



Computational Intelligence in Electrical Engineering
Vol. 16, No. 3, 2025
pp. 53-68
Research Paper

Investigating the inequity aversion model in multi-agent reinforcement learning approaches for enhancing power control in wireless cellular networks

Farinaz Alamiyan-Harandi ¹, Azam S Hosseinzadeh-Salaty ², Zeinab Maleki ³, Pouria Ramazi ⁴

¹ Department of Electrical and Computer Engineering, Isfahan University of Technology, Isfahan, Iran

² Electrical and Computer Engineering Group, Golpayegan College of Engineering, Isfahan University of Technology, Golpayegan, Iran

³ Department of Electrical and Computer Engineering, Isfahan University of Technology, Isfahan, Iran

⁴ Department of Biological Sciences and Dept. of Mathematics and Statistics, University of Calgary, Calgary, Canada

Abstract:

Recently, data-driven deep reinforcement learning algorithms have replaced traditional model-based algorithms in the problem of power allocation in wireless cellular networks. This article focuses on multi-agent reinforcement learning approaches that leverage centralized training and distributed execution. These approaches use an intrinsic reward function to solve the downlink power control problem in a multi-user wireless cellular network, with the aim of maximizing the sum rate. To enhance the quality of the agents' learning and their collaboration and coordination, the inequity aversion model is applied in the structure of multi-agent reinforcement learning methods. This model manipulates the agents' rewards based on the concepts of envy and guilt. The parameters of this model should first be selected properly. Various simulation results indicate that applying the inequity aversion model with appropriately tuned parameters can improve the performance of basic multi-agent reinforcement learning methods in the problem of power control in multi-user wireless cellular networks.

Keywords: Multi-Agent Reinforcement Learning, Wireless Cellular Networks, Resource Management.



This is an open access article under the CC BY-NC-ND/4.0/ License (<https://creativecommons.org/licenses/by-nc-nd/4.0/>).



<https://doi.org/10.22108/isee.2025.143537.1715>

مقاله پژوهشی

بررسی مدل تنفر از بی عدالتی در راهکارهای یادگیری تقویتی چندعاملی برای ارتقای کنترل توان در شبکه‌های سلولی بی سیم

فریناز اعلی‌یان هرندی^۱، اعظم سادات حسین‌زاده صلوتی^{۲*}، زینب مالکی^۳، پوریا رمضی^۴

۱- پژوهشگر پسادکتری، دانشکده مهندسی برق و کامپیوتر، دانشگاه صنعتی اصفهان، اصفهان، ایران

farinaz.alamiyan@gmail.com

۲- استادیار، گروه مهندسی برق و کامپیوتر، دانشکده فنی مهندسی گلپایگان، دانشگاه صنعتی اصفهان، گلپایگان، ایران

salaty@iut.ac.ir

۳- دانشیار، دانشکده مهندسی برق و کامپیوتر، دانشگاه صنعتی اصفهان، اصفهان، ایران

zmaleki@iut.ac.ir

۴- استادیار، گروه بیولوژی و گروه ریاضی و آمار، دانشگاه کلگری، کلگری، کانادا

p.ramazi@gmail.com

چکیده: به تازگی، الگوریتم‌های مبتنی بر داده یادگیری تقویتی عمیق جایگزین الگوریتم‌های مبتنی بر مدل سنتی در مسئله تخصیص توان در شبکه‌های سلولی بی سیم شده‌اند. تمرکز این مقاله بر راهکارهای یادگیری تقویتی چندعاملی است که از آموزش متمرکز و اجرای توزیع شده بهره می‌برند. در این راهکارها، از یک تابع پاداش داخلی برای حل مسئله کنترل توان لینک پایین‌رو در یک شبکه سلولی بی سیم چندکاربره، با هدف بیشینه‌سازی نرخ مجموع استفاده می‌شود. در این راستا، برای بهبود کیفیت یادگیری عامل‌ها و همکاری و هماهنگی میان آنها، مدل تنفر از بی عدالتی در ساختار یادگیری تقویتی چندعاملی اعمال می‌شود. این مدل بر اساس هر دو مفهوم حسادت و عذاب وجدان، پاداش‌های عامل‌ها را با هدف دستیابی به برابری دستکاری می‌کند. نتایج شبیه‌سازی‌های مختلف نشان می‌دهد اعمال مدل تنفر از بی عدالتی با پارامترهایی که به طور مناسب تنظیم شده‌اند، می‌تواند عملکرد روش‌های یادگیری تقویتی چندعاملی پایه را در مسئله تخصیص توان در شبکه‌های سلولی بی سیم بهبود بخشد.

واژه‌های کلیدی: یادگیری تقویتی چندعاملی، شبکه‌های سلولی بی سیم، مدیریت منابع.

۱- مقدمه

و برنامه‌ریزی کسری^۱ [۶، ۷] بهره‌برداری می‌کنند. با توجه به اینکه شبکه‌های بی سیم ماهیتی پویا دارند، هیچ تضمینی وجود ندارد که این الگوریتم‌ها برای مدیریت منابع رادیویی در تمام سناریوها به سطحی قابل قبول از عملکرد دست یابند. اما الگوریتم‌هایی که از طریق تعامل با محیط یک وظیفه را یاد می‌گیرند، قادر هستند به خوبی چنین پویایی‌هایی را مدیریت کنند. در این راستا، چارچوب‌هایی که فرایند تصمیم‌گیری خود را بر مبنای حجمی عظیم از

از آنجا که مدیریت منابع رادیویی یک مسئله غیرمحدب است و افزایش اندازه شبکه در پیچیدگی محاسبات آن نقشی به سزا دارد، الگوریتم‌های متمرکز و توزیع شده متنوعی برای مدیریت منابع رادیویی ارائه شده‌اند. این راهکارها از تکنیک‌های مطرح در زمینه‌های مختلف علمی همچون برنامه‌ریزی هندسی [۱]، بهینه‌سازی حداقل وزن دار میانگین مربعات خطا^۱ [۲]، نظریه بازی [۳]، نظریه اطلاعات [۴، ۵]،

نشانی نویسنده مسئول: گروه مهندسی برق و کامپیوتر، دانشکده فنی مهندسی گلپایگان، دانشگاه صنعتی اصفهان، گلپایگان، ایران

^۱ تاریخ ارسال مقاله : ۱۴۰۳/۰۹/۲۱

تاریخ پذیرش مقاله : ۱۴۰۴/۰۹/۰۹

نام نویسنده مسئول : اعظم سادات حسین‌زاده صلوتی

داده‌های فعلی شبکه‌های ارتباطی بی‌سیم قرار می‌دهند، برای مقابله با این چالش‌ها مناسب هستند. به همین دلیل، تکنیک‌های یادگیری ماشین در حل طیفی وسیع از مسائل مطرح در حوزه شبکه‌های بی‌سیم پیشنهاد شده‌اند.

روش‌های یادگیری تقویتی زیرمجموعه‌ای خاص از الگوریتم‌های یادگیری ماشین هستند که امکان یادگیری را از طریق تعامل با محیط برای یک عامل فراهم می‌کنند. در این روش‌ها، عامل در پاسخ به مشاهدات دریافتی خود، اقداماتی را با هدف بیشینه‌سازی پاداش در محیط مسئله اعمال می‌کند. اجرای عمل انتخابی عامل باعث می‌شود محیط به مرحله بعدی منتقل شود و عامل، علاوه بر دریافت یک مقدار پاداش عددی، مجموعه‌ای جدید از مشاهدات را درک کند. الگوریتم‌های یادگیری در حین این تعاملات، سیاست رفتاری عامل را اصلاح می‌کنند تا عامل پاداش خود را در طول زمان بیشینه کند.

پیچیدگی‌های فضای حالت و عمل در مسائل مختلف به ویژه مسائل دنیای واقعی، پژوهشگران را به سمت بهره‌برداری از راهکارهای یادگیری تقویتی عمیق سوق داده است. در این روش‌ها، از ساختار شبکه عصبی عمیق به عنوان یک تقریب‌زننده تابع جامع استفاده می‌شود تا تخمینی از احتمال انجام هر عمل از مجموعه عمل‌های مجاز عامل با توجه به هر مشاهده و نیز ارزش هر جفت مشاهده-عمل به دست آید. الگوریتم‌های یادگیری تقویتی عمیق در حل مسائل چالش‌برانگیز تصمیم‌گیری متوالی، به ویژه در انواع بازی‌ها، مانند Atari 2600 و Go به موفقیت‌هایی چشمگیر دست یافته‌اند [۸-۱۰].

به تازگی، از تکنیک‌های یادگیری تقویتی عمیق برای حل مسئله کنترل توان لینک پایین‌رو^۳ در شبکه‌های بی‌سیم سلولی به عنوان یک سیستم چندعاملی استفاده شده است و مکانیسم‌هایی برای برنامه‌ریزی ارسال‌ها پیشنهاد شده‌اند تا عملکرد شبکه برای همه کاربران در سراسر شبکه منصفانه باشد. برای نمونه، می‌توان به پژوهش‌های مینگ و همکاران [۱۱، ۱۲] اشاره کرد که در آن، یک شبکه عصبی با استفاده از الگوریتم DQN^۴ در یک محیط شبیه‌سازی با هدف تخصیص توان در شبکه‌های سلولی چندکاربره آموزش داده شده است. شبکه عصبی حاصل از این آموزش به صورت

پویا در سناریوهای واقعی با هدف انتقال یادگیری بهره‌برداری و تنظیم شده است. در کاربرد روش DQN، از نرخ مجموع برای تعریف تابع پاداش استفاده شده است. مقایسه عملکرد روش پیشنهادی مینگ و همکاران با الگوریتم‌های مبتنی بر مدل نشان داده است میانگین نرخ مجموع در این روش بیشتر و قابلیت تعمیم‌پذیری بهتر است.

در سال‌های اخیر، ادبیات کنترل توان مبتنی بر یادگیری تقویتی چندعاملی^۵ در شبکه‌های سلولی جهشی معنادار داشته است. از یک سو، رویکردهای غیرمتمرکز و مقیاس‌پذیر برای کنترل توان ارائه شده‌اند که فقط بر مشاهدات محلی تکیه دارند و همچنان به کارایی نزدیک روش‌های متمرکز دست می‌یابند [۱۳، ۱۴]. از سوی دیگر، چارچوب «آموزش متمرکز و اجرای توزیع‌شده»^۶ در قالب‌های تعاملی و مشارکتی برای کنترل توان به کار گرفته شده و در برخی موارد با شبکه عصبی حافظه طولانی کوتاه مدت^۷ نیز ترکیب شده است تا ناپایداری کانال را بهتر مدل کند [۱۵، ۱۶]. هم‌زمان، جریان روبه‌رشد انصاف‌محور در یادگیری تقویتی از توابع رفاه اجتماعی^۸ و شبکه‌های تفکیک ارزش^۹ برای موازنه کارایی-عدالت در تخصیص منابع بهره گرفته و نتایج امیدبخش در سناریوهای مخابراتی گزارش کرده است [۱۷، ۱۸]. این خط پژوهشی با ایده مطرح در این مقاله، یعنی به کارگیری پاداش داخلی «تنفر از بی‌عدالتی»^{۱۰} در یادگیری تقویتی چندعاملی هم‌راستاست و جایگاه نظری رویکرد این مقاله را تقویت می‌کند.

اثربخشی اعمال مدل تنفر از بی‌عدالتی در ساختار یادگیری تقویتی چندعاملی در بهبود کیفیت یادگیری عامل‌ها و همکاری و هماهنگی میان آنها در سناریوهایی همچون بازی‌های پاک‌سازی و برداشت^{۱۱} و کنترل ترافیک آزموده شده است [۱۹، ۲۰]. این مدل بر اساس هر دو مفهوم حسادت^{۱۲} و عذاب وجدان^{۱۳}، پاداش‌های عامل‌ها را با هدف دستیابی به برابری دستکاری می‌کند. مدل تنفر از بی‌عدالتی به جای تمرکز صرف روی بیشینه‌کردن کارایی کلی شبکه (مانند مجموع توان یا نرخ داده)، انصاف را به عنوان یک معیار اساسی در تابع پاداش عامل یادگیری تقویتی تعبیه می‌کند. این کار از نارضایتی کاربران با منابع ضعیف

کاربر در هر سلول به طور هم‌زمان توسط فرستنده ایستگاه پایه آن سلول سرویس‌دهی می‌شوند. در شیار زمانی t^m ، بهره کانال مستقل بین ایستگاه پایه n^m و کاربر k^m در سلول j^m به صورت $g_{n,j,k}^t$ در رابطه (۱) تعریف می‌شود:

$$g_{n,j,k}^t = |h_{n,j,k}^t|^2 \beta_{n,j,k} \quad (1)$$

که در آن، $h_{n,j,k}^t$ یک متغیر تصادفی گوسی مختلط است که پوش آن توزیع ریلی دارد و $\beta_{n,j,k}$ مؤلفه مقیاس بزرگ محوشدگی است و هر دو مورد تضعیف هندسی^{۱۸} و محوشدگی سایه^{۱۹} را شامل می‌شوند و فرض می‌شود مقدار آنها در بازه زمانی مدنظر ثابت است. مطابق مدل Jakes، محوشدگی سطح^{۲۰} مقیاس کوچک می‌تواند به صورت یک فرایند گوس-مارکوف مختلط مرتبه اول به صورت رابطه (۲) مدل شود:

$$h_{n,j,k}^t = \rho h_{n,j,k}^{t-1} + n_{n,j,k}^t \quad (2)$$

که در این رابطه $h_{n,j,k}^1 \sim \mathcal{CN}(0,1)$ و $n_{n,j,k}^t \sim \mathcal{CN}(0,1-\rho^2)$ و ضریب همبستگی ρ به صورت رابطه (۳) در نظر گرفته می‌شود:

$$\rho = J_0(2\pi f_d T_s) \quad (3)$$

به طوری که $J_0(\cdot)$ تابع بسل مرتبه صفر نوع اول، f_d فرکانس داپلر بیشینه و T_s دوره زمانی بین شیارهای متوالی است.

با فرض اینکه سیگنال‌های ارسالی از فرستنده‌های متفاوت مستقل از یکدیگر هستند، کانال‌ها در هر شیار زمانی ثابت در نظر گرفته می‌شوند. به این ترتیب، نسبت سیگنال به تداخل و نویز^{۲۱} برای لینک پایین‌رو از ایستگاه پایه n^m به کاربر k^m ($dl_{n,k}$) در شیار زمانی t به صورت رابطه (۴) تعریف می‌شود:

$$\gamma_{n,k}^t = \frac{g_{n,n,k}^t p_{n,k}^t}{\sum_{k' \neq k} g_{n,n,k'}^t p_{n,k'}^t + \sum_{n' \in D_n} g_{n',n,k}^t \sum_i p_{n',i}^t + \sigma^2} \quad (4)$$

که در آن، D_n مجموعه تمام سلول‌های تداخلی در اطراف سلول n^m ، $p_{n,k}^t$ توان انتشار فرستنده n^m به گیرنده k^m در شیار زمانی t^m و σ^2 توان نویز جمع‌شونده است.

جلوگیری می‌کند و پایداری و قابلیت اطمینان شبکه را در بلندمدت افزایش می‌دهد. در این مدل، اگر عامل یادگیری تقویتی دریابد یک کاربر برای مدت طولانی از سرویس محروم مانده است یا نرخ داده بسیار پایینی دارد، پاداش منفی دریافت می‌کند. این امر عامل را مجبور می‌کند تا برای بهبود وضعیت آن کاربر اقدام کند، حتی اگر این اقدام کمی از کارایی کلی شبکه بکاهد. ترافیک کاربران، تعداد کاربران و شرایط کانال در شبکه‌های مدرن به طور مداوم در حال تغییر است. عامل یادگیری تقویتی می‌تواند این شرایط پیچیده و متغیر را مستقیماً از داده‌ها بیاموزد و یک سیاست تطبیقی ایجاد کند که در شرایط مختلف، بی‌عدالتی را تشخیص و به آن واکنش نشان دهد. بهره‌برداری از مدل تنفر از بی‌عدالتی می‌تواند تعادل هوشمند و پویا بین کارایی و انصاف برقرار کند و با شرایط پیچیده و متغیر شبکه‌های مدرن سازگار شود.

ساختار مقاله به این شرح است: در بخش دوم، مسئله تخصیص پویای توان لینک پایین‌رو در یک شبکه بی‌سیم مطرح می‌شود. سپس، روش‌های یادگیری تقویتی پایه و مدل تنفر از بی‌عدالتی به ترتیب در بخش‌های سوم و چهارم معرفی می‌شوند. در بخش پنجم، عناصر یادگیری تقویتی در مسئله تحت بررسی شرح داده می‌شوند. آزمایش‌ها و نتایج حاصل از آنها در بخش ششم گزارش می‌شوند و در نهایت، در بخش هفتم، نتیجه‌گیری بیان می‌شود.

۲- بیان مسئله

در اینجا، مسئله تخصیص پویای توان لینک پایین‌رو در یک شبکه سلولی بی‌سیم با کانال دسترسی چندگانه تداخلی^{۲۲} با هدف بیشینه‌سازی نرخ مجموع^{۲۳} بررسی می‌شود. به عبارت دیگر، هدف اصلی مسئله انتخاب توان انتقال لینک پایین‌رو در پاسخ به شرایط فیزیکی کانال، تحت محدودیت‌های مقدار بیشینه توان است. محیط مسئله مطابق ساختار ارائه‌شده توسط منگ و همکاران [۱۱، ۱۲] یک شبکه سلولی بی‌سیم با یک کانال پخش تداخلی تک‌ورودی تک‌خروجی^{۲۴} است که از N سلول تشکیل و در مرکز هر سلول، یک ایستگاه پایه^{۲۵} مجهز به یک آنتن فرستنده مستقر شده است. با فرض باندهای فرکانسی مشترک، تعداد K

عمل می‌کنند. در این بین، عامل مرتبط با لینک پایین‌رو $dl_{n,k}$ فقط بخشی از اطلاعات وضعیت کانال $(g_{n,k}^t)$ را به عنوان ورودی دریافت می‌کند و توان انتقال لینک پایین‌رو مخصوص خود $(p_{n,k}^t)$ را به عنوان خروجی تولید می‌کند. اطلاعات جزئی وضعیت کانال $g_{n,k}^t$ به صورت رابطه (۱۰) تعریف می‌شود:

$$g_{n,k}^t = \{g_{n',n,k}^t | n' \in \{n, D_n\}\}. \quad (10)$$

به این ترتیب، مسئله بهینه‌سازی (۶) به صورت مسئله برنامه‌ریزی چندهدفه^{۲۴} رابطه (۱۱) در نظر گرفته می‌شود:

$$\begin{cases} \max_{p_{n,k}^t} C_{n,k}^t(g_{n,k}^t, p_{n,k}^t) | \forall n, k \\ s.t. \quad 0 \leq p_{n,k}^t \leq P_{max} \quad \forall n, k. \end{cases} \quad (11)$$

یادگیری چندعاملی برای حل این مسئله حتی به صورتی که در رابطه (۱۱) تعریف شده است، همچنان دشوار است؛ زیرا برای یادگیری تمام پارامترهای ساختار عامل‌ها که به صورت شبکه‌های عصبی عمیق هستند، به زمان یادگیری طولانی و تعداد زیادی داده آموزشی نیاز است. برای برطرف کردن این مشکل، از چارچوب آموزش متمرکز و اجرای توزیع‌شده بهره‌برداری می‌شود. در این چارچوب، تمام عامل‌ها در قالب یک عامل و با ساختاری مشابه در نظر گرفته می‌شوند و سیاستی یکسان برای تمام عامل‌ها به صورت اشتراکی و با استفاده از داده‌های جمع‌آوری‌شده از عملکرد تمام لینک‌های پایین‌رو یاد گرفته می‌شود. این چارچوب، با توجه به اینکه لینک‌ها از نظر مشخصات مکانی ثابت و ابعاد شبکه وسیع هستند، قابل اعمال است؛ زیرا لینک‌ها در فضاهای متمایز تقریباً مشابه یکدیگر هستند؛ بنابراین، سیاست یادگرفته‌شده می‌تواند با بهره‌برداری از مفاهیم انتقال یادگیری^{۲۵} بین عامل‌ها به اشتراک گذاشته شود. در واقع، هر عامل این سیاست را به صورت مستقل استفاده می‌کند. به عبارت دیگر، روش یادگیری متمرکز و روش اجرا توزیع‌شده است.

۳- روش‌های یادگیری تقویتی پایه

در این مقاله، سه الگوریتم یادگیری تقویتی عمیق از سه دسته اصلی معماری‌های یادگیری تقویتی شامل معماری

عبارت‌های $\sum_{k' \neq k} g_{n,n,k}^t p_{n,k'}^t$ و $\sum_{n' \in D_n} g_{n',n,k}^t \sum_i p_{n',i}^t$ نیز به ترتیب معرف تداخلات درون‌سلولی و بین‌سلولی هستند. با فرض پهنای بلند نرمال‌شده، نرخ لینک پایین‌رو $dl_{n,k}$ به صورت رابطه (۵) تعریف می‌شود:

$$C_{n,k}^t = \log_2(1 + \gamma_{n,k}^t). \quad (5)$$

با توجه به تعاریف بالا، مسئله بهینه‌سازی با در نظر گرفتن مجموعه توان غیرمنفی p که در آن، تمام عناصر محدودیت توان بیشینه را ارضا می‌کنند، به صورت بیشینه‌سازی نرخ مجموع در سرتاسر شبکه در رابطه (۶) تعریف می‌شود:

$$\begin{aligned} \max_{p^t} C(g^t, p^t) \\ s.t. \quad 0 \leq p_{n,k}^t \leq P_{max} \quad \forall n, k \end{aligned} \quad (6)$$

که در آن، P_{max} بیشترین توان انتشار سیگنال است و مجموعه توان p^t مجموعه بهره کانال g^t و نرخ مجموع $C(g^t, p^t)$ به صورت رابطه‌های (۷)، (۸) و (۹) تعریف می‌شوند:

$$p^t := \{p_{n,k}^t | \forall n, k\} \quad (7)$$

$$g^t := \{g_{n',n,k}^t | \forall n', n, k\} \quad (8)$$

$$C(g^t, p^t) := \sum_{n,k} C_{n,k}^t. \quad (9)$$

توجه شود که بهینه‌سازی بالا یک مسئله غیرمحدب و چندجمله‌ای غیرقطعی سخت^{۲۶} است. از آنجا که اطلاعات وضعیت کانال^{۲۳} در هر شیار زمانی، اطلاعات کافی درباره راه‌حل بهینه در اختیار قرار می‌دهد، تابعی برای نگاشت این اطلاعات به راه‌حل مسئله موجود و قابل دست‌یابی است. حل این مسئله بهینه‌سازی به صورت یک سیستم یادگیری تقویتی تک‌عاملی با مشکلاتی همچون افزایش نمایی ابعاد فضای حالت و عمل با افزایش تعداد سلول‌ها و عدم امکان انتقال تمام اطلاعات وضعیت کانال به عامل با تأخیری قابل تحمل مواجه است. به همین دلیل، تخصیص توان به صورت غیرمتمرکز و به عنوان یک مسئله یادگیری تقویتی چندعاملی مطرح می‌شود و برای کنترل توان انتقال هر لینک پایین‌رو از یک عامل یادگیری تقویتی مجزا بهره‌برداری می‌شود. همه عامل‌ها در شبکه ارتباطی به صورت هم‌زمان و توزیع‌شده

$$\theta_q^* = \arg \min_{\theta_q} \frac{1}{2} (Q(s, a; \theta_q) - r_s^a)^2. \quad (14)$$

گرایان θ_q به صورت رابطه (۱۵) محاسبه می‌شود:

$$\nabla \theta_q = (Q(s, a; \theta_q) - r_s^a) \nabla_{\theta_q} Q(s, a; \theta_q). \quad (15)$$

مطابق رابطه (۱۶)، عمل بهینه a^* به گونه‌ای انتخاب می‌شود که مقدار Q را بیشینه کند:

$$a^* = \arg \max_a Q(s, a; \theta_q). \quad (16)$$

برای مدیریت درجه کاوش الگوریتم در طول فریند یادگیری، یک سیاست شبه‌حریصانه (ϵ -greedy) تطبیق داده می‌شود که در آن، پارامتر ϵ به صورت رابطه (۱۷) تعریف می‌شود:

$$\epsilon_k := \epsilon_1 + \frac{k-1}{N_e-1} (\epsilon_{N_e} - \epsilon_1), k = 1, \dots, N_e \quad (17)$$

که در آن، N_e برابر تعداد کل گام‌ها در یک مرحله از فرایند یادگیری است و مقادیر ϵ_1 و ϵ_{N_e} به ترتیب نشان‌دهنده احتمال کاوش اولیه و نهایی هستند.

الگوریتم DDPG²⁹ [۲۲] منطبق با معماری عملگر-نقاد و مبتنی بر گرایان سیاست قطعی است که روی فضای عمل پیوسته عمل می‌کند. در این الگوریتم، یک عملگر حالت محیط S را مشاهده و از طریق نگاشت شبکه عصبی $A(S; \theta_a)$ عمل قطعی a را تولید می‌کند؛ به طوری که θ_a معرف پارامترهای عملگر است. نقاد مقدار Q متناظر با هر زوج حالت-عمل را از طریق شبکه عصبی $C(s_c, a; \theta_c)$ پیش‌بینی می‌کند که در آن، θ_c پارامترهای قلیل تنظیم نقاد و s_c حالت محیط است که توسط نقاد مشاهده شده است. نقاد و عملگر با یکدیگر همکاری دارند و سیاست قطعی بهینه را با حل مسئله بهینه‌سازی توأم رابطه‌های (۱۸) و (۱۹) به دست می‌آورند:

$$\theta_a^* = \arg \max_{\theta_a} C(s_c, a; \theta_c) |_{a=A(s; \theta_a)} \quad (18)$$

$$\theta_c^* = \arg \max_{\theta_c} \frac{1}{2} (C(s_c, a; \theta_c) |_{a=A(s; \theta_a)} - r_s^a)^2. \quad (19)$$

عملگر تلاش می‌کند تا مقادیر ارزش ارزیابی شده توسط نقاد را به بیشینه کند و هدف نقاد دقیق‌ترکردن این ارزیابی است. عملگر و نقاد هر دو مشتق‌پذیر هستند و طبق قاعده زنجیره، گرایان آنها به صورت رابطه‌های (۲۰) و (۲۱) محاسبه می‌شود:

عملگر-تنها، نقاد-تنها و عملگر-نقاد به عنوان روش‌های پایه در نظر گرفته شده‌اند. الگوریتم REINFORCE [۲۱] یک روش یادگیری گرایان سیاست از دسته معماری‌های عملگر-تنها و مبتنی بر روش مونت-کارلو²⁶ است. الگوریتم‌های مبتنی بر سیاست به صورت صریح سیاست تصادفی π را در قالب یک بردار احتمالاتی در خروجی یک شبکه عصبی با پارامترهای قابل یادگیری θ_π تولید می‌کنند. به عبارت دیگر، یادگیری سیاست $\pi(a|s; \theta_\pi)$ با استفاده از تنظیم پارامترهای یک شبکه عصبی برای هر یک از حالت‌های دریافتی s عمل مناسب a را تولید می‌کند. راهکار کلی در روش گرایان افزایشی تصادفی^{۲۷} به نمونه‌هایی نیاز دارد که امید گرایان آنها متناسب با گرایان واقعی معیار عملکرد به عنوان تابعی از پارامترها باشد. هدف الگوریتم REINFORCE مطابق رابطه (۱۲)، بیشینه‌سازی مقدار پاداش مورد انتظار تحت سیاست π است:

$$\theta_\pi^* = \arg \max_{\theta_\pi} \mathbb{E}_\pi [\sum_a \pi(a|s; \theta_\pi) r_s^a]. \quad (12)$$

گرایان نمونه‌برداری مونت-کارلو به صورت رابطه (۱۳) محاسبه می‌شود:

$$\nabla \theta_\pi = \mathbb{E}_\pi [\nabla_{\theta_\pi} \ln \pi(a|s; \theta_\pi) r_s^a | s = s^t, a = a^t] \quad (13)$$

به طوری که ∇ عملگر گرایان است. از آنجا که شبکه عصبی سیاست به صورت مستقیم سیاست تصادفی را تولید می‌کند، در طول یادگیری با استفاده از این شبکه، عمل‌های تصادفی تولید می‌شوند. اما زمانی که یادگیری انجام شد و سیاست بهینه یاد گرفته شد، این شبکه عصبی عمل بهینه را تولید خواهد کرد. برای کاهش حساسیت این الگوریتم به مقدار تابع پاداش، از مقادیر نرمال‌شده پاداش به صورت $\tilde{r} = \frac{r - \mu_r}{\sigma_r}$ استفاده می‌شود که در آن، μ_r و σ_r به ترتیب میانگین و انحراف معیار پاداش r هستند.

الگوریتم DQN²⁸ [۸] از معماری نقاد-تنها پیروی می‌کند. این الگوریتم مقدار تابع ارزش $Q(s, a; \theta_q)$ را تخمین می‌زند؛ به طوری که در آن، θ_q بردار پارامترهای قابل تنظیم شبکه عصبی است. انتخاب عمل‌های مناسب در این الگوریتم به دقت تخمین تابع ارزش بستگی دارد و هدف الگوریتم، مطابق رابطه (۱۴)، کمینه‌کردن تابع اتلاف l_2 با یافتن بردار پارامترهای بهینه θ_q^* است:

پژوهشگران پاداش‌های دریافتی از محیط را با پاداش‌های داخلی^{۳۰} غنی‌سازی می‌کنند تا جنبه‌هایی از رفتار عامل‌ها را در یادگیری لحاظ کنند که لزوماً توسط پاداش محیطی کدگذاری نمی‌شوند. به این ترتیب، برای هر عامل یادگیر در گام زمانی، پاداشی به صورت ترکیب خطی از پاداش محیطی r_{env} و پاداش داخلی $r_{intrinsic}$ به صورت رابطه (۲۵) تعریف می‌شود:

$$r = a r_{env} + b r_{intrinsic} \quad (25)$$

که در آن، a و b ضرایب ثابت هستند.

مدل تنفر از بی عدالتی [۲۳] برای بهبود عملکرد الگوریتم‌های یادگیری تقویتی چندعاملی بر اساس مفاهیم حسادت و عذاب وجدان، پاداش‌های عامل‌ها را دستکاری می‌کند. در این مدل، مطابق رابطه (۲۶)، هر عامل k پاداش خود (r_k) را با پاداش هر یک از دیگر عامل‌ها (r_j) مقایسه می‌کند. نتیجه هر مقایسه یکی از دو وضعیت حسادت یا عذاب وجدان را برای عامل در مقابل عامل همکارش مشخص می‌کند: وضعیت عذاب وجدان در نتیجه یک نابرابری مطلوب^{۳۱} رخ می‌دهد، یعنی زمانی که عامل در مقایسه پاداش خود با دیگری به این نتیجه برسد که درآمد بیشتری کسب کرده است ($r_k - r_j > 0$). حسادت بر اثر یک نابرابری نامطلوب^{۳۲} زمانی رخ می‌دهد که عامل نسبت به همکارش پاداش کمتری دریافت کرده باشد ($r_j - r_k > 0$). در این دو وضعیت، عامل مطابق مدل تنفر از بی عدالتی با دستکاری پاداش دریافتی خود، اقدام به تنبیه یا تشویق خود می‌کند. در رابطه (۲۶)، U_k مقدار پاداش نهایی عامل k ام، N تعداد کل عامل‌ها، α_k و β_k پارامترهای قابل تنظیم مدل هستند.

$$U_k(r_k, \dots, r_N) = r_k - \frac{\alpha_k}{N-1} \sum_{j \neq k} \max(r_j - r_k, 0) - \frac{\beta_k}{N-1} \sum_{j \neq k} \max(r_k - r_j, 0) \quad (26)$$

۵- عناصر یادگیری تقویتی در مسئله تحت

بررسی

انتخاب حالت محیط به عنوان آنچه از اطلاعات محیط توسط هر عامل درک می‌شود، یکی از موضوع‌های بسیار مهم و تأثیرگذار در روند یادگیری عامل هاست. بدیهی است، اطلاعات وضعیت فعلی کانال ($g_{n,k}^t$) که به صورت جزئی

$$\nabla \theta_a = \nabla_a C(s_c, a; \theta_c) |_{a=A(s; \theta_a)} \nabla_{\theta_a} A(s; \theta_a) \quad (20)$$

$$\nabla \theta_c = (C(s_c, a; \theta_c) - r_s^a) \nabla_{\theta_c} C(s_c, a; \theta_c) |_{a=A(s; \theta_a)} \quad (21)$$

مطابق رابطه (۲۲)، عمل قطعی به طور مستقیم از طریق عملگر به دست می‌آید:

$$a^* = A(s; \theta_a). \quad (22)$$

مشابه سیاست شبه‌حریصانه پویا، عمل حاصل از کاوش برای مرحله k در بازه $[0, P_{max}]$ و به صورت رابطه (۲۳) تعریف می‌شود:

$$a := [A(s; \theta_a) + n^k]_0^{P_{max}} \quad (23)$$

که در آن، n^k نویز اضافی است و از توزیع یکنواخت رابطه (۲۴) پیروی می‌کند:

$$n^k \sim u(-\frac{P_{max}}{k}, \frac{P_{max}}{k}). \quad (24)$$

نقاد $C(s_c, a; \theta_c)$ را می‌توان به عنوان یک شبکه کمکی برای انتقال گرادیان در یادگیری در نظر گرفت که پس از اتمام فرایند یادگیری در تولید عمل دخالتی ندارد. به همین دلیل، نقاد با وجود آنکه باید مشتق‌پذیر باشد، لزوماً آموزش‌پذیر نیست. در اینجا، یک نقاد تا حدودی مستقل از مدل پیشنهاد می‌شود که از هر دو امکان بهره‌برداری از دانش پیشین و انعطاف‌پذیری ساختار شبکه عصبی برخوردار است.

الگوریتم گرادیان سیاست REINFORCE بر اساس یک سیاست تصادفی توسعه یافته است، اما نمونه‌برداری در فضای عمل پیوسته یا با ابعاد بالا ناکارآمد است. در الگوریتم DDGP، از گرادیان سیاست قطعی برای غلبه بر این مشکل استفاده می‌شود. از سوی دیگر، با آنکه عملکرد نقاد $C(s_c, a; \theta_c)$ در الگوریتم DDGP با تخمین‌گر ارزش $Q(s, a; \theta_q)$ در الگوریتم DQN مشابه است، این دو ساختار از نظر ورودی‌ها با یکدیگر تفاوت دارند. در الگوریتم DQN، ارزش جفت حالت-عمل با دریافت حالت محیط تخمین زده می‌شود؛ اما در الگوریتم DDGP، این ارزش با دریافت حالت محیط و عمل عامل محاسبه می‌شود.

۴- مدل تنفر از بی عدالتی

برای بهبود الگوریتم‌های یادگیری تقویتی چندعاملی،

در اینجا، فضای حالت با در نظر گرفتن گروه ویژگی رابطه (۳۰) اجرا شده است:

$$f = \{\tilde{\mathbf{I}}_{n,k}^t, \tilde{\mathbf{p}}_{n,k}^{t-1}, \tilde{\mathbf{C}}_{n,k}^{t-1}\}. \quad (30)$$

مشابه مراحل که برای پیش پردازش حالت محیط S بیان شد، ورودی ساختار نقاد S_c در الگوریتم DDPG به صورت $\tilde{\mathbf{C}}_{n,k}^t$ مطابق رابطه (۲۹) تعریف می شود که برای ایجاد آن از تعداد I_c عنصر ابتدای لیست نزولی مرتبط شده $\mathbf{C}_{n,k}^t = \{C_{n,k}^t | \forall n, k\}$ استفاده می شود.

توان لینک پایین رو که مقداری پیوسته و غیر منفی است و با مقدار توان بیشینه P_{max} محدود شده است، به عنوان عمل هر عامل در نظر گرفته می شود. در الگوریتم های REINFORCE و DQN، بازه مجاز برای توان انتشار سیگنال به تعدادی سطح توان به صورت پلکانی گسسته سازی می شود و این سطوح توان انتشار به عنوان عمل های گسسته هر عامل یادگیری تقویتی به کار گرفته می شوند. از آنجا که توان های انتشار در بازه مجاز خود از نظر بزرگی مقدار تفاوت زیادی دارند، برای گسسته سازی بهتر توان انتقال از نرمال سازی لگاریتمی بهره برداری می شود؛ بنابراین، مجموعه عمل های مجاز یک عامل (A) به صورت رابطه (۳۱) تعریف می شود:

$$A := \left\{ 0, \left\{ P_{min} \left(\frac{P_{max}}{P_{min}} \right)^{\frac{b}{|A|-2}} \mid b = 0, \dots, |A| - 2 \right\} \right\} \quad (31)$$

که در آن، P_{min} کمینه غیر صفر توان انتشار است و $|A|$ تعداد عمل های مجاز عامل یادگیری تقویتی است.

یکی دیگر از روش های پایه که از ساختار عملگر-نقاد بهره برداری می کند الگوریتم DDPG است که در آن، عملگر می تواند مقادیر توان انتشار را از یک بازه پیوسته محدود شده تولید کند. این عملگر برای تولید خروجی از رابطه (۳۲) استفاده می کند و در آن، همان عمل انتخابی عامل و معادل با توان انتشار تولید شده است:

$$a := P_{max} \cdot \frac{1}{1 + \exp(-x)}, \quad (32)$$

به طوری که x خروجی لایه آخر شبکه عصبی عملگر در عامل یادگیرنده است.

تابع پاداش به گونه ای طراحی می شود که نرخ انتقال

توسط هر عامل درک می شود، یکی از مهم ترین ویژگی های تشکیل دهنده حالت محیط است. در اینجا، صورت لگاریتمی این اطلاعات پس از نرمال سازی آنها به شکل رابطه (۲۷) استفاده می شود:

$$\mathbf{\Gamma}_{n,k}^t = \log_2 \left(1 + \frac{g_{n,k}^t}{g_{n,n,k}^t} \right) \otimes \mathbf{1}_K \quad (27)$$

که در آن، $\otimes$ ضرب کرونیکر^{۳۳} و $\mathbf{1}_K$ یک بردار است که با تعداد K عدد یک پر شده است. برای نرمال سازی مقادیر $g_{n,k}^t$ از بهره کانال لینک پایین رو $dl_{n,k}$ استفاده می شود و ترجیح این است که از بازنمایی لگاریتمی این مقادیر استفاده شود؛ زیرا دامنه ها معمولاً از نظر مقدار بزرگی تفاوت زیادی دارند. تعداد عناصر $\mathbf{\Gamma}_{n,k}^t$ برابر $(|D_n| + 1)K$ است و با تغییر در تعداد کاربران هر سلول تغییر می کند.

برای آنکه ابعاد ورودی ساختارهای شبکه عصبی در عامل یادگیرنده کاهش داده شوند و متغیر بودن تعداد عناصر $\mathbf{\Gamma}_{n,k}^t$ مدیریت شود، مجموعه جدید $\tilde{\mathbf{\Gamma}}_{n,k}^t$ و اندیس های $I_{n,k}^t$ از طریق مرتب سازی نزولی عناصر $\mathbf{\Gamma}_{n,k}^t$ در یک لیست و انتخاب I_c عنصر ابتدای لیست به همراه اندیس آنها در $\mathbf{\Gamma}_{n,k}^t$ ایجاد می شوند. در واقع، از آنجا که تعداد تداخلات واقعی که در رابطه نسبت سیگنال دریافتی به تداخل و نویز اثرگذار هستند از تعداد عناصر $\mathbf{\Gamma}_{n,k}^t$ بسیار کمتر است و سیگنال تداخلی بسیاری از این عناصر به صفر نزدیک است، مرتب سازی عناصر $\mathbf{\Gamma}_{n,k}^t$ و انتخاب تعدادی از آنها نه فقط فضای حالت را کوچک تر می کند، بلکه تقریبی از عبارت مربوط به تداخل را در مخرج رابطه نسبت سیگنال دریافتی به تداخل و نویز (رابطه ۴) حاصل می کند.

با توجه به اینکه کلنال به صورت یک فرایند مارکوف مدل شده است و این فرایند به زمان وابسته است، آخرین حل، یعنی $p_{n,k}^{t-1}$ می تواند نقطه شروع بهبود یافته ای را برای ادامه یادگیری فراهم کند؛ بنابراین، مقادیر توان انتقال و نرخ لینک پایین رو که در زمان قبلی محاسبه شده اند، برای تمام عناصر مجموعه $\tilde{\mathbf{\Gamma}}_{n,k}^t$ به ورودی ساختار شبکه عصبی عامل یادگیرنده اضافه می شوند؛ به طوری که:

$$\tilde{\mathbf{p}}_{n,k}^{t-1} = \{p_{n,k}^{t-1} | (n, k) \in I_{n,k}^t\}, \quad (28)$$

$$\tilde{\mathbf{C}}_{n,k}^{t-1} = \{C_{n,k}^{t-1} | (n, k) \in I_{n,k}^t\}. \quad (29)$$

```

2. Initialization: Initialize DQN
    $Q(s, a; \theta_q)$  with random parameters  $\theta_q$ .
3. for  $k = 1$  to  $N_e$  do
4.   Update  $\varepsilon_k$  by (17).
5.   Receive initial state  $s^1$ .
6.   for  $t = 1$  to  $T$  do
7.     if  $\text{rand}() < \varepsilon_k$  then
8.       Randomly select action  $a^t \in A$ 
       with uniform probability.
9.     else
10.      Select action  $a^t$  by (16).
11.    end if
12.    Execute action  $a^t$ , achieve all
    agents' rewards  $r_1^t, \dots, r_N^t$ , and observe new state
     $s^{t+1}$ .
13.    Calculate  $U_k^t(r_1^t, \dots, r_N^t)$  by (26)
    as  $r^t$ .
14.    Calculate gradient  $\nabla \theta_q$  by (15),
15.    Update parameter along negative
    gradient direction:  $\theta_q \leftarrow \theta_q - \eta_q \nabla \theta_q$ .
16.     $s^t \leftarrow s^{t+1}$ 
17.  end for
18. end for
19. Output: learned DQN  $Q(s, a; \theta_q)$ .

```

۶- آزمایش‌ها و نتایج

برای انجام آزمایش‌ها و مقایسه الگوریتم‌های یادگیری تقویتی، از تنظیمات بیان‌شده در جدول (۱) برگرفته از [۱۱]، [۱۲] استفاده شده است. تنظیمات مربوط به ساختارهای شبکه عصبی عمیق از جمله توابع فعال‌سازی و تعداد نرون‌ها در هر لایه در جدول (۲) آورده شده است.

در مقایسه الگوریتم‌های یادگیری و ارزیابی عملکرد آنها، از معیار نرخ مجموع استفاده می‌شود. از آنجا که مقداردهی اولیه پارامترها و تولید نمونه‌های آموزشی در آزمایش‌ها به صورت تصادفی انجام می‌شود، در سنجش عملکرد الگوریتم از میانگین نتایج چندین آزمایش استفاده می‌شود. نمادهای σ_c^2 ، $\bar{C}$ و $\bar{C}^*$ به ترتیب به عنوان واریانس نرخ مجموع، میانگین نرخ جمع و میانگین نرخ مجموع، ۲۰ درصد برتر نسبت به تکرارهای مستقل آزمایش‌ها تعریف می‌شوند. معیار $\bar{C}^*$ یک شاخص برای اندازه‌گیری عملکرد الگوریتم‌هایی است که به‌خوبی آموزش دیده‌اند.

در اینجا، نتایج حاصل از آموزش سه الگوریتم یادگیری تقویتی عمیق DDGP، DQN و REINFORCE که به عنوان روش‌های پایه معرفی شدند، با استفاده از تابع پاداش حاصل از مدل تنفر از بی عدالتی گزارش می‌شود. برای هر

اطلاعات را بهبود بخشید و میزان تداخلات در لینک‌های همسایه را کاهش دهد. یکی از گزینه‌ها برای طراحی تابع پاداش استفاده از میانگین نرخ مجموع است و بر این اساس استفاده می‌شود که مجموع پاداش تمام عامل‌ها با نرخ مجموع شبکه برابر باشد. در تابع پاداش، همسایگی هر عامل لحاظ و تابع پاداش به صورت محلی برای هر عامل تعریف می‌شود:

$$r_{n,k}^t = C_{n,k}^t + \alpha (\sum_{n',k' \neq k} C_{n',k'}^t + \sum_{n' \in D_n, i} C_{n',i}^t), \quad (33)$$

که در آن، $\alpha \in \mathbb{R}^+$ ضریب وزنی برای اثر تداخل و $\mathbb{R}^+$ معرف مجموعه مقادیر حقیقی مثبت است. اگر تعداد سلول‌های شبکه (N) به قدر کافی بزرگ باشد، جمع پاداش‌های محلی با نرخ مجموع متناسب می‌شود:

$$\sum_{n,k} r_{n,k}^t \propto C(g^t, p^t). \quad (34)$$

در ترکیب مدل تنفر از بی عدالتی با روش‌های یادگیری تقویتی، مقادیر حاصل از رابطه (۳۳) برای هر کدام از عامل‌های k و j و j جایگزین مقادیر r_k و r_j در رابطه (۲۶) می‌شوند. با وجود آنکه در مدل تنفر از بی عدالتی، اثرگذاری نابرابری‌های مطلوب و نامطلوب در تغییرات مقدار پاداش عامل ممکن است به صورت تنبیه یا تشویق باشد، تا آنجا که بررسی‌ها نشان می‌دهند، تمام پژوهش‌های قبلی فقط از بخشی از مدل و آن هم با اثر تنبیهی بهره‌برداری می‌کنند. در این پژوهش، مجموعه‌ای از مقادیر به عنوان ضرایب مدل در نظر گرفته می‌شوند و اثر آنها در یادگیری عامل‌ها و بهبود نتایج بررسی می‌شود. مطابق نتایجی که در ادامه ارائه می‌شود، استفاده از ضرایب منفی و تشویق عامل به کمک افزودن ضریبی مثبت از تفاوت‌های محاسبه‌شده به پاداش عامل در برخی از روش‌های یادگیری تقویتی در مسئله تخصیص منابع شبکه‌های بی سیم می‌تواند به بهبود نتایج کمک کند.

برای وضوح بیشتر کاربرد مدل تنفر از بی عدالتی در روش‌های یادگیری تقویتی، گام‌های الگوریتم DQN برای عامل k به عنوان نمونه در زیر آورده شده‌اند:

1. Input: Episode times N_e , exploration times T , learning rate η_q , initial and final exploration probability $\varepsilon_1, \varepsilon_{N_e}$.

یک از روش‌های پایه، ترکیب‌هایی متفاوت از مقادیر $\alpha, \beta \in \{-1, -0.75, -0.5, -0.25, 0, 0.25, 0.5, 0.75, 1\}$ در مدل تنفر از بی‌عدالتی لحاظ می‌شوند و عامل‌ها به کمک آن روش پایه در محیط آموزش تحت یادگیری قرار می‌گیرند. در فرایند یادگیری، ۵۰۰۰ مرحله ۱۰ گامی در نظر گرفته می‌شوند، معیار متوسط نرخ مجموع پس از هر ۱۰۰ مرحله محاسبه و با بهترین نتیجه قبلی مقایسه می‌شود و در صورت بهبود نتیجه، ساختار یادگرفته‌شده به عنوان بهترین ساختار به همراه نتیجه حاصل‌شده برای مقایسه‌های بعدی ذخیره می‌شوند.

جدول (۱): پارامترهای شبیه‌سازی محیط در آزمایش‌ها

عنوان	متغیر	مقدار
تعداد سلول‌ها در شبکه	N	25
بیشترین تعداد کاربران در هر سلول	K	4
فضای درونی	R_{min}	km 0.01
نیم‌فاصله سلول به سلول	R_{max}	km 1
فرکانس داپلر	f_d	10 Hz
دوره زمانی*	T_s	20 ms
محوشدگی مقیاس بزرگ**	β	$-120.9 - 37.6 \log_{10} d + 10 \log_{10} z$
توان نویز گوسی سفید جمع‌شونده	σ^2	-114 dBm
کران پایین توان انتشار	P_{min}	5 dBm
کران بالای توان انتشار	P_{max}	38 dBm
بیشینه نسبت سیگنال به تداخل و نویز	$SINR$	30 dB
تعداد سلول‌های همسایه برای هر سلول	$ D_n , \forall n$	18
تعداد تداخلات انتخابی از ابتدای لیست‌های مرتب‌شده	I_c	16
تعداد سطوح توان (تعداد عمل ممکن برای هر عامل)	$ A $	10
ابعاد فضای حالت با در نظر گرفتن ویژگی‌های f_1 و f_2	$ S $	48, 32
ضریب وزن	α	1
تعداد مراحل ^۳ یادگیری	N_e	5000
تعداد شیارهای زمانی در هر مرحله	T	10

* برای شبیه‌سازی اثرات محوشدگی به کار گرفته شده است.

** در این رابطه، d طول لینک است و متغیر تصادفی z با توزیع لگاریتمی نرمال از $\ln z \sim N(0, \sigma_z^2)$ پیروی می‌کند که در آن، σ_z^2 برابر δdB است.

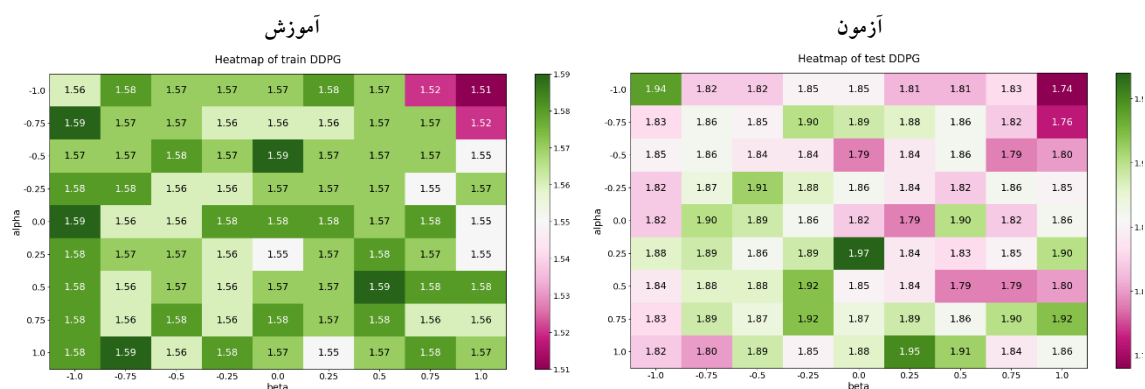
جدول (۲): تنظیمات یادگیری و پارامترهای ساختارهای شبکه عصبی عمیق در الگوریتم‌ها

الگوریتم				تنظیمات
DDPG		DQN	REINFORCE	
Critic	Actor			
$\eta_c = 1 \times 10^{-3}$	$\eta_a = 1 \times 10^{-4}$	$\eta_q = 1 \times 10^{-3}$	$\eta_\pi = 1 \times 10^{-4}$	نرخ یادگیری
-	رابطه (۲۳)	$\epsilon_1 = 0.2$ $\epsilon_{N_e} = 1 \times 10^{-4}$	رابطه (۱۷)	روش کاوش
linear, 1	رابطه (۳۲), 1	linear, A	softmax, A	لایه خروجی شبکه عصبی
RelU, 64	RelU, 64 softmax, 128	RelU, 64 softmax, 128	RelU, 64 softmax, 128	لایه پنهان شبکه عصبی
linear, l _c	linear, S	linear, S	linear, S	لایه ورودی شبکه عصبی
$RelU: f(x) = \max(0, x), \quad linear: f(x) = x, \quad softmax: f(x) = \frac{e^x}{\sum e^x}$				

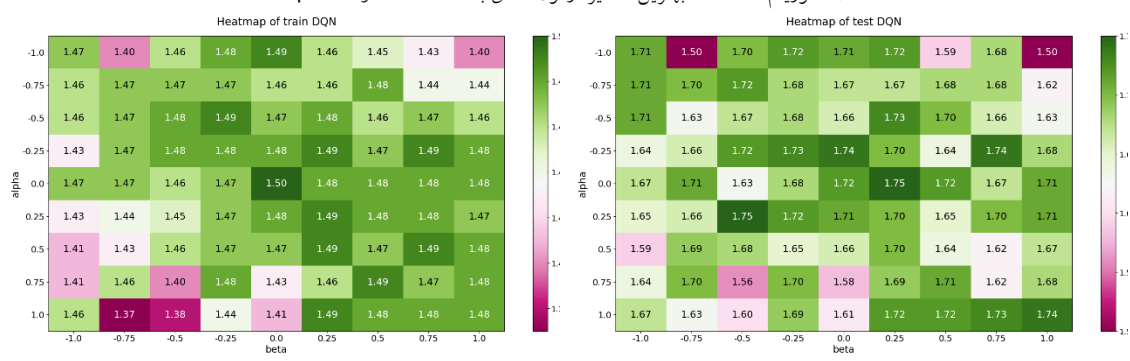
بررسی مدل تنفر از بی‌عدالتی در راهکارهای یادگیری تقویتی چندعاملی برای ارتقای کنترل توان در شبکه‌های سلولی بی‌سیم

نشان می‌دهد. بهبود نتایج $\bar{C}$ با استفاده از برخی از ترکیب‌های پارامترهای α و β برای هر یک از روش‌های پایه در این نمودارها مشهود است. با توجه به اینکه در برخی از پژوهش‌ها، علاوه بر معیار $\bar{C}$ از معیار $\bar{C}^*$ هم استفاده شده است، شکل (۲) مقایسه نتایج را با توجه به این دو معیار به صورت نمودارهای خطی برای ۸۱ ترکیب مختلف پارامترهای α و β نشان می‌دهد.

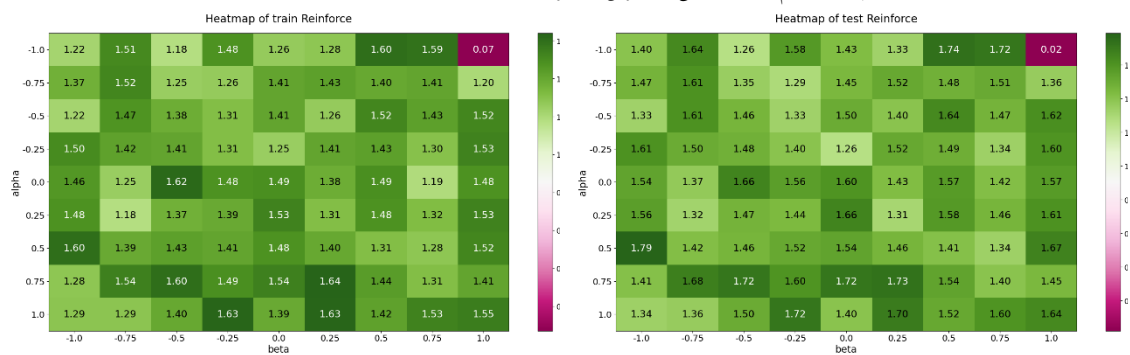
پس از پایان فرایند یادگیری، ساختار یادگرفته‌شده عامل‌ها در محیط آزمون ارزیابی می‌شود. برای فرایند آزمون، ۵۰۰ مرحله ۳۰۰ گامی در نظر گرفته شده‌اند که ارزیابی نتایج آنها با در نظر گرفتن معیار $\bar{C}$ انجام می‌شود. شکل (۱) این نتایج را برای فرایندهای آموزش و آزمون هر یک از روش‌های یادگیری تقویتی به صورت نمودارهای حرارتی و با توجه به مقادیر مختلف پارامترهای مدل تنفر از بی‌عدالتی



الف) الگوریتم DDPG، بهترین مقادیر آزمون متعلق به $\alpha = 0.25$ و $\beta = 0$ است.



ب) الگوریتم DQN، یکی از بهترین مقادیر آزمون متعلق به $\alpha = 1$ و $\beta = 1$ است.



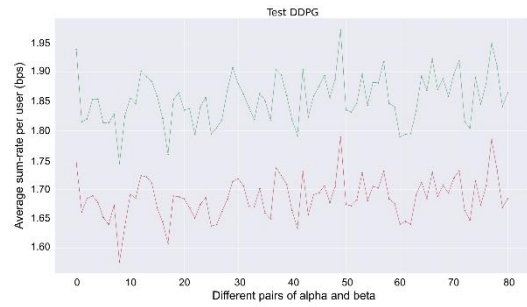
پ) الگوریتم REINFORCE، یکی از بهترین مقادیر آزمون متعلق به $\alpha = 0.75$ و $\beta = 0.25$ است.

شکل (۱): مقایسه معیار $\bar{C}$ در فرایند آموزش و آزمون عملکرد ساختار یادگرفته‌شده توسط هر یک از الگوریتم‌ها با در نظر گرفتن ترکیب‌هایی متفاوت از مقادیر پارامترهای α و β در مدل تنفر از بی‌عدالتی. در این مقایسه، بهترین مقادیر با رنگ سبز تیره قابل مشاهده هستند. ترکیب $\alpha = 0$ و $\beta = 0$ معرف نتیجه حاصل از اعمال روش پایه هر الگوریتم است.

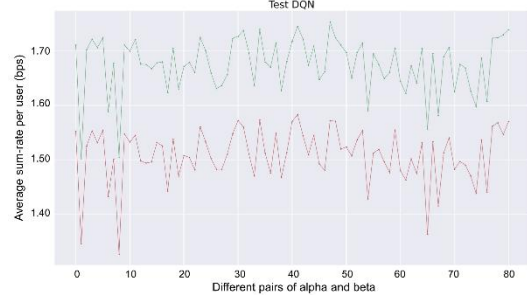
پارامترهای α و β در مدل تنفر از بی عدالتی در محیط آزمون. برای انجام آزمایش‌های تعمیم‌پذیری^{۳۰}، از میان ترکیب‌هایی از پارامترهای مدل تنفر از بی عدالتی که نتایج آزمون بهتری داشته‌اند، دو ترکیب انتخاب شده‌اند و عملکرد آنها در محیط‌های شبیه‌سازی با مشخصات متفاوت بررسی شده است. جدول (۳) مقادیر معیار $\bar{C}$ را برای هر یک از این ترکیب‌ها با روش پایه مقایسه می‌کند. در این آزمایش‌ها، از مشخصات بیان‌شده در جدول (۱) استفاده شده است.

ارزیابی تعمیم‌پذیری الگوریتم‌های یادگیری تقویتی با مدل تنفر از بی عدالتی در شکل (۳) که با توجه به مقادیر مختلف پارامترهای شبکه انجام شده است، نشان می‌دهد از میان ترکیب‌های بیان‌شده در جدول (۳)، بهترین نتایج مربوط به الگوریتم DDPG با پارامترهای $\alpha = 0.25$ و $\beta = 0$ است.

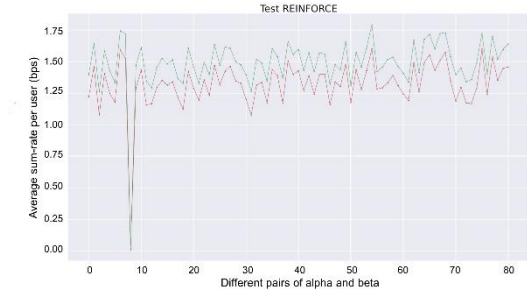
نتایج تعمیم‌پذیری بهترین ترکیبات پارامترهای مدل تنفر از بی عدالتی در هر روش یادگیری تقویتی با روش‌های پایه در شکل (۴) برتری بهره‌برداری از این مدل را در الگوریتم‌های DDPG و REINFORCE به‌وضوح نشان می‌دهد. در این شکل، نتایج استفاده از بیشترین توان و کاربرد توان تصادفی به عنوان دو روش اولیه و نتایج دو الگوریتم مبتنی بر مدل به نام‌های برنامه‌ریزی کسری [۷] و حداقل وزن‌دار میانگین مربعات خطا [۲] به عنوان روش‌های محک با نتایج الگوریتم‌های یادگیری تقویتی مورد بحث مقایسه شده است.



الف) الگوریتم DDPG



ب) الگوریتم DQN



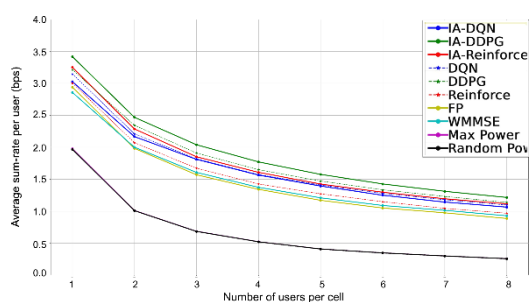
پ) الگوریتم REINFORCE

شکل (۲): مقایسه معیار $\bar{C}$ (قرمز رنگ) و C^* (سبز رنگ) در سنجش عملکرد ساختار یادگرفته‌شده به‌ازای مقادیر مختلف

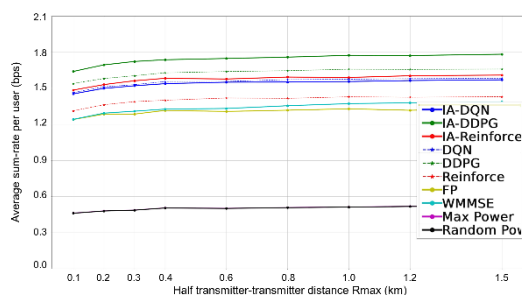
جدول (۳): مقایسه نتایج برخی از بهترین ترکیب‌های پارامترهای مدل تنفر از بی عدالتی در روش‌های پایه

IA-DDPG		DDPG	IA-DQN		DQN	IA-REINFORCE		REINFORCE	متغیر ارزیابی
$\alpha = 0.25$ $\beta = 0$	$\alpha = -0.5$ $\beta = -0.5$		$\alpha = 1$ $\beta = 1$	$\alpha = 0.25$ $\beta = 0$		$\alpha = 0.5$ $\beta = -1$	$\alpha = 0.75$ $\beta = 0.25$		
1.55	1.58	1.58	1.48	1.48	1.50	1.60	1.64	1.49	$\bar{C}$ در آموزش
1.97	1.84	1.82	1.74	1.71	1.72	1.79	1.73	1.60	$\bar{C}$ در آزمون

بررسی مدل تنفر از بی عدالتی در راهکارهای یادگیری تقویتی چندعاملی برای ارتقای کنترل توان در شبکه‌های سلولی بی سیم



(ب) بیشترین تعداد کاربران در هر سلول شبکه

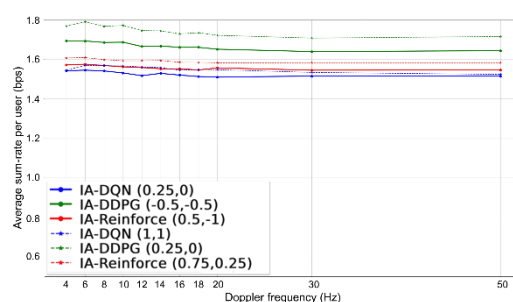


(پ) نیم فاصله سلول به سلول

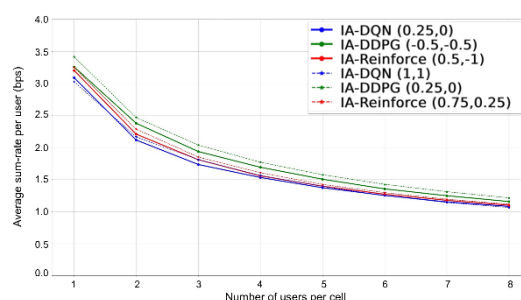
شکل (۴): بررسی تعمیم‌پذیری بهترین ترکیب پارامترهای مدل تنفر از بی عدالتی در هر یک از الگوریتم‌های یادگیری تقویتی و مقایسه نتایج با روش‌های پایه. این بررسی با توجه به مقادیر مختلف پارامترهای شبکه انجام شده است. در این شکل، منحنی‌های مربوط به الگوریتم‌های یادگیری تقویتی پایه به صورت خط چین نشان داده شده‌اند.

۷- نتیجه‌گیری

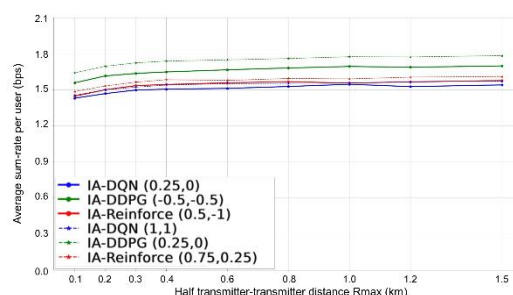
در این مقاله، عملکرد الگوریتم‌های یادگیری تقویتی عمیق برای حل مسئله تخصیص پویای توان توزیع شده در یک شبکه سلولی بی سیم با کانال دسترسی چندگانه تداخلی با هدف بیشینه‌سازی نرخ مجموع بررسی شد. در تنظیمات شبکه تحت بررسی، همکاری بین سلولی در نظر گرفته شد و برای مدیریت فرایند یادگیری، از روش آموزش متمرکز و اجرای توزیع شده بهره‌برداری شد. مدل تنفر از بی عدالتی به عنوان یک ساختار پاداش داخلی برای بهبود عملکرد الگوریتم‌های یادگیری تقویتی چندعاملی پیشنهاد و اثربخشی آن در بیشینه‌سازی نرخ مجموع مطالعه شد. نتایج آزمایش‌ها نشان می‌دهد تنظیم پارامترهای مدل تنفر از بی عدالتی تأثیری به‌سزا در بهبود عملکرد این الگوریتم‌ها به ویژه الگوریتم‌های DDPG و REINFORCE دارد. ارزیابی تعمیم‌پذیری نیز نشان داد بهره‌برداری از این پاداش داخلی با ضرایب منفی و



(الف) فرکانس داپلر

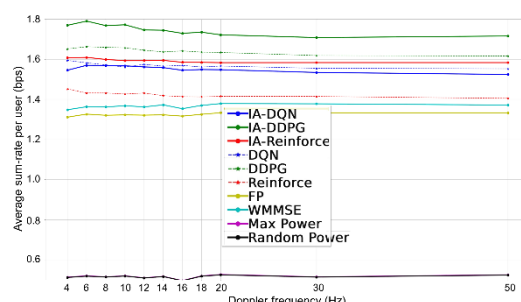


(ب) بیشترین تعداد کاربران در هر سلول شبکه



(پ) نیم فاصله سلول به سلول

شکل (۳): بررسی تعمیم‌پذیری الگوریتم‌های یادگیری تقویتی با توجه به پارامترهای انتخابی مدل تنفر از بی عدالتی. این بررسی با توجه به مقادیر مختلف پارامترهای شبکه انجام شده است.



(الف) فرکانس داپلر

- [6] K. Shen, W. Yu, "FPLinQ: A cooperative spectrum sharing strategy for device-to-device communications", IEEE International Symposium on Information Theory (ISIT), pp. 2323-2327, June 2017. <https://doi.org/10.1109/ISIT.2017.8006944>
- [7] K. Shen, W. Yu, "Fractional programming for communication systems—Part I: Power control and beamforming", IEEE Trans. on Signal Processing, Vol. 66, February 2018. <https://doi.org/10.1109/TSP.2018.2812733>
- [8] V. Mnih, K. Kavukcuoglu, D. Silver, A. A. Rusu, J. Veness, M. G. Bellemare, A. Graves, M. Riedmiller, A. K. Fidjeland, G. Ostrovski, "Human-level control through deep reinforcement learning", Nature, Vol. 518, February 2015. <https://doi.org/10.1038/nature14236>
- [9] D. Silver, A. Huang, C. J. Maddison, A. Guez, L. Sifre, G. Van Den Driessche, J. Schrittwieser, I. Antonoglou, V. Panneershelvam, M. Lanctot, "Mastering the game of Go with deep neural networks and tree search", Nature, Vol. 529, January 2016. <https://doi.org/10.1038/nature16961>
- [10] C. Berner, G. Brockman, B. Chan, V. Cheung, P. Debiak, C. Dennison, D. Farhi, Q. Fischer, S. Hashme, C. Hesse, "Dota 2 with large scale deep reinforcement learning", arXiv preprint, December 2019. <https://doi.org/10.48550/arXiv.1912.06680>
- [11] F. Meng, P. Chen, L. Wu, J. Cheng, "Power allocation in multi-user cellular networks: Deep reinforcement learning approaches", IEEE Trans. on Wireless Communications, Vol. 19, June 2020. <https://doi.org/10.1109/TWC.2020.3001736>
- [12] F. Meng, P. Chen, L. Wu, "Power allocation in multi-user cellular networks with deep Q learning approach", IEEE International Conference On Communications (ICC), pp. 1-6, December 2018. <https://doi.org/10.1109/ICC.2019.8761431>
- [13] Z. Wang, J. Zong, Y. Zhou, Y. Shi, V. W. Wong, "Decentralized multi-agent power control in wireless networks with frequency reuse", IEEE Trans. on Communications, Vol. 70, December 2021. <https://doi.org/10.1109/TCOMM.2021.3135540>
- [14] Y. Zhang, D. Guo, "Multi-Agent Reinforcement Learning for Multi-Cell Spectrum and Power Allocation", IEEE Trans. on Communications, Vol. 99, January 2025. <https://doi.org/10.1109/TCOMM.2025.3534565>
- [15] A. Kopic, E. Perenda, H. Gacanin, "A

تشویق عامل به کمک افزودن ضریبی مثبت از تفاوت‌های محاسبه‌شده به پاداش عامل در فرایند یادگیری الگوریتم DDPG بهترین عملکرد را دارد. علاوه بر این، همه رویکردهای مبتنی بر داده از روش‌های مبتنی بر مدل پیشرفته بهتر عمل می‌کنند و همچنین، عملکرد تعمیم‌پذیری بهتری دارند.

الگوریتم‌های یادگیری تقویتی عمیق به عنوان دسته‌ای از الگوریتم‌های مبتنی بر داده، راهکاری امیدوارکننده برای شبکه‌های هوشمند آینده است و الگوریتم DDPG ترکیب‌شده با مدل تنفر از بی‌عدالتی را می‌توان برای مسائلی متنوع با فضای حالت/عمل گسسته یا پیوسته و مسائل بهینه‌سازی توأم متغیرهای متعدد اعمال کرد. این الگوریتم را می‌توان برای بسیاری از مسائل مانند زمان‌بندی کاربر، مدیریت کانال و تخصیص توان در شبکه‌های ارتباطی مختلف اعمال کرد.

مراجع

- [1] A. Gjendemsjo, D. Gesbert, G. E. Oien, S. G. Kiani, "Binary power control for sum rate maximization over multiple interfering links", IEEE Trans. on Wireless Communications, Vol. 7, No. 8, August 2008. <https://doi.org/10.1109/TWC.2008.070227>
- [2] Q. Shi, M. Razaviyayn, Z. -Q. Luo, C. He, "An iteratively weighted MMSE approach to distributed sum-utility maximization for a MIMO interfering broadcast channel", IEEE Trans. on Signal Processing, Vol. 59, No. 9, September 2011. <https://doi.org/10.1109/TSP.2011.2147784>
- [3] L. Song, D. Niyato, Z. Han, E. Hossain, "Game-theoretic resource allocation methods for device-to-device communication", IEEE Wireless Communications, Vol. 21, No. 3, June 2014. <https://doi.org/10.1109/MWC.2014.6845058>
- [4] N. Naderializadeh, A. S. Avestimehr, "ITLinQ: A new approach for spectrum sharing in device-to-device communication systems", IEEE Journal on Selected Areas in Communications, Vol. 32, No. 6, June 2014. <https://doi.org/10.1109/jsac.2014.2328102>
- [5] X. Yi, G. Caire, "ITLinQ+: An improved spectrum sharing mechanism for device-to-device communications", 49th Asilomar Conference on Signals, Systems and Computers, pp. 1310-1314, November 2015. <https://doi.org/10.1109/ACSSC.2015.742135>

- Ramazi, "Inequity aversion reduces travel time in the traffic light control problem", arXiv preprint, 2023. <https://doi.org/10.48550/arXiv:2302.12053>
- [21] R. S. Sutton, D. McAllester, S. Singh, Y. Mansour, "Policy gradient methods for reinforcement learning with function approximation", *Advances in Neural Information Processing Systems*, Vol. 12, 1999. https://proceedings.neurips.cc/paper_files/paper/1999/file/464d828b85b0bed98e80ade0a5c43b0f-Paper.pdf
- [22] T. P. Lillicrap, J. J. Hunt, A. Pritzel, N. Heess, T. Erez, Y. Tassa, D. Silver, D. Wierstra, "Continuous control with deep reinforcement learning", arXiv preprint, 2015. <https://doi.org/10.48550/arXiv:1509.02971>
- [23] E. Hughes, J. Z. Leibo, M. Phillips, K. Tuyls, E. Dueñez-Guzman, A. García Castañeda, I. Dunning, T. Zhu, K. McKee, R. Koster, "Inequity aversion improves cooperation in intertemporal social dilemmas", *Advances in Neural Information Processing Systems*, Vol. 31, 2018. https://proceedings.neurips.cc/paper_files/paper/2018/file/7fea637fd6d02b8f0adf6f7dc36aed93-Paper.pdf
- collaborative multi-agent deep reinforcement learning-based wireless power allocation with centralized training and decentralized execution", *IEEE Trans. on Communications*, Vol. 72, November 2024. <https://doi.org/10.1109/TCOMM.2024.3409530>
- [16] H. Kim, J. So, "Distributed multi-agent deep reinforcement learning-based transmit power control in cellular networks", *Sensors*, Vol. 25, No. 13, June 2025. <https://doi.org/10.3390/s25134017>
- [17] Z. Huang, T. Li, C. Song, Z. Li, J. Wang, X. Liu, H. Chen, X. Zhao, Y. Cao, "Joint spectrum and power allocation scheme based on value decomposition networks in D2D communication networks", *EURASIP Journal on Wireless Communications and Networking*, Vol. 79, September 2024. <https://doi.org/10.1186/s13638-024-02393-1>
- [18] A. Reuel, D. Ma, "Fairness in reinforcement learning: A survey", *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, pp. 1218–1230, 7 February 2025. <https://doi.org/10.5555/3716662.3716769>
- [19] F. Alamiyan-Harandi, P. Ramazi, "Environmental-impact-based multi-agent reinforcement learning", *Applied Sciences*, Vol. 14, No. 15, July 2024. <https://doi.org/10.3390/app14156432>
- [20] M. Hassanjani, F. Alamiyan-Harandi, P.

¹ Weighted Minimum Mean Squared Error (WMMSE)

² Fractional Programming (FP)

³ Downlink

⁴ Deep Q Network

⁵ Multi Agent Reinforcement Learning (MARL)

⁶ Centralized Training Distributed Execution (CTDE)

⁷ Long Short Term Memory (LSTM)

⁸ Social welfare functions

⁹ Value Decomposition Networks

¹⁰ Inequity Aversion (IA)

¹¹ Cleanup and Harvest Games

¹² Envy

¹³ Guilt

¹⁴ Interfering Multiple-Access Channel (IMAC)

¹⁵ Sum-rate

¹⁶ Single-Input Single-Output (SISO) Interfering Broadcast Channel (IBC)

¹⁷ Base Station (BS)

¹⁸ Geometric attenuation

¹⁹ Shadow fading

²⁰ Flat fading

²¹ Signal-to-Interference-plus-Noise Ratio (SINR)

- ²² NP hard
- ²³ Channel State Information (CSI)
- ²⁴ Multi-objective programming
- ²⁵ Transfer Learning
- ²⁶ Monte-Carlo
- ²⁷ Stochastic gradient ascent
- ²⁸ Deep Q Network
- ²⁹ Deep Deterministic Policy Gradient
- ³⁰ Intrinsic Rewards
- ³¹ Advantageous Inequity
- ³² Disadvantageous Inequity
- ³³ Kronecker product
- ³⁴ Episode
- ³⁵ Generalization

