

Fall 2024, 5 (3), 101-113
DOR:

Received: 26 June 2024
Accepted: 6 Aug 2024

مقاله پژوهشی

Task Scheduling in Cloud Computing Using Deep Reinforcement Learning

Behnam Salari Hamzehkhani¹, Mehdi Akbari^{2*}, Faramarz Safi-Esfahani³, Behrang Barekatain⁴

1. Ph.D. Student, Faculty of Computer Engineering, Najafabad Branch, Islamic Azad University, Najafabad, Iran.
2. Assistant Professor, Faculty of Computer Engineering, Najafabad Branch, Islamic Azad University, Najafabad, Iran.
**Corresponding Author, mehdi.akbari@iaun.ac.ir*
3. Associate professor, Faculty of Computer Engineering, Najafabad Branch, Islamic Azad University, Najafabad, Iran.
4. Associate professor, Faculty of Computer Engineering, Najafabad Branch, Islamic Azad University, Najafabad, Iran.

Abstract

Introduction: Task scheduling and energy consumption are important issues in cloud computing. Failure to use an appropriate scheduling approach in cloud computing may result in high energy consumption and low resource efficiency. Due to the dynamics and limitations of cloud resources to execute diverse and time-varying requests from users, an effective scheduling mechanism is required that adapts to dynamic system conditions. Deep learning, an applied method for dealing with cloud computing resource management problems, has been considered as an innovative idea in recent years. Among the deep learning algorithms used in this field, we can mention DQL, ERDQL, and DRL-LSTM algorithms, which have achieved acceptable results in large datasets. However, these algorithms have problems such as unrealistic reward estimation, long-term training, random sampling of memory, and lack of proper use of state space, which affect the results obtained from their use in optimization problems.

Method: In this research, an improved deep reinforcement learning algorithm for task scheduling in cloud computing is proposed to improve the problems of unrealistic reward estimation, long-term training, and random sampling from memory. Also, the Markov decision process structure in the modeling part is defined in such a way that it can mitigate the problems of not using the state space properly.

Results: The proposed algorithm has been evaluated with two scenarios including a synthetic dataset and GOCJ (Google Cloud Jobs) dataset. The results of the evaluations in the simulated environment show that the proposed algorithm has achieved better results in less time in all evaluated parameters.

Discussion: In this paper, an improved algorithm based on PERDQN is presented. Then, with a proper definition of the cloud environment based on the Markov decision process, the proposed algorithm is used to schedule tasks on virtual machines in cloud computing. The results show that using the proposed algorithm, response time, waiting time, throughput, MakeSpan and the number of SLA violations are improved, which is more significant on a larger scale.

Keywords: Task scheduling, cloud computing, deep reinforcement learning, Markov decision process

زمانبندی کارها در محاسبات ابری با استفاده از یادگیری عمیق تقویتی

دوره پنجم، پاییز ۱۴۰۳
شماره سوم، صص: ۱۱۳-۱۰۱

تاریخ دریافت: ۱۴۰۳/۰۴/۰۶
تاریخ پذیرش: ۱۴۰۳/۰۵/۱۶

بهنام سالاری حمزه‌خانی^۱، مهدی اکبری^{۲*}، فرامرز صافی اصفهانی^۳، بهرنگ برکتین^۴

۱. دانشجوی دکتری، گروه مهندسی کامپیوتر، واحد نجف‌آباد، دانشگاه آزاد اسلامی، نجف‌آباد، ایران. salaribe@yahoo.com

۲. استادیار، گروه مهندسی کامپیوتر، واحد نجف‌آباد، دانشگاه آزاد اسلامی، نجف‌آباد، ایران. (نویسنده مسئول) mehdi.akbari@iaun.ac.ir

۳. دانشیار، گروه مهندسی کامپیوتر، واحد نجف‌آباد، دانشگاه آزاد اسلامی، نجف‌آباد، ایران. fsafi@iaun.ac.ir

۴. دانشیار، گروه مهندسی کامپیوتر، واحد نجف‌آباد، دانشگاه آزاد اسلامی، نجف‌آباد، ایران. Behrang_Barekatin@iaun.ac.ir

چکیده: زمان‌بندی وظایف و مصرف انرژی از جمله مسائل مهم در محاسبات ابری هستند. عدم استفاده از یک رویکرد زمان‌بندی مناسب در محاسبات ابری ممکن است باعث مصرف انرژی بالا و بهره‌وری پایین منابع شود. به علت پویایی و محدودیت منابع ابری برای اجرای درخواست‌های متنوع و متغیر با زمان کاربران، یک مکانیزم زمان‌بندی مؤثر مورد نیاز است که خود را با شرایط پویای سیستم وفق دهد. یادگیری عمیق، یک روش کاربردی برای مقابله با مشکلات مدیریت منابع رایانش ابری به عنوان یک ایده نوآورانه در سال‌های اخیر مورد توجه قرار گرفته است. از جمله الگوریتم‌های یادگیری عمیق استفاده شده در این زمینه می‌توان از الگوریتم‌های DQL، ERDQL و DRL-LSTM نام برد که در دیتاست‌های بزرگ به نتایج قابل قبولی رسیده‌اند. ولی این الگوریتم‌ها دارای مشکلاتی از جمله تخمین غیر واقعی پاداش، آموزش طولانی مدت، نمونه‌برداری تصادفی از حافظه و عدم استفاده مناسب از فضای حالت می‌باشند که بر نتایج حاصل از استفاده آن‌ها در مسائل بهینه‌سازی مؤثر است. در این تحقیق یک الگوریتم یادگیری عمیق تقویتی بهبود یافته برای زمان‌بندی کارها در محاسبات ابری ارائه شده است تا مشکلات تخمین غیرواقعی پاداش، آموزش طولانی مدت و نمونه‌برداری تصادفی از حافظه را بهبود دهد. همچنین ساختار فرآیند تصمیم‌گیری مارکوف در قسمت مدل‌سازی به گونه‌ای تعریف می‌شود که مشکلات عدم استفاده مناسب از فضای حالت را تعدیل نماید. لذا با استفاده از الگوریتم پیشنهادی نتایج مورد انتظار این است که منجر به بهبود زمان پاسخ، زمان انتظار، برون‌ده، زمان کل کارها و تعداد نقض SLA گردد. در این مقاله، الگوریتم توسعه یافته با مجموعه داده GOCJ (Google Cloud Jobs) مورد ارزیابی قرار می‌گیرد.

واژه‌های کلیدی: زمان‌بندی کارها، محاسبات ابری، یادگیری عمیق تقویتی، فرآیند تصمیم‌گیری مارکوف.

۱. مقدمه

هدف اصلی زمان‌بندی کار، به حداقل رساندن اتلاف زمان و به حداکثر رساندن عملکرد است به طوری که زمان پاسخ کار دریافت شده در محدوده قابل قبول بوده و بهره‌وری منابع نیز رعایت شده باشد [۱]. روش‌های رایج در زمان‌بندی دارای معایبی از قبیل پیچیدگی زمانی زیاد و افزایش زمان اجرای برنامه است. بنابراین، روشی مناسب است که تولید زمان اجرای کل کمینه گردد [۲-۴]. زمان‌بندی کارها با الگوریتم‌هایی مانند اکتشافی سنتی قادر به تطبیق سریع با تغییر و رسیدن کار تصادفی نیستند [۵]. ساختار پیچیده کارها با توالی پردازش انعطاف‌پذیر تصمیم‌گیری را مشکل نموده است. رویکردهای موجود در برآورد الزامات دینامیک و پاسخ سریع و یا برای مواجهه با وظایف با توالی‌های پردازش انعطاف‌پذیر مشکل دارند [۶].

با توجه به اینکه بارکاری برنامه‌های مبتنی بر ابر غیرقابل پیش‌بینی و متغیر هستند، انجام زمان‌بندی کار بصورت آنلاین برای استفاده در محیط ابری مناسب است و روش‌های مبتنی بر RL^۱ برای این سناریو مطلوب هستند، زیرا نیازی به دانستن حجم کاری کاربر، انتقال وضعیت و سایر اطلاعات سیستم اصلی ندارند. عامل RL در طول زمان تصمیمات مربوط به اختصاص کار به ماشین مجازی^۲ را دنبال می‌کند و بنابراین داده‌های آموزشی عظیمی را جمع‌آوری کرده تا یاد بگیرد که چگونه در آینده تصمیمات بهتری بگیرد و برای حل مسائل بهینه‌سازی پیچیده عملکرد مناسبی دارد و در طیف وسیعی از مسائل بهینه‌سازی قابل استفاده است [۷-۹]. به‌طور کلی استفاده از تکنیک یادگیری عمیق تقویتی در تحقیقات اخیر جهت زمان‌بندی کارها در محاسبات ابری مورد توجه قرار گرفته است [۱۰-۱۲]. الگوریتم‌های یادگیری عمیق تقویتی در زمان‌بندی ابری به‌علت اینکه پویایی محیط را در نظر گرفته و موجب بهبود عملکرد محاسبات ابری در بلندمدت می‌گردند لذا مبنای این تحقیق قرار گرفته است.

یادگیری تقویتی نوعی از روش‌های یادگیری ماشین است که یک عامل را قادر به یادگیری در محیطی تعاملی با استفاده از آزمون و خطا و استفاده از بازخوردهای اعمال و تجربیات خود می‌سازد و از پاداش‌ها و تنبیه‌ها به عنوان سیگنال‌هایی برای رفتار مثبت و منفی استفاده می‌شود. یادگیری عمیق تقویتی، ترکیبی از یادگیری تقویتی و یادگیری عمیق است که در آن عامل سعی در حداکثر کردن پاداش دارد. از جمله این الگوریتم‌ها SARSA و learning-Q است. learning-Q یک الگوریتم معروف و مستقل از مدل برای یادگیری تقویتی است. تمایز این الگوریتم با سایر الگوریتم‌ها در استراتژی‌های جست‌وجوی آن محسوب می‌شود. به‌دلیل اینکه اقدام بعدی با هدف حداکثر کردن و وضعیت بعدی value-Q انتخاب شده است، لذا learning-Q یک روش مستقل از سیاست است. این روش حول محور به‌روزرسانی ارزش‌های Q است که بر اساس معادله بلمن می‌باشد. [۱۳]

در این الگوریتم جدولی برای Qها تشکیل و بر اساس آن تصمیم‌گیری انجام می‌شود. با توسعه این الگوریتم برای رفع مشکلات آن، الگوریتم LQD^۳ مورد استفاده قرار گرفت. در این الگوریتم فرمول‌های QL وجود داشته و به‌جای جدول Qها از شبکه عصبی

استفاده شده است. ورودی شبکه عصبی حالت‌ها و خروجی آن مقادیر Q به ازای هر عمل می‌باشد. در این الگوریتم از TD^۴ برای محاسبه خطا استفاده می‌شود [۱۴، ۱۵]. الگوریتم DQN دارای ایراداتی از جمله تخمین غیرواقعی پاداش می‌باشد [۱۶].

سیس الگوریتم ERDQL^۵ توسعه یافت در این الگوریتم با ذخیره تجربیات قبلی در حافظه و نمونه‌برداری تصادفی از آن جهت آموزش شبکه عصبی، کارایی الگوریتم DQL بهبود یافته و پایدارتر می‌باشد. از دو شبکه عصبی یکی برای محاسبه مقادیر Q در حالت فعلی و دیگری برای حالت بعدی استفاده می‌شود. در این الگوریتم در هر اپیزود، تجربیات عامل در حافظه Replay ذخیره می‌شوند و از این حافظه به صورت تصادفی جهت آموزش شبکه نمونه‌برداری می‌شود. در ابتدای نمونه‌برداری محیط کاوش شده و یک عمل به صورت تصادفی انتخاب می‌شود و در مراحل بعدی بر اساس الگوریتم حریر صانه عمل با بالاترین مقدار Value-Q انتخاب می‌شود که در واقع یک off-trade بین استخراج و اکتشاف است [۱۷، ۱۸].

بنابراین، الگوریتم‌های یادگیری تقویتی عمیق که تاکنون برای زمان‌بندی وظایف در رایانش ابری استفاده شده‌اند، دارای مزایا و معایبی هستند. به عنوان مثال، در الگوریتم NDQ، تخمین بیش از حد پاداش و استفاده از شبکه یکسان برای انتخاب عمل و ارزیابی آن، ممکن است باعث نتایج غیربهرینه شود. الگوریتم ERDQN ممکن است به دلیل نمونه‌برداری تصادفی از حافظه، دوره آموزشی طولانی داشته باشد و الگوریتم‌های ترکیبی، پیچیدگی محاسباتی را در محیط ابری افزایش می‌دهند. برای رفع مشکل نمونه‌برداری تصادفی از حافظه در الگوریتم ERDQL که ممکن است تجربیات مفید کمتر یا اصلاً مورد استفاده قرار نگیرند الگوریتم PERDQL^۶ پیشنهاد گردید. لذا با اولویت‌بندی تجربیات توسط درخت باینری و قرارگیری اولویت‌ها در برگ‌های این درخت یک آرایه تشکیل می‌شود که به برگ‌های درخت اشاره نموده و حاوی تجربیات است. برای پیاده‌سازی این الگوریتم تجربیات قبلی در نوعی از حافظه با ظرفیت محدود ذخیره می‌گردند و در طول فرآیند یادگیری، تعدادی از رکوردهای ذخیره شده در حافظه بر اساس اولویت انتخاب می‌شوند [۱۹، ۲۰]. در این الگوریتم با وجود اینکه تا حدودی مشکلات الگوریتم‌های قبلی رفع شده است اما برای اطمینان از کارایی آن برای زمان‌بندی وظایف در محیط محاسبات ابری نیازمند بهبود می‌باشد. در الگوریتم پایه، وزن‌های شبکه هدف به صورت دوره‌ای تغییر می‌کنند. وقتی وزن‌های شبکه‌ها با یکدیگر مقایسه می‌شوند بدلیل به‌روزرسانی وزن‌های شبکه هدف در هر d تکرار الگوریتم، باعث افزایش تعداد تکرار مورد نیاز برای همگرایی و در نتیجه افزایش زمان یادگیری می‌شود. این مسئله برای محیط‌های پیچیده و بزرگ مانند محاسبات ابری اهمیت بیشتری دارد.

یادگیری تقویتی معمولاً در قالب فرآیند تصمیم‌گیری مارکوف مدل‌سازی^۷ می‌شود که یک چارچوب ریاضی است برای مدل‌سازی تصمیم‌گیری، در شرایطی که نتایج تا حدودی تصادفی و تحت کنترل

تصمیم‌گیر می‌باشد. این چارچوب برای مطالعه طیف گسترده‌ای از مسائل بهینه‌سازی که از طریق برنامه‌نویسی پویا حل می‌شوند مفید است. این فرآیند شامل پنج عنصر حالت^۸، عمل^۹، احتمال^{۱۰}، پاداش^{۱۱} و ضریب کاهش^{۱۲} می‌باشد [۲۱]. در این تحقیق برای تشریح محیط در یادگیری تقویتی مورد استفاده قرار گرفته‌است که هرچه خصوصیات بیشتری از کارها در نظر گرفته‌شود فضای حالت فرآیند تصمیم‌گیری مارکوف کامل‌تر می‌شود به‌عنوان نمونه یکی از دلایل شکست یک وظیفه، تخصیص نامناسب کارها و عبور از مهلت زمانی^{۱۳} می‌باشد [۲۲]. بنابراین مهلت زمانی کارها در این تحقیق در نظر گرفته‌شده‌است. سپس از طریق طراحی تابع پاداش، مدل را می‌توان برای اهداف مختلف و بارهای کاری بهینه‌کرد.

با توجه به مطالب عنوان شده، اگر بتوان به‌روزرسانی شبکه هدف را در هر تکرار الگوریتم یادگیری عمیق انجام‌داد آنگاه همگرایی در الگوریتم PERDQL افزایش یافته و در زمان کوتاه‌تری نسبت به الگوریتم پایه آموزش دیده و نتایج بهتری نیز به‌دست خواهد‌آمد. این موضوع برای زمانبندی کارها در محیط محاسبات ابری حائز اهمیت است. موضوع دیگری که در این تحقیق به آن پرداخته‌شده‌است چگونگی تعریف محیط با استفاده از MDP برای الگوریتم پیشنهادی است به‌طوری‌که ویژگی‌های بیشتری از کارها را پوشش‌دهد.

۲. کارهای مرتبط

در سال‌های اخیر تحقیقات متعددی در رابطه با زمانبندی کارها در محیط محاسبات ابری با استفاده از الگوریتم‌های یادگیری ماشین انجام شده‌است. از جمله الگوریتم‌های یادگیری ماشین که برای زمانبندی وظایف در ابر به جواب‌های قابل‌قبولی دست‌یافته‌اند می‌توان به یادگیری عمیق تقویتی اشاره نمود که در اغلب آن‌ها از مشتقات الگوریتم Learning_Q استفاده شده‌است [۲۳].

در مقاله ارائه‌شده توسط Wei و همکاران [۸] یک چارچوب تصمیم‌گیری مبتنی بر رویداد جهت زمانبندی وظایف بر روی ماشین‌های مجازی بیان‌شده‌است. بدین صورت که هر تصمیم زمانبندی پس از ورود درخواست کار جدید گرفته می‌شود. با توجه به این‌که دستیابی به مدیریت منابع و بهینه‌سازی عملکرد در محیط پویا و نامطمئن برای ارائه‌دهندگان برنامه‌های مبتنی بر ابر یک چالش بزرگ است لذا به تخصیص کارها به منابع مناسب و تأمین نیازمندی‌های QOS مدنظر کاربر پرداخته‌شده‌است. در این مقاله منظور از QOS زمان پاسخ کارها می‌باشد که از مجموع زمان مورد نیاز برای اجرای کار روی ماشین مجازی و زمان سپری شده در صف انتظار به‌دست می‌آید. در پایان نیز با سناریوهای مختلف از کارهایی که بیشتر محاسباتی یا ورودی و خروجی هستند نتایج مقایسه‌گردیده‌است. در این تحقیق علی‌رغم مدلسازی مناسب محیط در MDP، اما در قسمت الگوریتم، زمان آموزش زیادی مورد نیاز است.

در مقاله Swarup و همکاران [۲۴] از الگوریتم یادگیری عمیق learning q deep double جهت زمانبندی کارها در ابر استفاده شده‌است. برای پیاده‌سازی روش پیشنهادشده در این مقاله، محیط شبیه‌سازی کلودسیم و پایتون مورد استفاده قرار گرفته‌است. برای مجموعه اولیه تسک‌ها به صورت تصادفی زمانبندی انجام شده و برای تسک‌ها در مراحل بعدی سریع‌ترین ماشین مجازی انتخاب می‌شود برای این کار مدل Sequential با چهار لایه dense مورد استفاده قرار می‌گیرد و در نهایت سعی در کاهش زمان و هزینه اجرا دارد. با توجه به تحلیل نویسندگان این مقاله انجام زمانبندی سریع همچنان به‌عنوان چالش مطرح است.

در مقاله Rjoub و همکاران [۲۵] چهار الگوریتم RL, DQN, RNN-LSTM و LSTM-DRL برای زمانبندی در محاسبات ابری مورد بررسی قرار گرفته‌است. بعد از حدود ۱۰۰۰ تکرار، الگوریتم‌ها آموزش دیده و برای موارد جدید با داده تست، دقت الگوریتم بررسی شده‌است. در نهایت با نمودار مشخص شده‌است که LSTM-DRL عملکرد بهتری نسبت به بقیه دارد و در کل دقت آموزش در این چهار الگوریتم به ترتیب زیر بیان شده‌است:

$$RL < LSTM-RNN < DQN < LSTM-DRL$$

در حالت کلی با استفاده از روش‌های ترکیبی در شبکه‌های عصبی عمیق، دقت و عملکرد مدل بهبود می‌یابد [۲۶]. لذا با الگوریتم ترکیبی با DRL عملکرد بهتری جهت زمانبندی ابر نسبت به DQL به‌دست آمده‌است. اگرچه نتایج بسیار امیدوارکننده‌ای به‌دست آمده‌است، اما مسئله‌ای که باید مورد توجه قرار گیرد زمان زیاد مورد نیاز رویکرد LSTM-DRL است که ناشی از این واقعیت است که لایه LSTM باید به عقب برگردد و تاریخچه کامل حالت‌ها را بررسی نماید.

تعداد زیادی از مقالات نیز از الگوریتم‌های ترکیبی برای زمانبندی کارها در محاسبات ابری استفاده شده‌است. به‌عنوان نمونه در مقالات [۲۷] و [۲۸] الگوریتم A2C برای زمانبندی کارها پیشنهاد شده‌است این الگوریتم متشکل از دو عامل است یکی از آن‌ها مسئول یادگیری خودکار زمانبندی و دیگری خطای تخمین را کاهش می‌دهد. در واقع از دو شبکه عصبی یکی برای عمل و دیگری برای تخمین ارزش مورد استفاده قرار گرفته‌است. به‌طور کلی، استفاده از الگوریتم‌های ترکیبی برای زمانبندی وظایف، سربار محاسباتی در محیط ابری دارند. در مقاله Li و همکاران [۲۹] از الگوریتم DDQL جهت زمانبندی در محیط ابر استفاده شده‌است. زمانبندی یک وظیفه را به صف ماشین مناسب اختصاص می‌دهد. فضای حالت در این تحقیق بر اساس خصوصیات کار و ماشین مجازی که کار به آن اختصاص داده شده‌است تعیین می‌شود. عمل مورد نظر انتخاب ماشین مجازی است و محاسبه پاداش بر اساس دو هدف زمان کل^{۱۴} و هزینه تعریف شده‌است که این اهداف با هم تضاد دارند و می‌بایست بین این دو تعادل برقرار شود. به عنوان مثال، یک استراتژی کم هزینه، میزان استفاده از ماشین مجازی با قدرت پردازش پایین را افزایش می‌دهد که زمان

صرف شده را افزایش می دهد. اگر زمان کل هدف اصلی باشد میزان استفاده از ماشین مجازی با قدرت پردازش بالا افزایش خواهد یافت. بنابراین با استفاده از وزن ها تعادل بین این دو برقرار می شود. لذا پاداش به صورت رابطه (۱) می باشد [۲۰].

$$R = (r^1 * r^1_{rate} + r^2_i * r^2_{rate}) * \max(1 - r^3, 0.1) - r^3 \quad (1)$$

که در معادله فوق r^1 زمان کل، r^2 هزینه کار، r^3 زمان انتظار کار و r_{rate} وزن جهت کنترل پاداش می باشد. به طور خلاصه در این مقاله از الگوریتم DDQL که تکامل یافته DQL می باشد استفاده شده است. الگوریتم DDQL تجربیات قبلی را در حافظه ذخیره کرده و سپس با نمونه برداری تصادفی شبکه عصبی آموزش دیده و موجب بهبود عملکرد الگوریتم DQL می شود. اما شبکه عصبی را با نمونه گیری تصادفی آموزش می دهد که موجب افزایش زمان یادگیری می شود.

همچنین در مقاله های [۳۰] و [۳۱] از الگوریتم یادگیری عمیق تقویتی ERDQL برای بهینه سازی زمان کل و مصرف انرژی استفاده شده است. در مرجع [۳۰] زمان کل از مجموع دو زمان اجرا و انتظار محاسبه می شود، جهت محاسبه مصرف انرژی نیز ابتدا بهره وری بر اساس تقسیم تعداد ماشین های مجازی فعال روی سرور بر حداکثر تعداد ماشین مجازی قابل راه اندازی روی یک سرور به دست می آید و در نهایت مصرف انرژی در دو حالت استاتیک و دینامیک با استفاده از بهره وری، محاسبه شده و بر اساس تعداد سرورهای فیزیکی در دیتا سنتر با هم جمع می شوند. با توجه به اینکه دو هدف مطرح بوده لذا پاداش بر اساس زمان کل و میزان مصرف انرژی است و در پایان نرمال سازی جهت قرار دادن پاداش در بازه ۰ و ۱ انجام می شود. در این مقاله، فضای حالت تعریف شده برای زمان بندی وظایف در محیط ابری به صورت محدود تعریف شده است. در تحقیق [۳۱] نیز روشی جهت کاهش هزینه انرژی در دیتا سنترها و زمان بندی کارها ارائه شده است که موجب بهبود قابل ملاحظه ای در زمینه مصرف انرژی شده و همگرایی مناسبی دارد. دو مرحله یکی برای اختصاص کارها به یکی از نواحی سرورها و دیگری برای مشخص کردن سرور دارد. برای رسیدن به اهداف مشخص شده از DQN Replay Experience استفاده شده است که در آن اطلاعات در حافظه ذخیره شده و با نمونه برداری تصادفی از آن شبکه عصبی آموزش می بیند. از دو شبکه عصبی یکی برای محاسبه eQ_Value در حالت فعلی و دیگری برای محاسبه در حالت بعدی مورد استفاده قرار می گیرد و سپس به روز رسانی انجام می شود. که در نهایت پایداری این الگوریتم را نسبت به DQN پایه افزایش می دهد، اما به زمان آموزش بیشتری نیاز دارد. به طور کلی، استفاده از الگوریتم های یادگیری عمیق تقویتی برای دو هدف زمان بندی وظایف در محیط ابری و کاهش مصرف انرژی مرکز داده، نیازمند تعادل در تعریف تابع پاداش است.

با توجه به کارهای گذشته در خصوص استفاده از الگوریتم های یادگیری عمیق تقویتی برای زمان بندی کارها در محاسبات ابری که به نتایج قابل قبولی نسبت به سایر روش های زمان بندی رسیده اند. بنابراین مبنای این تحقیق قرار گرفته است ولی همچنان موارد یاد شده در قسمت مقدمه پابرجاست. لذا بهبود الگوریتم یادگیری عمیق تقویتی مورد

استفاده برای زمان بندی کارها در ابر با توجه به پارامترهای مورد نظر و همچنین تعریف مناسب محیط به طوری که ویژگی های بیشتری از کارها از جمله مهلت زمانی را شامل شود در این تحقیق مورد بررسی قرار گرفته است.

۳. الگوریتم پیشنهادی

الگوریتم DQN در محیط های دارای نویز تمایل به برآورد بیش از حد پاداش دارد. همچنین در این الگوریتم انتخاب عمل و ارزیابی عمل توسط یک شبکه یکسان انجام می شود در نتیجه آموزش شبکه عصبی پایدار و بهینه نیست. به دلیل اینکه در الگوریتم DQN انتخاب بهترین عمل برای حالت بعدی بر اساس حالت با بالاترین مقدار Value-Q قابل اطمینان نیست لذا با استفاده از الگوریتم DQN Double انتخاب عمل از ارزیابی عمل جدا شده است و از دو شبکه عصبی برای این کار استفاده می شود. شبکه اصلی DQN بهترین عمل برای حالت بعدی را انتخاب می کند و شبکه دیگر Value-Q target را محاسبه می کند. معادله نهایی به صورت رابطه (۲) است [۳۲، ۳۳].

$$Q_1(S, A) = Q_1(S, A) + \alpha (R + \gamma Q_2(S' + \operatorname{argmax}_a Q_1(S', a)) - Q_1(S, A)) \quad (2)$$

همچنین در برخی مواقع مقدار دقیق هر عمل مورد نیاز نیست بنابراین یادگیری تابع Value-State برای شبکه عصبی کافی می باشد لذا در الگوریتم DQN Dueling شبکه عصبی در لایه آخر به دو قسمت تقسیم می شود. یک قسمت برای تخمین value-state در حالت مورد نظر و دیگری برای تعیین مزایای هر عمل، سپس با هم ترکیب می شوند لذا معادله Value-Q مطابق رابطه (۳) است [۳۴].

$$Q(s, a) = V(s) + A(s, a) \quad (3)$$

که در معادله فوق $V(s)$ تابع ارزش و $A(s, a)$ مزیت یک عمل نسبت به سایر عمل ها را مشخص می نماید.

تکرار تجربه یک تکنیک کلیدی در یادگیری عمیق تقویتی است که عامل برای یادگیری از خاطرات قبلی استفاده کرده و سرعت یادگیری را با همبستگی مطلوب افزایش می دهد [۳۵]. در الگوریتم های مختلفی از جمله ERDQL و NTM^{۱۵} از ماژول حافظه جهت ذخیره تجربیات قبلی و استفاده مجدد از آن ها استفاده شده است [۳۶، ۳۷]. ایرادی که در الگوریتم ERDQN وجود دارد نمونه برداری تصادفی از حافظه می باشد که در الگوریتم PERDQL این مشکل برطرف شده است. این الگوریتم در حالت کلی مشابه الگوریتم ERDQL است با این تفاوت که در الگوریتم ERDQL انتخاب نمونه به صورت تصادفی از حافظه انجام می شود و ممکن است تجربیات مفید کمتر یا اصلاً مورد استفاده قرار نگیرند لذا با اولویت بندی تجربیات توسط درخت باینری mTreeSu و قرارگیری اولویت ها در برگ های این درخت یک آرایه تشکیل می شود

که به برگ‌های درخت اشاره نموده که حاوی تجربیات است. اولویت نمونه‌ها در الگوریتم PERDQL بر اساس معادله (۴) تعیین می‌شود که در این رابطه ϵ مقداری ثابت است که از صفر شدن احتمال نمونه‌ها جلوگیری می‌کند [۳۸، ۱۹].

$$P_i = |\delta_i| + \epsilon \quad (۴)$$

در معادله فوق δ_i خطای TD را نشان می‌دهد که بر اساس رابطه (۵) محاسبه می‌شود

$$\delta_t = R_t + \gamma Q(S_{t+1}, \arg\max(Q(S_{t+1}, a; \theta_t); \bar{\theta}_t) - Q(S_t, a_t; \theta_t) \quad (۵)$$

در معادله فوق، θ_t و $\bar{\theta}_t$ به ترتیب مربوط به وزن شبکه‌های اولیه و هدف هستند. احتمال انتخاب یک نمونه بر اساس معادله (۶) قابل محاسبه است که در این رابطه α مقداری بین ۰ و ۱ بوده و جهت کنترل میزان اولویت مورد استفاده قرار می‌گیرد. برای جلوگیری از over-fitting به نمونه‌ها با اولویت بالا که زیاد در شبکه استفاده می‌شوند وزن کوچک بر اساس معادله (۷) اختصاص داده شود که در آن N سایز بافر و β نشان‌دهنده بایاس می‌باشد.

$$P_i = p_i^\alpha / \sum_k p_k^\alpha \quad (۶)$$

$$W_i = (1/N * 1/p_i)^\beta \quad (۷)$$

برای درج در بافر با توجه به اولویت کارها به زمان $O(\log n)$ و جهت نمونه‌برداری از آن به زمان $O(n)$ نیاز می‌باشد لذا با توجه به زمان‌برد بودن این عملیات نمی‌توان تمامی بافر را بر اساس اولویت مرتب نمود بنابراین از درخت باینری SumTree استفاده شده و یک شیء حافظه برای این درخت باینری و داده‌ها ساخته می‌شود.

بنابراین با استفاده از Dueling Double Deep Q-Network و آنچه که در الگوریتم PERDQL برای ذخیره و بازیابی تجربیات قبلی از حافظه انجام شده است اکثر مشکلات الگوریتم‌های مورد استفاده تا کنون برای زمانبندی کارها در محاسبات ابری برطرف می‌گردد. اما در الگوریتم Dueling Double Deep Q-Network وزن‌های Network target در هر d تکرار به‌روزرسانی می‌شوند در صورتی که پیوند تغییر وزن‌های شبکه افزایش باید موجب همگرایی بهتر الگوریتم شده و به‌صورت تجربی نیز در زمان‌بندی کارها در محاسبات ابری نتایج بهتری به‌دست می‌آید. برای این منظور در این مقاله از معادله (۸) استفاده می‌شود.

$$w'_i = w'_{i-1} + \epsilon(w_{i-1} - w'_{i-1}) \quad (۸)$$

که در این معادله w_{i-1} وزن شبکه فعلی و w'_i وزن شبکه Network target در مرحله بعدی است. همچنین در الگوریتم پایه انتخاب action به‌صورت تصادفی با احتمال ثابت ϵ انجام می‌شود که می‌تواند بر کارایی الگوریتم اثرات منفی داشته باشد. لذا در الگوریتم پیشنهادی برای ϵ در ابتدا مقداری بین صفر و یک در نظر گرفته شده و سپس با معادله (۹) با توجه به تعداد تکرار لازم برای همگرایی مسأله، به‌روزرسانی انجام می‌شود.

$$\epsilon = \epsilon - ((\epsilon / \text{total episode}) * \text{number of episode}) \quad (۹)$$

با ایجاد تغییر فوق در مقدار ϵ در مراحل اولیه، کاوش بیشتری انجام می‌شود و عامل محیطی را به‌خوبی بررسی می‌کند، سپس به تدریج کاهش می‌یابد تا از تجربیات به دست آمده بیشتر استفاده شود و بهره‌برداری بیشتری حاصل شود. با نزدیک شدن به مراحل پایانی اجرای الگوریتم، معادله فوق به صفر همگرا شده و الگوریتم یادگیری عمیق تقویتی به‌صورت حریصانه‌تر عمل می‌کند.

بنابراین با اصلاح پیوند به‌روزرسانی وزن‌های target Network و مقدار ϵ برای انتخاب action و همچنین استفاده از Dueling Double Deep Q-Network شبه کد نهایی الگوریتم پیشنهادی بر اساس توسعه الگوریتم PERDQL به‌صورت الگوریتم ۱ است. در این الگوریتم، نمونه‌ها با استفاده از SumTree بر اساس اولویت ذخیره و بازیابی می‌شوند، این اولویت طبق رابطه (۴) تعیین می‌شود.

Algorithm 1: Extended Prioritized Experience Replay Deep Q Network (EPERDQN)

Input: minibatch k , replay period K , exponents α and β , update factor ϵ , budget T

1 Initialize replay memory D , Q-network parameters θ , $p_1=1$

2 For $t=1$ to T do:

3 Observe S_t and choose action A_t according to ϵ -greedy policy

4 Store transition (S_t, A_t, R_t, S_{t+1}) in D with maximal priority

$$P_t = \max_{i \in D} P_i$$

5 If $t \bmod K = 0$ do:

6 For $j=1$ to k do:

7 Sample transition $j \sim P(j) = p_j^\alpha / \sum_i p_i^\alpha$

8 Compute importance-sampling weight $w_j = (N * P(j))^\beta / \max_i w_i$

9 Compute TD-error $\delta_j = R_j + \gamma_j Q_{\text{target}}(S_j, \arg\max_a(Q(S_j, a)) - Q(S_{j-1}, A_{j-1})$

10 Update transition priority $p_j \leftarrow |\delta_j|$

11 End For

12 Update target network in each iteration $\theta'_i = \theta'_{i-1} + \epsilon(\theta_{i-1} - \theta'_{i-1})$

13 End If

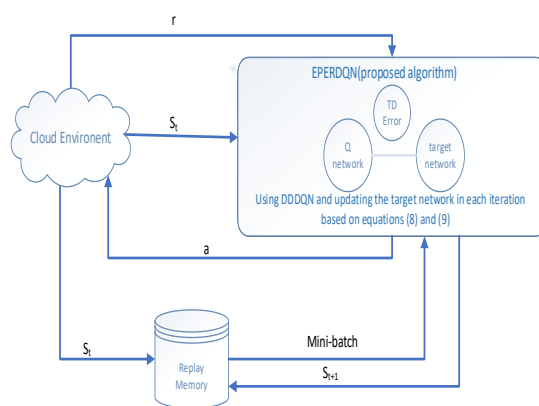
14 End For

در خط ۱۲ الگوریتم ۱، وزن‌های شبکه هدف بر اساس معادلات (۸) و (۹) به‌روز می‌شوند. با این روش به‌روزرسانی، پارامترهای شبکه هدف تغییر چندانی نمی‌کنند و الگوریتم پایدار است و در هر تکرار آموزش داده می‌شود، بنابراین زمان آموزش شبکه عصبی کاهش می‌یابد و سریعتر همگرا می‌شود که برای اهداف مورد نظر در این تحقیق مهم است. بنابراین این الگوریتم برای استفاده در محیط‌هایی با حالت‌های پیچیده و مقیاس‌های بزرگ مناسب است. در این تحقیق از آن برای زمان‌بندی وظایف در محیط محاسبات ابری استفاده شده است.

۴. پیاده‌سازی

در محیط محاسبات ابری با درخواست کار جدید، یک نمونه ماشین مجازی برای اجرای آن انتخاب می‌شود در صورتی که ماشین مجازی آزاد باشد کار بلافاصله اجرایی می‌شود در غیر این‌صورت در صف FIFO قرار

می گیرد. در این مدل زمانبندی به صورت آنلاین انجام می شود. چارچوب روش پیشنهادی در شکل ۱ بیان شده است. ابتدا محیط محاسبات ابری بر اساس کارها و ماشین های مجازی و خصوصیات مورد نظر IaaS^{۱۶} جهت زمان بندی، طبق فرآیند تصمیم گیری مارکوف پیاده سازی می شود. سپس بر اساس الگوریتم یادگیری عمیق تقویتی در زمان t مقدار S_t بعنوان ورودی شبکه عصبی به عامل داده می شود و عامل عمل a را انجام داده و پاداش r را دریافت می کند. سپس S_t و a_t و r_t و S_{t+1} به عنوان یک رکورد در حافظه ذخیره می شود و این کار ادامه می یابد تا زمانی که تعداد رکوردها در حافظه از مقدار آستانه بیشتر شود و سپس از این حافظه براساس اولویت درخت تصمیم دودویی برای آموزش مدل نمونه برداری می شود.



شکل ۱: چارچوب روش پیشنهادی

طراحی و پیاده سازی محیط یادگیری عمیق تقویتی معمولاً بر اساس فرآیند تصمیم گیری مارکوف انجام می شود. کارهایی که برای زمانبندی روی ماشین های مجازی وارد می شوند دارای خصوصیات از جمله اندازه، زمان اجرا و مهلت زمانی می باشند. همچنین مورد دیگری که برای حالت باید در نظر گرفت وضعیت فعلی ماشین های مجازی می باشد که بر اساس زمان تأخیر کار بر روی هر ماشین مجازی بر اساس بار کاری آن مشخص می شود. لذا در این مقاله حالت برای محیط از مجموع دو مورد ذکر شده به صورت زیر تعریف شده است. زمان اجرا بر اساس اندازه کار، تعداد پردازنده هر ماشین مجازی و سرعت آن ها محاسبه می شود.

$$State\ Space = \{job_size, job_ExecutionTime, job_constraints, VMs_WaitingTime\}$$

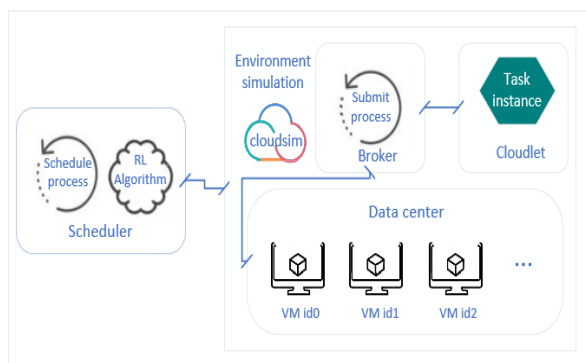
عامل می تواند هر یک از ماشین های مجازی را با توجه به حالت آن انتخاب نماید و هر ماشین مجازی محدودیتی در پذیرفتن کارها ندارد بنابراین فضای عمل به صورت زیر تعریف می شود.

$$Action\ Space = \{VM_{id^1}, VM_{id^2}, \dots, VM_{id^n}\}$$

عامل در یادگیری عمیق تقویتی سعی دارد که ماکزیمم پاداش را به دست بیاورد و به عنوان هدف از آن استفاده می کند. در این تحقیق پاداش در جهت دستیابی به اهداف مورد نظر به صورت $(d_i - T_i) / (1/n \sum_{i=1}^n |d_i - T_i| + \epsilon)$ تعریف شده است. در این معادله d_i نشان دهنده مهلت زمانی کار i و T_i نشان دهنده جمع زمان تأخیر کار i می باشد. کاربران محیط ابری انتظار دارند وظایف آن ها در مهلت مقرر

انجام شود، بنابراین الگوریتم زمان بندی باید مهلت زمانی را نیز در نظر بگیرد.

در این مقاله ابتدا مطابق چارچوب تعریف شده در شکل ۲ جهت انجام شبیه سازی محیط از کلودسیم استفاده شده است [۳۹]. همچنین از کراس که یک کتابخانه متن باز شبکه عصبی به زبان پایتون بوده و بر روی تنسورفلو قابل اجرا می باشد جهت پیاده سازی شبکه عصبی استفاده شده است [۴۰]. کلود سیم یک بسته نرم افزاری متن باز به زبان جاوا است که برای شبیه سازی محاسبات ابری مناسب بوده و امکانات متنوعی برای مقایسه نتایج اجرای الگوریتم های گوناگون بر محیط شبیه سازی شده با توجه به پارامترهای مورد نظر دارد. بعد از شبیه سازی محیط در کلود سیم، با فعال سازی listener جاوا، از طریق py4j ارتباط با محیط مورد استفاده برای پایتون برقرار می شود. Py4j امکان دسترسی پویا به اشیاء جاوا در یک ماشین مجازی را برای برنامه های پایتون با استفاده از شماره پورت و آدرس شبکه فراهم می کند.



شکل ۲: چارچوب ارزیابی الگوریتم EPERDQN در محیط شبیه سازی شده

در مرحله اول، آزمایش با یک مثال حاوی تعداد کمی از وظایف با داده های مصنوعی انجام می شود، سپس در مرحله دوم، الگوریتم EPERDQN با مجموعه داده GOCJ ارزیابی می شود. این مجموعه داده شامل چندین فایل متنی است که هر کدام شامل تعداد مشخصی از وظایف است و به محققان بینشی از رفتار واقعی کاربر می دهد [۴۱]. بنابراین، با توجه به ویژگی های ارائه شده از این مجموعه داده، حاوی داده های لازم برای آزمایش الگوریتم پیشنهادی است.

۵. ارزیابی

در مقالات مختلف، از جمله مراجع [۴۲-۴۴]، مقایسه ای بین الگوریتم های یادگیری عمیق تقویتی و سایر الگوریتم های مورد استفاده برای زمان بندی وظایف در فضای ابری بیان شده است و مزایای استفاده از الگوریتم DRL برای زمان بندی مورد بحث قرار گرفته است. بنابراین در این تحقیق مقایسه بین الگوریتم های یادگیری تقویتی انجام شده است.

جدول ۲: مقایسه پارامترهای خروجی بر اساس اهداف تعریف شده در سناریوی ۱

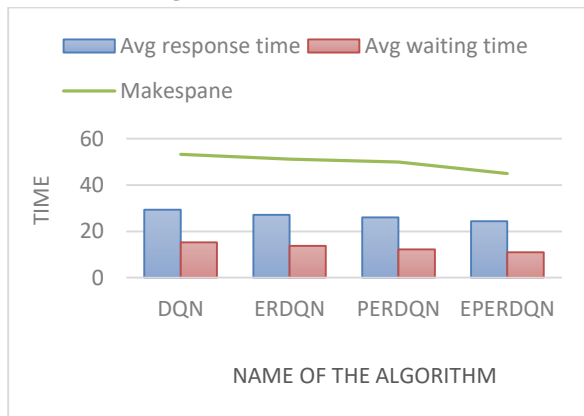
Algorithm	DQN	ERDQN	PERDQN	EPERDQN (proposed algorithm)
Avg response time	29.4	27.1	26	24.4
Avg waiting time	15.3	13.7	12.2	11
Makespan	53.24	51.24	50	45
Throughput	0.19	0.2	0.2	0.22

در این جدول زمان انتظار از تقسیم بار روی ماشین مجازی بر سرعت پردازنده آن به دست می آید. زمان پاسخ از مجموع دو زمان اجرا و انتظار محاسبه می شود. زمان کل نیز زمان اتمام آخرین کار بر روی آخرین پردازنده را نشان می دهد [۴۵]. زمان تکمیل نیز بر اساس معادلات (۱۰) و (۱۱) در محیط کلودسیم محاسبه می شود.

$$execTime = cloudlet.getCloudletLength() / (vm.getMips() * vm.getNumberOfPes()) \quad (10)$$

$$completionTime = execTime + waitingTime[vm_id] \quad (11)$$

با افزایش تعداد کارها درصد بهبود الگوریتم پیشنهادی افزایش می یابد. به طور کلی الگوریتم EPERDQN مطابق نمودار مقایسه ای شکل ۴ عملکرد بهتری در پارامترهای مورد ارزیابی در این تحقیق دارد.



شکل ۴: نمودار مقایسه ای پارامترهای خروجی در محیط کلودسیم در سناریوی ۱

توان عملیاتی^{۱۷} تعداد وظایف انجام شده در واحد زمان را نشان می دهد. در این تحقیق تعداد cloudlets اجرا شده در ثانیه به عنوان توان عملیاتی در نظر گرفته شده است. مقایسه بین الگوریتم های مورد استفاده در این تحقیق در شکل ۵ نشان داده شده است.

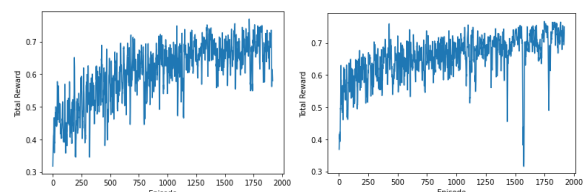
به منظور بررسی عملکرد الگوریتم پیشنهادی در این تحقیق و مقایسه آن با سایر الگوریتم های ذکر شده، دو سناریو مورد بررسی قرار گرفته است. برای شبیه سازی این دو سناریو، شبیه ساز کلودسیم با Eclipse IDE استفاده شده است.

سناریوی ۱: برای شبیه سازی در محیط کلودسیم با داده های مصنوعی، فرض شده یک دیتاسنتر با ۴ ماشین مجازی با خصوصیات مشخص از جمله تعداد و سرعت پردازنده، پهنای باند، ظرفیت حافظه و ... وجود دارد در این سناریو تعداد ۱۰ کار در نظر گرفته شده است. با تست الگوریتم ها با سناریوهای مختلف و بار کاری بالا نتایج مشابهی به دست آمده است. ابتدا محیط بر اساس اطلاعات فوق در کلودسیم ساخته می شود که پارامترها در جدول ۱ آورده شده است.

جدول ۱: محیط شبیه سازی برای سناریوی ۱

type	parameter	Value
Data Center	Number of DC	1
	VM-Scheduler	Time-Shared
Host	Host Number	1
	Processing Element	4000 MIPS
	RAM	2 GB
	Storage	1TB
	Bandwidth	10 GB/s
VM	Number of VMs	4
	Processing Element	220-450 MIPS
	Memory	512MB
	Bandwidth	1 GB/s
	VMM	Xen
Cloudlet	Task Number	10
	Length of Tasks	2000-10000

حال می توان الگوریتم EPERDQN را با زبان پایتون پیاده سازی نمود و ارتباط آن با کلودسیم را از طریق py4j برقرار نمود. پس از دو هزار بار تکرار الگوریتم یادگیری عمیق تقویتی نتایج به شرح شکل های ۳(a) و ۳(b) به دست آمده است.



شکل ۳: اجرای الگوریتم DQN (a) و اجرای الگوریتم EPERDQN (b)

در الگوریتم EPERDQN به نسبت الگوریتم های DQN، ERDQN و PERDQN پاداش بالاتری به دست آمده است. همچنین سریع تر از سایر الگوریتم های عنوان شده همگرا شده و پایدارتر می باشد. این نتایج با تست بارهای کاری مختلف تکرار شده است. اثرات اجرای الگوریتم ها بر پارامترهای مورد ارزیابی به شرح جدول ۲ می باشد.

جدول ۴: محیط شبیه سازی برای سناریوی ۲

Type	Parameter	value
Data Center	Number of DC	1
	VM-Scheduler	Time-Shared
Host	Host Number	20
	Processing Element	177730 MIPS
	RAM	16 GB
	Storage	2 TB
	Bandwidth	10 GB/s
VM	Number of VMs	5-100
	Processing Element	$(3.5-100) * 10^3$ MIPS
	Memory	1-4 GB
	Bandwidth	1-10 GB/s
	VMM	Xen
Cloudlet	Task Number	100-1000
	Length of Tasks	$(15-900) * 10^3$ MI

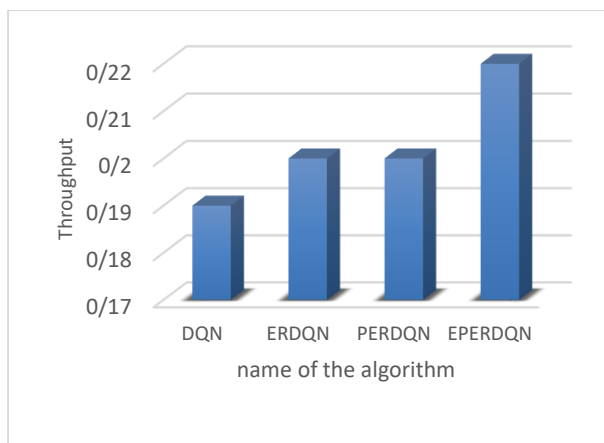
در این سناریو، زمان کل و توان عملیاتی برای دو گروه اندازه معمولی و اندازه بزرگ با مجموعه داده GOCJ ارزیابی می شود. نتایج با الگوریتم های DQN، ERDQN و PERDQN مقایسه شده اند. جدول ۵ و شکل ۶ نتایج زمان کل در الگوریتم EPERDQN را در مقایسه با سایر الگوریتم های ذکر شده نشان می دهد و جدول ۶ و شکل ۷ آنالیز عملکرد توان عملیاتی را برای مجموعه داده با اندازه معمولی نشان می دهد.

جدول ۵: زمان کل در سناریوی ۲ با استفاده از regular_size در مجموعه داده GOCJ

number of tasks	DQN	ERDQN	PERDQN	EPERDQN
100	87	74	72	60
200	180	162	153	132
300	227	208	188	157
400	265	247	232	204
500	326	304	301	259
600	402	360	317	274

جدول ۶: توان عملیاتی در سناریوی ۲ با استفاده از regular_size در مجموعه داده GOCJ

number of tasks	DQN	ERDQN	PERDQN	EPERDQN
100	1.1	1.4	1.4	1.7
200	1.1	1.2	1.3	1.5
300	1.3	1.4	1.6	1.9
400	1.5	1.6	1.7	2
500	1.5	1.6	1.7	1.9
600	1.5	1.7	1.9	2.2



شکل ۵: مقایسه توان عملیاتی در سناریوی ۱

بر اساس نتایج حاصل در پیاده سازی سناریوی ۱، الگوریتم EPERDQN در تمامی پارامترهای مورد ارزیابی نتایج بهتری به دست آورده است. همچنین در مدت زمان کمتری آموزش دیده و اهداف مورد انتظار را برآورده می نماید.

سناریو ۲: برای اطمینان از عملکرد الگوریتم پیشنهادی در محیط های واقعی و بزرگ از مجموعه داده GOCJ استفاده شده است که بر اساس ردیابی بار کاری خوشه گوگل^{۱۸} است. مجموعه داده GOCJ در مخزن داده مندی ذخیره می شود و شامل ۱۹ فایل متنی است که هر یک از آن ها GOCJ_dataset_xxx.text نام دارند که xxx تعداد کارهای موجود در آن فایل را نشان می دهد و هر ردیف اندازه کار را به MI نشان می دهد. وظایف موجود در این مجموعه داده بر اساس طول آن ها مطابق جدول ۳، به ۵ نوع تقسیم می شوند [۴۱].

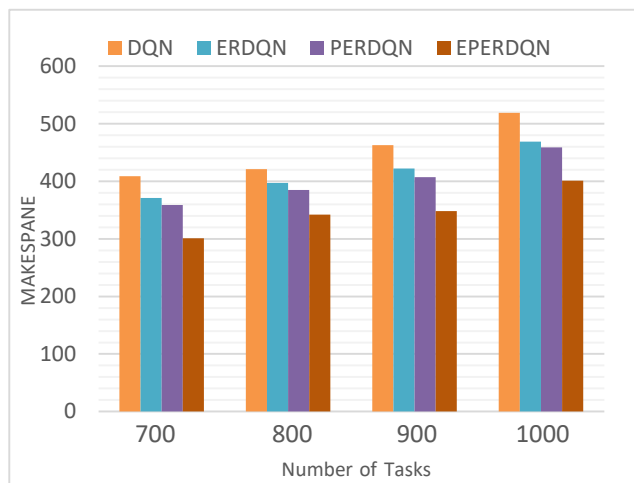
جدول ۳: اندازه کارها در مجموعه داده GOCJ

Job type	Small	Medium	Large	Extra-Large	Huge
Size of jobs(MI *10 ³)	15-55	59-99	101-135	150-525	525-900
Distribution	20%	40%	30%	6%	4%

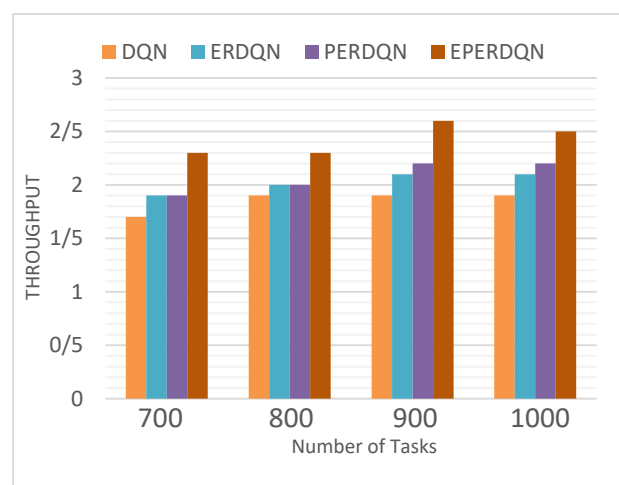
در این تحقیق فایل های موجود در مجموعه داده GOCJ به دو گروه تقسیم می شوند، یک گروه بین ۱۰۰ تا ۶۰۰ کار به عنوان مجموعه داده با اندازه معمولی^{۱۹} است و گروه دیگر شامل ۷۰۰ تا ۱۰۰۰ کار به عنوان مجموعه داده با اندازه بزرگ^{۲۰} است [۴۶]. ۵۰ ماشین مجازی برای برای مجموعه داده های اندازه معمولی و ۱۰۰ ماشین مجازی برای مجموعه داده های اندازه بزرگ استفاده می شود. سایر پارامترهای مورد استفاده در این سناریو برای شبیه سازی در کلودسیم مطابق جدول ۴ می باشد.

جدول ۸: توان عملیاتی در سناریوی ۲ با استفاده از larg_size در مجموعه داده GOCJ

number of tasks	DQN	ERDQN	PERDQN	EPERDQN
700	1.7	1.9	1.9	2.3
800	1.9	2	2	2.3
900	1.9	2.1	2.2	2.6
1000	1.9	2.1	2.2	2.5

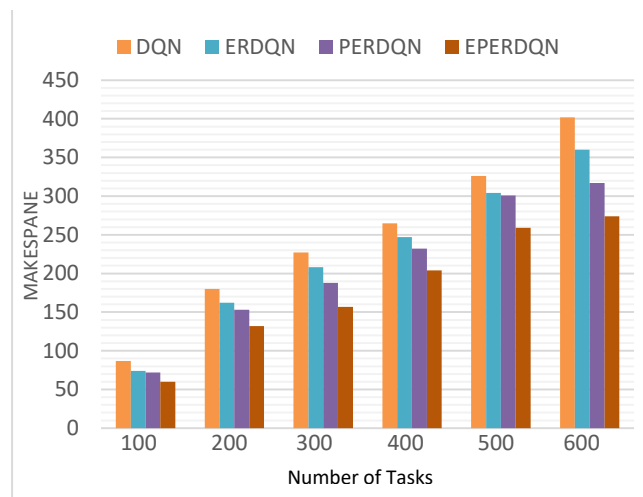


شکل ۸: مقایسه زمان کل در سناریوی ۲ با استفاده از larg_size در مجموعه داده GOCJ

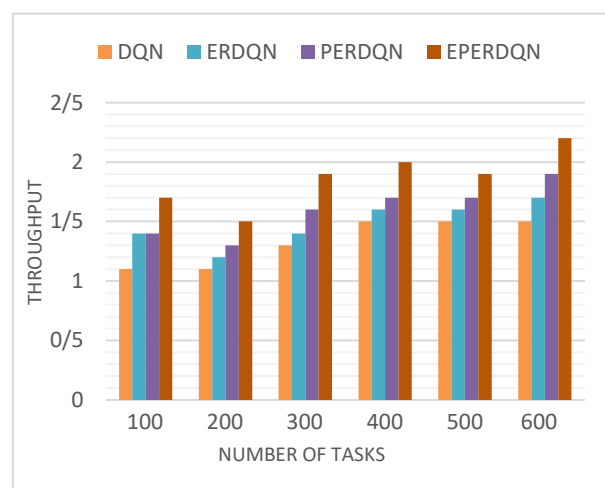


شکل ۹: مقایسه توان عملیاتی در سناریوی ۲ با استفاده از larg_size در مجموعه داده GOCJ

روند کلی ارزیابی پارامترهای مورد نظر با استفاده از مجموعه داده GOCJ مشابه سناریوی قبلی تکرار شده است و عملکرد الگوریتم EPERDQN در پارامترهای مورد نظر مطابق شکل‌های ۶-۹ بهتر است. نکته دیگری که از نتایج سناریوی ۲ می‌توان استنباط کرد این است که با افزایش تعداد کارها، الگوریتم پیشنهادی مقیاس‌پذیری بهتری دارد و با اندازه بزرگ در مقایسه با اندازه معمولی در مجموعه داده GOCJ، بهبود بیشتری نسبت به سایر الگوریتم‌های مورد ارزیابی در این تحقیق، حاصل می‌شود. نتایج آزمایش انجام شده در مجموعه داده GOCJ نشان



شکل ۶: مقایسه زمان کل در سناریوی ۲ با استفاده از regular_size در مجموعه داده GOCJ



شکل ۷: مقایسه توان عملیاتی در سناریوی ۲ با استفاده از regular_size در مجموعه داده GOCJ

تجزیه و تحلیل تجربی برای گروه اندازه بزرگ مجموعه داده GOCJ در این بخش ارائه شده است. مشابه آزمایش قبلی، عملکرد با استفاده از زمان کل و توان عملیاتی اندازه‌گیری می‌شود. تعداد وظایف بین ۷۰۰-۱۰۰۰ برای مجموعه داده با اندازه بزرگ در نظر گرفته شده است. این وظایف بر روی ۱۰ ماشین مجازی زمان‌بندی شده‌اند. جدول ۷ و شکل ۸ تجزیه و تحلیل مقایسه‌ای مجموعه داده‌های بزرگ را از نظر زمان کل نشان می‌دهد، همچنین جدول ۸ و شکل ۹ نتایج را برای توان عملیاتی نشان می‌دهد.

جدول ۷: زمان کل در سناریوی ۲ با استفاده از larg_size در مجموعه داده GOCJ

number of tasks	DQN	ERDQN	PERDQN	EPERDQN
700	409	371	359	301
800	421	397	385	342
900	463	422	407	348
1000	519	469	459	401

- survey," *Future Generation Computer Systems*, vol. 91, pp. 407-415, 2019.
- [4] C. Shetty, H. Sarojadevi, and S. Prabhu, "Machine learning approach to select optimal task scheduling algorithm in cloud," *Turkish Journal of Computer and Mathematics Education (TURCOMAT)*, vol. 12, no. 6, pp. 2565-2580, 2021.
- [5] J. Zhao, M. A. Rodríguez, and R. Buyya, "A deep reinforcement learning approach to resource management in hybrid clouds harnessing renewable energy and task scheduling," in *2021 IEEE 14th International Conference on Cloud Computing (CLOUD)*, 2021: IEEE, pp. 240-249.
- [6] X. Wang et al., "Dynamic scheduling of tasks in cloud manufacturing with multi-agent reinforcement learning," *Journal of Manufacturing Systems*, vol. 65, pp. 130-145, 2022.
- [7] K. Arulkumaran, M. P. Deisenroth, M. Brundage, and A. A. Bharath, "A brief survey of deep reinforcement learning," *arXiv preprint arXiv:1708.05866*, 2017.
- [8] Y. Wei, L. Pan, S. Liu, L. Wu, and X. Meng, "Drl-scheduling: An intelligent qos-aware job scheduling framework for applications in clouds," *IEEE Access*, vol. 6, pp. 55112-55125, 2018.
- [9] Z. Saadati And M. Zarei, "Solving the Multi-Objective Problem of IoT Service Placement in Fog Computing Using Reinforcement Learning Approaches," *Intelligent Multimedia Processing and Communication Systems (IMPCS)*, pp. 29-41, 4(3), 2023.
- [10] L. Zhang, C. Yang, Y. Yan, and Y. Hu, "Distributed Real-Time Scheduling in Cloud Manufacturing by Deep Reinforcement Learning," *IEEE Transactions on Industrial Informatics*, vol. 18, no. 12, pp. 8999-9007, 2022.
- [11] X. Wang, L. Zhang, Y. Liu, C. Zhao, and K. Wang, "Solving task scheduling problems in cloud manufacturing via attention mechanism and deep reinforcement learning," *Journal of Manufacturing Systems*, vol. 65, pp. 452-468, 2022.
- [12] K. Siddesha, G. Jayaramaiah, and C. Singh, "A novel deep reinforcement learning scheme for task scheduling in cloud computing," *Cluster Computing*, vol. 25, no. 6, pp. 4171-4188, 2022.
- [13] M. Naeem, S. T. H. Rizvi, and A. Coronato, "A gentle introduction to reinforcement learning and its application in different fields," *IEEE access*, vol. 8, pp. 209320-209344, 2020.
- [14] J. Liu, F. Gao, and X. Luo, "Survey of deep reinforcement learning based on value function and policy gradient," *Chinese Journal of Computers*, vol. 42, no. 6, pp. 1406-1438, 2019.
- [15] K. Sharma and S. Tripathy, "DEEP REINFORCEMENT LEARNING: A SURVEY."
- [16] W. Yuan, Y. Li, H. Zhuang, C. Wang, and M. Yang, "Prioritized experience replay-based deep q learning: Multiple-reward architecture for highway driving decision making," *IEEE Robotics & Automation Magazine*, vol. 28, no. 4, pp. 21-31, 2021.
- [17] W. Fedus et al., "Revisiting fundamentals of experience replay," in *International Conference on Machine Learning*, 2020: PMLR, pp. 3061-3071.
- [18] J. Fan, Z. Wang, Y. Xie, and Z. Yang, "A theoretical analysis of deep Q-learning," in *Learning for dynamics and control*, 2020: PMLR, pp. 486-489.
- [19] T. Schaul, J. Quan, I. Antonoglou, and D. Silver, "Prioritized experience replay," *arXiv preprint arXiv:1511.05952*, 2015.

می‌دهد که میانگین توان عملیاتی نسبت به DQN ۳۸ درصد، ERDQN ۲۳ درصد و PERDQN ۱۸ درصد بهبود یافته‌است. علاوه بر این، با در نظر گرفتن مهلت زمانی وظایف در تابع پاداش، ویژگی‌های بیشتری از نیازهای کاربران را برآورده می‌کند و تعداد نقض SLA را بهبود می‌بخشد. با توجه به نتایج سناریوهای ۱ و ۲ می‌توان استنباط کرد که الگوریتم PERDQN پایه در پارامترهای ارزیابی شده بهبود قابل توجهی نداشته‌است، حتی در برخی از اجراها نتایج کمتری نسبت به سایر الگوریتم‌ها داشته‌است، بنابراین الگوریتم پیشنهادی توسعه یافت. الگوریتم EPERDQN به درستی بار کاری را روی سرورها توزیع می‌کند و در تمام پارامترهای ارزیابی شده عملکرد بهتری دارد.

۶. نتیجه‌گیری و کارهای آینده

در این مقاله یک الگوریتم بهبود یافته بر اساس PERDQN برای غلبه بر مشکلات الگوریتم DQN شامل آموزش طولانی مدت، تخمین غیرواقعی پاداش و نمونه برداری تصادفی از حافظه ارائه شد. سپس با تعریف مناسب محیط ابری بر اساس فرآیند تصمیم‌گیری مارکوف از الگوریتم پیشنهادی برای زمانبندی کارها بر روی ماشین‌های مجازی در محاسبات ابری استفاده شده‌است. نتایج به دست آمده نشان می‌دهد که تمامی پارامترهای مورد نظر در الگوریتم EPERDQN بهبود می‌یابند که در مقیاس بزرگتر تفاوت بیشتر می‌شود. ولی در چندین حوزه امکان بهبود وجود دارد. از جمله اینکه می‌توان فضای حالت مسئله را توسعه داد به‌طوری‌که موارد بیشتری از خصوصیات کارها را شامل شود. همچنین این امکان وجود دارد که علاوه بر زمانبندی کارها برای اهداف دیگری به عنوان نمونه مصرف انرژی در مرکز داده از الگوریتم EPERDQN استفاده نمود. با وجود مزایای زیاد الگوریتم‌های یادگیری عمیق تقویتی، اما باید سربار محاسباتی آن‌ها برای استفاده در محیط ابری در نظر گرفته شود، مساله کاهش زمان یادگیری در الگوریتم پیشنهادی در نظر گرفته شده است، با این حال، همچنان می‌توان آن را به عنوان یک کار آینده بهبود بخشید. در قسمت شبکه عصبی نیز می‌توان از NTM استفاده نمود که در قسمت ذخیره و بازیابی، از یک حافظه خارجی استفاده نموده که موجب افزایش کارایی الگوریتم می‌شود.

مراجع

- [1] A. Keivani and J.-R. Tapamo, "Task scheduling in cloud computing: A review," in *2019 International Conference on Advances in Big Data, Computing and Data Communication Systems (icABCD)*, 2019: IEEE, pp. 1-6.
- [2] N. K. Gondhi and A. Gupta, "Survey on machine learning based scheduling in cloud computing," in *Proceedings of the 2017 International Conference on Intelligent Systems, Metaheuristics & Swarm Intelligence*, 2017, pp. 57-61.
- [3] A. Arunarani, D. Manjula, and V. Sugumaran, "Task scheduling techniques in cloud computing: A literature

- [35] R. Liu and J. Zou, "The effects of memory replay in reinforcement learning," in *2018 56th annual allerton conference on communication, control, and computing (Allerton)*, 2018: IEEE, pp. 478-485.
- [36] W. Zaremba and I. Sutskever, "Reinforcement learning neural turing machines-revised," *arXiv preprint arXiv:1505.00521*, 2015.
- [37] Y. Li, "Deep reinforcement learning: An overview," *arXiv preprint arXiv:1701.07274*, 2017.
- [38] D. Fährmann, N. Jorek, N. Damer, F. Kirchbuchner, and A. Kuijper, "Double deep q-learning with prioritized experience replay for anomaly detection in smart environments," *IEEE Access*, vol. 10, pp. 60836-60848, 2022.
- [39] T. Goyal, A. Singh, and A. Agrawal, "Cloudsim: simulator for cloud computing infrastructure and modeling," *Procedia Engineering*, vol. 38, pp. 3566-3572, 2012.
- [40] K. Ramasubramanian and A. Singh, "Deep learning using keras and tensorflow," in *Machine Learning Using R*: Springer, 2019, pp. 667-688.
- [41] A. Hussain and M. Aleem, "GoCJ: Google cloud jobs dataset for distributed and cloud computing infrastructures," *Data*, vol. 3, no. 4, p. 38, 2018.
- [42] P. Amini and A. Kalbasi, "An Adaptive Task Scheduling Approach for Cloud Computing Using Deep Reinforcement Learning," in *2024 Third International Conference on Distributed Computing and High Performance Computing (DCHPC)*, 2024: IEEE, pp. 1-9.
- [43] S. Mangalampalli, G. R. Karri, M. Kumar, O. I. Khalaf, C. A. T. Romero, and G. A. Sahib, "DRLBTSA: Deep reinforcement learning based task-scheduling algorithm in cloud computing," *Multimedia Tools and Applications*, vol. 83, no. 3, pp. 8359-8387, 2024.
- [44] S. Mangalampalli et al., "Multi Objective Prioritized workflow scheduling using Deep reinforcement based Learning in Cloud Computing," *IEEE Access*, 2024.
- [45] S. Mangalampalli, G. R. Karri, and U. Kose, "Multi Objective Trust aware task scheduling algorithm in cloud computing using Whale Optimization," *Journal of King Saud University-Computer and Information Sciences*, vol. 35, no. 2, pp. 791-809, 2023.
- [46] A. R. Kadhim and F. Rabee, "Deadline and Cost Aware Dynamic Task Scheduling in Cloud Computing Based on Stackelberg Game," *International Journal of Intelligent Engineering & Systems*, vol. 16, no. 3, 2023.
- [20] H. Zhang, C. Qu, J. Zhang, and J. Li, "Self-adaptive priority correction for prioritized experience replay," *Applied sciences*, vol. 10, no. 19, p. 6925, 2020.
- [21] Z. Tong, H. Chen, X. Deng, K. Li, and K. Li, "A scheduling scheme in the cloud computing environment using deep Q-learning," *Information Sciences*, vol. 512, pp. 1170-1191, 2020.
- [22] T. Zheng, J. Wan, J. Zhang, and C. Jiang, "Deep reinforcement learning-based workload scheduling for edge computing," *Journal of Cloud Computing*, vol. 11, no. 1, p. 3, 2022.
- [23] Y. Feng and F. Liu, "Resource Management in Cloud Computing Using Deep Reinforcement Learning: A Survey," in *Proceedings of the 10th Chinese Society of Aeronautics and Astronautics Youth Forum*, 2023: Springer, pp. 635-643.
- [24] S. Swarup, E. M. Shakshuki, and A. Yasar, "Task Scheduling in Cloud Using Deep Reinforcement Learning," *Procedia Computer Science*, vol. 184, pp. 42-51, 2021.
- [25] G. Rjoub, J. Bentahar, O. Abdel Wahab, and A. Saleh Bataineh, "Deep and reinforcement learning for automated task scheduling in large-scale cloud computing systems," *Concurrency and Computation: Practice and Experience*, p. e5919, 2020.
- [26] M. Roknaldini And E. Noroozi, "Presenting A Hybrid Method of Deep Neural Networks to Prevent Intrusion in Computer Networks," *Intelligent Multimedia Processing and Communication Systems(IMPCS)*, 4(4), pp. 57-65, 2023.
- [27] H. Che, Z. Bai, R. Zuo, and H. Li, "A deep reinforcement learning approach to the optimization of data center task scheduling," *Complexity*, vol. 2020, 2020.
- [28] Z. Chen, J. Hu, G. Min, C. Luo, and T. El-Ghazawi, "Adaptive and efficient resource allocation in cloud datacenters using actor-critic deep reinforcement learning," *IEEE Transactions on Parallel and Distributed Systems*, vol. 33, no. 8, pp. 1911-1923, 2021.
- [29] K. Li, Z. Peng, D. Cui, and Q. Li, "SLA-DQTS: SLA Constrained Adaptive Online Task Scheduling Based on DDQN in Cloud Computing," *Applied Sciences*, vol. 11, no. 20, p. 9360, 2021.
- [30] J. Lin, D. Cui, Z. Peng, Q. Li, and J. He, "A Two-Stage Framework for the Multi-User Multi-Data Center Job Scheduling and Resource Allocation," *IEEE Access*, vol. 8, pp. 197863-197874, 2020.
- [31] M. Cheng, J. Li, and S. Nazarian, "DRL-cloud: Deep reinforcement learning-based resource provisioning and task scheduling for cloud service providers," in *2018 23rd Asia and South pacific design automation conference (ASP-DAC)*, 2018: IEEE, pp. 129-134.
- [32] B. Ning, F. H. T. Lin, and S. Jaimungal, "Double deep q-learning for optimal execution," *Applied Mathematical Finance*, vol. 28, no. 4, pp. 361-380, 2021.
- [33] H. Van Hasselt, A. Guez, and D. Silver, "Deep reinforcement learning with double q-learning," in *Proceedings of the AAAI conference on artificial intelligence*, 2016, vol. 30, no. 1.
- [34] Z. Wang, T. Schaul, M. Hessel, H. Hasselt, M. Lanctot, and N. Freitas, "Dueling network architectures for deep reinforcement learning," in *International conference on machine learning*, 2016: PMLR, pp. 1995-2003.

- ¹² Discount factor
- ¹³ deadline
- ¹⁴ MakeSpan
- ¹⁵ Neural Turing Machine
- ¹⁶ Infrastructure as a Service
- ¹⁷ Throughput
- ¹⁸ google cluster traces
- ¹⁹ regular_size
- ²⁰ large_size

- ¹ reinforcement learning
- ² Virtual Machine (VM)
- ³ Deep Q Learning
- ⁴ Temporal difference
- ⁵ Experience Replay DQN
- ⁶ Prioritized Experience Replay DQL
- ⁷ Markov decision process (MDP)
- ⁸ State
- ⁹ Action
- ¹⁰ Probability
- ¹¹ Reward