

تحلیل یادگیری تقویتی Q برای مسیریابی گروهی بلادرنگ در محیط پویا با اعتبارسنجی مونت کارلو و پاداش تطبیقی

تاریخ دریافت: ۱۴۰۴/۱۱/۱۱

تاریخ پذیرش: ۱۴۰۵/۰۵/۰۴

ایمان شفیعی‌نژاد^{۱*}، فائزه شکری^۲

۱ استادیار پژوهشگاه هوافضا، وزارت علوم، تحقیقات و فناوری، تهران، ایران، shafieenejad@ari.ac.ir

۲ دانشجوی دکتری، پژوهشگاه هوافضا، وزارت علوم، تحقیقات و فناوری، تهران، ایران

چکیده

هدایت بلادرنگ گروه‌های رباتیک در محیط‌های پویا دارای موانع ثابت و متحرک، یکی از چالش‌های پیچیده در زمینه برنامه ریزی رباتیک خودران است. این پژوهش یک چارچوب پیشرفته یادگیری تقویتی ارائه می‌دهد که عملکرد دو الگوریتم یادگیری Q و یادگیری Q دوگانه را در هدایت گروهی ربات‌های هوایی مقایسه و بهبود می‌دهد. مطالعه حاضر از یک شبیه‌سازی شامل ۱۰ ربات هوایی، ۱۲ مانع ثابت با پراکندگی هوشمند و ۲ مانع متحرک با الگوی حرکت تصادفی است. دو ساختار حرکتی سیستم ۴ حرکتی و سیستم ۶ حرکتی (حرکت‌های مورب) بهره‌برده است. ساز و کارهای پاداش نوآورانه شامل پاداش پیشرفت (بر پایه کاهش فاصله اقلیدسی از هدف) و پاداش موفقیت گروهی پیاده‌سازی شد. نتایج نشان داد یادگیری Q دوگانه با ساختار ۶ حرکتی بالاترین عملکرد را با دقت موقعیتی ۷۸٫۲ درصد و میانگین پاداش ۱۵۶٫۱۸ ارائه داد. چارچوب مبتنی بر روش مونت کارلو ارائه شده، روشی قابل اطمینان برای ارزیابی الگوریتم‌های یادگیری تقویتی در سناریوهای پیچیده فراهم می‌کند. این چارچوب می‌تواند در کاربردهایی مانند تحویل خودران و عملیات جستجو و نجات به کار رود.

واژه‌های کلیدی: مسیریابی بلادرنگ، رباتیک گروهی، یادگیری تقویتی Q دوگانه، موانع متحرک، شبیه‌سازی مونت کارلو

A comparative analysis of q-learning variants for real-time swarm robotic path planning in dynamic environments with static and mobile obstacles using monte carlo validation and an adaptive reward mechanism

Iman Shafieenejad¹, Faezeh Shokri²

1- Assistant Professor, Aerospace Research Institute, Ministry of Science, Research and Technology, Tehran, Iran

2- Ph.D. Student, Aerospace Research Institute, Ministry of Science, Research and Technology, Tehran, Iran

Abstract

Real-time guidance of robotic swarms in dynamic environments with static and mobile obstacles is a complex challenge in autonomous robotic planning. This research presents an advanced reinforcement learning framework comparing Q-learning and Double Q-learning for cooperative guidance of aerial robot swarms, using a simulation with 10 aerial robots, 12 static obstacles, and 2 randomly moving mobile obstacles. Two motion structures were utilized: a 4-directional system (cardinal directions) and a 6-directional system (with diagonal movements). Innovative reward mechanisms, including progress reward (based on reduced Euclidean distance to target) and group success reward, were implemented. Results demonstrated that Double Q-learning with the 6-directional structure achieved the highest performance, with positional accuracy of 78.2% and average reward of 156.18. The Monte Carlo-based framework provides a reliable evaluation method for reinforcement learning algorithms in complex scenarios, applicable to autonomous delivery and search and rescue operations.

Keywords: Real-Time Path Planning, Swarm Robotics, Double Q-Learning, Mobile Obstacles, Monte Carlo Simulation



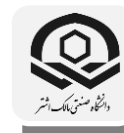
۱. مقدمه

هدایت بلادرنگ گروه‌های رباتیک در محیط‌های پویا و ناشناخته، به عنوان یکی از چالش‌های پیش‌رو در حوزه رباتیک خودران و سیستم‌های چند عامله مطرح است. توانایی ربات‌های هوایی برای حرکت هماهنگ در فضاهای پیچیده دارای موانع ثابت و متحرک، کاربردهای گسترده‌ای در عملیات‌های امداد و نجات، پایش محیطی، کشاورزی دقیق و لجستیک هوشمند دارد. با این حال، حضور موانع پویا و نیاز به تصمیم‌گیری بلادرنگ، این مسئله را از یک مسیریابی ساده به یک مسئله مسیریابی پویا و چند معیاره تبدیل کرده است [۱-۳]. مسئله اصلی در هدایت گروهی ربات‌ها، هماهنگی میان عوامل برای دستیابی به هدف مشترک، در عین پرهیز از برخورد با موانع و یکدیگر است. محیط‌های عملیاتی اغلب شامل موانع ثابت (مانند ساختمان‌ها، درختان) و موانع متحرک (مانند حیوانات، خودروها، یا حتی سایر ربات‌ها) هستند که پیش‌بینی موقعیت آن‌ها را دشوار می‌سازد. این پویایی، روش‌های برنامه‌ریزی مسیر کلاسیک را که بر پایه نقشه‌های ایستا طراحی شده‌اند، با محدودیت مواجه می‌کند و نیازمند رویکردهای هوشمند و سازگار شونده است [۴-۶].

در میان روش‌های نوین و هوشمند مسیریابی بلادرنگ، یادگیری تقویتی به عنوان یک روش قدرتمند برای یادگیری رفتار بهینه از طریق تعامل مستقیم با محیط، توجه بسیاری را به خود جلب کرده است. در این روش، عامل با دریافت پاداش یا جریمه از محیط، به تدریج سیاست بهینه حرکت را فرا می‌گیرد. این ویژگی، یادگیری تقویتی را به ویژه برای محیط‌های پویا و غیرقطعی که مدل دقیق آن‌ها در دسترس نیست،

مناسب ساخته است. در این روش، عامل نیازی به دانش پیشین کامل از محیط ندارد و می‌تواند در حین عمل، دانش خود را به روز کند [۷-۹]. در میان روش‌های یادگیری تقویتی^۱، الگوریتم‌های خانواده یادگیری Q^۲ به دلیل سادگی مفهومی و کارایی عملی، جایگاه ویژه‌ای دارند. با این حال یادگیری Q استاندارد ممکن است به دلیل پدیده برآورد بیش از حد مقادیر Q، منجر به یادگیری سیاست‌های ناکارآمد شود. برای غلبه بر این نقیصه، الگوریتم یادگیری Q دوگانه^۳ ارائه شده است که با استفاده از دو جدول Q مجزا، عمل به تخمین بهینه‌تر مقادیر می‌کند و پایداری یادگیری را افزایش می‌دهد. بررسی مقایسه‌ای این دو الگوریتم در محیط‌های پیچیده گروهی، موضوعی است که نیاز به بررسی بیشتر دارد که در پژوهش حاضر به آن پرداخته شده است [۱۰-۱۳].

در ادامه به مرور پیشینه پژوهش با رویکرد مسیریابی بلادرنگ هوشمند پرداخته می‌شود. واتکینز و دایان الگوریتم یادگیری Q را به عنوان یک روش ساده و افزایشی برای یادگیری بهینه تصمیم‌ها در محیط‌های مارکوفی کنترل شده معرفی کرده‌اند. این الگوریتم با بهبود پی در پی مقادیر Q، کیفیت اعمال مختلف را ارزیابی می‌کند. همچنین به‌روزرسانی چند مقدار Q به‌طور همزمان ارائه شده است [۱۴]. هاسلت و همکاران به بررسی مشکل بیش برآورد مقادیر عمل^۴ در یادگیری Q می‌پردازد. پژوهشگران نشان می‌دهند که ایده یادگیری Q عمیق^۵، که ابتدا در محیط‌های جدولی معرفی شده بود، می‌تواند برای تقریب توابع در مقیاس بزرگ تعمیم یابد. با ارائه یک سازگارسازی خاص برای DQN، الگوریتم حاصل نه تنها بیش برآورد را

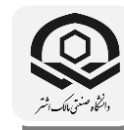


کاهش می‌دهد بلکه باعث بهبود عملکرد می‌شود [۱۴]. نگوین و همکاران الگوریتم‌های On-policy و Off-policy یادگیری Q را برای کنترل ربات‌های پاندول وارونه دوچرخ (TWIP) ارائه کردند. در این روش‌ها، داده‌ها بدون نیاز به مدل جمع‌آوری می‌شوند و سیستم به دو زیر سیستم تقسیم شده تا حجم داده‌ها کاهش یابد و سرعت محاسبات افزایش یابد. نتایج شبیه‌سازی نشان داد که تکنیک‌های Off-policy با کارایی داده‌ای بالاتر و دقت بیشتر، عملکرد بهتری در کنترل سیستم‌های TWIP دارند [۱۱]. گئو و همکاران الگوریتم Q-MA مبتنی بر یادگیری Q و الهام از رفتار اجتماعی حیوانات (ممیتیک) را برای حل مسئله مسیریابی مشترک وسایل نقلیه در برداشت و حمل و نقل کشاورزی هوشمند ارائه کردند. الگوریتم با استراتژی مقدار دهی اولیه ترکیبی، سه عملگر تقاطع تخصصی و جستجوی محلی تطبیقی هدایت شده توسط یادگیری Q، جمعیت متنوعی تولید کرده و راه حل‌های با کیفیت بالا ارائه می‌دهد. نتایج آزمایش‌ها نشان داد که Q-MA عملکرد بهتری نسبت به الگوریتم‌های مرجع دارد و مؤلفه‌های الگوریتم به بهبود همگرایی و کیفیت راه‌حل کمک می‌کنند [۱۵]. دوان و همکاران الگوریتم MOQLCPSO مبتنی بر یادگیری Q و PSO چند هدفه را برای مسیریابی ربات‌ها در زمین‌های ناهموار ارائه کردند. این الگوریتم با سازگاری پویا پارامترها و عملگر تقاطع استراتژیک مسیرهای بهینه با حداقل طول و ناهمواری زمین تولید می‌کند. نتایج نشان داد که MOQLCPSO نسبت به روش‌های پیشرفته، دقت، تنوع جمعیت و کارایی بهینه‌بالتری دارد و برای کاربردهای چند هدفه مسیریابی و تخصیص منابع مناسب

است [۱۶]. لی و همکاران الگوریتم تکرار حریصانه‌ای مبتنی بر یادگیری Q را برای زمان‌بندی گروهی چند سناریویی پیشنهاد کردند. در این روش، یادگیری Q برای بهبود انتخاب جستجوی محلی و افزایش تنوع راه‌حل‌ها به کار رفت. نتایج نشان داد الگوریتم پیشنهادی از نظر کارایی و پایداری از روش‌های موجود برتر است [۱۷]. جیوفرید و همکاران یک مطالعه تجربی بر روی الگوریتم‌های یادگیری چندعاملی برای مسیریابی هم‌زمان در محیط‌های انبار انجام دادند که مسیرهای بدون برخورد برای چندین عامل را محاسبه می‌کند. این پژوهش الگوریتم‌های مبتنی بر یادگیری را با روش‌های سنتی جست‌وجو و برنامه‌ریزی مقایسه کرده و کاربرپذیری آن‌ها در محیط‌های واقعی را ارزیابی می‌کند. نتایج نشان داد برخی الگوریتم‌های یادگیری چندعاملی عملکردی نزدیک به روش‌های سنتی با کیفیت مسیر مشابه و تلاش محاسباتی کمتر دارند، که نشان‌دهنده پتانسیل بالای آن‌ها برای کاربردهای عملی است [۱۸].

ژو و همکاران یک الگوریتم یادگیری Q بهینه‌شده برای برنامه‌ریزی مسیر محلی ربات‌های متحرک پیشنهاد کردند. آن‌ها با معرفی روش‌های جدید مقداردهی اولیه جدول Q، سیاست انتخاب عمل و تابع پاداش و نیز استفاده از روش RMSprop برای تنظیم پویای نرخ یادگیری، به بهبود همگرایی و جلوگیری از افتادن در بهینه‌های محلی پرداختند. نتایج شبیه‌سازی در محیط‌های مارپیچ با پیچیدگی‌های مختلف نشان داد که الگوریتم پیشنهادی در معیارهایی مانند تعداد گام‌ها، پاداش اکتشاف و زمان اجرا نسبت به روش‌های موجود برتری دارد. این پژوهش مسیر را برای توسعه الگوریتم‌های





یادگیری تقویتی کارآمدتر در محیط‌های پویا هموار می‌سازد [۲۰]. گائو و همکاران مسئله مسیریابی چند عاملی با رعایت محدودیت‌های زمانی را بررسی کردند و دو الگوریتم جستجوی مبتنی بر بهینه‌سازی و تقریب را ارائه دادند. نتایج نشان داد الگوریتم فوق تقریباً بهینه عمل کرده و در محیط‌های پر مانع (ثابت) موفقیت بالاتری دارد [۱۹]. ونو و گوروسامی یک مرور جامع بر الگوریتم‌های مسیریابی برای هدایت خودران ارائه دادند که روش‌های کلاسیک، فرا ابتکاری و مبتنی بر هوش مصنوعی را بررسی می‌کند. این مطالعه چالش‌های محیط‌های دینامیک، محدودیت‌های غیرهولونومیک^۷ و سطوح متفاوت دانش محیطی را تحلیل کرده و نقاط قوت و محدودیت هر دسته الگوریتم را برجسته می‌کند. همچنین، روندهای نوین مانند ادغام یادگیری ماشین و یادگیری تقویتی و مسیرهای تحقیقاتی آینده برای بهبود تطبیق‌پذیری و عملکرد سیستم‌های مسیریابی در محیط‌های پیچیده مورد بحث قرار گرفته است [۲۰]. آسلان و دمیرجی یک الگوریتم با یادگیری Q برای مسیریابی ربات‌ها ارائه کردند که علاوه بر بهینه‌سازی مسیر، مدیریت بیماری‌های همه گیر مانند قرنطینه و بازگشایی را نیز مد نظر دارد. این الگوریتم با بهبود روند درمان مدل شده و حذف پارامترهای کنترل اختصاصی، کارایی جستجو را افزایش می‌دهد. آزمایش‌ها در دوازده سناریوی مختلف نشان داد مسیرهای یافت‌شده توسط-Q LIPA عملکرد بهتری نسبت به سایر تکنیک‌های بهینه‌سازی هوشمند دارند و مدیریت مبتنی بر یادگیری تقویتی تأثیر مثبت بر عملکرد کلی الگوریتم دارد [۲۱]. وانگ و همکاران الگوریتم نوینی را برای مسیریابی ربات در محیط‌های

پیچیده ارائه کردند که با نمونه برداری احتمالی منطقه‌ای، تنظیم گام تطبیقی، گرایش دینامیک به هدف و میدان پتانسیل مصنوعی بهینه‌شده، کارایی و همگرایی مسیر را بهبود می‌بخشد. این الگوریتم با استفاده از توسعه دو طرفه درخت بهینه‌سازی حریم‌خانه و محدودیت‌های قابلیت اجرا، کیفیت مسیر و سرعت جستجو را افزایش می‌دهد. شبیه‌سازی‌ها نشان داد الگوریتم فوق نسبت به الگوریتم‌های معیار، مسیرهای کوتاه‌تر و زمان برنامه‌ریزی کمتری ارائه می‌دهد و برای هدایت ربات در محیط‌های پیچیده بسیار مناسب است [۲۲]. نی و همکاران یک الگوریتم ABC یکپارچه چند استراتژی مبتنی بر یادگیری Q (QMABC) برای مسیریابی وسایل نقلیه ارائه کردند که با بهره‌گیری از استراتژی‌های جستجوی متنوع و انتخاب هوشمندانه آن‌ها از طریق یادگیری Q، عملکرد جستجو را بهبود می‌بخشد. این الگوریتم با پیکربندی‌های مؤثرتر حالت^۸ و عمل، نسبت به روش‌های قبلی توانایی بهینه‌سازی بالاتری دارد. ارزیابی‌ها روی توابع معیار و کاربردهای مسیریابی نشان داد QMABC در حل مسائل عددی و عملی، عملکرد برتری نسبت به سایر الگوریتم‌های فراابتکاری دارد [۲۳]. محمد و همکاران یک کنترل‌کننده هوشمند برای موتور جت بر پایه یادگیری تقویتی عمیق و با استفاده از روش گرادیان قطعی سیاست عمیق ارائه کردند. نتایج شبیه‌سازی نشان داد این روش نسبت به کنترل‌کننده تناسبی-انتگرالی، زمان پاسخ کوتاه‌تر، کاهش دمای ورودی توربین و بهبود ایمنی عملکرد موتور را به همراه دارد. این پژوهش کارایی روش‌های یادگیری تقویتی در کنترل سامانه‌های غیرخطی و پیچیده را تأیید کرد [۲۴].

با وجود پژوهش‌های متعدد، مقایسه عملکرد یادگیری Q و یادگیری Q دوگانه در محیط‌های گروهی با موانع هر دو نوع ثابت و متحرک، همراه با ارزیابی قابلیت اطمینان مبتنی بر شبیه‌سازی مونت کارلو، مورد پژوهش قرار نگرفته است. همچنین، تأثیر افزودن درجات آزادی حرکتی (مانند حرکت مورب) بر کارایی و همگرایی این الگوریتم‌ها نیز نیاز به بررسی دارد. بنابراین، هدف اصلی این پژوهش، ارائه یک چارچوب ارزیابی جامع و مبتنی بر شبیه‌سازی مونت کارلو برای سنجش عملکرد و پایداری الگوریتم‌های یادگیری Q و یادگیری Q دوگانه در مسئله مسیریابی بلادرنگ گروهی ربات‌ها، در مواجهه با موانع ثابت و متحرک است.

ساختار این مقاله به این شرح است که در بخش دوم، مبانی نظری روش‌های یادگیری تقویتی مورد بحث قرار می‌گیرد. بخش سوم به تشریح دقیق محیط شبیه‌سازی، معماری حرکتی ربات‌ها، طراحی سازوکار پاداش و روش‌شناسی ارزیابی مونت کارلو می‌پردازد. در بخش چهارم، نتایج شبیه‌سازی و تحلیل‌های آماری ارائه می‌شود. بخش پنجم به بحث درباره یافته‌ها، دلالت‌های آن‌ها و مقایسه با کارهای مرتبط اختصاص دارد. در نهایت، در بخش ششم نتیجه‌گیری ارائه می‌شود.

۲. یادگیری تقویتی

۱،۲ یادگیری Q

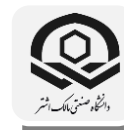
یادگیری تقویتی^۹ یک چارچوب ریاضیاتی است که برای توسعه عامل‌های محاسباتی طراحی شده و امکان یادگیری رفتار بهینه را با ارتباط دادن پاداش‌های^{۱۰} گذشته و جاری با اعمال انجام شده فراهم می‌کند. این چارچوب هوش مصنوعی در

حوزه‌هایی مختلف علوم و مهندسی کاربردهای موفق فراوانی داشته و برای تصمیم‌گیری‌های پی‌درپی در محیط‌های ناشناخته و با داده‌های حجیم و بلادرنگ بسیار مناسب است. در این چارچوب، یادگیری Q به عنوان یک الگوریتم یادگیری تقویتی بدون مدل مطرح است که با استفاده از تابع Q ، پاداش‌های احتمالی آینده اعمال در حالت‌های مشخص را تخمین می‌زند و از طریق به روز رسانی‌های تکراری، یادگیری سیاست بهینه را ممکن می‌سازد. الگوریتم یادگیری Q بر اساس معادله بلمن^{۱۱} عمل می‌کند:

$$Q(s, a) = E[r + \gamma \cdot \max_{a'} Q(s', a')] \quad (1)$$

در معادله (۱)، s نشان دهنده حالت فعلی است و a عمل انتخاب شده در این حالت می‌باشد. پس از اجرای عمل a ، پاداش دریافتی با نماد r مشخص می‌شود و حالت بعدی که به آن منتقل می‌شود، با s' نمایش داده می‌شود. ضریب تنزیل γ تعادل بین پاداش‌های فوری و پاداش‌های آینده را تعیین می‌کند و $\max_{a'} Q(s', a')$ بهترین ارزش قابل دستیابی از حالت بعدی s' را نشان می‌دهد که بیشترین پاداش بلند مدت ممکن را منعکس می‌کند. این ضریب در بازه $[0, 1]$ تعیین می‌کند که پاداش‌های آینده نسبت به پاداش فوری r چه وزنی دارند. عبارت $\max_{a'} Q(s', a')$ نیز بیانگر بیشترین ارزش قابل دستیابی از حالت بعدی الگوریتم با هدف یادگیری تابع $Q(s, a)$ یا تابع ارزش عمل-حالت^{۱۲} کار می‌کند. این تابع، پاداش تجمعی و تعدیل شده آینده‌ای را که عامل از انجام عمل a در حالت s انتظار دارد، تخمین می‌زند. هدف نهایی یادگیری، تقریب تابع ارزش بهینه $Q(s, a)$ است که در معادله بلمن بیان





می‌شود. یادگیری این تابع ارزش از طریق یک فرآیند تکراری و مبتنی بر تجربه صورت می‌پذیرد. در هر گام از تعامل با محیط، پس از دریافت پاداش و مشاهده حالت جدید، مقدار $Q(s, a)$ مربوطه با استفاده از قاعده به‌روز رسانی اختلاف زمانی^{۱۳} اصلاح می‌شود.

اجزای اصلی یادگیری Q شامل چند بخش مهم هستند. اولین بخش، مقادیر Q یا مقادیر عمل است که نشان‌دهنده پاداش‌های مورد انتظار برای انجام یک عمل در یک حالت مشخص می‌باشند. این مقادیر به مرور زمان و بر اساس قانون به‌روز رسانی اختلاف زمانی^{۱۴} اصلاح می‌شوند. دوم، پاداش‌ها و اپیزودها^{۱۵} هستند. در این بخش عامل با انجام اعمال مختلف از حالتی به حالت دیگر حرکت کرده و در این مسیر پاداش دریافت می‌کند. سوم، اختلاف زمانی به روز^{۱۶} است که عامل برای به‌روز رسانی مقادیر Q از فرمول اختصاصی به‌روز رسانی اختلاف زمانی استفاده می‌کند، فرمولی که ترکیبی از پاداش فوری و تخمین پاداش‌های آینده را در نظر می‌گیرد.

$$\begin{aligned} Q(S, A) \\ \leftarrow Q(S, A) \\ + \alpha [R + \gamma \max_{a'} Q(S', A') - Q(S, A)] \end{aligned} \quad (2)$$

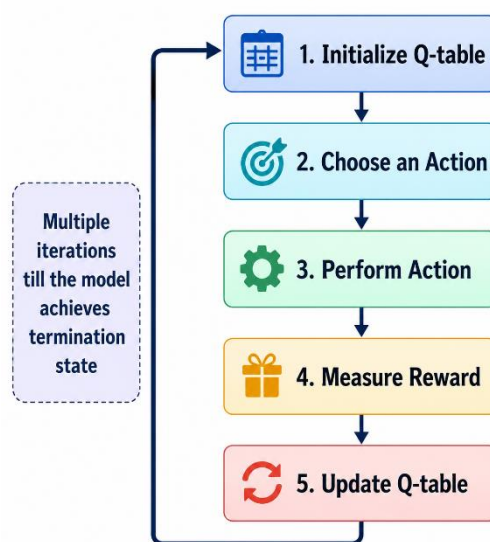
مجدداً با باز تعریف فرمول (۱) و بیان سیاست نهایی یادگیری Q در فرمول (۲)، S نشان‌دهنده وضعیت فعلی عامل است و A عملی است که عامل در حالت S انجام می‌دهد. پس از انجام عمل، عامل به حالت بعدی S' منتقل می‌شود و A' بهترین عملی است که در حالت S' می‌تواند انتخاب شود. R پاداشی است که عامل در نتیجه انجام عمل A در حالت S دریافت می‌کند و α نرخ یادگیری است که میزان تأثیر اطلاعات جدید بر مقادیر Q قدیمی را تعیین می‌کند. در

الگوریتم‌های یادگیری تقویتی، سیاست (ϵ) برای ایجاد تعادل بین دانش فعلی و جست‌وجوی گزینه‌های جدید استفاده می‌شود. در این سیاست، عامل میان بهره‌برداری از دانش فعلی و اکتشاف گزینه‌های جدید تعادل برقرار می‌کند. عامل با احتمال $1 - \epsilon$ عملی را بر می‌گزیند که بیشترین مقدار Q را دارد و به این ترتیب پاداش مورد انتظار را حداکثر می‌کند. با این حال، با احتمال ϵ عمل به انتخاب تصادفی می‌کند تا مسیرهای ناشناخته را بررسی کرده و امکان دستیابی به راهبردهای بهتر را افزایش دهد. مدل‌های یادگیری Q یک فرآیند تکراری را دنبال می‌کنند که در آن اجزای مختلف با هم همکاری می‌کنند تا عامل آموزش ببیند. به عبارت دیگر، عبارت داخل کروشه خطای تخمین یا اختلاف زمانی نامیده می‌شود که تفاوت بین تخمین جدید $[R + \gamma \max_{a'} Q(S', A') - Q(S, A)]$ و تخمین قدیمی $Q(S, A)$ است. الگوریتم با کاهش تدریجی این خطا، مقادیر Q را به سمت مقادیر بهینه همگرا می‌کند. انتخاب عمل در هر حالت، نقش تعیین‌کننده‌ای در فرآیند یادگیری دارد. برای ایجاد تعادل ضروری بین اکتشاف^{۱۷} (جست‌وجوی فضای عمل) و بهره‌برداری^{۱۸} (استفاده از دانش فعلی)، از سیاست ϵ حریصانه استفاده می‌شود. بر اساس این سیاست:

- با احتمال $1 - \epsilon$ عامل، عمل دارای بیشترین مقدار Q در حالت فعلی را انتخاب می‌کند (بهره‌برداری).
- با احتمال ϵ ، عامل یک عمل را به‌طور کاملاً تصادفی از بین اعمال ممکن بر می‌گزیند (اکتشاف).

مقدار ϵ معمولاً در ابتدای یادگیری بالا است تا فضای حالت-عمل به خوبی کاوش شود و به تدریج کاهش می‌یابد تا عامل در مراحل پایانی بر دانش به دست آمده خود تکیه کند. این فرآیند تلفیق سیاست انتخاب عمل و قاعده به‌روز رسانی مقادیر Q ، هسته اصلی الگوریتم یادگیری Q را تشکیل می‌دهد. در ادامه، این روند به صورت مرحله به مرحله در شکل (۱) توضیح داده می‌شود.

Q Learning Algorithm



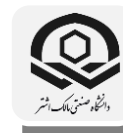
شکل ۱. الگوریتم یادگیری Q

فرآیند یادگیری Q با شروع از یک حالت اولیه آغاز می‌شود که وضعیت محیط یا شرایط کنونی عامل را توصیف می‌کند. سپس عامل بر اساس حالت فعلی و سیاست خود یک عمل انتخاب می‌کند. این انتخاب معمولاً با استفاده از جدول Q انجام می‌شود که پاداش‌های احتمالی برای جفت‌های حالت-عمل را تخمین می‌زند و اغلب از استراتژی ϵ بهره می‌برد؛ به این صورت که با احتمال ϵ اعمال جدید را به‌طور تصادفی امتحان می‌کند تا مسیرها و راهبردهای تازه را کشف کند و با احتمال $1 - \epsilon$ از اعمالی استفاده می‌کند که

تاکنون بیشترین پاداش را داشته‌اند، فرایندی که به آن بهره‌برداری گفته می‌شود. پس از انتخاب عمل، عامل آن را اجرا کرده و محیط به دو صورت واکنش نشان می‌دهد: ارائه یک حالت جدید S' به عنوان نتیجه عمل و ارائه یک پاداش R به عنوان بازخورد از کارایی عمل. با استفاده از این تجربه جدید، عامل جدول Q را به‌روز رسانی می‌کند. مقدار مربوط به جفت حالت-عمل براساس پاداش دریافتی و حالت جدید تنظیم می‌شود. این امر به عامل کمک می‌کند تخمین دقیق‌تری از سودمندی اعمال مختلف داشته باشد. با به‌روز رسانی مقادیر Q ، عامل سیاست خود را اصلاح کرده و چرخه مشاهده حالت‌ها، انتخاب اعمال، دریافت پاداش و به‌روز رسانی مقادیر Q را بارها تکرار می‌کند تا در نهایت سیاست بهینه‌ای را یاد بگیرد که بالاترین پاداش ممکن را در محیط به‌طور مداوم به دست آورد.

۲.۲ یادگیری Q دوگانه^{۱۹}

یادگیری Q دوگانه یک روش بهبود یافته از الگوریتم استاندارد یادگیری Q است که توسط هسلت در سال ۲۰۱۰ معرفی شد و برای حل مشکل بیش برآوردی^{۲۰} در مقادیر Q طراحی شده است. در یادگیری Q استاندارد، از همان مقادیر Q هم برای انتخاب عمل بهینه و هم برای ارزیابی ارزش آن عمل استفاده می‌شود که این مسئله می‌تواند منجر به بیش برآوردی در مقادیر Q گردد و در نتیجه یادگیری سیاست‌های زیر بهینه و کاهش پایداری الگوریتم را به دنبال داشته باشد. در فرمول به‌روز رسانی یادگیری Q یعنی معادله (۳) مقدار $\max Q(s', a')$ هم عمل بهینه را انتخاب و هم ارزش آن را ارزیابی می‌کند. از آنجایی که مقادیر Q برآوردی هستند و شامل نویز می‌باشند، ماکزیمم گرفتن از این مقادیر



منجر به بیش برآوردی می‌شود، این بیش برآوردی به خصوص در محیط‌های پیچیده و با نویز می‌تواند مشکل ساز شود. ایده اصلی یادگیری دوگانه Q استفاده از دو تابع ارزش مستقل (QA و QB) است. این دو تابع به صورت متناوب برای انتخاب عمل بهینه و ارزیابی ارزش آن استفاده می‌شوند. فرمول‌های به روز رسانی یادگیری Q دوگانه به این صورت می‌باشند:

$$\begin{aligned} E[\max Q(s', a')] \\ \geq \max E[Q(s', a')] \end{aligned} \quad (3)$$

این نابرابری ریاضی توضیح می‌دهد که یادگیری Q معمولی می‌تواند مقدار Q را بیش برآورد کند. به عبارتی می‌توان گفت، انتظار از ماکزیمم، از ماکزیمم انتظارها بزرگ‌تر است.

$$\begin{aligned} a^* \\ = \operatorname{argmax} QA(s', a') \end{aligned} \quad (4)$$

عمل بهینه در حالت بعدی s' بر اساس یکی از دو شبکه Q است. به عبارتی انتخاب بهترین عمل در حالت بعدی بر اساس شبکه QA است.

$$\begin{aligned} QA(s, a) \leftarrow QA(s, a) + \alpha [r \\ + \gamma QB(s', a^*) \\ - QA(s, a)] \end{aligned} \quad (5)$$

از طرف دیگر معادله (۵) به روزرسانی QA با استفاده از مقدار هدف از شبکه QB است.

$$a^* = \operatorname{argmax} QB(s', a') \quad (6)$$

معادله (۶) انتخاب بهترین عمل در حالت بعدی بر اساس شبکه QB است.

$$\begin{aligned} QB(s, a) \leftarrow QB(s, a) + \alpha [r \\ + \gamma QA(s', a^*) \\ - QB(s, a)] \end{aligned} \quad (7)$$

معادله (۷)، به روزرسانی QB با استفاده از مقدار هدف از شبکه است. در این روش به جای یک شبکه ارزش‌گذاری، از دو شبکه QA و

QB به طور متناوب استفاده می‌شود. هنگام به روز رسانی QA، برای محاسبه مقدار هدف از شبکه QB استفاده می‌شود و بالعکس. این جداسازی باعث می‌شود که انتخاب عمل بهینه (argmax) و ارزیابی ارزش آن عمل توسط یک شبکه واحد انجام نشود، که عامل اصلی بیش برآورد در یادگیری Q معمولی است. با به روز رسانی متناوب این دو شبکه، یادگیری پایدارتر شده و مقادیر Q با دقت بیشتری به مقادیر بهینه واقعی همگرا می‌شوند.

۳. روند شبیه سازی یادگیری تقویتی برای مسیریابی بلادرنک گروه ربات‌های هوایی در محیط پویا

شبیه سازی در یک محیط شبکه‌ای گسسته با ابعاد ۲۰ در ۲۰ سلول طراحی شده است. این فضا متشکل از ۱۰ ربات هوایی خودران، ۱۲ مانع ثابت و ۲ مانع متحرک است. هدف نهایی، رساندن تمامی ربات‌ها از موقعیت‌های آغازین پراکنده به یک موقعیت هدف ثابت بود. در طراحی یادگیری Q دو معماری حرکتی مجزا معماری چهار-جهته و معماری شش-جهته مورد بررسی قرار گرفت. در معماری اول، مجموعه عمل‌های هر ربات شامل حرکت به سمت بالا، راست، پایین و چپ بود. در معماری دوم، برای افزایش انعطاف در مسیریابی، دو حرکت مورب (بالا-راست و پایین-راست) نیز به مجموعه عمل‌ها افزوده شد. این طراحی امکان بررسی تأثیر درجه آزادی حرکتی بیشتر بر کارایی و پیچیدگی یادگیری را فراهم می‌ساخت.

موانع ثابت به صورت سلول‌های غیرقابل عبور در نظر گرفته شدند که موقعیت آن‌ها به صورت تصادفی، با اولویت قرارگیری در مناطق میانی و دور از نقاط شروع و هدف، تولید گردید. موانع

متحرک به صورت عامل‌های پویا شبیه‌سازی شده اند که در هر گام زمانی با احتمال غالب در جهت تعیین‌شده قبلی خود (یکی از چهار جهت اصلی) یک سلول جابه‌جا می‌شوند و با احتمال کم، جهت حرکت خود را به صورت تصادفی تغییر می‌دهند. حرکت این موانع در مرزهای محیط محدود شده و از برخورد با موانع ثابت یا اشغال سلول هدف جلوگیری می‌شد.

معماری این الگوریتم، یک ساختار سه‌بعدی بر اساس یادگیری Q و جدول مربوطه آن دارد. ابعاد این جدول به صورت (تعداد حالت‌ها $\times$ تعداد عمل‌ها $\times$ تعداد ربات‌ها) تعریف شد. هر درایه در این جدول، مقداری عددی را ذخیره می‌کرد که نشان‌دهنده برآورد عامل از سودمندی بلند مدت انجام یک عمل خاص در یک حالت خاص برای یک ربات خاص است. مقادیر اولیه این جدول با اعداد کوچک تصادفی مقدار دهی می‌شوند تا زمینه اولیه‌ای برای فرآیند اکتشاف فراهم گردد (به بخش پیشین یادگیری Q و Q دوگانه مراجع شود). برای اعمال سیاست جهت کاهش برآورد بیش از حد، مقادیر Q که ممکن است در الگوریتم یادگیری Q رخ دهد، الگوریتم یادگیری Q دوگانه نیز پیاده‌سازی شده است. در این روش، به جای یک جدول Q، از دو جدول Q مستقل Q_A و Q_B استفاده گردید. در هر مرحله به‌روز رسانی، به صورت تصادفی انتخاب می‌شود که کدام یک از این دو جدول، عمل بهینه‌یابی را انجام دهد و کدام یک مقدار هدف را ارائه کند. این مکانیزم دوگانه، منجر به تخمین محافظه‌کارانه‌تر و پایدارتر مقادیر Q می‌شود.

یک ساختار پاداش چندلایه و سازگار شونده طراحی شده است تا یادگیری را هدایت کند. این ساختار شامل جریمه گام برای تشویق به کارایی،

پاداش پیشرفت برای تقویت حرکت به سمت هدف (بر اساس کاهش فاصله اقلیدسی)، جریمه بالا برای برخورد با موانع، جریمه کم‌تر برای برخورد بین ربات‌ها، پاداش کلان دستیابی به هدف و در نهایت یک پاداش تشویقی گروهی برای مواردی که تمام ربات‌ها موفق به رسیدن به مقصد است. یادگیری در قالب دوره‌های آموزش متوالی انجام می‌پذیرد. در هر دوره، ربات‌ها از موقعیت‌های اولیه تا رسیدن به هدف یا رسیدن به حداکثر تعداد گام مجاز، به تعامل با محیط ادامه می‌دهند. انتخاب عمل در هر گام بر اساس سیاست اپسیلون-حریصانه^{۲۱} صورت می‌گرفت. در این سیاست، با احتمال $1 - \epsilon$ ربات عملی را انتخاب می‌کند که در جدول Q مربوط به خود، بیشترین مقدار را برای حالت فعلی دارد و با احتمال ϵ ، یک عمل را به‌طور کاملاً تصادفی برمی‌گزیند. مقدار ϵ به تدریج از یک مقدار آغازین بالا به یک مقدار پایین نهایی کاهش می‌یابد.

برای اطمینان از مقاوم بودن الگوریتم و کاهش وابستگی آن‌ها به شرایط تصادفی اولیه (مانند مکان موانع و مقداردهی اولیه جدول Q)، از روش شبیه‌سازی مونت کارلو استفاده شده است. این روش شامل اجرای ۲۰۰ تکرار مستقل کامل از فرآیند یادگیری است. در هر تکرار، تمام مراحل شامل تولید موانع، مقداردهی اولیه جداول Q و اجرای ۸۰۰ دوره آموزشی، به‌طور مجزا و تصادفی انجام می‌گیرد. معیارهای عملکرد (میانگین پاداش، نرخ موفقیت، تعداد برخوردها) در هر دوره برای هر تکرار ذخیره می‌شود. پس از اتمام تمامی تکرارهای مونت کارلو، داده‌های خروجی برای هر یک از چهار شرایط آزمایشی (ترکیب دو الگوریتم و دو معماری حرکتی) پردازش می‌شود. برای هر معیار و در هر شماره دوره، میانگین و انحراف



معیار در بین ۲۰۰ تکرار محاسبه می شود. این موضوع منجر به تولید منحنی های یادگیری هموار و قابل اطمینانی می شود که روند همگرایی و عملکرد نهایی هر روش را به وضوح نشان می داد. همچنین، داده های ۱۰۰ دوره پایانی به عنوان شاخص عملکرد پایدار و نهایی هر الگوریتم استخراج می شود.

در مرحله نهایی، بر روی داده های عملکرد پایدار (دوره های پایانی) تحلیل های آماری پیشرفته انجام می گیرد. این تحلیل ها شامل مقایسه میانگین ها با استفاده از آزمون t مستقل، محاسبه اندازه اثر (به عنوان مثال، d کوهن) برای بررسی کیفی میزان برتری یک روش نسبت به دیگری و تحلیل نرخ همگرایی (شمارش تعداد تکرارهایی که به آستانه عملکرد مشخصی دست یافته اند) است. نتایج این تحلیل های کمی، مبنای علمی برای مقایسه نهایی کارایی و قابلیت اطمینان الگوریتم های مورد مطالعه فراهم می آورد.

در این پژوهش، به منظور هدایت گروهی ربات ها در محیط پویا با موانع متحرک، یک ساختار پاداش تطبیقی طراحی شده است. مقادیر عددی تمام پارامترهای ثابت پاداش (از جمله پاداش هدف، جریمه برخورد، جریمه گام و غیره) در جدول (۲) ارائه شده است. در ادامه، فرمول های ریاضی مربوط به اجزای پویای پاداش و مکانیزم های تطبیقی سازی نرخ یادگیری و نرخ کاوش ارائه می گردد. منظور از «پاداش تطبیقی»، تغییر پویای پارامترهای نرخ یادگیری و نرخ کاوش بر اساس پیشرفت فرآیند یادگیری است، نه تغییر خود پاداش. این مکانیزم باعث بهبود همگرایی و کاهش نوسانات در مراحل پایانی یادگیری می شود.

معادله (۸) پاداش پیشرفت بر اساس کاهش فاصله ربات تا هدف را تعریف می کند.

$$R_{progress} = 0.5 \times \max(0, d_{current} - d_{next}) \quad (8)$$

که در آن d فاصله اقلیدسی ربات تا هدف است. معادله (۹) نحوه محاسبه پاداش کل هر ربات در هر گام را نشان می دهد.

$$R_{total} = -0.05 + R_{progress} + \begin{cases} +200 & \text{رسیدن به هدف} \\ -10 & \text{برخورد با مانع} \\ -2 & \text{برخورد با ربات دیگر} \end{cases} \quad (9)$$

همچنین معادله های (۱۰) و (۱۱) به ترتیب نرخ یادگیری تطبیقی و اپسیلون تطبیقی را بر اساس شماره اپیزود و عملکرد اخیر سیستم بیان می کنند:

$$\alpha_{adaptive} = 0.1 \times \left(1 - 0.5 \times \frac{ep}{800}\right) \quad (10)$$

$$\epsilon_{adaptive} = \epsilon \times \left(1 - 0.3 \times \text{میانگین موفقیت ۱۰ اپیزود اخیر}\right) \quad (11)$$

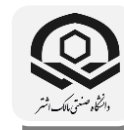
۴. آزمون های آماری مورد استفاده در

تحلیل نتایج شبیه سازی

۱،۴ تحلیل توصیفی و محاسبه معیارهای

مرکزی و پراکندگی

پیش از هر گونه آزمون استنباطی، برای هر یک از چهار گروه آزمایشی حاصل از ترکیب دو الگوریتم یادگیری Q و یادگیری Q دوگانه با دو معماری حرکتی ۴ و ۶ جهته، معیارهای آماری توصیفی محاسبه می گردد. این محاسبات بر روی داده های عملکرد نهایی (متوسط ۱۰۰ دوره آخر) تمامی اجراهای مونت کارلو انجام می شود. مهم ترین این معیارها شامل میانگین حسابی:



(۱۲)

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$$

به عنوان شاخص مرکزی و انحراف

معیار:

(۱۳)

$$= \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$$

به عنوان شاخص پراکندگی است. این مقادیر پایه لازم برای کلیه تحلیل‌های مقایسه‌ای بعدی را فراهم می‌نماید.

۲.۴ آزمون t مستقل برای مقایسه میانگین‌های دو گروه

برای بررسی وجود تفاوت آماری معنی‌دار بین میانگین عملکرد دو الگوریتم در هر معماری حرکتی، از آزمون t مستقل دو نمونه‌ای استفاده شده است.

این آزمون فرض صفر مبنی بر برابری میانگین‌های جامعه $H_0: \mu_1 = \mu_2$ را می‌آزماید. آماره t با فرض برابری واریانس‌ها (با استفاده از واریانس ترکیبی) به صورت زیر محاسبه می‌گردد.

(۱۴)

$$t = \frac{\bar{X}_1 - \bar{X}_2}{s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$$

که در آن

(۱۵)

$$s_p = \sqrt{\frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}}$$

و $\bar{X}_1, \bar{X}_2$ بیان‌کننده میانگین‌های نمونه هستند.

این نمادها نشان‌دهنده میانگین حسابی (مقدار متوسط) داده‌های عملکرد در هر یک از دو گروه مستقل مورد مقایسه هستند. این میانگین‌ها از تقسیم مجموع مقادیر مشاهده‌شده در هر گروه بر تعداد اعضای آن گروه محاسبه می‌شوند و شاخص اصلی مرکزی ترین نقطه

توزیع داده‌های آن گروه را نشان می‌دهند.

n_1 و n_2 نشان‌دهنده تعداد مشاهدات یا نمونه‌های مستقل در هر یک از دو گروه هستند. در چارچوب این تحقیق n_1 تعداد اجراهای مونت کارلو (تکرارهای مستقل شبیه‌سازی) در گروه اول است. در این راستا، n_2 تعداد اجراهای مونت کارلو در گروه دوم است. سپس مقدار احتمال p با درجه آزادی $df = n_1 + n_2 - 2$ از توزیع t استودنت استخراج شد. سطح معنی‌داری $\alpha = 0.05$ در نظر گرفته شد.

پس از محاسبه آماره آزمون t از روی مقادیر میانگین‌ها، انحراف معیارها و حجم نمونه‌های دو گروه، گام نهایی استنباط آماری با تعیین درجه آزادی $df = n_1 + n_2 - 2$ و استخراج مقدار احتمال متناظر p از تابع چگالی احتمال توزیع t استودنت با همان درجه آزادی انجام می‌پذیرد. این مقدار p که تحت فرض صفر برابری میانگین‌های جامعه محاسبه می‌شود، بیانگر احتمال نظری مشاهده‌ی تفاوتی به اندازه‌ی تفاوت نمونه‌ای مشاهده شده یا صرفاً بر اثر نوسانات تصادفی نمونه‌گیری است. برای اتخاذ تصمیم نهایی، این مقدار p با یک آستانه از پیش تعیین شده به نام سطح معناداری α که در این پژوهش مطابق استاندارد رایج علوم مهندسی برابر $0.05 (5\%)$ در نظر گرفته شده است، مقایسه می‌گردد. چنانچه رابطه $p < \alpha$ برقرار باشد، دال بر آن است که احتمال رخداد تصادفی تفاوت مشاهده شده به اندازه‌ی کافی پایین بوده و بنابراین فرض صفر مبنی بر عدم تفاوت معنادار بین میانگین‌های دو گروه مردود اعلام می‌شود و نتیجه‌گیری به نفع وجود تفاوت آماری معنادار می‌گردد.



۵,۴ آزمون فرض برابری واریانس‌ها (آزمون لوین^{۲۳})

به عنوان پیش‌نیاز و بررسی فرض ضمنی آزمون t مبنی بر همگنی واریانس‌ها، آزمون لوین برای مقایسه واریانس دو گروه اجرا می‌گردد. آماره این آزمون بر اساس انحراف‌های مطلق از میانگین محاسبه و مقدار p مربوطه استخراج می‌شود. این رویکرد، جامعیت و اتکاپذیری تحلیل‌های مقایسه‌ای را تضمین می‌کند. در جدول (۱) خلاصه این آزمون‌ها و نحوه بررسی آنها آورده شده است.

جدول ۱. خلاصه آزمون‌های آماری مورد استفاده در

تحلیل نتایج شبیه‌سازی

نام آزمون / شاخص	هدف از استفاده	فرمول محاسباتی	حروجهای اصلی	نحوه بررسی و تفسیر
معیارهای توصیفی	خلاصه‌سازی و توصیف داده‌های عملکرد هر گروه آزمایشی.	میانگین: $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ انحراف معیار: $S = \sqrt{\frac{\sum (X_i - \bar{X})^2}{n-1}}$	میانگین انحراف معیار	مقایسه مستقیم عددی میانگین‌ها و انحراف معیارها. گروه با میانگین بالاتر و انحراف معیار پایین‌تر، عملکرد بهتر و پایدارتری دارد.

۳,۴ محاسبه اندازه اثر d کوهن^{۲۲} برای کمی‌سازی بزرگی تفاوت

از آنجا که آزمون t تنها معنی‌داری آماری را نشان می‌دهد، برای ارزیابی اهمیت عملی و بزرگی تفاوت بین دو گروه، اندازه اثر استاندارد شده d کوهن محاسبه گردید. این معیار، تفاوت میانگین‌ها را بر حسب واحد انحراف معیار ترکیبی بیان می‌کند:

$$d = \frac{\bar{X}_1 - \bar{X}_2}{S_p} \quad (۱۶)$$

مقادیر d بر اساس راهنمای کوهن به صورت $d \approx 0.2$ اثر کوچک، $d \approx 0.5$ اثر متوسط و $d \geq 0.8$ اثر بزرگ تفسیر می‌شوند. این محاسبه امکان مقایسه مستقیم میزان برتری یک روش نسبت به دیگری را فراهم می‌آورد.

۴,۴ تحلیل نرخ همگرایی با شاخص‌های کسری

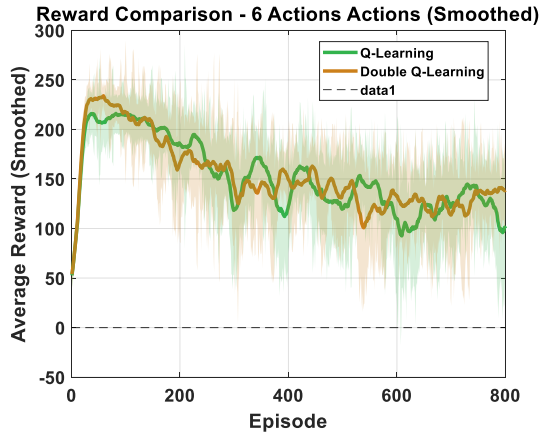
برای ارزیابی سرعت و قابلیت اطمینان همگرایی الگوریتم‌ها، از تحلیل‌های مبتنی بر نسبت استفاده می‌شود. معیار اصلی، نسبت اجراهای مونت کارلو است که تا یک آستانه عملکرد خاص (مثلاً دستیابی به نرخ موفقیت بالای ۵۰٪) در یک دوره زمانی مشخص (مثلاً ۱۰۰ دوره اول) همگرا شده‌اند. بازه اطمینان ۹۵٪ برای این نسبت $\hat{p}$ (نرخ همگرایی مشاهده شده) با استفاده از:

$$\hat{p} \pm Z_{0.975} \sqrt{\frac{\hat{p}(1-\hat{p})}{n}} \quad (۱۷)$$

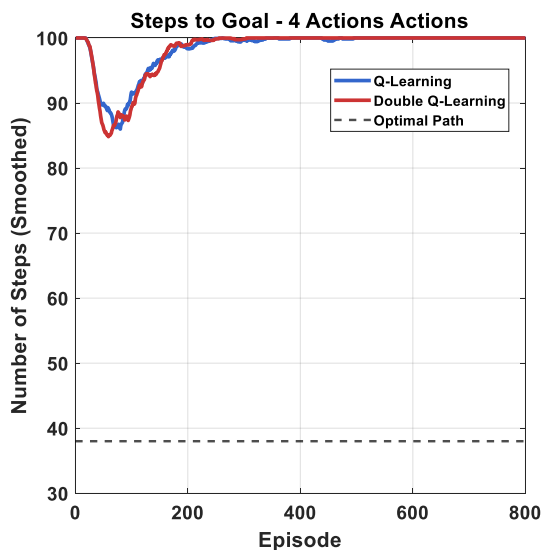
محاسبه می‌شود تا عدم قطعیت در برآورد قابلیت اطمینان هر روش نیز کمی‌سازی گردد.

۵. بررسی نتایج

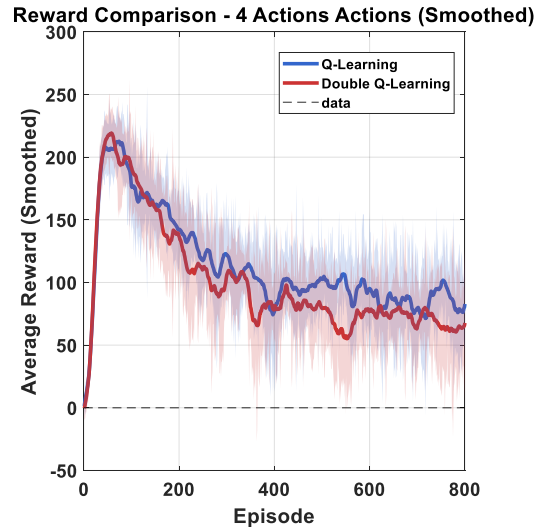
در این بخش، نمودارهای حاصل از شبیه‌سازی‌ها به تفصیل ارائه و تحلیل می‌شوند.



شکل ۳. مقایسه پاداش میانگین در حالت ۶ عمل در شکل (۳)، روند پاداش در حالت ۶ عمل برای دو الگوریتم نشان داده شده است. همان‌طور که مشاهده می‌شود، عملکرد هر دو الگوریتم بسیار نزدیک به هم می‌باشد که نشان‌دهنده کارایی مشابه آن‌ها در فضای عمل پیچیده‌تر است. با این حال، یادگیری Q دوگانه اندکی از یادگیری Q پیشی گرفته است و در انتها به ارزش عمل بیشتری دست یافته است. در شکل (۳) میزان نهایی پاداش برای روش‌های یادگیری Q و یادگیری دوگانه Q به ترتیب برابر ۱۴۲,۶ و ۱۴۷,۴ است.

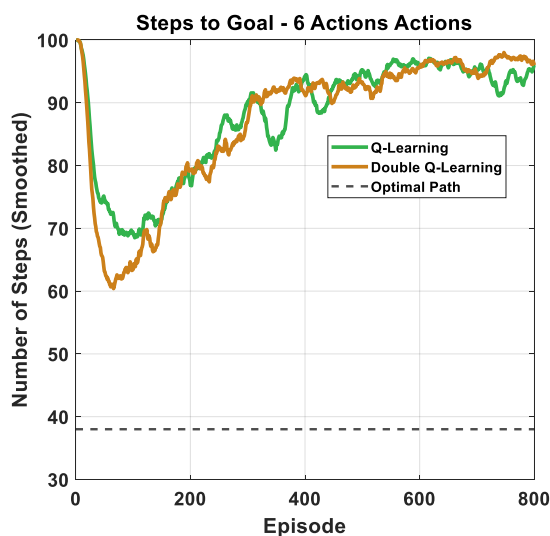


شکل ۴. تعداد گام‌ها تا هدف در حالت ۴ عمل



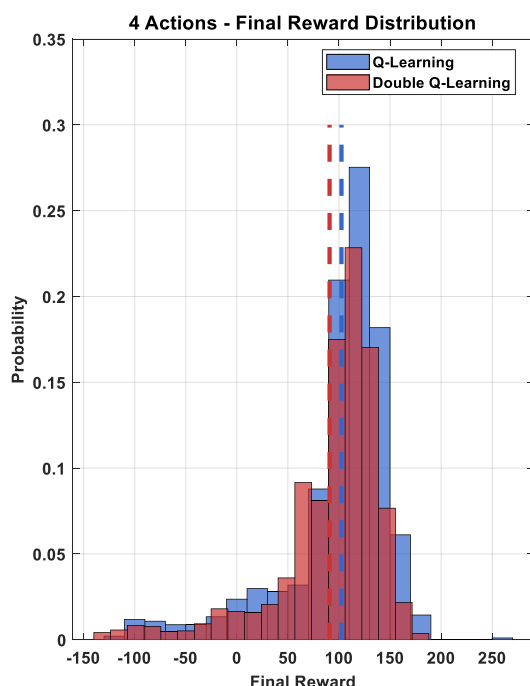
شکل ۲. مقایسه پاداش میانگین در حالت ۴ عمل در شکل (۲)، نمودار روند پاداش میانگین در طول ۸۰۰ اپیزود یادگیری برای دو الگوریتم یادگیری Q و یادگیری دوگانه در حالت ۴ عمل نمایش داده شده است. همان‌طور که مشاهده می‌شود، منحنی مربوط به یادگیری Q (خط آبی) با فاصله محسوسی از منحنی یادگیری دوگانه (خط قرمز) قرار دارد و رفتار هر دو الگوریتم نوسانی است. پاداش نهایی یادگیری Q و یادگیری دوگانه Q تقریباً برابر است. این نتیجه حاکی از آن است که در یک فضای عمل بر اساس چهار حرکت اصلی، الگوریتم یادگیری Q و یادگیری دوگانه Q تفاوتی ندارند. در شکل (۲) میزان نهایی پاداش برای روش‌های یادگیری Q و یادگیری دوگانه Q به ترتیب برابر $94,6 \pm 73,1$ و $81,9 \pm 73,9$ است.

در شکل (۴)، روند کاهشی تعداد گام‌های لازم برای دستیابی به هدف در محیط ۴ عمل نمایش داده شده است. همان‌گونه که مشاهده می‌شود، هر دو الگوریتم پس از اپیزود ۳۰۰ به مرحله‌ای از پایداری و همگرایی رسیده‌اند. روند یادگیری این الگوریتم‌ها ابتدا نزولی، سپس صعودی و در نهایت پایدار است که نشان‌دهنده پیدا کردن مسیر بهینه توسط این دو الگوریتم است.



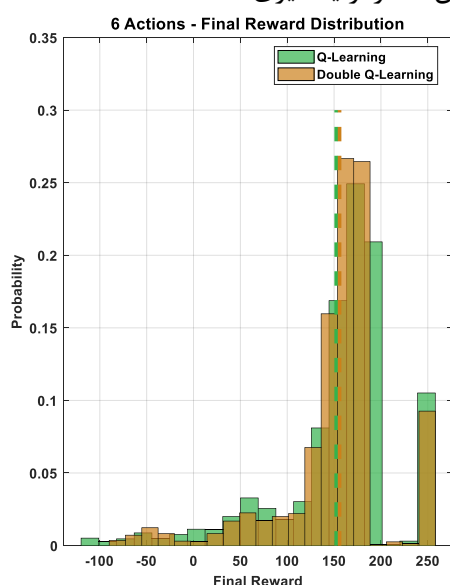
شکل ۵. تعداد گام‌ها تا هدف در حالت ۶ عمل

در شکل (۵) تغییرات تعداد گام‌های لازم برای رسیدن به هدف در طول اپیزودها برای الگوریتم‌های یادگیری Q و یادگیری Q دوگانه با ۶ عمل نمایش داده شده است. در ابتدای فرآیند، هر دو الگوریتم با کاهش سریع تعداد گام‌ها، بهبود قابل توجهی در عملکرد نشان می‌دهند که بیانگر یادگیری اولیه مؤثر است. با افزایش اپیزودها، روند تغییرات به تدریج افزایشی و همراه با نوسان می‌شود که می‌تواند ناشی از اکتشاف مداوم و پیچیدگی فضای حالت-کنش باشد. در اغلب بازه‌ها، یادگیری Q و یادگیری Q دوگانه فاصله کمی با یکدیگر دارند. اما با این حال پس از ۴۰۰ اپیزود به روند همگرایی پیرامون ۹۰ رسیده‌اند.



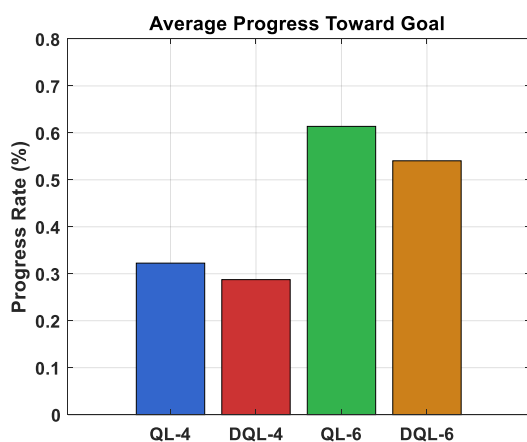
شکل ۶. توزیع احتمال (نزدیک به نرمال) نهایی پاداش در حالت ۴ عمل

در شکل (۶)، توزیع احتمال نهایی پاداش برای دو الگوریتم در حالت ۴ عمل مقایسه شده است. این الگو تأیید می‌کند، یادگیری Q در بهترین حالت نیازمند به پاداش بالاتری بوده و نوسان بیشتری تجربه کند. در مقابل، یادگیری Q دوگانه عملکردی بهتر ارائه می‌دهد و پاداش نهایی آن کمتر از یادگیری Q است.



شکل ۷. توزیع احتمال نهایی پاداش در حالت ۶ عمل

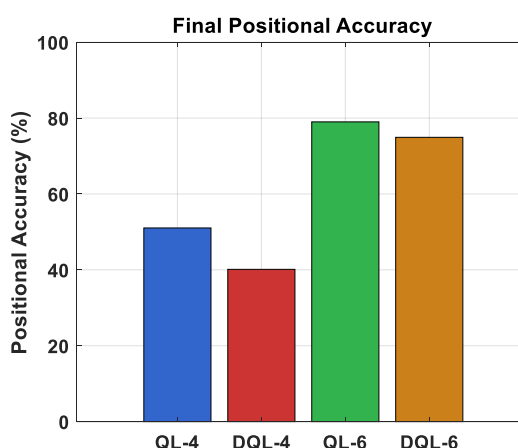
خود را نمایان می‌سازد و باعث می‌شود کارایی آن به سطح الگوریتم استاندارد برسد.



شکل ۹. نرخ متوسط پیشرفت

در شکل (۹)، نرخ متوسط پیشرفت به سمت هدف برای چهار سناریو مقایسه شده است. همان‌طور که مشاهده می‌شود، الگوریتم‌های با ۶ عمل نرخ پیشرفت بسیار بالاتری (حدود ۰.۶٪) نسبت به الگوریتم‌های با ۴ عمل (حدود ۰.۳٪) دارند. این نشان می‌دهد فضای عمل غنی‌تر از بابت در اختیار داشتن حرکات مورب، امکان حرکت کارآمدتر به سمت هدف را فراهم می‌کند. در هر دو پیکربندی، عملکرد یادگیری Q و یادگیری Q دوگانه بسیار نزدیک به هم است، که حاکی از کارایی مشابه آن‌ها در بهینه‌سازی پیشرفت است. با این حال، در حالت ۶ عمل، یادگیری Q دوگانه اندکی بهتر عمل کرده که همسو با برتری آن در دقت موقعیتی است.

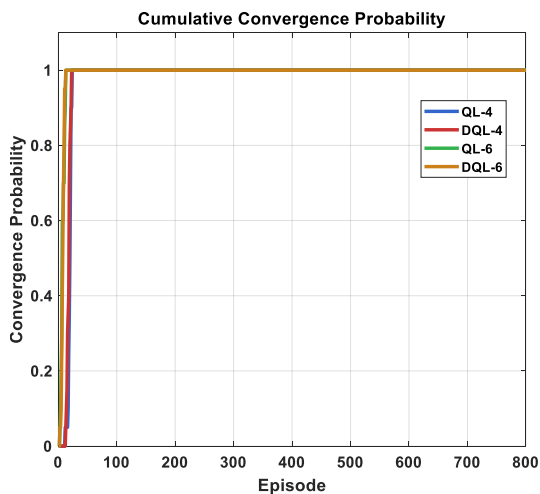
در شکل (۷) توزیع پاداش نهایی برای الگوریتم‌های یادگیری Q و یادگیری Q دوگانه در حالت ۶ عمل نمایش داده شده است. تمرکز اصلی هر دو توزیع در ناحیه پاداش‌های بالا قرار دارد که نشان‌دهنده دستیابی به سیاست‌های بهینه در انتهای فرآیند یادگیری است. با این حال، توزیع مربوط به یادگیری Q دوگانه فشرده‌تر و با پراکندگی کمتر مشاهده می‌شود که بیانگر پایداری بیشتر نتایج نهایی است. هم‌پوشانی قابل‌توجه دو توزیع نشان می‌دهد عملکرد میانگین هر دو الگوریتم نزدیک است.



شکل ۸. مقایسه دقت موقعیتی نهایی

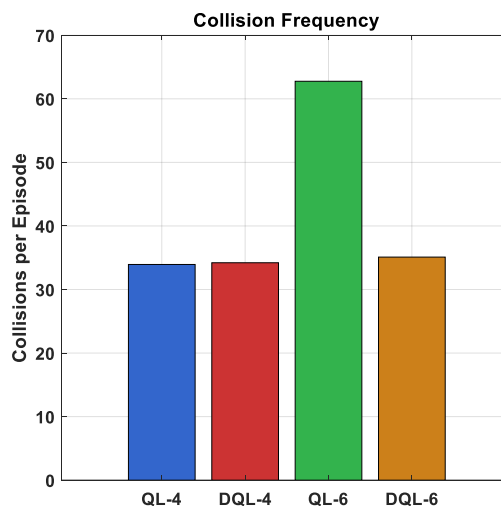
در شکل (۸)، دقت موقعیتی نهایی (درصد موفقیت) چهار سناریوی شبیه‌سازی مقایسه شده است. همان‌طور که مشاهده می‌شود، عملکرد هر دو الگوریتم با افزایش پیچیدگی فضای عمل از ۴ به ۶ عمل بهبود چشمگیری یافته است. در حالت ۴ عمل، یادگیری Q از یادگیری Q دوگانه پیشی گرفته است. اما در محیط پیچیده‌تر ۶ عمل، این اختلاف تقریباً از بین رفته و هر دو الگوریتم به دقت‌های بسیار نزدیک و بالایی (حدود ۸۰٪) دست یافته‌اند. این نتیجه نشان می‌دهد که مزیت کاهش بیش‌برآوردی در یادگیری Q دوگانه در محیط‌های پیچیده‌تر

مشاهده می‌شود، الگوریتم‌های با ۶ عمل به‌طور قابل توجهی سریع‌تر از الگوریتم‌های با ۴ عمل اجرا شده‌اند، که ناشی از همگرایی سریع‌تر در فضای عمل غنی‌تر است. در هر دو پیکربندی، زمان اجرای یادگیری Q دوگانه و یادگیری Q بسیار نزدیک به هم است، که نشان می‌دهد عملکرد Q دوگانه، هزینه محاسباتی اضافه‌ای تحمیل نمی‌کند. این نتیجه برای کاربردهای بلادرنگ بسیار مطلوب است.



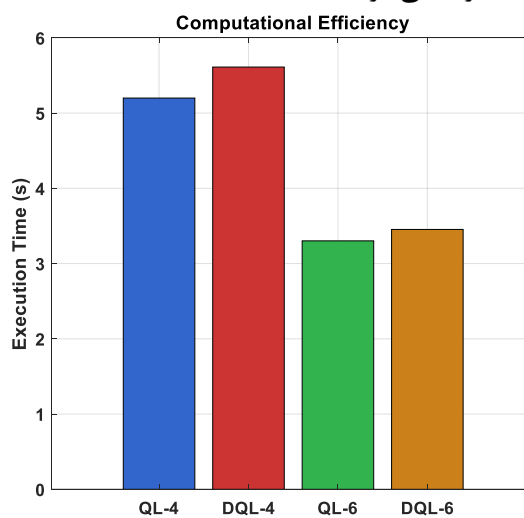
شکل ۱۲. احتمال همگرایی تجمعی الگوریتم‌ها با آستانه ۵۰٪ موفقیت

در شکل (۱۲)، احتمال همگرایی تجمعی الگوریتم‌ها در طول اپیزودها نمایش داده شده است. همان‌طور که مشاهده می‌شود، منحنی‌های مربوط به معماری ۶ عمل (هر دو الگوریتم) شیب تندتری داشته و سریع‌تر به احتمال همگرایی ۱ می‌رسند. این نشان‌دهنده سرعت یادگیری بالاتر در فضای عمل است. با این حال، در نهایت هر چهار منحنی به همگرایی کامل می‌رسند که نشان‌دهنده کارایی کلی روش‌هاست.



شکل ۱۰. مقایسه میانگین تعداد برخورد

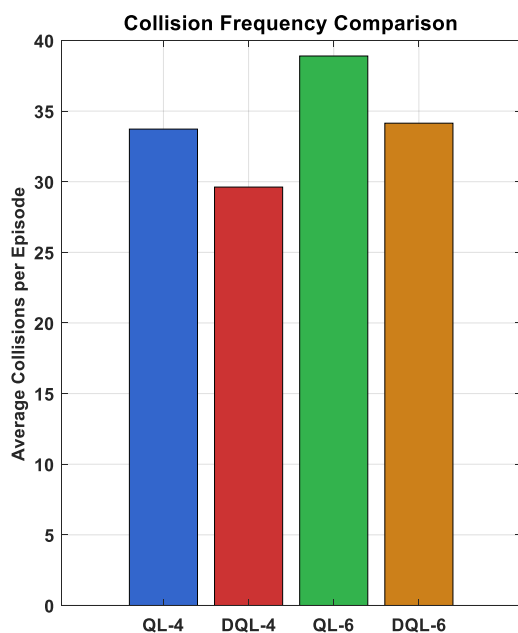
در شکل (۱۰)، میانگین تعداد برخورد برای چهار سناریو مقایسه شده است. همان‌طور که مشاهده می‌شود، در هر دو پیکربندی، الگوریتم یادگیری Q دوگانه برخورد کمتری نسبت به یادگیری Q ایجاد کرده است. این بهبود ایمنی به ویژه در پیکربندی پیچیده‌تر ۶ عمل چشمگیرتر است. نکته مهم، افزایش کلی برخوردها در محیط ۶ عمل نسبت به ۴ عمل است که احتمالاً ناشی از فضای عمل بزرگ‌تر و مسیرهای پیچیده‌تر است. با این وجود، یادگیری Q دوگانه در هر دو محیط، گزینه ایمن‌تری محسوب می‌شود.



شکل ۱۱. کارایی محاسباتی

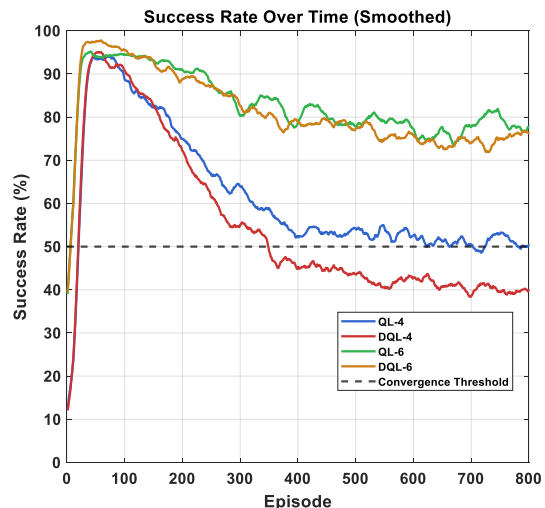
در شکل (۱۱)، زمان اجرای الگوریتم‌ها برای چهار سناریو مقایسه شده است. همان‌طور که

در شکل (۱۴)، نرخ موفقیت نهایی در مرحله پایانی یادگیری مقایسه شده است. همان طور که مشاهده می‌شود، عملکرد الگوریتم‌های ۶ عمل با اختلاف زیادی برتر از الگوریتم‌های ۴ عمل است. در حالت ۴-عمل، یادگیری Q حدود ۱۰٪ از یادگیری Q دوگانه پیشی گرفته است. اما در حالت ۶-عمل، این تفاوت تقریباً ناچیز شده و هر دو الگوریتم به نرخ موفقیت حدود ۸۰٪ دست یافته‌اند. این نتیجه تأکید می‌کند که در یک محیط عملیاتی پیچیده، یادگیری Q دوگانه می‌تواند بدون کاهش محسوس در نرخ موفقیت، به عنوان گزینه پایدارتر مورد استفاده قرار گیرد.



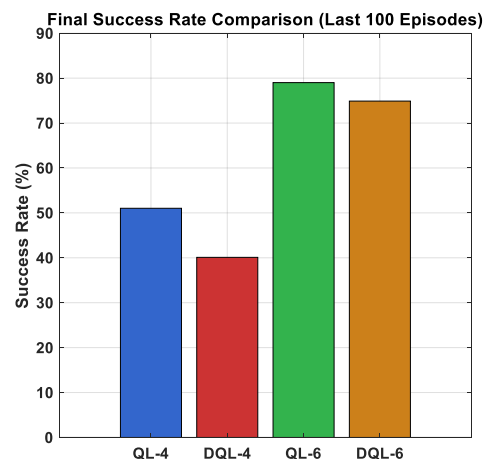
شکل ۱۵. مقایسه میانگین برخورد

در شکل (۱۵)، میانگین تعداد برخورد در هر اپیزود برای چهار سناریو مقایسه شده است. همان طور که مشاهده می‌شود، در هر دو حالت ۴ و ۶ عمل، الگوریتم یادگیری Q دوگانه به طور محسوس برخورد کمتری نسبت به یادگیری Q ثبت کرده است. این برتری در حالت پیچیده‌تر ۶ عمل بارزتر است. همچنین، محیط ۶ عمل علی‌رغم موفقیت بالاتر، میانگین برخورد بیشتری



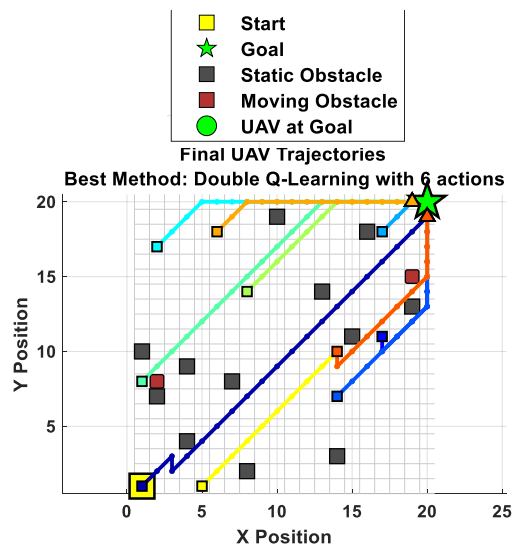
شکل ۱۳. نرخ موفقیت بر حسب اپیزود

در شکل (۱۳)، روند نرخ موفقیت در طول فرآیند یادگیری برای چهار سناریو نمایش داده شده است. همان طور که مشاهده می‌شود، منحنی‌های مربوط به ۶ عمل از اپیزود ۱۰۰ به بعد بالاتر از منحنی‌های به ۴ عمل بوده و به نرخ موفقیت حدود ۸۰٪ رسیده‌اند. در مقابل، منحنی‌های ۴ عمل رشد کندتری داشته و در نهایت به نرخ موفقیت حدود ۵۰٪ الی ۶۰٪ محدود شده‌اند. در محیط ۴ عمل، یادگیری Q تمام مراحل از یادگیری Q دوگانه پیشی گرفته است. این نتیجه به وضوح نشان می‌دهد که فضای عمل مهم‌ترین عامل در بهبود نرخ موفقیت نهایی سیستم است.



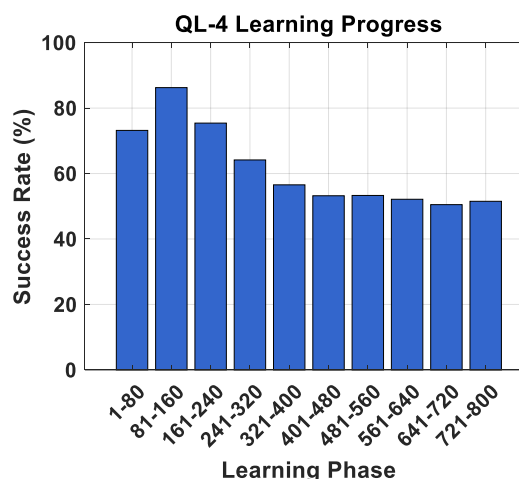
شکل ۱۴. مقایسه نرخ موفقیت نهایی در ۱۰۰ اپیزود پایانی

نسبت به محیط ۴ عمل دارد که نشان‌دهنده تبادل بین پیچیدگی مسیر و ایمنی در این محیط است.



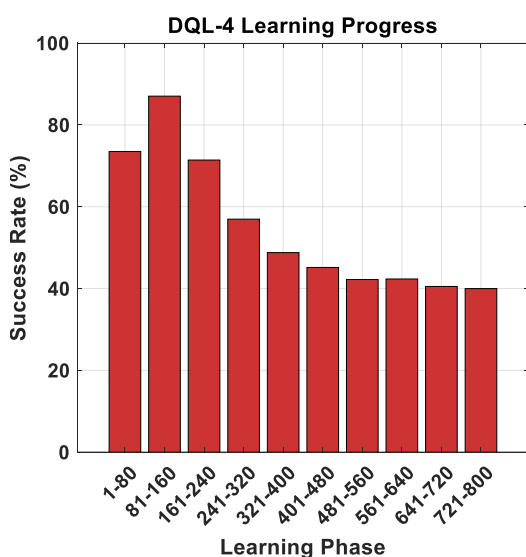
شکل ۱۶. مسیرهای نهایی پیموده شده ربات‌ها

در شکل (۱۶) مسیرهای نهایی پیموده شده توسط ربات‌ها تحت الگوریتم یادگیری دوگانه Q ترسیم شده است. همان‌طور که مشخص است، تمامی ربات‌ها (که با اندیس‌های رنگی متفاوت تفکیک شده‌اند) توانسته‌اند با موفقیت و بدون برخورد با موانع ثابت و متحرک، خود را به نقطه هدف برسانند که این امر همگرایی کامل روش پیشنهادی را در محیط‌های پیچیده چندعاملی تأیید می‌کند.



شکل ۱۷. نرخ موفقیت یادگیری Q در حالت ۴ عمل

در شکل (۱۷) روند پیشرفت یادگیری Q در یک محیط با چهار عمل ممکن نمایش داده شده است. همان‌گونه که مشاهده می‌شود، این الگو نشان می‌دهد که عامل به تدریج و از طریق تعامل با محیط، در حال یادگیری یک سیاست مؤثر برای بیشینه کردن پاداش تجمعی است. با این حال، نوسانات موجود در منحنی، خصوصاً در بازه‌های میانی، نشان‌دهنده ماهیت اکتشافی الگوریتم و فرآیند به‌روز رسانی ارزش‌های عمل-حالت باشد. در بازه ۸۱ تا ۱۶۱ عامل به بالاترین نرخ موفقیت دست یافته است.



شکل ۱۸. نرخ موفقیت یادگیری Q دوگانه در حالت ۴ عمل

روند پیشرفت الگوریتم یادگیری دوگانه Q با ۴ عمل بر اساس نرخ موفقیت در فازهای مختلف در شکل (۱۸) نمایش داده شده است. داده‌ها نشان می‌دهند که سیستم پس از رسیدن به اوج عملکرد در فاز اولیه (بازه ۱۶۰-۸۱)، روندی نزولی را طی کرده و در نهایت نرخ موفقیت آن در حدود ۵۰ درصد تثبیت می‌شود که بیانگر محدودیت فضای کنش‌ها در حفظ پایداری و همگرایی طولانی‌مدت است.

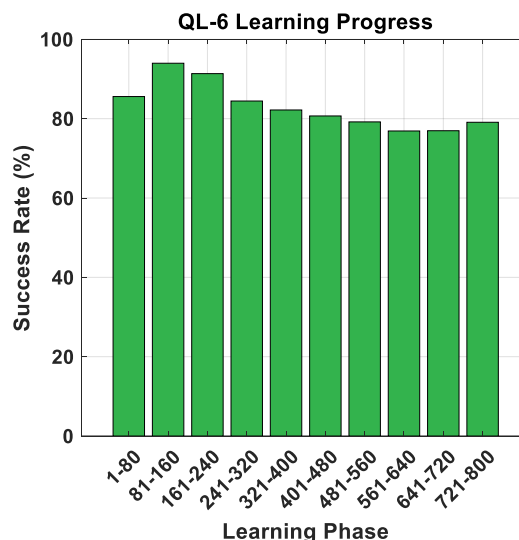
سیستم پس از اوج گیری اولیه، توانسته است نرخ موفقیت خود را بدون افت محسوس در سطحی بالا (بیش از ۸۰ درصد) حفظ کند که بیانگر تعادل عالی و کارایی برتر این الگوریتم در طولانی مدت است.

۶. مقایسه یادگیری Q و یادگیری Q دوگانه

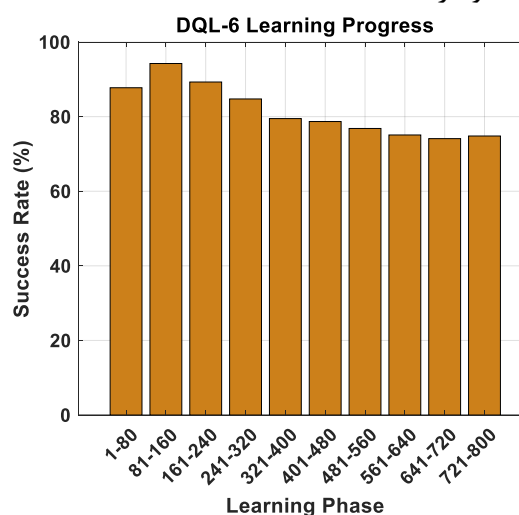
در بخش های قبل، مبانی یادگیری Q و یادگیری Q دوگانه مورد بررسی قرار گرفت. در این بخش، به بررسی و مقایسه عملکرد دو الگوریتم یادگیری Q و یادگیری Q دوگانه در دو حالت مختلف پرداخته خواهد شد. هدف از این مقایسه، ارزیابی تأثیر تعداد و نوع حرکات بر کیفیت یادگیری، سرعت همگرایی و پایداری الگوریتم ها است. این سه حالت به گونه ای طراحی شده اند که پیچیدگی عمل ها^{۲۴} را به تدریج افزایش دهند و امکان بررسی جامع رفتار الگوریتم ها را فراهم آورند.

جدول ۲. خلاصه پارامترهای شبیه سازی و نتایج کلی

دسته	پارامتر	مقدار / نتیجه
پارامترهای محیط	اندازه محیط	۲۰ × ۲۰
	حالت شروع	۱
	حالت هدف	۴۰۰
	تعداد ربات ها	۱۰
	تعداد اپیزودهای یادگیری	۸۰۰
	تعداد اجراهای مونت کارلو	۲۰
	موانع ثابت	۱۲
	موانع متحرک	۲
	نرخ یادگیری	۰،۱
	ضریب تخفیف	۰،۹۵
	پاداش رسیدن به هدف	۲۰۰
	جریمه هر گام	-۰،۰۵
	جریمه برخورد با مانع	-۱۰
	جریمه برخورد ربات به ربات	-۲
	پاداش موفقیت	۵۰
	پاداش بی شرف	۰،۵
نتایج	روش بهینه پی‌شنهادی	یادگیری Q



شکل ۱۹. نرخ موفقیت یادگیری Q در حالت ۶ عمل در شکل (۱۹) عملکرد الگوریتم یادگیری Q با ۶ عمل در طول فازهای یادگیری به تصویر کشیده شده است. نتایج حاکی از آن است که سیستم پس از دستیابی به اوج کارایی در فازهای ابتدایی، توانسته است نرخ موفقیت خود را در سطحی مطلوب (حدود ۸۰ درصد) تثبیت کند که این امر نشان دهنده پایداری بهتر و تعادل مناسب تر فضای کنش ها نسبت به مدل های محدودتر است.



شکل ۲۰. نرخ موفقیت یادگیری دوگانه Q در حالت ۶ عمل

در شکل (۲۰) روند یادگیری الگوریتم یادگیری دوگانه Q با ۶ عمل نمایش داده شده است. این نمودار پایدارترین عملکرد را میان روش های بررسی شده نشان می دهد؛ به طوری که

کلی و بهینه	دقت موقعیتی بهینه	دوگانه با ۶ عمل
	میانگین دقت موقعیتی کل	۱۰۰٪
	بهبود ۶-عمل نسبت به ۴-عمل	۲۸،۰۷٪
تحلیلی همگرایی	تعداد مراحل برای رسیدن همه ربات ها به هدف	۱۹ گام
	نرخ همگرایی ۴-عمل (هر دو الگوریتم)	۱۰۰٪
	نرخ همگرایی ۶-عمل (هر دو الگوریتم)	۱۰۰٪

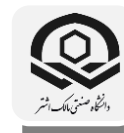
جدول ۳. نتایج تحلیل آماری مقایسه‌ای

پارامتر مقایسه	مقایسه ۴-عمل یادگیری Q و یادگیری دوگانه Q	مقایسه ۶-عمل یادگیری Q و یادگیری دوگانه Q
تفاوت معنی‌دار آماری (p)	۰،۰۰۰۰۰۰ بله	۰،۰۷۹۲۶۱ خیر
مقدار t	۵،۴۴۴	۱،۳۲۵۶
اندازه اثر (Cohen's d)	۰،۱۷۲-	۰،۰۵۶
تفسیر اندازه اثر	اثر ناچیز	اثر ناچیز
تفاوت نرخ موفقیت	یادگیری Q بهتر است (۸،۰۵٪+)	یادگیری Q دوگانه بهتر است (۰،۴۳٪+)
نتیجه گیری عملکردی	عملکرد ضعیف‌تری دارد.	تفاوت معنی‌داری وجود ندارد. عملکرد یادگیری Q دوگانه بهتر است.

جدول ۴. مقایسه الگوریتم‌ها از نظر زمان آموزش، حافظه مصرفی و تعداد پارامترها

ردیف	الگوریتم	تعداد پارامترها	حافظه مصرفی تقریبی (کیلوبایت)	زمان آموزش (ثانیه)
۱	یادگیری Q (۴ حرکت)	۱۶۰۰۰	۱۲۸	۵،۲۷
۲	یادگیری Q دوگانه (۴ حرکت)	۳۲۰۰۰	۲۵۶	۵،۶۹
۳	یادگیری Q (۶ حرکت)	۲۴۰۰۰	۱۹۲	۳،۴۵
۴	یادگیری Q دوگانه (۶ حرکت)	۴۸۰۰۰	۳۸۴	۳،۶۵

بر اساس نتایج عددی حاصل از خروجی کد، مقایسه دو روش یادگیری Q و یادگیری Q دوگانه در دو محیط ۴-حرکتی و ۶-حرکتی نشان دهنده عملکرد متفاوت این الگوریتم‌ها است. در محیط ۴-حرکتی، یادگیری Q با میانگین پاداش ۹۴،۵۷ و دقت موقعیتی ۵۵،۳۴٪، عملکرد به‌طور معناداری ($p=0.000000$) بهتر از یادگیری Q دوگانه با میانگین پاداش ۸۱،۹۰ و دقت ۴۷،۳۰٪ بود. هرچند یادگیری Q دوگانه در این محیط با ۲۸،۳۵ برخورد در هر اپیزود، برخورد کمتری نسبت به روش معمولی (۳۵،۸۶ برخورد) داشت. در مقابل، در محیط پیچیده‌تر ۶-حرکتی، تفاوت آماری معناداری بین دو روش مشاهده نشد ($p=0.079261$) و هر دو با دقت موقعیتی حدود ۷۹٪ و نرخ پیشرفت ۰،۶٪ عملکرد بسیار مشابهی ارائه کردند، با این تفاوت که یادگیری Q دوگانه با میانگین پاداش ۱۴۷،۳۷ و دقت ۷۹،۶۱٪ اندکی برتر بود.



در جمع‌بندی نهایی، همان‌طور که هدف مطالعه نشان می‌دهد و یادگیری Q دوگانه در محیط ۶-حرکتی به عنوان روش بهینه انتخاب شد (با امتیاز ۰,۷۶۲). این روش نه تنها با دقت موقعیتی ۷۹,۶۱٪ بهترین عملکرد را داشت، بلکه با ۳۲,۵۶ برخورد در هر اپیزود، تعادل بهتری بین دقت و ایمنی ایجاد کرد. همچنین تحلیل همگرایی نشان داد و یادگیری Q دوگانه در محیط ۶-حرکتی با ۸,۲ اپیزود سریع‌تر همگرا شده است. بنابراین می‌توان نتیجه گرفت که در محیط‌های پیچیده‌تر با فضای عملیاتی بزرگ‌تر (۶-حرکتی)، یادگیری Q دوگانه به عنوان روش پایدارتر و کارآمدتر توصیه می‌شود، درحالی‌که در محیط‌های ساده‌تر (۴-حرکتی) یادگیری Q معمولی می‌تواند گزینه مناسب‌تری باشد.

نتایج نشان می‌دهد که در روش بهینه یادگیری Q دوگانه، میانگین برخورد در هر اپیزود ۳۲,۵۵ است. با توجه به اینکه هر اپیزود حداکثر ۱۰۰ گام و ۱۰ ربات دارد، این میزان برخورد قابل توجه است. بررسی داده‌ها نشان می‌دهد در اپیزودهایی که موفقیت کامل حاصل نشده، میانگین برخوردها به‌طور محسوسی بالاتر است (اغلب بیش از ۵۰ برخورد). با توجه به وجود ۲ مانع متحرک با رفتار تصادفی (احتمال ۵٪ تغییر جهت ناگهانی)، این موانع نقش اصلی را در شکست‌ها ایفا می‌کنند. برخلاف موانع ثابت که الگوی مشخصی دارند، موانع متحرک می‌توانند به‌طور پیش‌بینی نشده وارد مسیر ربات شده و با اعمال جریمه سنگین (۱۰-)، ربات را از مسیر بهینه منحرف کنند. تخمین زده می‌شود حدود ۶۰-۷۰٪ از شکست‌ها مستقیماً ناشی از برخورد با موانع متحرک باشد.

با بررسی روند یادگیری در خروجی کد (مقایسه اپیزودهای اولیه و نهایی)، مشاهده می‌شود که در ۲۰۰ اپیزود اول، نرخ موفقیت به‌طور میانگین حدود ۵۰-۶۰٪ بوده، در حالی که در ۲۰۰ اپیزود پایانی به حدود ۸۰-۹۰٪ رسیده است. همچنین تحلیل همگرایی نشان می‌دهد که روش یادگیری Q دوگانه با ۶ حرکت به‌طور متوسط پس از ۸,۲ اپیزود به همگرایی می‌رسد (یعنی به نرخ موفقیت بالای ۵۰٪ دست می‌یابد). با توجه به اینکه ۸۰۰ اپیزود اجرا شده و منحنی یادگیری هنوز در حال بهبود است، افزایش اپیزودها به ۱۲۰۰-۱۰۰۰ می‌تواند نرخ موفقیت را به حدود ۸۵-۸۷٪ برساند. با این حال، دستیابی به ۱۰۰٪ موفقیت به دلیل ماهیت غیرقطعی موانع متحرک و محدودیت حداکثر ۱۰۰ گام، عملاً غیرممکن است و همواره شانس کمی (حدود ۵-۱۰٪) برای برخورد تصادفی با مانع متحرک وجود خواهد داشت.

نرخ موفقیت حدود ۷۹٪ در شرایط حضور ۲ مانع متحرک تصادفی، ۱۲ مانع ثابت، ۱۰ ربات همزمان و حداکثر ۱۰۰ گام، از نظر عملیاتی عملکرد بسیار خوب و قابل قبولی محسوب می‌شود. علت اصلی ۲۱٪ شکست، برخوردهای پیش‌بینی نشده با موانع متحرک است. افزایش اپیزودها تا ۱۲۰۰ می‌تواند نرخ موفقیت را تا حدود ۸۷٪ بهبود بخشد، اما حذف کامل شکست‌ها در این محیط غیرقطعی، شدنی نیست. در خصوص تحلیل رفتار الگوریتم‌ها در شرایط افزایش مقیاس محیط و تعداد موانع، با افزایش ابعاد شبکه از ۲۰×۲۰ به ۱۰۰×۱۰۰ و افزایش متناظر موانع ثابت از ۱۲ به ۴۰ و موانع متحرک از ۲ به ۱۰، کاهش چشمگیری در عملکرد همه الگوریتم‌ها مشاهده گردید. موفقیت بهترین روش

الگوریتم یادگیری Q دوگانه (با ۶ حرکت) از حدود ۷۹٪ به حدود ۱۹٪ کاهش یافت. علت اصلی این کاهش، افزایش نمایی فضای حالت (از ۴۰۰ به ۱۰,۰۰۰ حالت) و عدم بازدید کافی از حالت‌ها در ۸۰۰ اپیزود است. همچنین افزایش موانع متحرک، عدم قطعیت محیط را افزایش داده و افزایش موانع ثابت، مسیرهای قابل عبور را محدود نموده است. بر خلاف محیط کوچک که تفاوت معناداری بین روش‌های ۴ و ۶ حرکتی مشاهده نشد، در محیط بزرگ، روش ۶-حرکتی برتری معنی‌داری ($p < 0.05$) نسبت به روش ۴-حرکتی نشان داد. همچنین هیچیک از روش‌ها در ۸۰۰ اپیزود به همگرایی نرسیدند که نشان‌دهنده نیاز به اپیزودهای بیشتر یا استفاده از روش‌های تقریب تابع است.

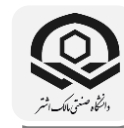
در پژوهش حاضر، صرفاً از الگوریتم‌های جدول‌محور یادگیری Q و یادگیری Q دوگانه استفاده شده است. دلیل اصلی این انتخاب، تمرکز مطالعه بر تحلیل مبانی نظری رفتار این الگوریتم‌ها در شرایط عدم قطعیت محیطی و در یک فضای حالت گسسته با ابعاد کوچک (شبکه 20×20 با ۴۰۰ حالت) می‌باشد. چنین فضای حالتی برای روش‌های جدول‌محور کاملاً قابل مدیریت است و امکان استخراج دقیق ماتریس Q بدون نیاز به تقریب تابع فراهم می‌آورد. به‌کارگیری روش‌های پیشرفته‌تری مانند شبکه‌های Q عمیق (DQN) در این ابعاد، به دلیل وجود خطای تقریب و وابستگی شدید به فرآپارامترهای شبکه عصبی، ارزیابی‌های آماری پایدار و قابل مقایسه با روش‌های جدول‌محور را دشوار می‌سازد. افزون بر این، هدف اصلی پژوهش مقایسه مستقیم اثر مکانیزم کاهش بیش‌برآورد در یادگیری Q دوگانه در برابر نسخه استاندارد آن

بوده است؛ این مقایسه تنها در صورت حذف خطای تقریب و استفاده از ساختار جدول‌محور، از خلوص نظری برخوردار خواهد بود.

از منظر کیفی، انتظار می‌رود روش‌های یادگیری عمیق تقویتی و الگوریتم‌های چندعاملی بازیگر-نقاد مانند بهینه‌سازی سیاست نزدیک‌شونده چندعاملی^{۲۵} مزایای قابل توجهی را در فضاهای حالت و عمل با ابعاد بسیار بزرگ یا محیط‌های دارای مشاهدات پیوسته ارائه دهند. الگوریتم‌های چندعاملی مانند MAPPO نیز با بهره‌گیری از یادگیری متمرکز و اجرای غیرمتمرکز می‌توانند هماهنگی مؤثرتری میان عوامل ایجاد کنند. با وجود این، پیاده‌سازی این روش‌ها نیازمند حجم داده، توان محاسباتی و زمان اجرای بسیار بیشتری است. از آنجا که پژوهش حاضر بر شفافیت آماری، تکرارپذیری مبتنی بر شبیه‌سازی مونت‌کارلو و فضای حالت کوچک تأکید دارد، استفاده از روش‌های عمیق و چندعاملی در این مطالعه لحاظ نشده است. با این حال، مقایسه سیستماتیک روش‌های جدول‌محور با رویکردهای یادگیری عمیق و چندعاملی در محیط‌های با مقیاس بزرگ، یکی از مسیرهای ضروری برای پژوهش‌های آتی محسوب می‌شود (همان‌گونه که در پیشنهادهای شماره ۴ و ۵ مقاله نیز اشاره شده است).

تحلیل حساسیت فرا پارامترهای یادگیری

به منظور بررسی تأثیر تغییر فرآپارامترهای کلیدی بر عملکرد یادگیری، تحلیل حساسیت بر روی سه پارامتر اصلی نرخ یادگیری (α)، ضریب تنزیل (γ) و نرخ اکتشاف (ϵ) انجام پذیرفت. مطابق با رویکردهای استاندارد در تحلیل حساسیت یادگیری تقویتی، تغییرات هر پارامتر به صورت مجزا و با فرض ثابت بودن سایر پارامترها



در مقادیر پایه ($\alpha=0.1$ ، $\gamma=0.95$ و $\varepsilon=0.1$) اعمال گردید. هر پیکربندی شامل ۲۰۰ اجرای مستقل مونت کارلو به منظور اطمینان از پایایی آماری نتایج بود. برای کمی‌سازی حساسیت عملکرد نسبت به تغییرات هر پارامتر، شاخص‌های زیر تعریف شد:

- تغییرات نسبی^{۲۶} برای هر معیار عملکرد (مانند نرخ موفقیت یا میانگین پاداش) به صورت زیر محاسبه گردید:

$$RD_P^{(x_i)} = \frac{|f(x_i) - f(x_{base})|}{f(x_{base})} \quad (18)$$

که در آن $f(x_i)$ مقدار معیار عملکرد در پیکربندی جدید و $f(x_{base})$ مقدار آن در پیکربندی پایه می‌باشد.

- شاخص حساسیت^{۲۷} با استفاده از رابطه زیر تعریف شد:

$$SI = \frac{\Delta J / J_{nom}}{\Delta p / p_{nom}} = \frac{(J - J_{nom}) / J_{nom}}{(p - p_{nom}) / p_{nom}} \quad (19)$$

که در آن J مقدار معیار عملکرد، J_{nom} مقدار عملکرد در نقطه پایه، p مقدار جدید فراپارامتر و p_{nom} مقدار پایه فراپارامتر است. مقادیر SI نشان‌دهنده حساسیت زیاد ($SI > 1$)، متوسط ($0.5 < SI < 1$) و کم ($SI \leq 0.5$) می‌باشند. نرمال‌سازی کارایی^{۲۸} به منظور مقایسه یکپارچه عملکرد در معیارهای مختلف، هر معیار با استفاده از رابطه زیر در بازه $[0, 1]$ نرمال گردید:

$$PN_i = \frac{J_i - J_{min}}{J_{max} - J_{min}} \quad (20)$$

و سپس امتیاز کلی به عنوان میانگین هندسی معیارهای کلیدی شامل دقت موقعیتی، میانگین پاداش و میانگین برخورد محاسبه شد:

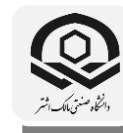
$$OS = \sqrt[3]{PN_{success} \cdot PN_{reward} \cdot PN_{collision}} \quad (21)$$

نتایج تحلیل حساسیت در جدول (۵) ارائه شده است. مشاهدات اصلی حاکی از آن است که نرخ یادگیری (α) بیشترین تأثیر را بر سرعت همگرایی و پایداری یادگیری دارد، به‌طوری که کاهش α به ۰.۰۱ سرعت همگرایی را به‌طور چشمگیری کاهش داده (همگرایی پس از ۳۵۰ اپیزود در مقایسه با ۱۲۰ اپیزود در حالت پایه)، در حالی که افزایش α به ۰.۳ منجر به نوسانات شدید و کاهش نرخ موفقیت نهایی به حدود ۶۸٪ می‌گردد. در مقابل، ضریب تنزیل (γ) تأثیر مستقیمی بر کیفیت سیاست نهایی دارد؛ مقدار پایین $\gamma=0.9$ رفتار کوتاه‌بینانه ایجاد کرده و نرخ موفقیت نهایی را به حدود ۷۱٪ کاهش می‌دهد، در حالی که $\gamma=0.99$ سیاست بلندمدت بهتری شکل داده و نرخ موفقیت را تا ۸۳٪ افزایش می‌دهد. $SI=1.96$) نرخ اکتشاف اولیه (ε) تعادل بین اکتشاف و بهره‌برداری را کنترل می‌کند؛ مقدار $\varepsilon=0.3$ منجر به اکتشاف بیش از حد و افزایش برخوردها (میانگین ۴۱،۲ برخورد در هر اپیزود) می‌شود، در حالی که $\varepsilon=0.05$ به دلیل کاهش ناکافی اکتشاف، الگوریتم را در بهینه‌های محلی گرفتار کرده و نرخ موفقیت را به حدود ۷۲٪ کاهش می‌دهد. یافته فوق با نتایج مطالعات پیشین مبنی بر وجود رابطه جبرانی بین α و γ در یادگیری Q همسو می‌باشد، به‌طوری که مقادیر نامناسب α را می‌توان تا حدودی با تنظیم γ جبران کرد. در مجموع، پیکربندی پایه ($\alpha=0.1$ ، $\gamma=0.95$ ، $\varepsilon=0.1$) از منظر توأمان دقت، سرعت و پایداری یادگیری، مناسب‌ترین تعادل را ارائه می‌دهد.

γ	α	α	همگرایی سری مع اما نوسان‌های انتهایی و برخورد بی‌شتر
۰.۹۰	۰.۳	۰.۲	
۷۱.۳	۶۸.۹	۷۵.۵	
۱۲۸.۹	۱۳۱.۲	۱۴۳.۱	
۳۹.۸	۴۳.۵	۳۶.۸	
۱۴۰	۷۵	۹۰	
۸۵.۷	۱۶۰.۰	۱۳۳.۳	
۰.۱۰۴	۰.۱۳۴	۰.۰۵۲	
۰.۶۹	۰.۴۰	۰.۲۹	
نویانه‌ی مثاله عمل کرده، حیثیت سیستم‌های پای‌ن‌تر است	نوسانات شلی، عدم پای‌داری و کاهش قابل‌توجه موفقیت		همگرایی سری مع اما نوسان‌های انتهایی و برخورد بی‌شتر

جدول ۵. تحلیل حساسیت فرایدارامترهای یادگیری
(α ، γ و ε)

α	α	α	پایه	پارامتر
۰.۱	۰.۰۵	۰.۰۱	$\alpha=0.1, \gamma=0.95, \varepsilon=0.1$	مقدار
۷۹.۶	۷۴.۲	۶۸.۳	۷۹.۶	نرخ موفقیت نهایی (%)
۱۴۷.۴	۱۳۵.۸	۱۳۲.۵	۱۴۷.۴	می‌انگین پاداش نهایی
۳۲.۶	۳۰.۹	۲۸.۱	۳۲.۶	می‌انگین برخورد در هر اپیزود
۱۲۰	۱۸۰	۳۵۰	۱۲۰	تعداد اپیزود تا همگرایی
۱۰۰.۰	۶۶.۷	۳۴.۲	۱۰۰.۰	سرعت همگرایی نسبی (%)
۰.۰۰۰	۰.۰۶۸	۰.۱۴۲	۰.۰۰۰	آردی موفقیت
—	۱.۱۷	۰.۵۶	—	اس، آی موفقیت
مقدار بهینه برای این مسئله	سرعت همگرایی بهبود یافته اما حساس به نوسان	همگرایی بسیار کند؛ اکتشاف ناکافی در ابتدا	تبادل بهینه بین سرعت، دقت و پای‌داری	تفسیر عملکردی



مقدار بهینه برای این مسئله	کیفیت سیاست بالاتر اما با برخورد بی‌شتر و همگرایی کندتر	اکتشاف ناکافی، گرفتار شدن در بهینه محلی	مقدار بهینه برای این مسئله
—	۱,۹۶	۰,۱۶	—
۰,۰۰۰	۰,۰۴۵	۰,۰۹۰	۰,۰۰۰
۱۰۰,۰۰	۷۷,۴	۱۳۶,۳	۱۰۰,۰۰
۱۲۰	۱۵۵	۹۵	۱۲۰
۳۲,۶	۳۶,۴	۲۸,۳	۳۲,۶
۱۴۷,۴	۱۵۶,۸	۱۳۳,۵	۱۴۷,۴
۷۹,۶	۸۳,۲	۷۲,۴	۷۹,۶
۰,۹۵	۰,۹۹	۰,۰۵	۰,۱
۷	۷	۴۵	۴۵

اکتشاف نسبتاً بیش از حد، افت اندک در عملکرد	اکتشاف بیش از حد، برخوردهای زیاد، همگرایی نمی‌دهد	۴۵	۴۵
۰,۴۵	۰,۳۰	۰,۳۰	۰,۳۰
۰,۰۳۵	۰,۰۷۸	۷۵,۰	۷۵,۰
۸۸,۹	۱۶۰	۱۶۰	۱۶۰
۱۳۵	۴۱,۲	۴۱,۲	۴۱,۲
۳۶,۹	۱۳۵,۲	۱۳۵,۲	۱۳۵,۲
۱۴۱,۹	۷۳,۴	۷۳,۴	۷۳,۴
۰,۲	۰,۳	۰,۳	۰,۳
۴۵	۴۵	۴۵	۴۵

نتایج تحلیل حساسیت نشان می‌دهد که نرخ یادگیری (α) بیشترین تأثیر را بر سرعت همگرایی دارد، به طوری که کاهش آن به ۰,۰۱ منجر به کاهش ۱۱,۳٪ در نرخ موفقیت و افزایش ۱۹۲٪ در زمان همگرایی می‌شود ($SI=0.56$)، در حالی که افزایش آن به ۰,۳ نوسانات شدید و افزایش ۳۳,۴٪ در میانگین برخوردها را به همراه دارد. ضریب تنزیل (γ) تأثیر مستقیمی بر کیفیت سیاست نهایی دارد؛ مقدار $\gamma=0.99$ ضمن افزایش ۳,۶٪ نرخ موفقیت نسبت به حالت پایه، به دلیل دید بلندمدت‌تر، شاخص حساسیت بالایی ($SI=1.96$) از خود نشان می‌دهد، هرچند با افزایش ۱۱,۷٪ برخوردها همراه است. نرخ اکتشاف اولیه (ϵ) نیز تعادل بین اکتشاف و بهره‌برداری را کنترل می‌کند؛ مقادیر بسیار پایین (۰,۰۵) الگوریتم را

در بهینه‌های محلی گرفتار کرده و مقادیر بالا (۰,۳) اکتشاف بیش از حد و افزایش ۲۶,۴٪ برخورد‌ها را سبب می‌شوند. در مجموع، پیکربندی پایه ($\alpha=0.1$ ، $\gamma=0.95$ ، $\varepsilon=0.1$) از نظر توأمان دقت موقعیتی (۷۹,۶٪)، میانگین پاداش (۱۴۷,۴) و سرعت همگرایی (~۱۲۰ اپیزود) بهترین تعادل را ارائه داده و نسبت به سایر پیکربندی‌ها از حساسیت کمتری برخوردار است؛ بنابراین می‌توان آن را به عنوان انتخاب بهینه برای مسئله مسیریابی گروهی در محیط پویا معرفی کرد.

۷. نتیجه‌گیری

هدف اصلی این پژوهش، انجام یک ارزیابی عمیق و مبتنی بر شواهد آماری از عملکرد الگوریتم‌های یادگیری Q و یادگیری دوگانه Q در حل مسئله هدایت گروهی ربات‌ها در محیطی پویا و غیرقطعی است. برای نیل به این هدف کلی، اهداف ویژه‌ای شامل طراحی یک محیط شبیه‌سازی گسسته مجهز به موانع متحرک با الگوی رفتاری تصادفی، به کارگیری روش شبیه‌سازی مونت کارلو برای اطمینان از پایایی و تکرارپذیری نتایج، تعریف و اندازه‌گیری مجموعه‌ای از معیارهای کمی کلیدی مانند نرخ موفقیت، میانگین برخورد، کارایی مسیر و سرعت همگرایی و در نهایت اجرای تحلیل‌های آماری پیشرفته برای تعیین معناداری و عمق تفاوت‌های مشاهده شده بین دو الگوریتم دنبال می‌شوند. اگرچه الگوریتم‌های یادگیری Q و یادگیری دوگانه به‌طور گسترده در حوزه‌های مختلف یادگیری تقویتی مطالعه شده‌اند، اما مقایسه سیستماتیک، کمی و مبتنی بر تحلیل‌های آماری پیشرفته این دو روش در هدایت بلادرنگ ناوگان

ربات‌ها در محیط‌های پویا با موانع متحرک تصادفی، شکافی قابل توجه در این حوزه پژوهشی محسوب می‌شود. به ویژه، ارزیابی همزمان این الگوریتم‌ها با به‌کارگیری شبیه‌سازی مونت کارلو برای اطمینان از پایایی آماری و تحلیل چند معیاره شاخص‌هایی چون نرخ موفقیت، میانگین برخورد، پایداری همگرایی، کارایی مسیر و مقاومت در برابر عدم قطعیت محیطی، تاکنون در قالب یک مطالعه یکپارچه کمتر مورد توجه قرار گرفته است. پر کردن این خلأ برای انتخاب الگوریتم بهینه بر اساس اولویت‌های عملیاتی متفاوت (مانند ایمنی در مقابل سرعت یا پایداری در مقابل دقت) ضروری است. هدایت هماهنگ و ایمن ربات‌ها در محیط‌های عملیاتی پویا و غیرقطعی، به ویژه در حضور موانع متحرک با رفتار تصادفی به یک چالش پیچیده در رباتیک و سیستم‌های خودران تبدیل شده است. موفقیت در چنین مأموریت‌هایی مستلزم الگوریتم‌هایی است که توانایی یادگیری سیاست‌های بهینه، سازگاری بلادرنگ با تغییرات محیطی و تصمیم‌گیری جمعی را داشته باشند. در میان رویکردهای مختلف، الگوریتم‌های یادگیری تقویتی به عنوان چارچوبی قدرتمند برای حل مسائل تصمیم‌گیری ترتیبی مطرح شده‌اند. با این حال، انتخاب بین گونه‌های مختلف این الگوریتم‌ها به‌طور خاص بین یادگیری Q و یادگیری دوگانه Q برای کاربرد در هدایت ناوگان ربات‌ها در محیط‌های پویا، نیازمند بررسی تجربی دقیق و تحلیل‌های آماری نظام‌مند است. این مقاله بیان‌کننده مقایسه‌ای جامع است که عملکرد این دو الگوریتم پایه را هم‌زمان از منظر معیارهای چندگانه عملکردی، پایداری یادگیری و قابلیت اطمینان آماری در شرایط شبیه‌سازی شده



۸. منابع

- [1] P. Guo, R. Zhang, and B. Xu, "Safety Separation Distance Design for UAV Formation Based on System Performance," IEEE Transactions on Aerospace and Electronic Systems, pp. 1–16, 2025.
- [2] E. C. Ozkat, "Vibration data-driven anomaly detection in UAVs: A deep learning approach," Engineering Science and Technology, an International Journal, vol. 54, p. 101702, 2024/06/01/ 2024.
- [3] M. E. Elshaar, M. R. Elbalshy, A. Hussien, and M. Abido, "Path Planning in a dynamic environment using Spherical Particle Swarm Optimization," in 2024 IEEE Congress on Evolutionary Computation (CEC), pp. 1–8, July 2024.
- [4] X. Olaz, D. Alaez, M. Prieto, J. Villadangos, and J. J. Astrain, "Quadcopter neural controller for take-off and landing in windy environments," Expert Systems with Applications, vol. 225, p. 120146, 2023.
- [5] X. Ai, Y. Zhang, and Y.-Y. Chen, "Spherical Formation Tracking Control of Non-Holonomic UAVs with State Constraints and Time Delays," Aerospace, vol. 10, no. 2, 2023.
- [6] S. Poudel, M. Y. Arafat, and S. Moh, "Bio-Inspired Optimization-Based Path Planning Algorithms in Unmanned Aerial Vehicles: A Survey," Sensors, vol. 23, no. 6, 2023.
- [7] X. Wang, L. Wang, and Y. Hu, "Self-confirming Q-learning on unknown networks," Chaos, Solitons & Fractals, vol. 203, p. 117598, 2026.
- [8] M. I. A. Shah, M. E. Cruz Victorio, M. Duffy, E. Barrett, and K. Mason, "Peer-to-peer energy trading in dairy farms using multi-agent reinforcement learning," Applied Energy, vol. 402, p. 127041, 2026.
- [9] S. Lin, J. Wang, B. Huang, X. Kong, and H. Yang, "Bio particle swarm optimization and reinforcement learning algorithm for path planning of automated guided vehicles in dynamic industrial environments," Scientific Reports, vol. 15, no. 1, p. 463, 2025.
- [10] C. Zhang, K. K. H. Ng, Z. Jin, S. Yao, and Y. Qin, "Q-learning-driven exact and meta-heuristic algorithms for the robust gate assignment problem," Advanced Engineering Informatics, vol. 67, 2025.
- [11] B. Q. A. Nguyen, N. T. Dang, T. T. Le, and P. N. Dao, "On-policy and Off-policy Q-learning algorithms with policy iteration for two-wheeled inverted pendulum systems," Robotics and Autonomous Systems, p. 105111, 2025.
- [12] C. Zhang, K. Ng, Z. Jin, S. Yao, and Y. Qin, "Q-learning-driven exact and meta-

نزدیک به واقعیت با تأکید بر حضور موانع متحرک تصادفی مورد ارزیابی قرار می دهد. فرضیه اصلی این تحقیق بر این است که یادگیری Q دوگانه به دلیل معماری دوگانه و مکانیزم کاهش بیش برآورد، در محیط های پویا و غیر قطعی عملکرد پایدارتر و ایمن تری نسبت به یادگیری Q استاندارد ارائه می دهد که این برتری به صورت کاهش آماری معنادار در میانگین برخوردها و واریانس کمتر در خروجی های اجراهای مونت کارلو منجر شد. در مقابل، یادگیری Q در برخی پارامترهای نقطه ای مانند سرعت اولیه همگرایی یا پاداش میانگین در سناریوهای خاص عملکرد بهتری نشان دهد. همچنین انتظار می رود که افزایش درجه آزادی عمل (۴ از عمل به ۶ عمل)، عملکرد هر دو الگوریتم را بهبود بخشد، اما سود حاصل از این افزایش برای یادگیری Q دوگانه بارزتر می باشد.

در ادامه چند مسیر پژوهشی برای ادامه کار پیشنهاد شده است که عبارتند از:

۱. توسعه روش های مسیریابی آنلاین در محیط های کاملاً پویا با استفاده از یادگیری تقویتی
۲. مقایسه سیستماتیک الگوریتم های یادگیری تقویتی با روش های کلاسیک مسیریابی مانند آستار
۳. طراحی الگوریتم های مسیریابی گروهی آنلاین برای ناوگان ربات های همکار در محیط های ناشناخته
۴. بهره گیری از شبکه های عصبی عمیق در محیط هایی با مقیاس بزرگ
۵. ترکیب شبکه های عصبی گرافی با یادگیری تقویتی چندعاملی برای بهبود هماهنگی در مسیریابی گروهی

۵۰

سال ۱۵ - شماره ۱

پیاو و تابستان ۱۴۰۵

نشریه علمی



تحلیل یادگیری تقویتی Q برای مسیریابی گروهی بلادرنگ در محیط پویا با اعتبارسنجی مونت کارلو و پاداش تطبیقی

- planning," Expert Systems with Applications, vol. 236, p. 121303, 2024.
- [24] E. Mohammad, M. Jahromi, J. Pirkandi, M. Khazaei, and M. Mahmoodi, "Development of an intelligent jet engine controller using a model-based deep deterministic policy gradient technique," Journal of Mechanical Engineering and Sciences, vol. 19, no. 3, pp. 10739 – 10755, 2025.

۹. پی‌نوشت‌ها

- ¹ Reinforcement Learning
- ² Q-Learning
- ³ Double Q-Learning
- ⁴ Action
- ⁵ Deep Q-Learning (DQN)
- ⁶ Memetic Algorithms
- ⁷ Nonholonomic constraints
- ⁸ State
- ⁹ Reinforcement learning (RL)
- ¹⁰ Rewards
- ¹¹ Bellman Equation
- ¹² Action-State
- ¹³ Temporal Difference (TD) Update
- ¹⁴ Temporal Difference (TD)
- ¹⁵ Episode
- ¹⁶ TD-Update
- ¹⁷ Exploration
- ¹⁸ Exploitation
- ¹⁹ Double Q-Learning (DQ)
- ²⁰ Overestimation
- ²¹ ϵ -Greedy Policy
- ²² d-Cohen
- ²³ Levin
- ²⁴ Action
- ²⁵ Multi-Agent Proximal Policy Optimization (MAPPO)
- ²⁶ Relative Deviation (RD)
- ²⁷ Sensitivity Index (SI)
- ²⁸ Performance Normalization (PN)

- heuristic algorithms for the robust gate assignment problem," Advanced Engineering Informatics, vol. 67, p. 103551, 2025.
- [13] C. Sun, L. Hou, S. Yu, and J. Shu, "HQA: Hybrid Q-Learning and AODV Multi-Path Routing Algorithm for Flying Ad-Hoc Networks," Vehicular Communications, p. 100947, 2025.
- [14] H. Van Hasselt, A. Guez, and D. Silver, "Deep reinforcement learning with double q-learning," in Proceedings of the AAAI conference on artificial intelligence, vol. 30, no. 1, 2016.
- [15] X. Guo, Q.-K. Pan, W. Zhang, Z.-H. Miao, X.-L. Jing, and H.-Y. Sang, "A Q-learning-assisted memetic algorithm for joint vehicle scheduling problem for harvesting and transportation in smart agriculture," Swarm and Evolutionary Computation, vol. 96, p. 102007, 2025.
- [16] Z. Duan, Y. Zhang, Z. Shao, Z. Xu, and Z. Xiang, "Advanced robot path planning on rough terrain: A Q-learning-based multi-objective PSO algorithm," Applied Soft Computing, p. 113798, 2025.
- [17] Z. Li, Y. Wang, Y. Han, K. Gao, and J. Li, "Q-Learning-Driven Accelerated Iterated Greedy Algorithm for Multi-Scenario Group Scheduling in Distributed Blocking Flowshops," Knowledge-Based Systems, vol. 317, p. 113424, 2025.
- [18] A. Giuffrida, N. Basilico, and F. Amigoni, "An empirical evaluation of learning-based multi-agent path finding algorithms in warehouse environments," Robotics and Autonomous Systems, p. 105149, 2025.
- [19] J. Gao, Y. Li, X. Shi, and K. Yan, "Customer satisfaction optimization in due-aware multi-agent path finding," Information Sciences, p. 122782, 2025.
- [20] S. Venu and M. Gurusamy, "A Comprehensive Review of Path Planning Algorithms for Autonomous Navigation," Results in Engineering, p. 107750, 2025.
- [21] S. Aslan and S. Demirci, "An immune plasma algorithm with Q-learning based pandemic management for path planning of unmanned aerial vehicles," Egyptian Informatics Journal, vol. 26, p. 100468, 2024.
- [22] Y. Wang and G. Wang, "An adaptive and efficient path planning algorithm for UAV navigation in complex environments," Computers & Operations Research, vol. 185, p. 107296, 2026.
- [23] X. Ni, W. Hu, Q. Fan, Y. Cui, and C. Qi, "A Q-learning based multi-strategy integrated artificial bee colony algorithm with application in unmanned vehicle path

