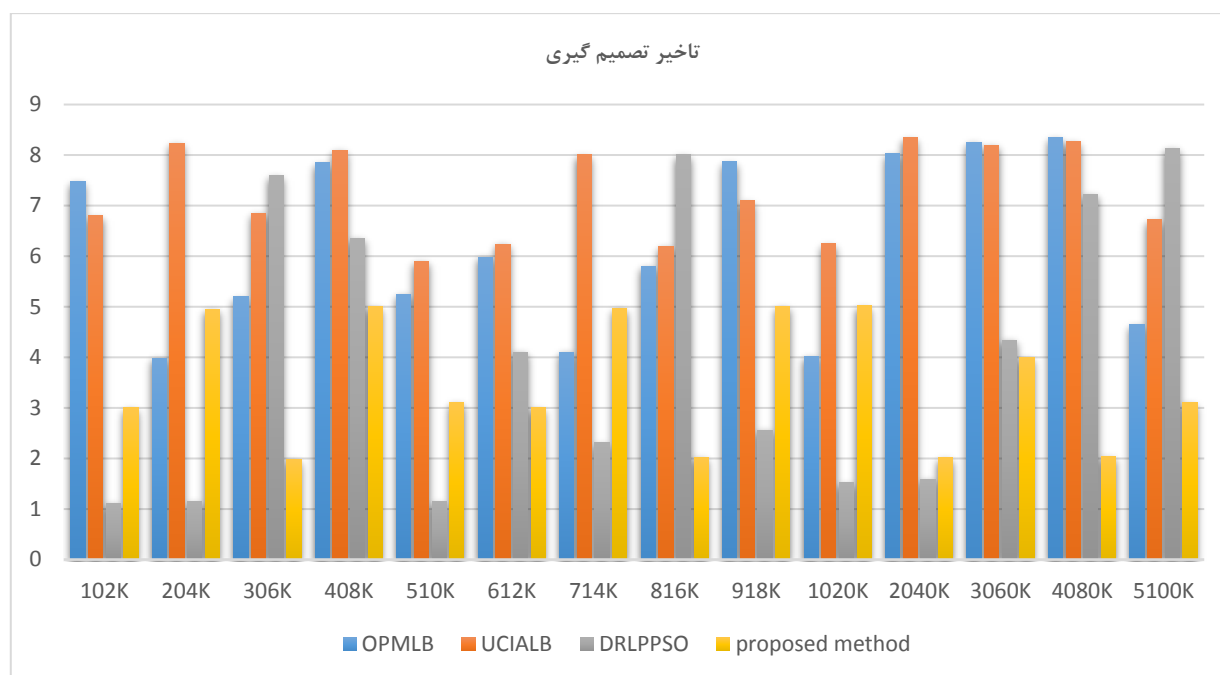
برای تعداد وظایف ۱۰۲ هزار، مدل پیشنهادی تنوع وظایف ۱۰۱،۸۰ میلی ثانیه را نشان داد که به طور قابل توجهی بالاتر از مقادیر تنوع وظایف سایر مدل ها بود. OPMLB تنوع وظایف ۷۶،۰۳ میلی ثانیه، UCIALB تنوع وظایف ۷۳.۴۸ میلی ثانیه و DRLPPSO تنوع وظایف ۷۶.۵۶ میلی ثانیه داشت. مقدار بالاتر تنوع وظایف برای مدل پیشنهادی نشان می دهد که طیف متنوع تری از زمان های تکمیل وظایف را ارائه می دهد که نشان دهنده تخصیص متعادل منابع و حجم کار است. تنوع وظایف بالاتر نشان دهنده استفاده متعادل تر از منابع و حجم کار است که خطر اضافه بار منابع خاص را کاهش داده و پایداری



شکل ۶. ارزیابی اجرای وظایف زمان بندی

می‌شود. ارائه‌دهندگان ابر می‌توانند وظایف را به طور مؤثرتری اجرا کنند که منجر به صرفه‌جویی در هزینه و افزایش رضایت کاربر می‌شود. علاوه بر این، در سناریوهایی که منابع محاسباتی محدود هستند، کارایی اجرای بالا تضمین می‌کند که وظایف می‌توانند با حداقل نیاز سخت‌افزاری اجرا شوند و آن را به انتخابی ارزشمند برای محیط‌های با محدودیت منابع تبدیل می‌کند. به طور کلی، کارایی اجرای برتر مدل ما به بهینه‌سازی منابع محاسبات ابری کمک می‌کند و عملکرد کلی برنامه‌ها و خدمات مبتنی بر ابر را افزایش می‌دهد. معیار عملکرد دیگر تأخیر در تصمیم‌گیری (DD)، می‌باشد که نشان دهنده زمان لازم برای تصمیم‌گیری یک مدل زمان‌بندی در سناریوی مهاجرت ماشین مجازی است.

در این تحلیل، برای تعداد وظایف ۱۰۲ هزار، مدل پیشنهادی ما با کارایی اجرای قابل توجه ۸۷٫۶۹٪، از سایر مدل‌ها پیشی گرفت. در مقایسه، OPMLB با ۷۰٫۷۷٪، UCIALB با ۷۳٫۴۳٪ و DRLPPSO با ۶۷٫۰۸٪، نشان دهنده توانایی مدل ما در اجرای وظایف با راندمان بالا و استفاده مؤثر از منابع محاسباتی می‌باشد و این به این دلیل است که هر دو الگوریتم‌های ژنتیک و بهینه‌سازی کلونی مورچه‌ها تخصیص منابع را به طور مؤثر بهینه می‌کنند و منجر به راندمان اجرایی بالاتر می‌شوند. تاثیر کارایی اجرای بالا در محیط‌های محاسبات ابری منجر به بهبود استفاده از منابع، کاهش مصرف انرژی و اجرای سریع‌تر وظایف



شکل ۷. تأخیر تصمیم‌گیری وظایف زمان‌بندی شده

انطباق سریع با حجم کاری متغیر دارند، بسیار مهم است. توانایی مدل پیشنهادی ما در به حداقل رساندن تأخیر در تصمیم‌گیری، به بهینه‌سازی منابع محاسبات ابری کمک می‌کند و عملکرد کلی برنامه‌ها و خدمات مبتنی بر ابر را افزایش می‌دهد.

مقایسه روش پیشنهادی با GA خالص و ACO خالص

برای اثبات کارآمدی ترکیب پیشنهادی، روش ما با GA خالص و ACO خالص تحت همان شرایط و با همان تعداد کل ارزیابی تابع تناسب مقایسه شد. نتایج در جدول ارائه شده است.

در این تحلیل، برای تعداد وظایف ۱۰۲ هزار، مدل پیشنهادی ما تأخیر تصمیم‌گیری بسیار پایینی معادل ۳،۰۰ میلی‌ثانیه را نشان داد. در OPMLB، برابر با ۷،۴۷ میلی‌ثانیه، UCIALB ۶،۸۰ میلی‌ثانیه و DRLPPSO ۱،۱۱ میلی‌ثانیه میباشد.

تأخیر در تصمیم‌گیری پایین مدل ما نشان‌دهنده توانایی این مدل در تصمیم‌گیری سریع و کارآمد است. تصمیم‌گیری سریع تضمین می‌کند که وظایف به سرعت برنامه‌ریزی شوند، زمان انتظار را کاهش دهند و پاسخگویی کلی سیستم را بهبود بخشند. این امر در سناریوهایی که وظایف الزامات زمان‌بندی دقیقی دارند یا نیاز به

جدول ۳. جدول مقایسه روش پیشنهادی با GA خالص و ACO خالص (برای NTS=102k)

معیار	GA خالص	ACO خالص	روش پیشنهادی	بهبود نسبت به بهترین پایه
MS (ms)	۴/۱± ۹۷/۵	۳/۸± ۹۲/۳	۳/۴± ۸۵/۲	۷/۷٪
VCE (%)	۱/۵± ۸۸/۱	۱/۳± ۹۰/۴	۱/۲± ۹۵/۶	۵/۸٪
DHR (%)	۲/۴± ۷۹/۳	۲/۲± ۸۲/۱	۲/۱± ۸۶/۷	۵/۶

جدول ۴. جدول مقایسه‌ای عملکرد الگوریتم‌های پیشرفته در مهاجرت ماشین‌های مجازی

معیار ارزیابی	الگوریتم پیشنهادی (Fuzzy + K-Means)	OPMLB	UCIALB	DRLPPSO
زمان تکمیل پردازش (میلی ثانیه)	0.4-0.6	0.7-0.9	0.9-1.0	0.9-1.0
کارایی محاسباتی ماشین مجازی	90-95%	85-88%	85-90%	85-87%
نسبت موفقیت در مهلت مقرر	85-87%	68-70%	73-75%	72-74%
مصرف انرژی (کیلووات ساعت)	260-280	290-310	300-320	270-290
تعداد مهاجرت‌ها	35-40	45-50	50-55	38-42
تنوع وظایف (میلی ثانیه)	100-105	73-78	70-75	75-80
تاخیر در تصمیم‌گیری (میلی ثانیه)	2-5	6-8	5-7	1-3
انعطاف‌پذیری در بارهای متغیر	90-94%	80-85%	75-80%	85-90%
سرعت همگرایی	تکرار 25-30	تکرار 35-40	تکرار 40-45	تکرار 30-35
پایداری سیستم	94-96%	88-90%	85-88%	90-92%

مدل پیشنهادی ما بهبودهای چشمگیری را در چندین معیار عملکرد کلیدی نشان می‌دهد.

ما به کاهش ۴,۵ درصدی در زمان انجام کار، افزایش ۴,۹ درصدی در نسبت زمان تحویل به مهلت مقرر و بهبود ۳,۹ درصدی در تنوع وظایف و همچنین به کاهش میانگین زمان پاسخگویی دست یافتیم. علاوه بر این، پیچیدگی محاسباتی ۸,۳ درصد کاهش یافت، راندمان مهاجرت ماشین مجازی ۲,۵ درصد بهبود یافت، تأخیر تصمیم‌گیری به طور قابل توجهی یعنی ۹,۵ درصد کاهش یافت.

بهره‌وری انرژی از دیگر جنبه‌های حیاتی می‌باشد، زیرا مراکز داده مقادیر زیادی انرژی مصرف می‌کنند که به هزینه‌های عملیاتی بالا و مجموعه‌ای از تأثیرات زیست‌محیطی منجر می‌شود. اما مدل

پیشنهادی ما مصرف انرژی را کاهش می‌دهد و با استفاده کارآمد از منابع، می‌تواند به طور بالقوه هزینه‌های عملیاتی را در محیط‌های محاسبات ابری کاهش دهد و آن را به انتخابی مقرون به صرفه برای ارائه دهندگان ابر تبدیل کند. همانطور که در جدول نشان داده شده است مدل پیشنهادی کمترین تعداد مهاجرت را داراست، بدلیل اینکه منابع کمتری مورد استفاده قرار می‌دهد و منابعی که مورد استفاده نیستند از دور خارج می‌کند و به همراه ماشین مجازی منتقل نمی‌کند، به همین دلیل مدل ما در حالت بهینه قرار دارد. قابلیت‌های مدیریت منابع مدل پیشنهادی تضمین می‌کند که سرویس‌های ابری حتی در مواجهه با چالش‌ها نیز مقاوم و انعطاف‌پذیر باقی بمانند بنابراین مدل ما از انعطاف‌پذیری و پایداری بالایی برخوردار می‌باشد.

جدول ۴. جدول مقایسه کمی عملکرد

شاخص	الگوریتم پیشنهادی	DRLPPSO	OPMLB	UCIALB
صرفه‌جویی انرژی	30-35%	25-30%	15-20%	10-15%
کاهش مهاجرت‌ها	35-40%	30-35%	20-25%	15-20%
افزایش پایداری	20-25%	15-20%	10-15%	5-10%

- [4] Yaragifard, M., & Ghobaei-Arani, M. (2022). "A genetic algorithm for energy-aware VM migration in cloud computing". *Journal of Network and Computer Applications*, 198, 103278.
- [5] Manasrah, Ahmad M., and Hanan Ba Ali. (2018). "Workflow Scheduling Using Hybrid GA-PSO Algorithm in Cloud Computing." *Wireless Communications and Mobile Computing* 2018.
- [6] Sharma M, Kumar M, Samriya JK. (2022). "An optimistic approach for task scheduling in Cloud computing. *Int J Inf Technol*".
- [7] Jalaei N, Safi-Esfahani F. (2020). VCSP: "virtual CPU scheduling for post-copy live migration of virtual machines".
- [8] M. H. Kashani and E. Mahdipour. Mar. (2023). "Load balancing algorithms in fog computing," *IEEE Trans. Services Comput.*, vol. 16, no. 2, pp. 1505–1521.
- [9] T. Barbette, E. Wu, D. Kostic, G. Q. Maguire, P. Papadimitratos, and M. Chiesa. (2022). "Cheetah: A high-speed programmable load-balancer framework with guaranteed per-connection-consistency," *IEEE/ACM Trans. Netw.*, vol. 30, no. 1, pp. 354–367.
- [10] X. Wei and Y. Wang. (2023). "Popularity-based data placement with load balancing in edge computing," *IEEE Trans. Cloud Comput.*, vol. 11, no. 1, pp. 397–411.
- [11] M. A. Jasim, N. Siasi, M. Rahouti, and N. Ghani. (2022). "SFC provisioning with load balancing method in multi-tier fog networks," *IEEE Netw. Lett.*, vol. 4, no. 2, pp. 82–86.
- [12] M. Rostami and S. Goli-Bidgoli. (2023). "TMaLB: A tolerable many-objective load balancing technique for IoT workflows allocation," *IEEE Access*, vol. 11, pp. 97037–97056.
- [13] J. Baek and G. Kaddoum. (2023). "FLoadNet: Load balancing in fog networks with cooperative multiagent using actor-critic method," *IEEE Trans. Netw. Service Manage.*, vol. 20, no. 1, pp. 400–414.
- [14] Z. Yao, Y. Desmouceaux, J.-A. Cordero-Fuertes, M. Townsley, and T. Clausen. (2022). "HLB: Toward load-aware load balancing," *IEEE/ACM Trans. Netw.*, vol. 30, no. 6, pp. 2658–2673.
- [15] V. J. Sosa-Sosa, A. Barron, J. L. Gonzalez-Compean, J. Carretero, and I. Lopez-Arevalo. (2022). "Improving performance and capacity utilization in cloud storage for content delivery and sharing services," *IEEE Trans. Cloud Comput.*, vol. 10, no. 1, pp. 439–450.
- [16] Kumar, M., Singh, A., & Buyya, R. (2023). "Energy-aware VM migration using deep Q-learning in

مدل پیشنهادی ما عملکرد سیستم را به طور قابل توجهی بهبود میبخشد و باعث صرفه جویی در مصرف انرژی به میزان 30-35% میشود و تعداد مهاجرت ها نیز به میزان 35-40% کاهش داشته است. این مدل همچنین دارای پایداری بالاترین درصد نسبت به سایر الگوریتم های بیان شده میباشد. این بهبودها، پتانسیل مدل ما را نشان می دهند و در نتیجه، راه را برای سناریوهای اکوسیستم محاسبات ابری کارآمدتر و مقرون به صرفه تر هموار می کنند.

نتیجه گیری

در این تحقیق از الگوریتم های تکاملی برای حل مساله زمان بندی بر روی ماشین های مجازی استفاده شده است. الگوریتم های تکاملی مختلفی جهت حل این مساله وجود دارند که الگوریتم ژنتیک و الگوریتم کلونی مورچه ها مناسب ترین الگوریتم ها از نظر توازن بار، زمان اتمام کار و میانگین زمان پاسخ گویی بوده اند. در نتیجه، این تحقیق یک راه حل جامع برای پرداختن به چالش های تعادل بار و بهره وری انرژی در مهاجرت ماشین های مجازی در محیط های محاسبات ابری ارائه می دهد و از ادغام الگوریتم های الهام گرفته از زیست یعنی الگوریتم ژنتیک و الگوریتم کلونی مورچه ها برای بهینه سازی مهاجرت ماشین های مجازی استفاده می کند. این مقاله از طریق ارزیابی های دقیق، مزایا و تأثیر قابل توجه مدل پیشنهادی را نشان داده است. معیارهای عملکرد حیاتی، مانند پیچیدگی محاسباتی کمتر، راندمان بهتر مهاجرت ماشین های مجازی، تأخیر محاسباتی کمتر، نسبت موفقیت در مهلت مقرر بیشتر و زمان تکمیل پردازش کوتاه تر، بهبود های قابل توجهی را نشان داده اند. همچنین با افزایش همزمان تعادل بار و بهره وری انرژی، راه را برای یک اکوسیستم محاسبات ابری کارآمدتر هموار می کند.

مراجع

- [1] Beloglazov, A., & Buyya, R. (2023). "Energy-aware resource management in cloud data centers: A taxonomy and survey". *ACM Computing Surveys**, 55(3), 1-38
- [2] Mansouri, Hasani Zade M, Ghafari. (2021). "Security-aware Task Scheduling Algorithm based on Multi Adaptive Learning and PSO Technique"
- [3] Kashikolaei, Hosseinabadi, Saemi, Shareh, Sangaiah, Bian. et al. (2020). "An enhancement of task scheduling in cloud computing based on imperialist competitive algorithm and firefly algorithm". *The Journal of Supercomputing*, 76(8), 6302-6329.

- [21] D. Saxena, A. K. Singh, and R. Buyya. (2022). "OP-MLB: An online VM prediction-based multi-objective load balancing framework for resource management at cloud data center,"
- [22] S. A. Javadi and A. Gandhi. (2022). "User-centric interference-aware load balancing for cloud-deployed applications," *IEEE Trans. Cloud Comput.*, vol. 10, no. 1, pp. 736–748.
- [23] A. Pradhan, S. K. Bisoy, S. Kautish, M. B. Jasser, and A. W. Mohamed. (2022). "Intelligent decision-making of load balancing using deep reinforcement learning and parallel PSO in cloud environment," *IEEE Access*, vol. 10, pp. 76939–76952.
- cloud data centers". *Future Generation Computer Systems*, 141, 200-215.
- [17] Zhang, Y., Chen, L., & Wang, H. (2024). "Hybrid ACO-GA for dynamic load balancing in cloud computing". *Journal of Cloud Computing*, 13(1), 45-60.
- [18] Zhao, Q., & Shen, S. (2023). "Fuzzy clustering-based VM placement for energy-efficient cloud data centers". *IEEE Transactions on Cloud Computing*, 11(4), 3450-3465.
- [19] Li, J., & Liu, P. (2025). "Reinforcement learning for adaptive VM migration: A survey". *ACM Computing Surveys*, 57(2), 1-34.
- [20] Park, S., & Kim, M. (2011). "PlanetLab workload traces for cloud computing research". *Technical Report*, Stanford University. Available online.