

Research Article

Multi-Agent Task Offloading Approach based on Software Defined-Networking in Vehicular Fog Networks

Kobra Behravan ¹ | Seyed Amin Hosseini Seno ² | Nazbanoo Farzaneh ^{*3} | Mohsen Jahanshahi ⁴

¹ Department of Computer Engineering, CT.C., Islamic Azad University, Tehran, Iran, k.behravan@iau.ac.ir

² Department of Computer Engineering, Ferdowsi University of Mashhad, Mashhad, Iran, hosseini@um.ac.ir

³ Department of Computer Engineering, Imam Reza International University, Mashhad, Iran, nazbanou.farzaneh@imamreza.ac.ir

⁴ Department of Computer Engineering, CT.C., Islamic Azad University, Tehran, Iran, mjahanshahi@iauctb.ac.ir

Correspondence

*Nazbanoo Farzaneh, Assistant professor, Department of Computer Engineering, Imam Reza International University, Mashhad, Iran, nazbanou.farzaneh@imamreza.ac.ir

Main Subjects: Finding optimal task offloading policy

Received: 15 November 2025

Revised: 4 December 2025

Accepted: 19 December 2025

<https://doi.org/10.82195/ntds.2025.1224422>

Abstract

Vehicular fog computing (VFC) been recognized as an effective architecture to address rising demands in smart vehicles. Fog servers deployed on moving or parked vehicles can provide spatio-temporal heterogeneity of computational resources that capable of accomplishing computation and deadline intensive-tasks beyond the capacity which is embedded inside the vehicles. In multi-tier VFC architecture, vehicles can share idle low cost resources to increases acceptance tasks. However, the main issue is choosing optimal destination fog server for executing hard deadline tasks at each time slot. Therefore, in this work proposed a federated multi-agent deep Q-learning task offloading approach to provide collaborative learning and fast convergence. This approach improves privacy of data for agents and reduces average response time, energy consumption and processing cost in vehicular networks. Software-defined networking (SDN)-based architecture can provide flexibility, scalability, programmability, and overall network knowledge. With the help of SDN control plane configuration, the network can not only adapt to dynamic network changes, but also respond to emergency situations. Therefore, SDN technology integrated with

federated reinforcement learning (FRL) method increases programmability, centralized network management, and dynamic configuration. The results show that the proposed method reduces the average response time, average energy consumption, and average economic cost of performing tasks in the network and will increase the successful performance of high-priority tasks in vehicular fog networks.

Keywords: Vehicular Fog Computing (VFC), Task Offloading, Software-Defined Networking (SDN), Federated Reinforcement Learning (FRL)

پژوهشی

رویکرد برون سپاری وظیفه چند عامله مبتنی بر شبکه نرم افزار محور در شبکه های مه خودرویی

کبری بهروان^۱ | سید امین حسینی سنو^۲ | نازبانو فرزانه^۳ | محسن جهانشاهی^۴

چکیده:

رایانش مه خودرویی (VFC) به عنوان یک معماری مؤثر برای رسیدگی به تقاضاهای رو به افزایش در خودروهای هوشمند شناخته شده است. سرورهای مه مستقر در خودروهای در حال حرکت یا پارک شده می توانند ناهمگونی فضایی-زمانی منابع محاسباتی را فراهم کنند که قادر به انجام محاسبات وظایف فشرده با مهلت فراتر از ظرفیت تعبیه شده در داخل خودروها هستند. در معماری VFC چند لایه، خودروها می توانند منابع کم هزینه و بیکار را برای افزایش پذیرش وظایف به اشتراک بگذارند. با این حال، مسئله اصلی انتخاب سرور مه مقصد بهینه برای اجرای وظایف با مهلت سخت در هر برش زمانی است. بنابراین، در این کار یک رویکرد برون سپاری وظیفه Q-learning عمیق چند عاملی فدرال برای ارائه یادگیری مشارکتی و همگرایی سریع پیشنهاد شده است. این رویکرد حریم خصوصی داده ها را برای عامل ها بهبود می بخشد و میانگین زمان پاسخ، مصرف انرژی و هزینه پردازش را در شبکه های خودرویی کاهش می دهد. معماری مبتنی بر شبکه نرم افزار محور (SDN) می تواند انعطاف پذیری، مقیاس پذیری، قابلیت برنامه ریزی و دانش کلی شبکه را فراهم نماید. با کمک پیکربندی صفحه ی کنترلی SDN، شبکه می تواند نه فقط با تغییرات پویای شبکه، وفق پذیر باشد بلکه می تواند به وضعیت های اورژانسی پاسخ دهد. بنابراین، SDN که با روش یادگیری تقویتی فدرال (FRL) ادغام شده است، قابلیت برنامه ریزی، مدیریت شبکه متمرکز و پیکربندی پویا را افزایش می دهد. نتایج، نشان می دهند که روش پیشنهادی میانگین زمان پاسخدهی، میانگین مصرف انرژی و میانگین هزینه اقتصادی جهت انجام وظایف در شبکه را کاهش می دهد و باعث افزایش انجام موفقیت آمیز وظایف با اولویت بالا در شبکه های مه خودرویی خواهد شد.

^۱ گروه مهندسی کامپیوتر، واحد تهران مرکزی، دانشگاه آزاد اسلامی، تهران، ایران،

k.behravan@iau.ac.ir

^۲ گروه مهندسی کامپیوتر، دانشگاه فردوسی مشهد، مشهد، ایران،

hosseini@um.ac.ir

^۳ گروه مهندسی کامپیوتر، دانشگاه بین المللی امام رضا، مشهد، ایران،

nazbanou.farzaneh@imamreza.ac.ir

^۴ گروه مهندسی کامپیوتر، واحد تهران مرکزی، دانشگاه آزاد اسلامی، تهران، ایران،

mjahanshahi@iauctb.ac.ir

نویسنده مسئول

* نازبانو فرزانه، استادیار، گروه مهندسی کامپیوتر، دانشگاه بین المللی امام رضا، مشهد، ایران،

nazbanou.farzaneh@imamreza.ac.ir

موضوع اصلی: یافتن سیاست برون سپاری وظیفه بهینه

تاریخ دریافت: ۲۴ آبان ۱۴۰۴

تاریخ بازنگری: ۱۳ آذر ۱۴۰۴

تاریخ پذیرش: ۲۸ آذر ۱۴۰۴

<https://doi.org/10.82195/ntds.2025.1224422>

کلید واژه ها: برون سپاری وظیفه، رایانش مه خودرویی، شبکه نرم افزار محور، یادگیری تقویتی عمیق

۱-مقدمه

آینده همیشه با ماشین های مدرن و تخیلی تصویر می شود که در مقایسه با وسایل حمل و نقل کلاسیک بیشتر شبیه ربات هستند. این تصویر امروزه به نوعی با وسایل نقلیه خودران محقق شده است و ما بیش از هر زمان دیگری به سبک حمل و نقل آینده نگر نزدیک هستیم. بازار اینترنت خودرویی^۱ در سال ۲۰۲۱ معادل ۹۶ میلیارد دلار تخمین زده شده و تقریباً پیش بینی می شود که در سال ۲۰۲۸ به ۳۷۰ میلیارد دلار برسد [۱]. با تکامل هوش مصنوعی^۲، سرویس های حمل و نقل هوشمندتر و بیشتری در وسایل نقلیه هوشمند، مانند موقعیت یابی با دقت بالا، رانندگی هوشمند و سرگرمی های چندرسانه ای، پدیدار می شوند که تجربه رانندگی کاربران را تا حد زیادی بهبود می بخشد. با این حال، این سرویس های هوش مصنوعی بلادرنگ اغلب محاسباتی و حساس به زمان هستند که به دلیل منابع محاسباتی محدود موجود در خودروها، مدیریت کامل آنها توسط خودروها دشوار است. با بهره گیری از توسعه سریع فناوری ارتباطات بی سیم، برون سپاری وظایف از خودروها به سرورها به یک الگوی امیدوارکننده برای کمک به خودروها با منابع محاسباتی خارجی تبدیل شده است [۲]. رویکرد سنتی برون سپاری وظیفه مبتنی بر ابر، از منابع محاسباتی عظیم روی یک ابر متمرکز استفاده می کند که معمولاً در یک مکان دورافتاده قرار دارد. با این حال، رویکرد رایانش ابری^۳ اغلب از تأخیرهای بالای انتقال و اتصالات بی سیم نوسان دار رنج می برد که نمی تواند الزامات وظایف حساس به زمان را برآورده کند [۳].

برای برآورده کردن نیاز به تأخیر کم در برون سپاری وظیفه، رایانش لبه موبایل^۴ در شبکه های خودرویی، یعنی رایانش لبه خودرویی^۵، معرفی شده است که در آن واحد کنار جاده^۶ یا ایستگاه پایه^۷ به مقدار مشخصی از منابع محاسباتی و ذخیره سازی برای ارائه خدمات محاسباتی برای خودروها مجهز شده اند. با این حال، با افزایش تراکم ترافیک، تضمین کیفیت خدمات^۸ همه برنامه های خودرویی برای منابع محاسباتی محدود در واحد کنار جاده یا ایستگاه پایه دشوار است. علاوه بر این، سرورهای معمولاً به صورت پراکنده مستقر می شوند و عملکرد آنها توسط پوشش رادیویی محدود می شود [۴].

برای غلبه بر محدودیت های ذکر شده، رایانش مه^۹ به عنوان یک راه حل امیدوارکننده برای وظایف حساس به تأخیر در برنامه های کاربردی وسایل نقلیه در شبکه های وسایل نقلیه ظهور کرد [۵]. کنسرسیوم OpenFog معماری مرجع خود را برای ساخت یک سکوی رایانش مه سازگار و مقیاس پذیر در سال ۲۰۱۷ منتشر کرده است. با افزایش ترافیک سرویس های خودرویی، گره های مه احتمالاً با مشکل مواجه خواهند شد. یک سوال طبیعی این است که چگونه می توان منابع محاسباتی گره های مه را افزایش داد. رایانش مه خودرویی^{۱۰} به عنوان یک مورد استفاده از رایانش مه، در زمینه سیستم های حمل و نقل هوشمند پدیدار شده است. رایانش مه خودرویی از منابع خودرویی برای ارتقای قابلیت محاسباتی و همچنین کاهش بیشتر تأخیر محاسبات مه استفاده می کند. بنابراین، منابع محاسباتی استفاده نشده خودروها می توانند به عنوان اجزای گره های مه، مانند وسایل نقلیه در پارکینگ ها یا مراکز خرید، تکمیل شوند [۶].

با این حال، فناوری اینترنت با بهترین تلاش فعلی، انعطاف پذیری تصمیم گیری پویا در مورد برون سپاری وظایف را محدود می کند. برای رفع این محدودیت ها، محققان طرح های برون سپاری وظیفه را در شبکه های نرم افزار محور^{۱۱} پیشنهاد کردند که در آن یک کنترل کننده متمرکز تصمیمات برون سپاری وظیفه را می گیرد. علاوه بر این، رویکرد شبکه نرم افزار محور به

¹ Internet of Vehicles

² Artificial Intelligence

³ Cloud Computing

⁴ Mobile Edge Computing

⁵ Vehicular Edge Computing

⁶ Road Side Unit

⁷ Base Station

⁸ Quality of Service

⁹ Fog Computing

¹⁰ Vehicular Fog Computing

¹¹ Software Defined Networks

کنترل‌کننده اجازه می‌دهد تا شبکه را به شیوه‌ای انعطاف‌پذیر و ساده مدیریت کند [۷]. در نتیجه، در این کار، ما بر برون‌سپاری وظیفه مبتنی بر رویکرد نرم‌افزار محور در حوزه رایانش مه خودرویی تمرکز می‌کنیم.

شبکه نرم‌افزار محور برای رفع عدم انعطاف‌پذیری در شبکه قدیمی از طریق جداسازی صفحه ارسال از صفحه کنترل و فراهم کردن قابلیت برنامه‌ریزی در صفحه کنترل طراحی شده است. همچنین، کنترلر شبکه نرم‌افزار محور، مسئول جمع‌آوری تمام اطلاعات در صفحه داده به صورت دوره‌ای است. این اطلاعات مانند قابلیت محاسبه، تأخیر ارتباط، میانگین زمان انتظار و هزینه واحد برای حفظ یک نمای کلی از شبکه است. بنابراین، می‌تواند تمام منابع موجود در شبکه را به روشی کارآمد و سریع مدیریت، پیکربندی و هماهنگ کند. با آوردن این ویژگی به معماری پیشنهادی، یک سیاست برون‌سپاری وظیفه مبتنی بر شبکه نرم‌افزار محور برای تصمیم‌گیری صحیح برون‌سپاری به گره مه بهینه پیشنهاد شد [۸].

در نتیجه، اخیراً برون‌سپاری وظیفه در شبکه مه خودرویی توجه فزاینده‌ای را به خود جلب کرده است. با این حال، به دلیل ویژگی‌های منحصر به فرد شبکه‌های خودرویی، به ویژه تحرک بالای گره‌ها و ویژگی‌های کانال پویا، طراحی یک طرح برون‌سپاری وظیفه مبتنی بر لبه کارآمد بسیار چالش برانگیز است. بنابراین، چگونگی طراحی یک استراتژی کارآمد برای برون‌سپاری وظیفه که بتواند با محیط پویای شبکه سازگار شود، از منابع محاسباتی لبه به طور کامل استفاده کند و به تعادل بار بین دستگاه‌های مختلف دست یابد، به یک مسئله تحقیقاتی فوری تبدیل شده است.

چالش‌های زیادی در رایانش مه خودرویی وجود دارند. اولین چالش، انرژی محدود خودروهای کاربران است. اگر دستگاه‌های کاربران، انرژی کافی برای اجرای برنامه نداشته باشند، برنامه می‌تواند به طور کامل^۱ و جزئی^۲ به سرورهای دیگر در شبکه منتقل شود. علاوه بر این، اکثر خودروهای الکتریکی با قابلیت‌های محاسباتی، ذخیره‌سازی و انرژی بالا بیش از ۹۰ درصد از وقت خود را در پارکینگ‌ها می‌گذرانند. بنابراین، یک ایده می‌تواند استفاده از منابع محاسباتی خودروهای الکتریکی پارک شده برای ارائه خدمات محاسباتی بر اساس تقاضا در سراسر شبکه باشد [۹].

چالش دیگر، حفظ کیفیت سرویس در سرورهای مه است. کیفیت سرویس را می‌توان با تقسیم بار محاسباتی به طور مساوی بین سرورهای مه بهبود بخشید. بار محاسباتی در یک سرور مه به نرخ‌های ورود درخواست مرتبط با آن بستگی دارد. سرورهای مه ممکن است نرخ‌های ورود درخواست متفاوتی داشته باشند، این ممکن است وضعیتی را ایجاد کند که یک سرور مه بیش از بقیه سرورها بارگذاری شود. از این رو، ارتباط بین سرورهای مه باید به گونه‌ای انجام شود که عدم تعادل بار در سرورهای مه به حداقل برسد.

امروزه، از تکنیک‌های یادگیری تقویتی^۳ برای یافتن یک استراتژی بهینه برون‌سپاری وظیفه استفاده می‌شود [۱۰]. یادگیری تقویتی به طور هوشمندانه سیاست بهینه را برای به حداکثر رساندن پاداش بر اساس بازخورد دائمی از محیط یاد می‌گیرد. با این حال، زمان مورد نیاز برای کاوش کامل محیط پیچیده و در نهایت همگرایی به یک سیاست بهینه، یادگیری تقویتی سنتی را غیرممکن می‌سازد. این امر به ویژه هنگام مقیاس‌بندی یادگیری تقویتی برای محیط‌های بزرگ و پویا با فضای حالت-عمل گسترده و پیوسته چالش برانگیز است. برای پرداختن به چنین فضاهایی در مقیاس بزرگ، یادگیری تقویتی عمیق^۴ معرفی شده است تا تعمیم و سازگاری بهتری با محیط‌های ناشناخته ارائه دهد [۱۱].

در این مقاله، ما یک تکنیک برون‌سپاری وظیفه فدرال بر خط^۵ مبتنی بر یادگیری عمیق^۶ را برای رایانش مه خودرویی پیشنهاد می‌کنیم که به طور مشترک مصرف انرژی و تعادل بار را در سراسر سرورهای مه بهینه می‌کند. ما از یک شبکه مه خودرویی مبتنی بر منطقه استفاده می‌کنیم که در آن هر منطقه دارای یک سرور مه ایستا، سرور مه پویا و چندین واحد کنار جاده ای است. در هر منطقه مه، یک شبکه یادگیری عمیق توسط یک عامل آموزش داده می‌شود که یک استراتژی برون‌سپاری وظیفه محلی بهینه برای سرورهای مه مربوط به خود تعیین می‌کند. یک استراتژی برون‌سپاری وظیفه بهینه، مصرف انرژی را

¹ Binary

² Partial

³ Reinforcement learning

⁴ Deep Reinforcement Learning

⁵ Online federated task offloading

⁶ Deep Learning

در سرورهای مه به حداقل می‌رساند و بار را به طور یکنواخت در سرورهای مه در یک منطقه مه توزیع می‌کند. یک رویکرد برون‌سپاری وظیفه سراسری با جمع کردن پارامترهای یادگیری محلی مناطق مختلف مه به دست می‌آید. دستاوردهای مقاله ما به شرح زیر است:

- ما مسئله برون‌سپاری وظیفه را به عنوان یک فرآیند تصمیم‌گیری مارکوف^۱ مدل‌سازی می‌کنیم و در عین حال ظرفیت محاسباتی متغیر با زمان یک سرور مه پویا و صف‌های اولویت دار برای زمانبندی وظایف در وسایل نقلیه را در نظر می‌گیریم. مدل مارکوف هم بهینه‌سازی انرژی و هم تعادل بار در سرورهای مه را در نظر می‌گیرد.
 - ما یک تکنیک برون‌سپاری وظیفه فدرال بر خط چند عامله مبتنی بر یادگیری عمیق را برای رایانش مه خودرویی پیشنهاد می‌کنیم. یک تکنیک برون‌سپاری وظیفه فدرال^۲، با تجمیع مدل‌های برون‌سپاری وظیفه محلی مناطق مختلف مه، یک مدل برون‌سپاری وظیفه سراسری توسعه می‌دهد.
 - ما شبیه‌سازی‌های گسترده‌ای را برای ارزیابی عملکرد رویکرد پیشنهادی خود را با استفاده از داده‌های واقعی حرکت خودروها انجام می‌دهیم. نتایج شبیه‌سازی ما نشان می‌دهد که روش پیشنهادی با نرخ همگرایی مشابهی با رویکرد متمرکز همگرا می‌شود و در مقایسه با سایر تکنیک‌های پیشرفته برون‌سپاری وظیفه، عملکرد بهتری دارد.
- ساختار این مقاله به شرح زیر است. بخش ۲ کارهای مرتبط را ارائه می‌دهد. بخش ۳ مدل سیستم و فرمول‌بندی مسئله را شرح می‌دهد. در بخش ۴، تکنیک پیشنهادی ارائه می‌شود. ارزیابی عملکرد در بخش ۵ نشان داده شده است. در نهایت، نتیجه‌گیری در بخش ۶ ارائه شده است.

۲-مروری بر کارهای انجام شده

ماهیت پویای رایانش مه خودرویی، تحقیقات گسترده‌ای را در مورد استراتژی‌های برون‌سپاری وظیفه، با تمرکز اصلی بر کاهش زمان پاسخ و مصرف انرژی، برانگیخته است. محدودیت‌های مهم مطالعات موجود شامل عدم قطعیت در تحرک و تقاضاهای درخواست، پوشش محدود سرورهای مه و پیچیدگی محاسباتی وظایف مختلف است. در این بخش، ما یک دسته بندی برای کارهای مرتبط بر اساس رویکردهای برون‌سپاری وظیفه بهینه مبتنی بر یادگیری تقویتی تک عامله و رویکردهای برون‌سپاری وظیفه بهینه مبتنی بر یادگیری تقویتی چند عامله در محیط رایانش مه خودرویی را ارائه می‌دهیم. جدول ۱ تحقیقات مرتبط و روش پیشنهادی را با توجه به پارامترهای مختلف و روش‌های یادگیری ماشین مقایسه می‌کند.

۲-۱- رویکردهای برون‌سپاری وظیفه بهینه مبتنی بر یادگیری تقویتی تک عامله

با پیشرفت یادگیری ماشین، بسیاری از تحقیقات موجود یادگیری تقویتی تک عاملی را برای حل مشکلات تخصیص منابع محاسباتی در رایانش مه خودرویی به کار گرفته اند. در [۴]، یک سیاست برون‌سپاری وظیفه ارائه شده است که پویایی شبکه خودرو، اولویت وظیفه و در دسترس بودن سرویس را در نظر می‌گیرد. مسئله برون‌سپاری وظیفه با استفاده از فرآیند تصمیم‌گیری مارکوف فرموله شده و با استفاده از الگوریتم یادگیری تقویتی عمیق مبتنی بر مدل منتقد بازیگر نرم (SAC)^۳ حل شده است. هدف این مقاله به حداکثر رساندن میانگین مطلوبیت آگاه از تأخیر است. یادگیری تقویتی عمیق در [۱۲] برای یافتن بهترین راه‌حل بلندمدت در یک چارچوب تصمیم‌گیری برون‌سپاری وظیفه مبتنی بر دانش، برای اینترنت خودرویی استفاده می‌شود. هنگامی که سرویس در حال کار است، خودروها یک مدل جدید را در گره محاسبات لبه به‌روزرسانی می‌کنند و یادگیری برخط ناهمزمان را انجام می‌دهند. در [۱۳]، یک معماری شبکه رایانش لبه خودرویی با استفاده از فناوری شبکه نرم‌افزار محور پیشنهاد شد. سپس، یک طرح برون‌سپاری وظیفه مبتنی بر یادگیری عمیق برای سازگاری با نرخ بالای تغییر شبکه ارائه گردید. در [۱۴]، بهینه‌سازی تحویل مبتنی بر یادگیری ارائه شده است. با استفاده از شبکه عصبی پیش‌خور، گره مه در زمان و مکان داده شده پیش‌بینی می‌شود و شبکه عصبی بازگشتی نیز برای پیش‌بینی کارایی و هزینه پیشنهاد شده است. برون‌سپاری وظیفه جزئی

¹ Markov Decision Process

² Federated

³ Soft Actor-Critic

دودویی خودرو به خودرو (V2V) در [۱۵] برای افزایش دسترسی به سرویس و استفاده از منابع پیشنهاد شده است. مدل‌سازی ریاضی توسط الگوریتم یادگیری عمیق حل شده است تا تأخیر و هزینه اجرای وظیفه را کاهش دهد. بوبکر و همکاران [۱۶] پیشنهاد دادند که یک سیاست بهینه برای برون‌سپاری جزئی وظایف گردش کار و مهاجرت پویا در محیط رایانش مه خودرویی ناهمگن و محاسبات ابری پیدا کنند. این مقاله بر رعایت محدودیت‌های مهلت و وابستگی وظایف تمرکز دارد. گردش کار به عنوان یک گراف جهتدار غیر مدور (DAG) برای نشان دادن مجموعه وابستگی‌های متوالی بین جفت وظایف نشان داده می‌شود. در [۱۷]، یک مکانیسم تشویقی مبتنی بر نظریه قرارداد پیشنهاد شده است که از روش یادگیری تقویتی عمیق به صورت توزیع‌شده برای کاهش پیچیدگی پیاده‌سازی استفاده می‌کند و تصمیمات برون‌سپاری وظیفه می‌تواند از طریق شبکه عصبی عمیق سریع‌تر انجام شود. خودروها منابع خود را به اشتراک می‌گذارند و بر اساس میزان منابعی که ارائه می‌دهند، پاداش دریافت می‌کنند. همچنین، خودرو از پاداش انباشته خود برای تبادل پهنای باند انتقال بی‌سیم استفاده می‌کند. این روش مصرف انرژی و نسبت تکمیل وظیفه را بهبود می‌بخشد.

۲-۲- رویکردهای برون‌سپاری وظیفه بهینه مبتنی بر یادگیری تقویتی چند عامله

برون‌سپاری وظیفه و تخصیص منابع، یک نقطه کانونی تحقیقاتی در محاسبات مشارکتی رایانش مه خودرویی است. یادگیری تقویتی تک عاملی، دارای برخی نقص‌ها مانند استحکام پایین^۲ و نادیده گرفتن مقیاس‌پذیری^۳ است. مرجع [۱۸] یک یادگیری تقویتی عمیق چندعاملی فدرال را برای تصمیمات بهتر برون‌سپاری وظیفه پیشنهاد می‌دهد. این روش با یادگیری تصمیمات برون‌سپاری در چندین لایه، حریم خصوصی داده‌ها و همگرایی را در مدل‌های یادگیری مشارکتی بهبود می‌بخشد. نتایج نشان می‌دهد که رویکرد پیشنهادی زمان اقامت، تأخیر انتها به انتها و هزینه کلی سیستم را کاهش می‌دهد. در مرجع [۱۹]، یک استراتژی مهاجرت و همکاری برای انتقال نتایج محاسبات بین واحدهای کنار جاده ای پیشنهاد شده است که می‌تواند میانگین شکست در تحویل نتایج محاسبات، تأخیر و مصرف انرژی را کاهش دهد. رویکرد یادگیری عمیق چندعاملی، تصمیمات برون‌سپاری وظیفه و تخصیص منابع را از طریق آموزش متمرکز بهینه می‌کند. دای و همکارانش [۲۰]، سناریوهای مختلف برون‌سپاری چند وظیفه‌ای (MTO)^۴ را با یادگیری تقویتی متا^۵ برای پیش‌بینی زمان اجرای وظیفه و سازگاری با سناریوهای پویا پیشنهاد کردند. وظایف کاربر به صورت گراف مبتنی بر گراف جهتدار غیر مدور مدل‌سازی می‌شوند تا زمان اجرا به حداقل برسد. مرجع [۲۱] یک چارچوب مدیریت لبه ارائه می‌دهد که کارایی یادگیری چندین عامل را بهبود می‌بخشد. طرح یادگیری چندعاملی توزیع‌شده می‌تواند هزینه برون‌سپاری وظیفه خودرو را تحت محدودیت تأخیر کاهش داده و استفاده از منابع را افزایش دهد. همچنین رایانش لبه‌ای نرم‌افزار محور ارائه شد که از روش یادگیری ماشین توزیع‌شده برای ایجاد یک معماری مقیاس‌پذیر، مدیریت منابع و هماهنگی استفاده می‌شود [۲۲]. مرجع [۲۳] معماری سلسله مراتبی رایانش مه خودرویی را در نظر می‌گیرد و وظایف را بر اساس یادگیری تقویتی چندعاملی، هم در داخل و هم در داخل ناحیه توزیع می‌کند. رویکرد یادگیری تقویتی چندعاملی به نام COMA^۶ یک مبنای ارزیابی عملکرد برای هر عامل ارائه می‌دهد. نتایج نشان می‌دهد که این روش می‌تواند از قابلیت محاسباتی سیستم بهتر استفاده کند. در [۱۰]، یک تکنیک برون‌سپاری وظیفه مبتنی بر یادگیری Q عمیق فدرال برای رایانش مه خودرویی برای بهبود مصرف انرژی و تعادل بار در واحدهای کنار جاده پیشنهاد شده است. برخلاف کارهای تحقیقاتی مذکور، روش پیشنهادی ما یک تکنیک برون‌سپاری وظیفه پویا است که جنبه‌های بهبود کیفیت سرویس در شبکه‌های مه خودرویی مانند کاهش هزینه، تأخیر و مصرف انرژی، متعادل سازی بار در سرورهای مه و استفاده از اطلاعات سراسری در توسعه یک مدل برون‌سپاری وظیفه را پوشش می‌دهد.

¹ Vehicle to Vehicle

² Directed Acyclic Graph

³ Low robustness

⁴ Scalability

⁵ Multi Task Offloading

⁶ Meta reinforcement learning

⁷ Counterfactual multi-agent

جدول ۱: مقایسه کارهای انجام شده براساس پارامترهای تاثیر گذار روی برون سپاری وظیفه
Table 1. Comparison of completed tasks based on parameters affecting task offloading

محیط	هزینه	اولویت وظایف	مهلت زمانی	انرژی	زمان انتظار در صف	مرجع
خودرو-مه	×	✓	✓	×	✓	[۴]
خودرو-مه	×	×	✓	✓	×	[۱۰]
خودرو-لبه	×	×	×	×	×	[۱۲]
خودرو-MEC-ابر	×	×	×	×	×	[۱۳]
خودرو-مه-ابر	×	×	×	×	×	[۱۴]
خودرو-مه	✓	×	✓	×	×	[۱۵]
خودرو-مه-ابر	×	✓	✓	✓	×	[۱۶]
خودرو-مه	×	×	✓	×	×	[۱۷]
خودرو-مه	×	✓	×	×	×	[۱۸]
خودرو-لبه-ابر	×	✓	✓	✓	×	[۱۹]
خودرو-MEC-ابر	×	×	×	×	✓	[۲۰]
خودرو-لبه	✓	×	✓	×	×	[۲۱]
خودرو-مه-ابر	×	×	×	×	×	[۲۲]
خودرو-مه	×	✓	✓	×	×	[۲۳]
خودرو-مه-ابر	✓	✓	✓	✓	✓	روش
پیشنهادی						

یادداشت: ✓=پارامترهایی که در نظر گرفته شده اند ، ×=پارامترهایی که در نظر گرفته نشده اند

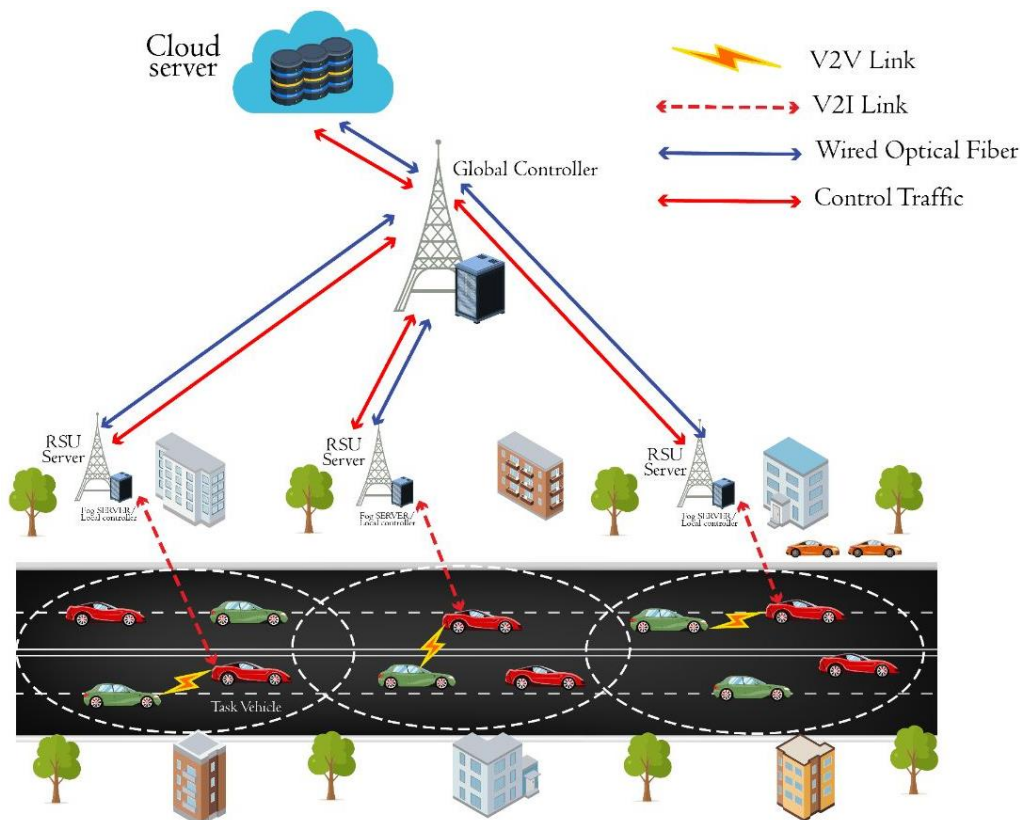
۲-روش پیشنهادی

معماری سلسله مراتبی مبتنی بر شبکه نرم افزار محور ، برای محیط رایانش مه خودرویی در شکل ۱ نشان داده شده است. هر منطقه شامل چند خوشه است. در هر خوشه، یک واحد کنار جاده ای، خودروهای پارک شده کنار جاده (RSPVs)^۱ و یک خودروی در حال حرکت با قابلیت محاسباتی بالا مانند اتوبوس به عنوان گره های مه عمل می کنند. ما فرض می کنیم که یک کنترل کننده محلی که عامل روی آن قرار دارد، یک خوشه را مدیریت می کند. کنترل محلی روی هر واحد کنار جاده مستقر شده است. در لایه اول کنترل کننده های محلی قرار دارند و در لایه دوم کنترل کننده سراسری قرار دارد. همانطور که قبلا گفته شد، ما از یک تکنیک برون سپاری وظیفه فدرال بر خط چند عامله مبتنی بر یادگیری عمیق برای پیاده سازی برون سپاری وظیفه استفاده می کنیم. اعمال رویکردهای یادگیری فدرال بر روی شبکه های نرم افزار محور می تواند مزایایی را داشته باشد. یک صفحه کنترل شبکه نرم افزار محور می تواند چندین کنترل کننده داشته باشد که در لایه ها سازماندهی شده اند؛ کنترل کننده های لایه پایینی می توانند در لبه شبکه قرار گیرند تا با عدم تقارن در سوئیچ ها و میزبان های متصل مقابله کنند و کنترل کننده در لایه بالایی می تواند کل شبکه را به صورت مرکزی و سراسری نظارت کند. اعمال یادگیری فدرال در شبکه های نرم افزار محور با یک صفحه کنترل توزیع شده لایه ای، می تواند ضمن بهبود همکاری بین شرکت کنندگان برای تولید مدل های با کیفیت بالاتر بر روی شبکه های نامتقارن، از حریم خصوصی داده های هر شرکت کننده محافظت کند [۲۴].

هنگامی که یک خودروی در حال حرکت نمی تواند وظایف را به صورت محلی اجرا کند، مشخصات وظایف را به کنترلر محلی در خوشه خود ارسال می کند. پس از آن، کنترل کننده محلی سیاست بهینه را برای اختصاص برنامه ها به گره های مه در خوشه خود انتخاب می کند تا مهلت انجام وظایف رعایت شود. یک کنترل کننده سراسری شبکه نرم افزار محور از تمامی خوشه های آن

¹ Road Side Parked vehicles

منطقه آگاهی دارد. در پایان هر قسمت^۱ از یادگیری، وزن‌های مدل برون‌سپاری وظیفه که به صورت محلی آموخته شده است، با کنترلر سراسری به اشتراک گذاشته می‌شوند. کنترلر سراسری، وزن‌های مدل محلی را تجمیع کرده و یک مدل سراسری بروز شده را برای کنترل‌های محلی به اشتراک می‌گذارد. در نتیجه، مسئله بهینه‌سازی در لایه پایین‌تر حل می‌شود که در آن کنترلرها یا عامل‌های محلی، مدل برون‌سپاری وظیفه را به صورت توزیع‌شده یاد می‌گیرند.



شکل ۱: معماری توزیع شده پیشنهادی شامل کنترلر سراسری و کنترلرهای محلی است. کنترلرهای محلی منابع محلی را تخصیص می‌دهند و کنترلر سراسری دانش کلی شبکه را دارد.

Figure 1. The proposed distributed architecture consists of a global controller and local controllers. Local controllers allocate local resources and the global controller has global knowledge of the network.

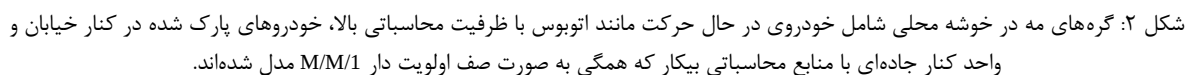
۲-۱- مدل وظیفه

وظایف در خودروها را می‌توان به سه دسته وظایف حیاتی، وظایف با اولویت بالا و اولویت پایین تقسیم‌بندی کرد. وظایف حیاتی، وظایف هسته‌ای و مرتبط با ایمنی هستند و باید به صورت محلی در خودروها اجرا شوند. بنابراین، خودروها باید مقداری از منابع محاسباتی را برای وظایف حیاتی محلی رزرو کنند. وظایف با اولویت بالا نشان‌دهنده دسته‌ای از وظایف هستند که محدودیت تأخیر شدیدی دارند، مانند ناوبری، حسگر جاده و غیره. اگر یک وظیفه با اولویت بالا نتواند در حداکثر تأخیر قابل تحمل به پایان برسد، وظیفه با شکست مواجه می‌شود و ممکن است ضررهایی را برای خودرو به همراه داشته باشد. در مقایسه با وظایف با اولویت پایین، وسایل نقلیه باید اطمینان حاصل کنند که وظایف با اولویت بالا ابتدا اجرا می‌شوند. وظایف با اولویت پایین دسته‌ای از وظایف مقاوم در برابر تأخیر هستند، مانند برنامه‌های سرگرمی خودرویی، خدمات ارزش افزوده و غیره. اگر زمان تکمیل یک وظیفه با اولویت پایین از تأخیر مرجع بیشتر شود، نتیجه همچنان قابل استفاده خواهد بود، اما با افزایش زمان اجرا، در دسترس بودن نتیجه کاهش می‌یابد [۴]. در روش پیشنهادی فرض می‌شود که وظایف با اولویت بالا و اولویت پایین

¹ Episode

خودروها وظایف محاسباتی برنامه‌ها را تولید می‌کنند که هر وظیفه i مربوط به خودروی v به صورت چندگانه $(x_{i,v}, y_{i,v}, \lambda_{i,v}, d_{i,v}, \tau_{i,v}, L_{i,v}, b_{i,v}, B_{i,v}, K_{i,v})$ تعریف می‌شوند. که در آن $x_{i,v}$ و $y_{i,v}$ مختصات خودرو v شامل وظیفه i را نشان می‌دهند، $\lambda_{i,v}$ ، نرخ تولید وظیفه محاسباتی، $d_{i,v}$ ، اندازه وظیفه (کد برنامه و پارامترهای مربوطه) برای انتقال، $\tau_{i,v}$ ، حداکثر تأخیر قابل تحمل وظیفه، $L_{i,v}$ ، مقدار منبع محاسباتی مورد نیاز برای اجرای وظیفه، $b_{i,v}$ ، بودجه وظیفه، $B_{i,v}$ ، پهنای باند مورد استفاده وظیفه و $K_{i,v}$ ، اولویت وظیفه را نشان می‌دهد و مقدار آن صفر یا یک است.

این بخش، مدل ارتباطی، مدل محاسباتی، مدل‌سازی ریاضی و الگوریتم پیشنهادی در سیستم مه خودروپی را معرفی می‌کند. از سیستم صف جهت مدل‌سازی گره‌های مه مختلف استفاده شده است که در شکل ۲ نشان داده شده است.



۲-۲-۱- بیرون‌سیاری وظیفه خودرو به خودرو^۱

به منظور کاهش بار محاسباتی واحدهای کنار جاده ای، هنگامی که دستگاه‌های اطراف خودرو منابع محاسباتی آزادی دارند که می‌توانند مورد استفاده قرار گیرند، خودروها می‌توانند از طریق ارتباط خودرو به خودرو با یکدیگر ارتباط برقرار کنند تا وظایف را واگذار کرده و فشار محاسباتی را کاهش دهند. این مقاله استفاده از خودروهای در حال حرکت و پارک شده را پیشنهاد می‌دهد. به‌طوری که آنها می‌توانند به عنوان گره مه عمل کنند. نرخ انتقال داده بین خودروی v و خودروی v' برای وظیفه i را می‌توان به صورت رابطه ۱ بیان کرد:

$$\text{tr}_{v,v'} = B_{i,v} \log_2 \left(1 + \frac{|h_{v,v'}|^2 p_{v,v'}}{\sigma(\text{dis}_{v,v'})^{\bar{\omega}}} \right) \quad (1)$$

¹ Vehicle-to-Vehicle task offloading

که در آن B_i ، نشان دهنده پهنای باند کانال، $h_{v,v'}$ ، ضریب محوشدگی کانال، $p_{v,v'}$ ، نشان دهنده توان انتقال خودروی v برای ارسال به خودروی v' ، σ ، نشان دهنده توان نویز سفید گاوسی، $\bar{w}$ ، نشان دهنده توان تلفات مسیر و $dis_{v,v'}$ ، فاصله بین خودروی v و خودروی v' است. این فاصله را می توان با رابطه ۲ بیان کرد:

$$dis_{v,v'} = \sqrt{(x_v - x_{v'})^2 + (y_v - y_{v'})^2} \quad (2)$$

که در آن x_v و y_v مختصات خودروی v را نشان می دهند در حالی که $x_{v'}$ و $y_{v'}$ مختصات خودروی v' در حال حرکت v' را نشان می دهند. زمان انتقال وظیفه i (زمان بارگذاری) به گره v' را می توان طبق رابطه ۳ محاسبه کرد:

$$T_{v,v'}^u = \frac{d_{i,v}}{tr_{v,v'}} \quad (3)$$

در روش پیشنهادی، خودروهای پارک شده و در حال حرکت توسط صف های اولویت دار $M/M/1$ مدل سازی می شوند. کنترل کننده محلی در خوشه، نرخ رسیدن دستورالعمل های $\Delta_{v'}^1$ و $\Delta_{v'}^2$ را به ترتیب برای کلاس ۱ و کلاس ۲ به خودروی v' اختصاص می دهد. برای یک وظیفه با نرخ ورود $\lambda_{i,v}$ و میزان پردازش $L_{i,v}$ ، آنگاه $\Delta_{v'}^1$ و $\Delta_{v'}^2$ را می توان به ترتیب به صورت رابطه ۴ و رابطه ۵ محاسبه کرد:

$$\Delta_{v'}^1 = \sum_i L_{i,v} \lambda_{i,v} K_{i,v} x_{i,v'} \quad (4)$$

$$\Delta_{v'}^2 = \sum_i L_{i,v} \lambda_{i,v} (1 - K_{i,v}) x_{i,v'} \quad (5)$$

مجموع کل نرخ رسیدن دستورالعمل ها برای خودروی v' در خوشه به صورت رابطه ۶ محاسبه می شود:

$$\Delta_{v'} = \Delta_{v'}^1 + \Delta_{v'}^2 \quad (6)$$

میانگین زمان انتظار در صف برای وظیفه با اولویت بالا و اولویت پایین را می توان به ترتیب به صورت معادلات ۷ و ۸ محاسبه کرد:

$$W_{i,v'}^1 = \frac{\Delta_{v'}}{\mu_{v'}(\mu_{v'} - \Delta_{v'}^1)} \quad (7)$$

$$W_{i,v'}^2 = \frac{\Delta_{v'}}{(\mu_{v'} - \Delta_{v'}^1)(\mu_{v'} - \Delta_{v'}^1 - \Delta_{v'}^2)} \quad (8)$$

$T_{i,v'}^w$ را می توان به صورت رابطه ۹ محاسبه کرد:

$$T_{i,v'}^w = \begin{cases} W_{i,v'}^1 & \text{if } K_{i,v} = 1 \\ W_{i,v'}^2 & \text{if } K_{i,v} = 0 \end{cases} \quad (9)$$

بنابراین، میانگین زمان مورد نیاز خودروی v' برای تکمیل وظیفه i را می توان به صورت رابطه ۱۰ بیان کرد:

$$T_{i,v'}^p = \frac{1}{\mu_{v'}} \quad (10)$$

$\mu_{v'}$ ، ظرفیت محاسباتی گره v' است. ما فرض می کنیم هزینه واحد به ازای هر MIPS^۱ برای گره v' برابر با $C_{v'}$ است. بنابراین، میانگین هزینه برون سپاری محاسباتی را می توان با رابطه ۱۱ محاسبه کرد:

$$C_{i,v'} = C_{v'} \Delta_{v'} \quad (11)$$

میانگین مصرف انرژی مورد نیاز ارتباطی و محاسباتی برای تکمیل وظیفه با استفاده از رابطه ۱۲ محاسبه می شود:

$$E_{i,v'} = E_{v'} \Delta_{v'} + p_{v,v'} T_{v,v'}^u \quad (12)$$

¹ Million Instruction Per second

که در آن $E_{v'}$ نشان دهنده مصرف انرژی به ازای هر MIPS برای گره مه v' است. میانگین زمان برون سپاری یا میانگین زمان پاسخ وظیفه i که توسط گره مه خودرویی v' پردازش می شود، می تواند در کنترل کننده محلی محاسبه شود که عموماً به صورت رابطه ۱۳ تعریف می شود:

$$tf_{i,v'} = T_{i,v'}^u + T_{i,v'}^w + T_{i,v'}^p + T_{i,v'}^d \quad (13)$$

$T_{i,v'}^d$ و $T_{i,v'}^u$ به ترتیب زمان های بارگذاری و بازگشت پاسخ در خوشه محلی توسط گره مه خودرویی v' هستند. برای سادگی، زمان بارگذاری و بازگشت پاسخ از گره مه با یکدیگر برابر در نظر گرفته می شوند.

۲-۲-۲- برون سپاری وظیفه خودرو به زیرساخت^۱

در این روش، خودروها می توانند وظایف خود را به واحدهای کنار جاده ای (RSUs)^۲ یا ابر برون سپاری نمایند. ارتباط بین خودروها به واحد کنار جاده ای از طریق ارتباط بیسیم و از واحد کنار جاده ای به ابر از طریق فیبر نوری است. نرخ انتقال داده بین خودروی v و واحد کنار جاده ای طبق معادله ۱۴ محاسبه می شود:

$$tr_{v,r} = B_{i,v} \log_2 \left(1 + \frac{|h_{v,r}|^2 p_{v,r}}{\sigma (dis_{v,r})^{\bar{\omega}}} \right) \quad (14)$$

که در آن $h_{v,r}$ ضریب محوشدگی کانال، $p_{v,r}$ ، نشان دهنده توان انتقال خودرو، σ ، نشان دهنده توان نویز سفید گاوسی، ω ، نشان دهنده توان تلفات مسیر و $dis_{v,r}$ ، فاصله بین خودروی v و واحد کنار جاده ای r است. این فاصله را می توان با رابطه ۱۵ بیان کرد:

$$dis_{v,r} = \sqrt{(x_v - x_r)^2 + (y_v - y_r)^2} \quad (15)$$

زمان انتقال وظیفه i در گره مه خودرویی v به واحد کنار جاده ای r را می توان بر طبق رابطه ۱۶ محاسبه کرد:

$$T_{v,r}^u = \frac{d_{i,v}}{tr_{v,r}} \quad (16)$$

مشابه با خودروهای مه، واحدهای کنار جاده ای نیز به صورت صف $M/M/1$ در نظر گرفته می شود. بنابراین، نرخ رسیدن دستورالعمل های Δ_r^1 و Δ_r^2 ، به ترتیب برای کلاس ۱ و کلاس ۲ به واحد کنار جاده ای به صورت زیر محاسبه می شوند:

$$\Delta_r^1 = \sum_i L_{i,v} \lambda_{i,v} K_{i,v} x_{i,r} \quad (17)$$

$$\Delta_r^2 = \sum_i L_{i,v} \lambda_{i,v} (1 - K_{i,v}) x_{i,r} \quad (18)$$

که در این صورت کل نرخ رسیدن دستورالعمل ها به واحد کنار جاده ای مجموع نرخ دستورالعمل های کلاس ۱ و کلاس ۲ است که می توان مشابه با فرمول های ۷ و ۸، میانگین زمان انتظار با اولویت بالا و پایین را بدست آورد که به صورت $W_{i,r}^1$ و $W_{i,r}^2$ نمادگذاری شده اند. همچنین میانگین زمان مورد نیاز برای پردازش وظیفه i ، به صورت رابطه ۱۹ محاسبه می گردد:

$$T_{i,r}^p = \frac{1}{\mu_r} \quad (19)$$

μ_r ، ظرفیت محاسباتی گره r است. همچنین مشابه خودروهای پارک شده و در حال حرکت، میانگین مصرف انرژی، میانگین هزینه و میانگین زمان برون سپاری وظیفه از خودرو به واحد کنار جاده ای به ترتیب با استفاده از رابطه های ۲۰ تا ۲۲ بدست می آیند:

$$E_{i,r} = E_r \Delta_r + p_{v,r} T_{v,r}^u \quad (20)$$

$$C_{i,r} = C_r \Delta_r \quad (21)$$

¹ Vehicle-to-Infrastructure task offloading

² Road Side Units (RSUs)

$$tf_{i,r} = T_{i,r}^u + T_{i,r}^w + T_{i,r}^p + T_{i,r}^d \quad (22)$$

$T_{i,r}^d$ و $T_{i,r}^u$ به ترتیب زمان های بارگذاری و بازگشت پاسخ در خوشه محلی توسط واحد کنار جاده ای r هستند. $C_{i,r}$ ، $E_{i,r}$ و $tf_{i,r}$ به ترتیب میانگین مصرف انرژی، میانگین هزینه و میانگین زمان برون سپاری وظیفه خودرو به واحد کنار جاده ای هستند.

در صورتیکه منابع مه محلی پاسخگوی نیازمندی های وظیفه نباشند، کنترلر محلی وظیفه را به ابر منتقل می کند. از آنجا که ارتباط بین واحدهای کنار جاده ای با مرکز ابر از طریق فیبر نوری است، بنابراین سرعت ارسال داده ای ثابت در نظر گرفته می شود. بنابراین زمان انتقال داده ها از واحد کنار جاده ای به ابر از طریق رابطه زیر بدست می آید:

$$T_{r,c}^u = \frac{d_{i,v}}{tr_{r,c}} \quad (23)$$

در این مقاله، فرض می شود که مرکز ابر با C ماشین مجازی به صورت صف $M/M/C$ مدل سازی می شوند که میانگین زمان انتظار در صف می تواند بر اساس فرمول های زیر محاسبه گردد:

$$T_{i,c}^w = \frac{C(c,\rho)}{(c\mu_c - \Delta_c)} \quad (24)$$

$$C(c,\rho) = \frac{\left(\frac{c\rho^c}{c!}\right)\left(\frac{1}{1-\rho}\right)}{\sum_{k=0}^{c-1} \left(\frac{k\rho^k}{k!}\right) + \left(\frac{c\rho^c}{c!}\right)\left(\frac{1}{1-\rho}\right)} \quad (25)$$

$C(c,\rho)$ ، فرمول C ارلنگ [۲۵] و $\rho = \frac{\Delta_c}{c\mu_c}$ نرخ شدت ترافیک در سرور مجازی C است. همچنین میانگین زمان مورد نیاز برای

پردازش وظیفه t در ابر توسط سرور مجازی C ، به صورت معادله ۲۶ محاسبه می گردد:

$$T_{i,c}^p = \frac{1}{\mu_c} \quad (26)$$

μ_r ، ظرفیت محاسباتی گره r است. با توجه به اینکه داده ها ابتدا باید به واحد کنار جاده ای و سپس به ابر منتقل شوند، بنابراین میانگین مصرف انرژی، میانگین هزینه پردازش و میانگین زمان برون سپاری وظیفه از خودرو به ابر به ترتیب با استفاده از رابطه های ۲۷ تا ۲۹ بدست می آیند:

$$E_{i,c} = p_{v,r} T_{v,r}^u + p_{r,c} T_{r,c}^u + E_c \Delta_c \quad (27)$$

$$C_{i,c} = C_c \Delta_c \quad (28)$$

$$tf_{i,c} = T_{i,r}^u + T_{r,c}^u + T_{i,r}^w + T_{i,c}^p + T_{i,c}^d \quad (29)$$

$T_{i,c}^d$ و $T_{i,c}^u$ به ترتیب زمان های بارگذاری و بازگشت پاسخ از ابر هستند. $C_{i,c}$ ، $E_{i,c}$ و $tf_{i,c}$ به ترتیب میانگین مصرف انرژی، میانگین هزینه و میانگین زمان برون سپاری وظیفه خودرو به ابر هستند. در روش پیشنهادی، فرض می شود که محدودیتی برای مصرف انرژی و همچنین پهنای باند، برای ارسال وظیفه به ابر وجود ندارد.

۲-۲-۳-مدلسازی ریاضی

این زیر بخش، تابع هدف و محدودیت های مسئله برون سپاری وظیفه را بیان می کند. تابع هدف براساس نرخ وظایف انجام شده با اولویت بالا و پایین را با توجه به محدودیت های تاخیر، انرژی، هزینه و پهنای باند خطوط ارتباطی است. با توجه به مدل های برون سپاری مختلف ذکر شده، سه متغیر تصمیم باینری نحوه انتخاب سرور مقصد را در برش زمانی t بیان می کنند که به صورت زیر است:

$$X_{i,v',t}, X_{i,r,t}, X_{i,c,t} \in \{0,1\} \quad (30)$$

$X_{i,v',t}$ ، متغیر تصمیم باینری برای انتخاب خودروی مه محلی، $X_{i,r,t}$ ، متغیر تصمیم باینری برای انتخاب واحد کنار جاده ای و $X_{i,c,t}$ ، متغیر تصمیم باینری برای انتخاب ماشین مجازی در ابر است. مجموعه ی خودروها، خودروهای مه، ماشین های مجازی، و وظایف با استفاده از رابطه های ۳۱ تا ۳۴ بیان می شود:

$$V=\{vehicle_1, vehicle_2,..., vehicle_v\} \quad (31)$$

$$V'=\{moving vehicle_1, parked vehicle_2,..., parked vehicle_n\} \quad (32)$$

$$C=\{virtual machine_1, virtual machine_2,..., virtual machine_c\} \quad (33)$$

$$I=\{task_1, task_2,..., task_m\} \quad (34)$$

تابع هدف براساس رابطه ۳۵ بیان می شود که در هر برش زمانی t ، ماکزیمم نرخ وظایف پذیرفته شده مربوط به خودروی v به عنوان تابع هدف در نظر گرفته می شود:

$$\max Acc_v^t = \frac{\sum_{i \in I(t)} \psi_{i,v,t}}{I(t)} \quad (35)$$

$$s.t. (c1) \quad X_{i,v',t}, X_{i,r,t}, X_{i,c,t} \in \{0,1\}, \quad \forall i \in I, \forall v' \in V', \forall c \in C, \forall t \in T$$

$$(c2) \quad X_{i,v',t} + X_{i,r,t} + X_{i,c,t} \leq 1, \quad \forall i \in I, \forall v' \in V', \forall c \in C, \forall t \in T$$

$$(c3) \quad \psi_{i,v,t} \in \{0,1\}, \quad \forall i \in I, \forall v \in V, \forall t \in T$$

$$(c4) \quad \sum_{v' \in V'} X_{i,v',t} tf_{i,v'} + X_{i,r,t} tf_{i,r} + \sum_{c \in C} X_{i,c,t} tf_{i,c} \leq \tau_{i,v}, \quad \forall i \in I, \forall v \in V, \forall v' \in V', \forall c \in C, \forall t \in T$$

$$(c5) \quad \sum_{v' \in V'} X_{i,v',t} C_{i,v'} + X_{i,r,t} C_{i,r} + \sum_{c \in C} X_{i,c,t} C_{i,c} \leq b_{i,v}, \quad \forall i \in I, \forall v' \in V', \forall c \in C, \forall t \in T$$

$$(c6) \quad B_{i,v} \leq \sum_{v' \in V'} X_{i,v',t} B_{i,v'} + X_{i,r,t} B_{i,r}, \quad \forall i \in I, \forall v \in V, \forall v' \in V', \forall t \in T$$

$$(c7) \quad \sum_{v' \in V'} X_{i,v',t} E_{v'} \Delta_{v'} + X_{i,r,t} E_r \Delta_r \leq \sum_{v' \in V'} X_{i,v',t} E_{v'}^{total} + X_{i,r,t} E_r^{total},$$

$$\forall i \in I, \forall v \in V, \forall v' \in V', \forall t \in T$$

$$(c8) \quad \sum_{v' \in V'} X_{i,v',t} \Delta_{v'} + X_{i,r,t} \Delta_r + \sum_{c \in C} X_{i,c,t} \Delta_{i,c} \leq \sum_{v' \in V'} X_{i,v',t} \mu_{v'} + X_{i,r,t} \mu_r + \sum_{c \in C} X_{i,c,t} \mu_{i,c},$$

$$\forall i \in I, \forall v' \in V', \forall c \in C, \forall t \in T$$

$$(c9) \quad \sum_{v' \in V'} X_{i,v',t} = 0, \quad \forall i \in I, \forall v' \in V', \forall t \in T$$

$$(c10) \quad \sum_{c \in C} X_{i,c,t} = 0, \quad \forall i \in I, \forall c \in C, \forall t \in T$$

محدودیت (c1)، نشان می دهد که سه نوع اجرای وظیفه در نظر گرفته می شوند. محدودیت (c2)، بیان می کند که تنها یک نوع اجرای وظیفه باید برای هر وظیفه در نظر گرفته شود. محدودیت (c3)، بیان می کند که متغیر تصمیم $\psi_{i,v,t}$ مقدار یک را می گیرد وقتی که وظیفه i در خودروی v به طور موفق آمیز اجرا شود و در غیر صورت مقدارش صفر می گردد. محدودیت (c4)، توضیح می دهد که زمان اجرای وظیفه باید کمتر یا مساوی از مقدار تعیین شده باشد. محدودیت (c5)، بیان می کند که هزینه انجام وظیفه باید کمتر یا مساوی از بودجه تعریف شده برای وظیفه باشد. محدودیت (c6)، نشان می دهد که پهنای باند مورد استفاده برای ارسال وظیفه به گره مه باید کمتر یا مساوی از پهنای باند تخصیص داده شده به گره مه باشد. محدودیت (c7)، بیان می کند که انرژی مصرف شده برای اجرای وظیفه باید کمتر یا مساوی از انرژی گره مه باشد. محدودیت (c8)، توضیح می دهد که منابع محاسباتی مورد نیاز وظایف باید کمتر یا مساوی ظرفیت محاسباتی گره مه یا سرور ابر باشد. محدودیت (c9) و محدودیت (c10) به ترتیب بیان می کنند که یک و فقط یک خودروی مه و یا سرور ابر به وظیفه تخصیص داده می شود. بر اساس جدول ۲، نمادگذاری های متداول در مدل سازی ریاضی را خلاصه می کند.

جدول ۲: نمادهای متداول استفاده شده در مدل سازی

Table 2. Common symbols used in modeling

نماد	تعریف
$X_{i,v',t}$	متغیر تصمیم گیری باینری برای انتخاب خودروی مه v' در برش زمانی t
$X_{i,r,t}$	متغیر تصمیم گیری باینری برای انتخاب واحد کنار جاده ای r در برش زمانی t
$X_{i,c,t}$	متغیر تصمیم گیری باینری برای انتخاب سرور ابر c در برش زمانی t
$\psi_{i,v,t}$	متغیر تصمیم گیری باینری برای پذیرش وظیفه i در خودروی v

$tf_{i,v'}$	زمان برون‌سپاری وظیفه i به خودروی مه v' (ثانیه)
$tf_{i,r}$	زمان برون‌سپاری وظیفه i به واحد کنار جاده ای r (ثانیه)
$tf_{i,c}$	زمان برون‌سپاری وظیفه i به سرور ابر c (ثانیه)
$\tau_{i,v}$	مهلت زمانی وظیفه i در خودروی v (ثانیه)
$C_{i,v'}$	هزینه برون‌سپاری وظیفه i به خودروی مه v' (واحد هزینه)
$C_{i,r}$	هزینه برون‌سپاری وظیفه i به واحد کنار جاده ای r (واحد هزینه)
$C_{i,c}$	هزینه برون‌سپاری وظیفه i به سرور ابر c (واحد هزینه)
$b_{i,v}$	بودجه وظیفه i در خودروی v (واحد هزینه)
$B_{i,v}$	پهنای باند مورد نیاز وظیفه i در خودروی v (مگا هرتز)
$B_{v'}$	پهنای باند تخصیص داده شده به خودروی مه v' (مگا هرتز)
B_r	پهنای باند تخصیص داده شده به واحد کنار جاده ای r (مگا هرتز)
$E_{v'}$	انرژی مورد نیاز به ازای هر MIPS برای خودروی مه v' (ژول بر ثانیه)
E_r	انرژی مورد نیاز به ازای هر MIPS برای واحد کنار جاده ای r (ژول بر ثانیه)
$\Delta_{v'}$	نرخ رسیدن کل دستورات به خودروی مه v' (MIPS)
Δ_r	نرخ رسیدن کل دستورات به واحد کنار جاده ای r (MIPS)
Δ_c	نرخ رسیدن کل دستورات به سرور ابر c (MIPS)
$E_{v'}^{total}$	انرژی تخصیص داده شده به خودروی مه v' (ژول)
E_r^{total}	انرژی تخصیص داده شده به واحد کنار جاده ای r (ژول)
$\mu_{v'}$	ظرفیت محاسباتی خودروی مه v' (MIPS)
μ_r	ظرفیت محاسباتی واحد کنار جاده ای r (MIPS)
μ_c	ظرفیت محاسباتی سرور ابر c (MIPS)

۳- فرمول‌بندی یادگیری و راه‌حل پیشنهادی

مسئله بهینه‌سازی مطرح شده، یک مسئله بهینه‌سازی عدد صحیح مختلط^۱ است و چنین مسائلی معمولاً، مسائل NP سخت^۲ در نظر گرفته می‌شوند [۲۶]. بنابراین، ما از یادگیری تقویتی برای حل مسئله بهینه‌سازی فرموله شده استفاده می‌کنیم. در این بخش، ابتدا پیاده‌سازی یادگیری تقویتی را شرح می‌دهیم و سپس راه‌حل یادگیری تقویتی عمیق را بررسی می‌کنیم.

۳-۱- پیاده‌سازی یادگیری تقویتی

ما از یک مدل فرآیند تصمیم‌گیری مارکوف متناهی (MDP) برای ثبت دینامیک انتقال حالت شبکه استفاده می‌کنیم. یک MDP را می‌توان با مجموعه‌ای از حالت‌های سیستم، مجموعه‌ای از اقدامات و مجموعه‌ای از توابع پاداش با مقدار حقیقی تعریف کرد [۲۷]. عامل‌های خودرویی با محیط خودرویی تعامل می‌کنند تا تصمیمات برون‌سپاری خود را بر اساس تجربیات قبلی بهبود بخشند. ما ارتباط بین کنترلرهای محلی و سرورهای مه و ابر را به عنوان یک فرآیند MDP به صورت $(S, A, R, \Pi_{SS'})$ فرموله می‌کنیم، که در آن S ، نشان دهنده فضای حالت، A ، نشان دهنده فضای اقدام، R ، نشان دهنده فضای پاداش و $\Pi_{SS'}$ ، نشان دهنده احتمال انتقال است. نمادگذاری‌ها و توابع مختلف مورد استفاده به شرح زیر است:

- **فضای حالت:** یک فضای حالت شامل تمام پیکربندی‌های شبکه در یک سیستم است. مجموعه وظایف در برش زمانی t به صورت $\Phi^t = \{\Phi_1^t, \Phi_2^t, \dots, \Phi_{I'}^t\}$ نمادگذاری می‌شوند. مشخصات وظیفه i ، در برش زمانی t ، را به صورت $\Phi_i^t = \{x_{i,v}^t, y_{i,v}^t, \lambda_{i,v}^t, d_{i,v}^t, \tau_{i,v}^t, L_{i,v}^t, b_{i,v}^t, B_{i,v}^t, K_{i,v}^t\}$ نمادگذاری می‌کنیم. گره‌های مه و سرور ابری تنظیم شده در بازه زمانی t به صورت $F^t = \{F_1^t, F_2^t, \dots, F_{V'+1}^t\}$ و $C^t = \{C_1^t, C_2^t, \dots, C_C^t\}$ نشان داده می‌شوند. مشخصات گره مه با اندیس j به صورت $F_j^t = \{x_j^t, y_j^t, \mu_j^t, E_j^t, C_j^t, B_j^t\}$ تعریف می‌شود که به ترتیب x_j^t و y_j^t ، مختصات گره مه، μ_j^t ، ظرفیت محاسباتی گره مه، E_j^t ، مقدار انرژی گره مه، C_j^t ، هزینه پردازش وظیفه و B_j^t ، پهنای باند تخصیص داده شده به گره مه می‌باشند. همچنین مشخصات سرور ابر با اندیس c به صورت $C_c^t = \{\mu_c, E_c, C_c\}$ نمادگذاری شود که به ترتیب μ_c ، مشخص کننده ظرفیت

¹ Mixed-integer linear optimization problem

² NP-hard problems

محاسباتی ماشین مجازی، E_c ، مقدار انرژی ماشین مجازی و C_c ، هزینه استفاده از ماشین مجازی است. بنابراین، بردار حالت سیستم در بازه زمانی t را می توان به صورت زیر تعریف کرد:

$$S^{(t)} = [\phi^t, F^t, C^t] \quad (36)$$

- **فضای عمل:** در رویکرد پیشنهادی، یک عامل در هر لحظه زمانی t ، استراتژی برون سپاری وظیفه را تعیین می کند. بنابراین، یک عمل در لحظه زمانی t به صورت زیر تعریف می شود:

$$a^{(t)} = [X_{i,1,t}, X_{i,2,t}, \dots, X_{i,v',t}, X_{i,v'+1,t}, \dots, X_{i,v'+1+c,t}] \quad (37)$$

- **تابع پاداش:** تابع پاداش بر اساس نرخ پذیرش وظایف، مصرف انرژی، تاخیر و هزینه است. به طور خاص، پاداش $r^{(t)}$ برای برنامه i که از سیاست $\pi(a^{(t)}|r^{(t)})$ پیروی می کند، می تواند به صورت معادله ۳۸ تا ۴۰ بیان شود:

$$r^{(t)} = Acc_v^t + penalty_v^t \quad (38)$$

$$penalty_v^t = \sum_{i \in I(t)} penalty_{i,v,t} \quad (39)$$

$$penalty_{i,v,t} = \begin{cases} 0 & \text{if } \sum_{v' \in V'} X_{i,v',t} tf_{i,v'} + X_{i,r,t} tf_{i,r} + \sum_{c \in C} X_{i,c,t} tf_{i,c} \leq \tau_{i,v} \\ -\omega & \text{if } \sum_{v' \in V'} X_{i,v',t} tf_{i,v'} + X_{i,r,t} tf_{i,r} + \sum_{c \in C} X_{i,c,t} tf_{i,c} > \tau_{i,v} \text{ and } \kappa_{i,v} = 1 \end{cases} \quad (40)$$

در اینجا $-\omega$ ، یک عدد صحیح منفی است. در واقع در روش پیشنهادی سعی شده که وظایف با اولویت بالا بیشتر پذیرش شوند.

۳-۲-مقدمات

با توجه به دانش محدود در مورد احتمال انتقال بین حالت ها و فضای قابل توجه حالت-عمل در شبکه، برنامه نویسی پویا^۱ سنتی نمی تواند سیاست بهینه را به طور موثر پیدا کند. بنابراین، ما از یادگیری تقویتی عمیق (DRL) را برای حل مسئله انتخاب سرور پیشنهادی اتخاذ می کنیم [۲۸]. ما یادگیری Q عمیق را به دلایل مختلف به عنوان یک رویکرد مبتنی بر ارزش اتخاذ می کنیم. اولاً، سیستم زمان گسسته ای دارد و در هر بازه زمانی در نظر گرفته می شود. دوماً، تمام منابع، مانند قابلیت های محاسباتی، بر اساس مقادیر گسسته مدل سازی شده و به سطوح گسسته تقسیم می شوند. سوماً، راه حل یادگیری ما باید فقط یک سرور را برای هر بازه زمانی به عنوان نتیجه یادگیری عمیق انتخاب کند. چهارماً، یکی از اهداف استفاده از یک شبکه Q عمیق^۲، به حداکثر رساندن پاداش در مرحله ها است. رویکرد یادگیری Q می تواند برای یافتن یک سیاست تصمیم گیری بهینه مورد استفاده قرار گیرد. در رویکرد یادگیری Q ، یک تابع Q تعریف می شود که عامل را به سمت راه حل بهینه هدایت می کند. تابع Q به صورت بازگشتی با استفاده از معادله بلمن^۴ محاسبه می شود:

$$Q(s^t, a^t) = Q(s^t, a^t) + \alpha [r^t + \gamma \max_{a^{t+1}} Q(s^{t+1}, a^{t+1}) - Q(s^t, a^t)] \quad (41)$$

هر عامل می تواند یک جدول Q را نگهداری کند که توسط جدول $Q(s^t, a^t)$ داده شده است، که در آن s^t وضعیت فعلی در بازه زمانی t و a^t عملی است که توسط عامل در بازه زمانی t انجام می شود. در فرمول بالا، α نرخ یادگیری، γ ضریب تخفیف، s^{t+1} حالت بعدی و a^{t+1} عملی در حالت s^{t+1} است که مقدار Q را به حداکثر می رساند. یادگیری Q یک راه حل امیدوارکننده است که می تواند به طور موثر به سیاست بهینه برسد، زمانی که فضاهای حالت و عمل کوچک هستند. با این حال، در اکثر محیط های دنیای واقعی، فضاهای حالت و عمل بسیار بزرگ هستند و بنابراین نگهداری جدول Q به فضای ذخیره سازی عظیمی

¹ Dynamic Programming

² Q-learning

³ Q-network

⁴ Bellman equation

نیاز دارد. رویکرد یادگیری Q در الگوریتم ۱ شرح داده شده است. در ابتدا، مقادیر Q برای تمام لحظات زمانی ۰ در نظر گرفته می‌شوند. یک عامل اطلاعات وضعیت را از واحدهای کنار جاده‌ای دریافت می‌کند و با استفاده از سیاست ϵ -حریصانه^۱ یک عمل را اجرا می‌کند. در سیاست ϵ -حریصانه، عامل یک اقدام بهینه با احتمال $1 - \epsilon$ و یک اقدام تصادفی را با احتمال ϵ انجام می‌دهد. جدول مقدار Q برای T برش زمانی در هر مرحله به‌روزرسانی می‌شود. پاداش، به عنوان میانگین پاداش‌های فوری در داخل یک مرحله محاسبه می‌شود. بنابراین، رویکرد یادگیری Q عمدتاً برای MDPهایی که فضاهای حالت یا عمل کوچکی دارند، استفاده می‌شود.

Algorithm 1: Q-Learning Approach

Input: Set of all tasks and set of all possible associations between controllers and predefined servers.

Output: Efficient offloading decisions

- 1: Initialize $Q(s^t, a^t)$ to 0 for all time instants ;
 - 2: **for** each episode **do**
 - 3: **for** $t=1$ to T **do**
 - 4: At time instant t , an agent acquires the state s^t ;
 - 5: An agent chooses an action $a^t = \arg\max_a Q(s^t, a^t)$ with probability $1-\epsilon$ or randomly choose action from action space with probability ϵ ;
 - 6: An agent updates $Q(s^t, a^t)$ from Eq. 41 ;
 - 7: **end for**
 - 8: **end for**
 - 9: return $Q(S, A)$;
-

از این رو، از یک رویکرد شبکه یادگیری Q عمیق (DQN)^۲ برای بدست آوردن راه حل بهینه استفاده می‌شود. در DQN، یک تابع مقدار Q، با استفاده از یک شبکه عصبی عمیق (DNN)^۳ با مجموعه پارامتر θ تخمین زده می‌شود که حالت ورودی را به یک عمل نگاشت می‌دهد. جدا از شبکه عصبی اصلی، یک شبکه عصبی هدف برای تثبیت شبکه وجود دارد. پس از تعداد ثابتی از تکرارها، پارامترهای شبکه عصبی هدف از شبکه عصبی اصلی کپی می‌شوند. پارامترهای شبکه عصبی اصلی به طور مکرر برای به حداقل رساندن تابع اتلاف به‌روزرسانی می‌شوند. تابع اتلاف DQN به صورت زیر تعریف می‌شود:

$$L^t(\theta^t) = (y^t - Q(s^t, a^t | \theta^t))^2 \quad (42)$$

که در آن y^t مقدار هدف شبکه هدف برای بازه زمانی t است، و $Q(s^t, a^t | \theta^t)$ مقدار تخمینی شبکه اصلی با پارامتر وزنی θ^t است. مقدار هدف y^t به صورت معادله ۴۳ تعریف می‌شود:

$$y^t = r^t + \gamma \max_{a'} Q(s^t, a' | \theta^t) \quad (43)$$

که در آن θ^t و θ'^t به ترتیب پارامترهای وزنی شبکه‌های اصلی و هدف هستند. یک رویکرد DQN در الگوریتم ۲ شرح داده شده است.

Algorithm 2: DQN Approach

Input: Set of all tasks and set of all possible associations between controllers and predefined servers.

Output: Efficient offloading decisions

- 1: Initialize Replay Buffer, Paramter set (θ and θ') of main and target networks, Update interval U ;
 - 2: **for** each episode **do**
 - 3: **for** $t=1$ to T **do**
 - 4: At time instant t , an DQN agent acquires the state s^t ;
 - 5: An agent chooses an action based on the ϵ -greedy policy;
-

¹ ϵ -greedy

² Deep Q-learning Network (DQN)

³ Deep Neural Network (DNN)

```

6:   Execute action  $a_i^{(t)}$ , and observe the reward  $r_i^{(t)}$  and the next state  $S^{t+1}$ 
7:   Process  $S_i^{(t+1)}$  to be the next state;
8:   Store transition  $(S_i^{(t)}, a_i^{(t)}, r_i^{(t)}, S_i^{(t+1)})$  into replay memory;
9:   Sample random mini-batch of transitions  $(S_i^{(t)}, a_i^{(t)}, r_i^{(t)}, S_i^{(t+1)})$  from replay memory;
10:  Compute  $y^t$  from Eq. 43;
11:  The DQN agent uses gradient descent to minimize the loss function as shown in Eq. 42 ;
12:  After every U steps, update the parameters of target network ;
7:  end for
8:  end for
9:  return  $Q(S,A; \theta)$ ;

```

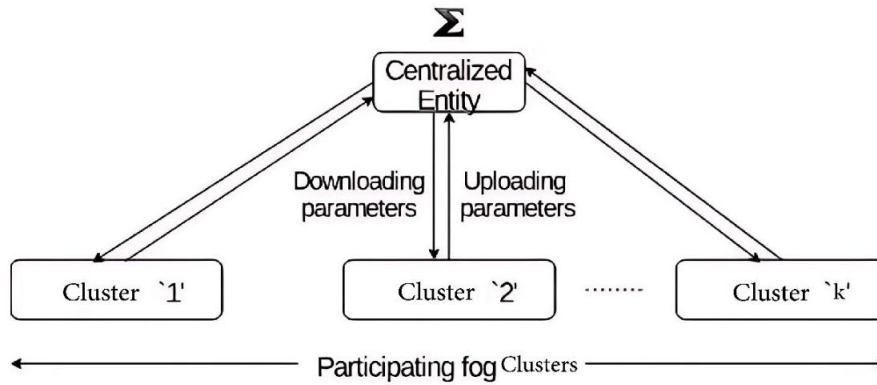
در مرحله اول، پارامترهای DQN شبکه‌ی اصلی و هدف مقداردهی اولیه می شوند. هر قسمت دارای T مرحله است. در هر مرحله، عامل DQN وضعیت فعلی s^t را به دست می آورد. پس از به دست آوردن حالت، یک عامل DQN، براساس سیاست ϵ -حریصانه عملی را انتخاب می کند. چنگانه^۱ انتقال در یک بافر پخش مجدد^۲ ذخیره می شود. سپس از روش گرادیان نزولی برای بروز سانی وزن ها و کم کردن تابع اتلاف استفاده می شود. پس از تعداد ثابتی از مراحل U، پارامترهای θ شبکه هدف از طریق شبکه اصلی بروزرسانی می شود.

۳-۳- الگوریتم پیشنهادی مبتنی بر یادگیری فدرال

در این کار، ما یک چارچوب برون سپاری وظیفه فدرال مبتنی بر DRL چند عاملی ارائه می دهیم. راه حل پیشنهادی یک معماری چند لایه را ارائه می دهد که در آن مدل های محلی و سراسری به طور متقابل دانش را به اشتراک می گذارند تا تصمیمات برون سپاری وظیفه آگاهانه ای را بگیرند. مدل ها به صورت محلی یاد گرفته می گیرند و در کنترلرهای محلی نگهداری می شوند. یادگیری فدرال با آموزش محلی در کنترلرهای محلی به کاهش تأخیرهای ارتباطی کمک می کند. همچنین حریم خصوصی خودروها و داده ها را بهبود می بخشد، زیرا مقدار کمی از داده های حیاتی به اشتراک گذاشته می شود. همچنین احتمال تغییر داده ها کمتر است زیرا فقط گرادیان ها در شبکه ها به اشتراک گذاشته می شوند. گردش کار یادگیری فدرال در شکل ۳ نشان داده شده است که شامل دو لایه است: لایه های پایین و بالا. لایه بالایی شامل کنترلر سراسری و مبتنی بر فدراسیون کنترلرهای محلی هستند، در حالی که لایه پایینی شامل کنترلرهای محلی هستند. عامل ها در لایه پایینی، حالت را به عنوان ورودی خود در نظر می گیرند و استراتژی بهینه را با استفاده از تغییر وزن ها (پارامترها) در شبکه اصلی یاد می گیرند. سپس بروزرسانی ها به عامل در لایه بالایی به اشتراک گذاشته می شوند که بروزرسانی های پارامترهای همه عامل ها در محدوده انتقال را جمع می کند. این اشتراک گذاری متقابل بروزرسانی ها بین لایه پایینی و بالایی، مدل تصمیم گیری برون سپاری وظیفه را بهبود می بخشد. شبکه سراسری، وزن های مدل محلی را جمع کرده و یک مدل سراسری بروز شده را با عامل های لایه پایینی به اشتراک می گذارد.

¹ Tuple

² Replay buffer



شکل ۳: فرآیند یادگیری فدرال برای k خوشه و ارتباط بین کنترلرهای محلی و سراسری را نشان می‌دهد.

Figure 3. Shows the federated learning process for k clusters and the connection between local and global controllers.

مسئله بهینه‌سازی در لایه پایین‌تر حل می‌شود که در آن کنترلرهای محلی مدل بروزرسانی وظیفه را به صورت توزیع‌شده یاد می‌گیرند. هر کنترلر محلی با محیط تعامل دارد، به صورت بر خط تجربه جمع‌آوری می‌کند و تصمیمات مربوط به برون‌سپاری خود را بهینه می‌کند. لایه بالایی مسئول تجمیع پارامترهای دریافتی از لایه پایینی و بازگرداندن بروزرسانی‌ها به لایه پایینی است. لایه بالایی با جمع‌آوری بروزرسانی‌ها برای همه کنترلرها در یک صحنه شبیه‌سازی، به عامل‌ها کمک می‌کند تا تجربه آموزشی متنوع‌تری داشته باشند. رویکرد فدرال توسط الگوریتم ۳ بیان شده است.

Algorithm 3: DQN-based Federated Learning Approach

- 1: **Input:** The Initial set of local agents K ;
- 2: **Output:** Reward $r^{(t)}$ and offloading action $a^{(t)}$;
- 3: **Repeat**
- 4: each local controller agent in cluster $k \in K$ updates its parameters set θ_k^t using Algorithm 2;
- 5: each local controller agent in cluster $k \in K$ sends its parameters set θ_k^t to the centralized controller agent;
- 6: The centralized controller agent aggregates the local parameters
- 7: sets of k number of fog clusters as $\theta_k^{t+1} = \sum_{k \in K} \frac{D_k}{D} * \theta_k^t$
- 8: $t \leftarrow t + 1$;
- 9: **until** $t > T$

در ابتدا، تعداد عامل‌های اولیه مقدار دهی می‌شود. یادگیری فدرال شامل تعداد T دور است. در هر دور، کنترلرهای محلی، مجموعه پارامترهای خود را بروزرسانی کرده و برای تجمیع به کنترلر متمرکز ارسال می‌کنند، همانطور که در مراحل ۴ و ۵ نشان داده شده است. در مرحله ۷، کنترلر متمرکز، مجموعه پارامترهای سراسری θ_k^{t+1} را در هر دور t با تجمیع مجموعه پارامترهای محلی θ_k^t هر کنترلر محلی شرکت‌کننده $k \in K$ به شرح زیر پیدا می‌کند [۲۹]:

$$\theta_k^{t+1} = \sum_{k \in K} \frac{D_k}{D} * \theta_k^t \quad (44)$$

در اینجا، D_k تعداد وظایف در k امین خوشه است و D اندازه کل وظایف است که به صورت $D = \sum_{k \in K} D_k$ محاسبه می‌شود.

۴- نتایج شبیه‌سازی

این بخش با اتخاذ یک مجموعه داده دنیای واقعی، تنظیمات تجربی را ارائه می‌دهد و کارهای مورد مقایسه را معرفی می‌کند. سپس، نتایج تجربی و بحث‌ها را نشان می‌دهد.

۴-۱- تنظیمات شبیه‌سازی و کارهای مورد مقایسه

از جریان ترافیک واقعی خودروها برای ارزیابی عملکرد رویکرد پیشنهادی مدیریت منابع تصمیم‌گیری سرورهای مه در یک شبکه خودرویی استفاده می‌شود. بنابراین، از یک مجموعه داده واقعی خودروها با استفاده از ردیابی‌های حرکتی تاکسی‌ها در

شهر رم^۱، ایتالیا استفاده شد [۳۰]. این مجموعه داده، شامل مختصات GPS تقریباً ۳۲۰ تاکسی است که طی ۳۰ روز، از ۱ فوریه تا ۲ مارس ۲۰۱۴، جمع آوری شده اند که شامل موقعیت مکانی رانندگان است که هر ۷ ثانیه جمع آوری شده اند. مکان تاکسی ها از عرض جغرافیایی ۳۹/۳۶ تا عرض جغرافیایی ۴۵/۵۱ و از طول جغرافیایی ۱۲/۰۰ تا طول جغرافیایی ۱۲/۸۹ متغیر است. ما به طور تصادفی با استفاده از مجموعه داده ها، مسیر هر تاکسی در طول یک روز در نظر می گیریم. از آنجایی که سرورهای RSU در جاده نزدیک به خودرو قرار دارند، ما فاصله حداکثر و حداقل محدوده مکان خودروها را بر تعداد سرورهای RSU تقسیم می کنیم تا مکان هر RSU را پیدا کنیم. علاوه بر این، برای اختصاص یک موقعیت جغرافیایی برای سرور مرکزی، یک مقدار بالا به حداکثر مقدار عرض و طول جغرافیایی موقعیت تاکسی ها اضافه می کنیم. همچنین، برای ارائه بهینه الگوریتم و مقایسه آن، از سه مدل به عنوان مبنا استفاده می کنیم.

روش پایه اول، روش حریصانه^۲ برای تخصیص هر وظیفه به یک سرور است. در راه حل حریصانه، وظایف به سرورهایی با قابلیت های منابع بالاتر اختصاص داده می شوند. روش پایه دوم، یادگیری تقویتی جدولی است، که RL^۳ نامیده می شود و وظایف را به یک سرور مناسب اختصاص می دهد. در روش RL، یک جدول بزرگ، برای حفظ تمام شرایط حالت پیاده سازی می شود. روش پایه سوم، سیاست برون سپاری مبتنی بر SDN ارائه شده در [۸] است. این روش، یک سیاست بهینه سعی می کند که با دانش سراسری از شبکه خودروری مه توسط SDN سراسری، تفاوت زمانی بین زمان پاسخ سرور مقصد و مهلت درخواست را به حداکثر برساند.

ما از کتابخانه PyTorch [۳۱]، برای پیاده سازی یادگیری تقویتی عمیق پیشنهادی خود استفاده می کنیم. علاوه بر این، ما از چهار لایه از یک شبکه عصبی کاملاً متصل با لایه های پنهان با اندازه ۱۲۸ نورون استفاده کرده ایم. همچنین از تابع ReLU [۳۲] که یک تابع فعال سازی است به لایه های مرتبط اضافه شده است که ویژگی های غیرخطی را به شبکه عصبی معرفی می کند و عملکرد شبکه عصبی عمیق را بهبود می بخشد. برخی از پارامترهای روش مدیریت منابع پیشنهادی و فرآیند یادگیری در جدول ۳ آمده است.

جدول ۳: پارامترهای شبیه سازی

Table 3. Simulation parameters

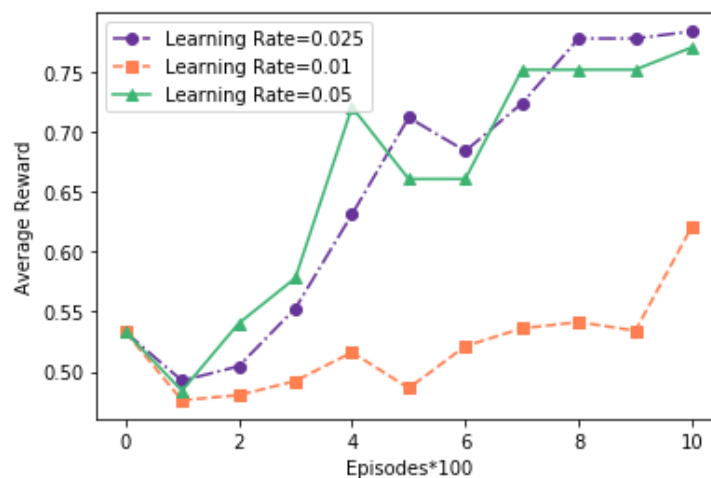
مقدارها	پارامترها
۵	خوشه ها
۱۲۸	تعداد نورون ها
Relu	تابع فعال سازی
۶۰۰۰	اندازه بافر
Adam	بهینه ساز
۱۲۸	Mini-batch
۰/۰۲۵	نرخ یادگیری
۵۰۰-۴۰۰۰ (MIPS)	نرخ ورود دستورالعمل های وظایف
-۱۲۰dB	قدرت نویز
۱-۲ (GHz)	پهنای باند خودرو
۱۰۰-۲۰۰ (MHz)	پهنای باند RSU
۱-۱۰۰ (MHz)	پهنای باند وظیفه
۱۰ (GB/s)	نرخ ارسال داده ها از RSU به ابر
۲۰۰۰-۳۰۰۰/۵۰۰۰/۱۰۰۰۰ (MIPS)	ظرفیت محاسباتی خودرو/RSU/ابر
۰/۵ watts	توان انتقال خودرو
۰/۰۰۰۸-۰/۰۰۰۱۲ / ۰/۰۰۰۴/ ۰/۰۰۰۱ (joule/s)	انرژی پردازش خودرو/RSU/ابر

¹ Rome² Greedy³ Reinforcement Learning(RL)

هزینه پردازش خودرو/RSU/بر	۰/۰۰۲۲و۰/۰۰۳۳ / ۰/۰۰۶۱/ ۰/۰۰۸۲
(unit cost/s)	
هزینه ارتباط بین خودرو و خودرو/RSU/بر	۰/۰۰۴/۰/۰۰۳۳/ ۰/۰۰۱ (unit costs)
تعداد خودروها	۴۰۰-۵۰۰
تعداد گره‌های مه	۳۵

۲-۴- نتایج تجربی

با توجه به اینکه نرخ یادگیری یک پارامتر اساسی در یادگیری عمیق است که ممکن است بر عملکرد مدل تأثیر بگذارد، این مقاله نرخ یادگیری را روی ۰/۰۵، ۰/۰۱ و ۰/۰۲۵ تنظیم می‌کند. این مقاله بررسی می‌کند که چگونه آنها می‌توانند عملکرد الگوریتم پیشنهادی را افزایش دهند و شکل ۴ یافته‌ها را نشان می‌دهد. به خاطر نتایج زیاد، برای نشان دادن بهتر نتایج، هر ۱۰۰ قسمت^۱، میانگین گرفته می‌شود. همانطور که شکل نشان می‌دهد، الگوریتم پیشنهادی با سرعت همگرایی سریع‌تر و مقدار همگرایی بهتر در طول مرحله آموزش، زمانی که نرخ یادگیری روی ۰/۰۲۵ تنظیم شده است، عملکرد بهتری دارد. مقادیر همگرایی ضعیف مدل زمانی که نرخ یادگیری روی ۰/۰۵ و ۰/۰۱ تنظیم شده است، نشان داده شده است. نتایج بیان می‌کند که افزایش نرخ یادگیری می‌تواند منجر به همگرایی سریع‌تر شود. با این حال، افزایش بیشتر نرخ یادگیری هیچ بهبود قابل توجهی در مقادیر نشان نمی‌دهد که احتمال افتادن در راه‌حل بهینه محلی و شکست برای رسیدن به راه‌حل بهینه را افزایش می‌دهد. اما نرخ یادگیری پایین برای همگرایی زمان زیادی خواهد برد. نتیجه نشان می‌دهد که نرخ یادگیری بالاتر ممکن است منجر به یک مرحله به‌روزرسانی وزن کوچک شود که زمان آموزش را افزایش می‌دهد و برای دستیابی به حالت همگرایی، تکرارهای بیشتری را ضروری می‌کند. بنابراین، در تمام آزمایش‌ها، مقدار نرخ یادگیری ۰/۰۲۵ است.



شکل ۴: مقایسه نرخ‌های یادگیری
Figure 4. Comparison of learning rates

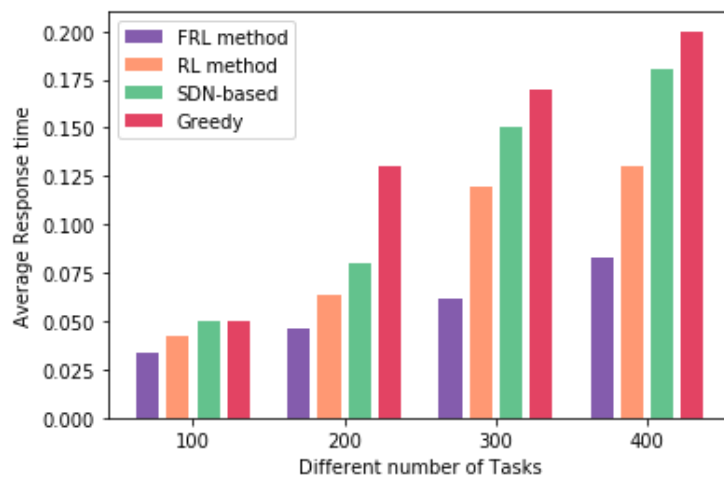
نتایج آزمایش در شکل ۵ نشان داده شده است، که در آن محور طولی نشان‌دهنده تعداد وظایف و محور عرضی نشان‌دهنده مقدار میانگین زمان پاسخدهی (زمان برون سپاری وظیفه) است. زمان پاسخدهی شامل زمان انتقال وظیفه، زمان انتظار در صف، و زمان پردازش وظیفه است. همانطور که در شکل ۵ نشان داده شده است، زمان پاسخدهی در طرح پیشنهادی ما برای تعداد وظایف مختلف کمترین مقدار را دارد. تحلیل نتایج نشان می‌دهد که FRL^۲ یک سکوی^۳ مقیاس پذیر را برای یادگیری سیاست بهینه فراهم می‌کند. در بهینه‌سازی فدرال، چندین کنترلر توزیع شده برای بهینه‌سازی یک تابع هدف مشترک همکاری می‌کنند،

¹ Episode

² Federated Reinforcement Learning (FRL)

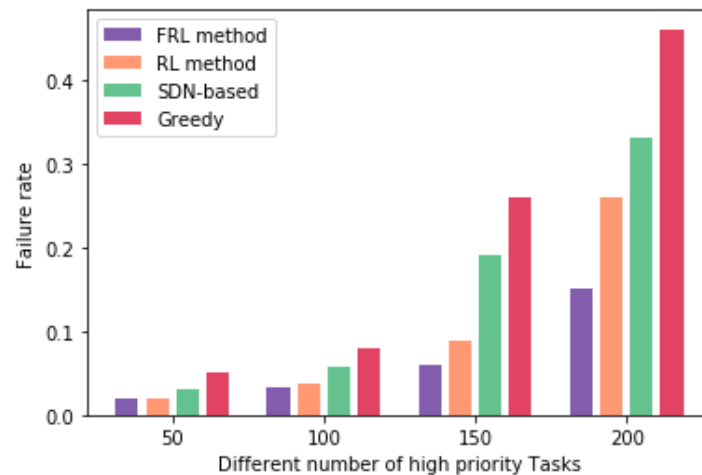
³ Platform

در حالی که عامل ها یاد می گیرند که در یک محیط غیرمتمرکز و بدون هماهنگی مرکزی عمل کنند. FRL با آموزش الگوریتم های مشارکتی بر روی داده های محلی، به اشتراک گذاری به روزرسانی های مدل سراسری به جای داده های خام، کار می کند و در نتیجه به ایجاد مدل با ایجاد وزن های بهینه کمک می کند. همچنین، از شکل ۵ می توانیم دریابیم که راه حل های RL و FRL در حال یادگیری نحوه تخصیص وظایف به سرورهای مناسب بر اساس منابع موجود هستند، در حالی که تقاضاهای تأخیر وظیفه را برآورده می کنند. روش RL در محیط های بزرگ خوب کار نمی کند و استراتژی پیشنهادی می تواند یک راه حل خوب باشد. یک جدول بزرگ در روش RL پیاده سازی شده است تا تمام شرایط حالت را حفظ کند و از کمبود مقیاس پذیری رنج می برد. رویکرد پیشنهادی از خودروهای پارک شده و در حال حرکت جهت انجام وظایف استفاده می کند بنابراین می تواند میانگین زمان پاسخدهی را کاهش دهد. روش SDN-based ارائه شده، در درون خوشه تنها از منابع RSU ها استفاده می کند. از آنجا که منابع RSU ها محدود هستند، بنابراین وظیفه نمی تواند به صورت محلی اجرا شود و در نتیجه باید براساس تصمیم کنترلر سراسری و از طریق سوییچ ها به سرور مقصد برون سپاری گردد. در راه حل حریصانه، هر وظیفه به سروری با ظرفیت منابع بالا اختصاص داده می شود و از هیچ سیاست یادگیری پیروی نمی شود که در نتیجه زمان پاسخدهی در روش حریصانه برای استراتژی بهینه در نظر گرفته نمی شود.



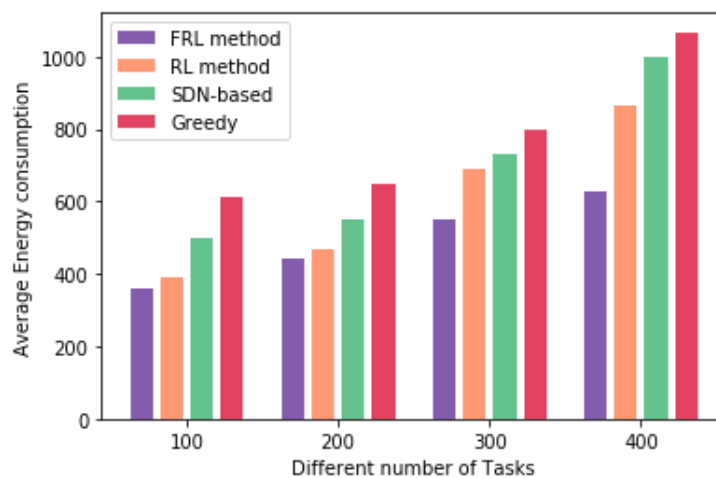
شکل ۵: میانگین زمان پاسخدهی برای چهار الگوریتم با تعداد وظایف مختلف
Figure 5. Average response time for four algorithms with different numbers of tasks

همانطور که قبلاً بیان شد، وظایف با اولویت بالا دارای مهلت زمانی از پیش تعریف شده اند و باید مهلت زمانی رعایت شود. روش حریصانه، رویکرد مناسبی برای انتخاب سیاست بهینه برای برنامه های بلادرنگ نیست. از آنجا که الگوریتم حریصانه معمولاً مانند برنامه نویسی پویا به طور کامل روی همه داده ها عمل نمی کند، معمولاً (اما نه همیشه) در کشف یک راه حل بهینه شکست می خورد. از طرف دیگر، تکنیک های یادگیری تقویتی می توانند با استفاده از قابلیت های شبکه های عصبی، سرورهای مقصد را برای برون سپاری وظایف به طور مؤثرتری انتخاب کنند. روش SDN-based، اولویت وظایف را در نظر نمی گیرد و در نتیجه نمی تواند اجرای وظایف بلادرنگ را تضمین نماید.



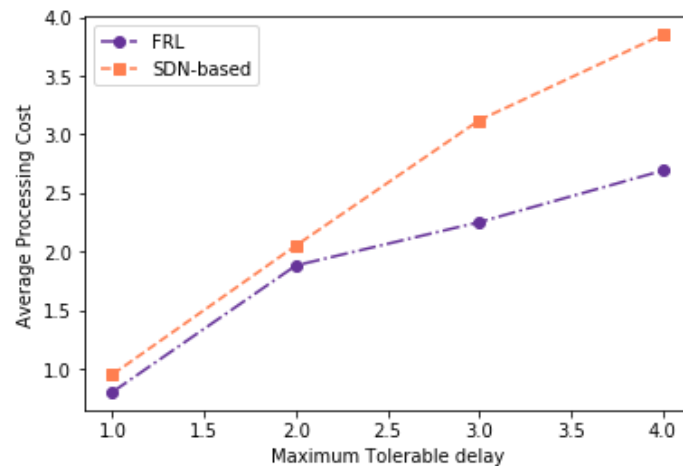
شکل ۶: نرخ شکست وظیفه برای چهار الگوریتم با تعداد وظایف اولویت بالای مختلف
Figure 6. Task failure rates for four algorithms with different numbers of high priority tasks

شکل ۷، میانگین مصرف انرژی برای چهار الگوریتم مختلف را نشان می‌دهد که الگوریتم‌های FRL و RL با یادگیری استراتژی بهینه می‌توانند تصمیمات بهتری را با توجه به در نظر گرفتن مصرف انرژی به عنوان محدودیت اتخاذ کنند. همانطور که شکل نشان می‌دهد الگوریتم RL نمی‌تواند در شبکه‌های با مقیاس بالا خوب کار کند. الگوریتم SDN-based و حریصانه مصرف انرژی را در نمی‌گیرند و تنها بر روی انجام وظیفه در مهلت زمانی خود تاکید دارند که با کوتاه‌ترین زمان وظیفه انجام شود و بنابراین کارایی مناسبی از لحاظ انرژی ندارند.



شکل ۷: میانگین مصرف انرژی برای چهار الگوریتم با تعداد وظایف اولویت بالای مختلف
Figure 7. Average energy consumption for four algorithms with different numbers of high priority tasks

همانطور که بیان شد، روش پیشنهادی از یادگیری فدرال همراه با فناوری SDN استفاده می‌کند که می‌تواند انعطاف‌پذیری و مقیاس‌پذیری را در شبکه افزایش دهد. در شکل ۸، دو الگوریتم مبتنی بر FRL و SDN-based با هم مقایسه شده‌اند. در الگوریتم FRL، تابع هدف افزایش نرخ پذیرش وظایف و همچنین در الگوریتم SDN-based، تابع هدف براساس کاهش زمان برون‌سپاری وظیفه است. بنابراین با افزایش مهلت زمانی، وظایف می‌توانند روی RSU و ابر اجرا شوند که می‌تواند، هزینه اجرا را افزایش دهد. می‌توان گفت که هزینه پردازش در مدل هر دو الگوریتم به عنوان محدودیت است که لزوماً هزینه کمینه نمی‌شود. نتایج نشان می‌دهد که الگوریتم پیشنهادی با یادگیری استراتژی بهینه بر اساس الگوی ترافیکی خودروها می‌تواند هزینه را نیز کاهش دهد.



شکل ۸: میانگین هزینه در مقابل بیشینه مهلت زمانی برای وظایف مختلف

Figure 8. Average cost vs. maximum deadline for different tasks

۵- ارزیابی و نتیجه‌گیری

در این مقاله، ما یک تکنیک برون‌سپاری وظیفه مبتنی بر Q-learning عمیق فدرال را برای توسعه یک مدل برون‌سپاری سراسری برای محاسبات مه خودرویی پیشنهاد کرده‌ایم. روش پیشنهادی هم مصرف انرژی و هم عدم تعادل بار در سرورهای مه را به حداقل می‌رساند. نتایج تجربی نشان می‌دهند که تکنیک برون‌سپاری پیشنهادی ما از تکنیک‌های برون‌سپاری پایه از لحاظ میانگین زمان پاسخدهی و هم میانگین هزینه پردازش کارآمدتر است. در آینده، ما قصد داریم این کار را گسترش دهیم تا تکنیک‌های یادگیری تقویتی عمیق پیچیده‌تری را برای حل مسئله برون‌سپاری وظیفه همراه با دیگر مجموعه داده‌های محیط واقعی نظر بگیریم.

مراجع

- [1] M. K. Farimani, S. Karimian-Aliabadi, R. Entezari-Maleki, B. Egger, and L. Sousa, "Deadline-aware task offloading in vehicular networks using deep reinforcement learning," *Expert Systems with Applications*, vol. 249, p. 123622, 2024.
- [2] O. Akyıldız, F. Y. Okay, İ. Kök, and S. Özdemir, "Road to efficiency: Mobility-driven joint task offloading and resource utilization protocol for connected vehicle networks," *Future Generation Computer Systems*, vol. 156, pp. 157-167, 2024.
- [3] X. Zhang, Y. Zhu, C. Wang, J. Cao, Y. Chen, and J. Wang, "Multi-agent reinforcement learning for vehicular task offloading with multi-step trajectory prediction," *CCF Transactions on Pervasive Computing and Interaction*, vol. 6, no. 2, pp. 101-114, 2024.
- [4] J. Shi, J. Du, J. Wang, J. Wang, and J. Yuan, "Priority-aware task offloading in vehicular fog computing based on deep reinforcement learning," *IEEE Transactions on Vehicular Technology*, vol. 69, no. 12, pp. 16067-16081, 2020.
- [5] S. Zhou, Y. Sun, Z. Jiang, and Z. Niu, "Exploiting moving intelligence: Delay-optimized computation offloading in vehicular fog networks," *IEEE Communications Magazine*, vol. 57, no. 5, pp. 49-55, 2019.
- [6] Z. Zhou, H. Liao, X. Wang, S. Mumtaz, and J. Rodriguez, "When vehicular fog computing meets autonomous driving: Computational resource management and task offloading," *IEEE Network*, vol. 34, no. 6, pp. 70-76, 2020.
- [7] S. Misra and S. Bera, "Soft-VAN: Mobility-aware task offloading in software-defined vehicular network," *IEEE Transactions on Vehicular Technology*, vol. 69, no. 2, pp. 2071-2078, 2019.
- [8] A. A. Khadir and S. A. H. Seno, "SDN-based offloading policy to reduce the delay in fog-vehicular networks," *Peer-to-Peer Networking and Applications*, vol. 14, no. 3, pp. 1261-1275, 2021.

- [9] H. M. Birhanie, M. A. Messous, S.-M. Senouci, E.-H. Aglzim, and A. M. Ahmed, "MDP-based resource allocation scheme towards a vehicular fog computing with energy constraints," in *2018 IEEE Global Communications Conference (GLOBECOM)*, 2018, pp. 1-6: IEEE.
- [10] V. Sethi and S. Pal, "FedDOVe: A Federated Deep Q-learning-based Offloading for Vehicular fog computing," *Future Generation Computer Systems*, vol. 141, pp. 96-105, 2023.
- [11] F. L. Lewis and D. Vrabie, "Reinforcement learning and adaptive dynamic programming for feedback control," *IEEE circuits and systems magazine*, vol. 9, no. 3, pp. 32-50, 2009.
- [12] Q. Qi *et al.*, "Knowledge-driven service offloading decision for vehicular edge computing: A deep reinforcement learning approach," *IEEE Transactions on Vehicular Technology*, vol. 68, no. 5, pp. 4192-4203, 2019.
- [13] H. Guo, J. Liu, J. Ren, and Y. Zhang, "Intelligent task offloading in vehicular edge computing networks," *IEEE Wireless Communications*, vol. 27, no. 4, pp. 126-132, 2020.
- [14] S. Memon and M. Maheswaran, "Using machine learning for handover optimization in vehicular fog computing," in *Proceedings of the 34th ACM/SIGAPP Symposium on Applied Computing*, 2019, pp. 182-190.
- [15] J. Shi, J. Du, J. Wang, and J. Yuan, "Deep Reinforcement Learning-Based V2V Partial Computation Offloading in Vehicular Fog Computing," in *2021 IEEE Wireless Communications and Networking Conference (WCNC)*, 2021, pp. 1-6: IEEE.
- [16] N. E. H. Boubaker, K. Zarour, N. Guermouche, and D. Benmerzoug, "Optimizing Workflow Offloading and Migration under Timed Constraints in Fog and Cloud Computing," *Journal of Grid Computing*, vol. 23, no. 1, p. 10, 2025.
- [17] J. Zhao, M. Kong, Q. Li, and X. Sun, "Contract-based computing resource management via deep reinforcement learning in vehicular fog computing," *IEEE Access*, vol. 8, pp. 3319-3329, 2019.
- [18] B. Shabir, A. U. Rahman, A. W. Malik, R. Buyya, and M. A. Khan, "A federated multi-agent deep reinforcement learning for vehicular fog computing," *The Journal of Supercomputing*, pp. 1-27, 2022.
- [19] J. Xue, L. Wang, Q. Yu, and P. Mao, "Multi-agent deep reinforcement learning-based partial offloading and resource allocation in vehicular edge computing networks," *Computer Communications*, vol. 234, p. 108081, 2025.
- [20] P. Dai, Y. Huang, K. Hu, X. Wu, H. Xing, and Z. Yu, "Meta Reinforcement Learning for Multi-task Offloading in Vehicular Edge Computing," *IEEE Transactions on Mobile Computing*, 2023.
- [21] K. Zhang, J. Cao, and Y. Zhang, "Adaptive digital twin and multiagent deep reinforcement learning for vehicular edge computing and networks," *IEEE Transactions on Industrial Informatics*, vol. 18, no. 2, pp. 1405-1413, 2021.
- [22] S.-C. Lin, K.-C. Chen, and A. Karimoddini, "SDVEC: software-defined vehicular edge computing with ultra-low latency," *IEEE Communications Magazine*, vol. 59, no. 12, pp. 66-72, 2021.
- [23] Y. Hou, Z. Wei, R. Zhang, X. Cheng, and L. Yang, "Hierarchical task offloading for vehicular fog computing based on multi-agent deep reinforcement learning," *IEEE Transactions on Wireless Communications*, 2023.
- [24] X. Ma, L. Liao, Z. Li, R. X. Lai, and M. Zhang, "Applying federated learning in software-defined networks: A survey," *Symmetry*, vol. 14, no. 2, p. 195, 2022.
- [25] M. U. Thomas, "Queueing systems. volume 1: Theory (leonard kleinrock)," *SIAM Review*, vol. 18, no. 3, pp. 512-514, 1976.
- [26] B. Yang *et al.*, "Edge intelligence for autonomous driving in 6G wireless system: Design challenges and solutions," *IEEE Wireless Communications*, vol. 28, no. 2, pp. 40-47, 2021.
- [27] M. Li, J. Gao, L. Zhao, and X. Shen, "Deep reinforcement learning for collaborative edge computing in vehicular networks," *IEEE Transactions on Cognitive Communications and Networking*, vol. 6, no. 4, pp. 1122-1135, 2020.

- [28] Z. Ning and L. Xie, "A survey on multi-agent reinforcement learning and its application," *Journal of Automation and Intelligence*, vol. 3, no. 2, pp. 73-91, 2024.
- [29] B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas, "Communication-efficient learning of deep networks from decentralized data," in *Artificial intelligence and statistics*, 2017, pp. 1273-1282: PMLR.
- [30] L. Bracciale, M. Bonola, P. Loreti, G. Bianchi, R. Amici, and A. Rabuffi, "CRAWDAD dataset roma/taxi," ed: Version 2014-07-17, 2014.
- [31] S. Imambi, K. B. Prakash, and G. Kanagachidambaresan, "PyTorch," in *Programming with TensorFlow: solution for edge computing applications*: Springer, 2021, pp. 87-104.
- [32] I. Daubechies, R. DeVore, S. Foucart, B. Hanin, and G. Petrova, "Nonlinear approximation and (deep) ReLU networks," *Constructive Approximation*, vol. 55, no. 1, pp. 127-172, 2022.