

# انتقال دانش تنظیم شده برای یادگیری تقویتی



## چندعاملی

نیلوفر علوی و جعفر طهمورث نژاد\*

دانشگاه صنعتی ارومیه دانشکده مهندسی فناوری اطلاعات و کامپیوتر،

دانشگاه صنعتی ارومیه، ارومیه، ایران

### چکیده

یادگیری تقویتی به آموزش مدل‌های یادگیری ماشین برای اتخاذ تصمیمات متوالی اشاره می‌کند که در آن یک عامل از طریق تعامل با محیط، آموزش دیده، نتایج این تعامل را مشاهده کرده و بر این اساس، پاداش مثبت یا منفی دریافت می‌کند. یادگیری تقویتی کاربردهای زیادی برای سیستم‌های چندعاملی به خصوص در محیط‌های پویا و ناشناخته دارد. با این حال، بیش‌تر الگوریتم‌های یادگیری تقویتی چندعاملی با مشکلاتی همچون پیچیدگی محاسباتی نمایی برای محاسبه فضای حالت مشترک مواجه هستند که منجر به عدم مقیاس‌پذیری الگوریتم‌ها در مسائل چندعاملی واقعی می‌شود. کاربردهای یادگیری تقویتی چندعاملی را می‌توان از فوتبال ربات‌ها، شبکه‌ها، محاسبات ابری، زمانبندی شغل تا اعزام نیروی واکنشی دسته‌بندی کرد. در این مقاله یک الگوریتم جدید به نام انتقال دانش تنظیم شده برای یادگیری تقویتی چندعاملی (RKT-MARL) معرفی می‌شود که براساس مدل تصمیم‌گیری مارکوف کار می‌کند. این الگوریتم برخلاف روش‌های یادگیری تقویتی سنتی، مفاهیم تعاملات پراکنده و انتقال دانش را برای رسیدن به تعادل بین عامل‌ها استفاده می‌کند. علاوه بر این، RKT-MARL از سازوکار مذاکره برای یافتن مجموعه تعادل و از روش کمینه واریانس برای انتخاب بهترین عمل در مجموعه تعادل به دست آمده استفاده می‌کند. همچنین الگوریتم پیشنهادی، دانش مقادیر حالت-عمل را در میان عامل‌های مختلف انتقال می‌دهد. از طرفی، الگوریتم RKT-MARL مقادیر  $Q$  را در حالت‌های هماهنگی به عنوان ضریبی از اطلاعات محیطی جاری و دانش قبلی مقداردهی می‌کند. به منظور ارزیابی عملکرد روش پیشنهادی، یک گروه از آزمایش‌ها بر روی پنج بازی جهانی انجام شده و نتایج حاصل بیانگر همگرایی سریع و مقیاس‌پذیری بالا در RKT-MARL است.

واژگان کلیدی: یادگیری تقویتی چندعاملی، انتقال دانش، تعادل‌های متناوب، تنظیم‌پذیری، تعاملات پراکنده، مذاکره عامل‌ها.

## Regularized Knowledge Transfer for Multi-Agent Reinforcement Learning

Niloofer Alavi & Jafar Tahmoresnezhad\*

Faculty of IT & Computer Engineering, Urmia University of Technology, Urmia, Iran

### Abstract

Reinforcement learning (RL) refers to the training of machine learning models to make a sequence of decisions on which an agent learns by interacting with its environment, observing the results of interactions and receiving a positive or negative reward, accordingly. RL has many applications for multi-agent systems, especially in dynamic and unknown environments. However, most multi-agent reinforcement learning (MARL) algorithms suffer from some problems specifically the exponential computational complexity to calculate the joint state-action space, which leads to the lack of scalability of algorithms in realistic multi-agent problems. Applications of MARL can be categorized from robot soccer, networks, cloud computing, job scheduling, and to optimal reactive power dispatch.

In the area of reinforcement learning algorithms, there are serious challenges such as the lack of application of equilibrium-based algorithms in practice and high computational complexity to find

\* Corresponding author

\* نویسنده عهده‌دار مکاتبات

equilibrium. On the other hand, since agents have no concept of equilibrium policies, they tend to act aggressively toward their goals, which it results the high probability of collisions.

Consequently, in this paper, a novel algorithm called Regularized Knowledge Transfer for Multi-Agent Reinforcement Learning (RKT-MARL) is presented that relies on Markov decision process (MDP) model. RKT-MARL unlike the traditional reinforcement learning methods exploits the sparse interactions and knowledge transfer to achieve an equilibrium across agents. Moreover, RKT-MARL benefits from negotiation to find the equilibrium set. RKT-MARL uses the minimum variance method to select the best action in the equilibrium set, and transfers the knowledge of state-action values across various agents. Also, RKT-MARL initializes the Q-values in coordinate states as coefficients of current environmental information and previous knowledge. In order to evaluate the performance of our proposed method, groups of experiments are conducted on five grid world games and the results show the fast convergence and high scalability of RKT-MARL. Therefore, the fast convergence of our proposed method indicates that the agents quickly solve the problem of reinforcement learning and approach to their goal.

**Keyword:** Multi-agent reinforcement learning, Knowledge transfer, Meta and Nash equilibriums, Regularization, Sparse interactions, Agents negotiations.

## ۱- مقدمه

دانشمندان و پژوهشگران در حوزه سیستم‌های چندعاملی و یادگیری ماشین علاقه بسیاری به یادگیری‌های چندعاملی نشان داده‌اند. یادگیری چندعاملی یک تکنیک کاربردی برای عامل‌ها است تا بتوانند رفتارهای هماهنگ و مؤثری را در سیستم‌های چندعاملی بیاموزند. در یادگیری چندعاملی، فرآیندهای یادگیری توزیع شده و هم‌زمان، محیط یادگیری را برای عامل‌ها پویا می‌کنند. توسعه یک رویکرد یادگیری مؤثر برای هماهنگ کردن رفتارهای عامل‌ها در یک محیط پویا، به‌خصوص زمانی که عامل‌ها ساختار محیط را نمی‌دانند و فقط مشاهدات محلی از محیط دارند، مسئله دشواری است [۱].

یکی از پرکاربردترین روش‌های یادگیری در سیستم‌های چندعاملی، یادگیری تقویتی است. در فرآیند یادگیری تقویتی، یک عامل، فضای استراتژی‌های ممکن را بررسی کرده و بازخوردی را در مورد نتیجه انتخاب‌هایی که دارد، دریافت می‌کند [۲]. بنابراین عامل یاد می‌گیرد که برای دریافت پاداش بیشتر از محیط، چه عمل‌هایی را باید انتخاب کند. نکته قابل توجه در یادگیری تقویتی این است که انجام یک عمل توسط عامل علاوه بر دریافت پاداش فوری، بر حالت‌های بعدی و همچنین مجموع پاداش‌ها نیز تأثیر می‌گذارد [۳، ۴ و ۵].

یکی از تکنیک‌های کلیدی در یادگیری چندعاملی، یادگیری تقویتی چندعاملی<sup>۱</sup> (MARL) است که توسعه‌ای از یادگیری تقویتی در محیط‌های چندعاملی است. یادگیری تقویتی چندعاملی، روشی برای عامل‌های خودمختار فراهم می‌کند تا بتوانند مسائل تصمیم‌گیری متوالی را در سیستم‌های چندعاملی حل کنند [۵ و ۶].

<sup>1</sup> Multi-agent reinforcement learning

چندین مدل ریاضی به‌عنوان چارچوب یادگیری تقویتی چندعاملی ساخته شده‌اند، که می‌توان به بازی‌های مارکوف<sup>۲</sup> (MG) [۷] و فرآیندهای تصمیم‌گیری مارکوف با تعاملات پراکنده غیرمتمرکز<sup>۳</sup> (Dec-SIMDP) [۸] اشاره کرد. فرض بازی‌های مارکوف این است که عامل‌ها در تمام فضای مشترک حالت-عمل<sup>۴</sup>، کاملاً مشاهده‌پذیر باشند. درحالی‌که فرض الگوریتم‌های مبتنی بر فرآیند تصمیم‌گیری مارکوف با تعاملات پراکنده غیرمتمرکز، برپایه مشاهدات محلی عامل‌ها است [۹]. به دلیل اینکه عامل‌ها در فرآیند تصمیم‌گیری مارکوف با تعاملات پراکنده غیرمتمرکز، خارج از مناطق تعاملی قرار دارند، با فرآیند تصمیم‌گیری مارکوف تک‌عاملی مدل‌سازی می‌شوند، ولی در صورتی که عامل‌ها در داخل مناطق تعاملی قرار داشته باشند، از بازی مارکوف برای مدل‌سازی آن‌ها استفاده می‌شود [۱۰].

چالش‌های موجود در حوزه یادگیری تقویتی شامل موارد زیر می‌باشند: (۱) برخلاف پیشرفت سریع در تئوری‌ها و الگوریتم‌های یادگیری تقویتی چندعاملی، کاربردهای عملی این یادگیری با توجه به محدودیت‌های روش‌های موجود به تلاش‌های بسیاری نیاز دارد، (۲) الگوریتم‌های یادگیری تقویتی چندعاملی مبتنی بر تعادل، متکی بر فرآیندهای یادگیری با اتصالات محکم<sup>۵</sup> هستند که مانع از کاربرد آن‌ها در عمل می‌شوند، (۳) به دست آوردن تعادل (به‌عنوان نمونه، تعادل نش) در هر گام زمانی، حتی برای محیط‌های نسبتاً کوچک با دو یا سه عامل، پیچیدگی محاسباتی زیادی دارد، (۴) از طرفی در الگوریتم‌های یادگیری تقویتی چندعاملی با تعاملات

<sup>2</sup> Markov games

<sup>3</sup> Decentralized sparse interaction markov decision processes

<sup>4</sup> State - action

<sup>5</sup> Tightly coupled learning process

دیگر کار کرده و برای جلوگیری از برخورد هماهنگی برقرار کنند [۱۳، ۱۴]. به کارگیری سازوکار انتقال دانش در الگوریتم‌های یادگیری تقویتی، سرعت یادگیری عامل‌ها را بالا برده [۱۵] و عملکرد الگوریتم‌ها را بهبود می‌بخشد و عامل‌ها را در برقراری هماهنگی با یکدیگر توانمندتر می‌کند. نکته قابل توجه این است که انتقال چه نوع دانشی می‌تواند برای یادگیری سیاست‌های بهتر در سیستم‌های چندعاملی، سودمند باشد. در بعضی سیستم‌های چندعاملی، عامل‌ها دانش تک‌عاملی (به عنوان نمونه، تابع مقدار محلی<sup>۵</sup> یا سیاست محلی) را قبل از فرآیند یادگیری چندعاملی می‌آموزند و استفاده از چنین دانشی در طول یادگیری ممکن است، مفید باشد.

به منظور ارزیابی کارایی الگوریتم پیشنهادی در این مقاله، چندین مجموعه داده برای نشان دادن عملکرد الگوریتم از لحاظ همگرایی، مقیاس‌پذیری و قابلیت هماهنگی مورد آزمایش قرار گرفته و نتایج آن‌ها، با دیگر الگوریتم‌های یادگیری تقویتی چندعاملی مقایسه شده‌است. نتایج حاصل، نشان‌دهنده برتری قابل ملاحظه الگوریتم پیشنهادی نسبت به الگوریتم‌های دیگر در این حوزه‌است.

ساختار ادامه مقاله به صورت زیر است. در بخش ۲ پژوهش‌های مرتبط مورد بررسی قرار می‌گیرند. در بخش ۳ تعاریف پایه‌ای و روش پیشنهادی گنجانده شده‌است. بخش ۴ شامل مجموعه داده‌ها و الگوریتم‌های مورد ارزیابی و فرضیات آزمایش‌ها است. در بخش ۵ نتایج حاصل از عملکرد الگوریتم پیشنهادی و مقایسه آن با دیگر الگوریتم‌های موجود، قرار دارد؛ و در نهایت در بخش ۶ نتیجه‌گیری و کارهای پیشنهادی آینده ارائه شده‌است.

## ۲- کارهای پیشین

در حوزه یادگیری ماشین، روش‌های بسیاری در زمینه یادگیری تقویتی پیشنهاد شده، که اکثر آن‌ها از الگوریتم یادگیری  $Q$ <sup>۶</sup> [۱۶] برای برقراری هماهنگی بین عامل‌ها استفاده می‌کنند. به طور کلی، روش‌های پیشنهاد شده در زمینه یادگیری تقویتی چندعاملی، به سه دسته تقسیم می‌شوند: روش‌های مبتنی بر مفهوم تعادل<sup>۷</sup>، روش‌های مبتنی بر سازوکار انتقال دانش<sup>۸</sup> و روش‌های مبتنی بر تعاملات پراکنده<sup>۹</sup>. به طور خلاصه، انواع روش‌های یادگیری تقویتی چندعاملی در جدول (۱) نشان داده شده‌است.

پراکنده<sup>۱</sup>، عامل‌ها هیچ مفهومی از سیاست تعادلی ندارند و به صورت مهاجمانه به سمت هدف خود حرکت می‌کنند که به احتمال زیاد منجر به برخورد می‌شود [۱۱]. در نتیجه برای مواجهه با چالش‌های ذکر شده، در این مقاله تمرکز اصلی بر این است که چگونه سازوکار تعادل و انتقال دانش، در الگوریتم‌های مبتنی بر تعاملات پراکنده و الگوریتم‌های یادگیری تقویتی چندعاملی مبتنی بر مذاکرات و تعاملات پراکنده استفاده شود [۱۲].

الگوریتم پیشنهادی در این مقاله با عنوان RKT-MARL<sup>۲</sup> شامل چهار مرحله اصلی زیر است: (۱) ایجاد یک چارچوب مبتنی بر مذاکره برای تعاملات پراکنده، (۲) مذاکره برای به دست آوردن مجموعه تعادل، (۳) روش کمینه واریانس<sup>۳</sup> برای انتخاب یک عمل مشترک، و (۴) انتقال دانش مقادیر  $Q$  محلی<sup>۴</sup> به صورت تنظیم شده. الگوریتم RKT-MARL ابتدا بر پایه مدل فرآیند تصمیم‌گیری مارکوف شروع شده و فرض اولیه بر این است که عامل‌ها سیاست بهینه تک‌عاملی را قبل از یادگیری در محیط چندعاملی به دست می‌آورند؛ سپس عامل‌ها برای به دست آوردن مجموعه تعادل در حالت هماهنگی، باهم مذاکره کرده و از روش کمینه واریانس برای انتخاب بهترین عمل مشترک استفاده می‌کنند. در مرحله بعد، عامل‌ها طبق این عمل مشترک انتخاب شده به حالت بعدی حرکت و پاداش فوری خود را دریافت می‌کنند. در نهایت، مقدار  $Q$  با توجه به سود تعادلی و پاداش عامل‌ها، به روز می‌شود؛ همچنین نکته قابل توجه این است که در مرحله چهارم که انتقال دانش صورت می‌گیرد،  $Q$  در حالت هماهنگی مشترک به صورت تنظیم شده، توسط مجموعی از اطلاعات محیطی جاری و دانش آموزش داده شده قبلی، مقداردهی اولیه می‌شود. این ضرایب تنظیم باعث می‌شوند که بتوان درصد اثرگذاری اطلاعات جاری و دانش قبلی را در مقداردهی اولیه  $Q$  کنترل کرد؛ بنابراین، مقدار بالاتر ضرایب نشان‌دهنده تاثیر بیشتر اطلاعات جاری و دانش قبلی است.

در سال‌های اخیر، از سازوکار انتقال دانش برای بهبود سرعت در مسائل یادگیری تقویتی چندعاملی استفاده شده‌است. انتقال دانش به عامل‌ها اجازه می‌دهد که ابتدا بر روی یک محیط منبع ساده، یادگیری انجام داده و سپس اطلاعات به دست آمده از آن را در صورت لزوم به محیط‌های پیچیده‌تری انتقال دهند، تا بتوانند با عامل‌های

<sup>۵</sup> Local value function

<sup>۶</sup> Q-Learning

<sup>۷</sup> Equilibrium-based methods

<sup>۸</sup> Knowledge transfer-based methods

<sup>۹</sup> Sparse-interaction-based methods

<sup>۱</sup> Sparse interaction

<sup>۲</sup> Regularized knowledge transfer for multi-agent reinforcement learning

<sup>۳</sup> Minimum variance method

<sup>۴</sup> Local Q-value

(جدول - ۱): انواع روش‌های یادگیری تقویتی چندعاملی

(Table-1). Types of multi-agent reinforcement learning methods

| ویژگی هر الگوریتم                                                                                                                                           | نمونه الگوریتم    | اساس کار هر روش                                                                         |                                      |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------|-----------------------------------------------------------------------------------------|--------------------------------------|
| • استفاده از سه نوع استراتژی پروفایل غیراحتمالی به جای استراتژی پروفایل احتمالی، به عنوان مفهوم تعادل                                                       | الگوریتم NegoQ    | • استفاده از مدل بازی مارکوف به عنوان چارچوب یادگیری تقویتی                             | (۱) روش‌های مبتنی بر تعادل           |
| • استفاده از یک فرآیند مذاکره چندمرحله‌ای برای یافتن سه استراتژی پروفایل                                                                                    |                   | • استفاده از مفاهیم تعادلی به عنوان سیاست بهینه                                         |                                      |
| • استفاده از سه سازوکار انتقال دانش (۱) انتقال تابع مقدار، (۲) انتقال تابع مقدار انتخابی، و (۳) انتزاع بازی مبتنی بر انتقال مدل برای یادگیری سیاست‌های بهتر | الگوریتم NegoQVFT | • بهره گرفتن از دانش تک‌عاملی از قبل آموخته‌شده در طول یادگیری چندعاملی                 | (۲) روش‌های مبتنی بر انتقال دانش     |
|                                                                                                                                                             |                   | • انتقال دانش و تجربیات مورد نیاز به عامل از یک محیط ساده به یک محیط پیچیده             |                                      |
| • بهره بردن از سه نوع سازوکار انتقال دانش                                                                                                                   | الگوریتم NegoQVFT | • استفاده از تعاملات پراکنده در سیستم‌های چندعاملی به منظور بهبود عملکرد یادگیری تقویتی | (۳) روش‌های مبتنی بر تعاملات پراکنده |
| • استفاده از تعاملات پراکنده                                                                                                                                |                   |                                                                                         |                                      |

به اشتراک‌گذاری این تابع غیرمنطقی بوده و همچنین، به دلیل بالا بودن پیچیدگی محاسباتی استراتژی احتمالی، این الگوریتم‌ها کاربرد گسترده‌ای در عمل ندارند. الگوریتم NegoQ از جمله روش‌های مبتنی بر تعادل است، که برای حل دو محدودیت بالا پیشنهاد شده و دو ایده کلی دارد: (۱) به جای استراتژی پروفایل احتمالی، از سه نوع استراتژی پروفایل غیر احتمالی<sup>۴</sup> به عنوان مفهوم تعادل استفاده می‌کند، و (۲) به دلیل اینکه عامل‌ها نمی‌توانند تابع مقدار عامل‌های دیگر را رونوشت کنند، یک فرایند مذاکره چندمرحله‌ای برای یافتن سه استراتژی پروفایل غیر احتمالی استفاده می‌کند.

در روش‌های مبتنی بر انتقال دانش NegoSI<sup>۵</sup> [۱۲، ۱۵، ۱۹] و NegoQ with VFT<sup>۶</sup> [۲۰]، عامل‌ها دانش تک‌عاملی (به عنوان نمونه، تابع مقدار محلی یا سیاست محلی) را قبل از فرآیند یادگیری چندعاملی می‌آموزند و استفاده از چنین دانشی در طول یادگیری، می‌تواند مفید باشد. در این روش‌ها می‌توان انواع دانش و تجربیات مورد نیاز، مثل تابع مقدار، تابع مقدار محلی، سیاست [۱۹] و تعادل را به عامل‌ها انتقال داد که ایده کلیدی انتقال تعادل [۱۵]، استفاده مجدد از تعادل‌هایی است که در قبل محاسبه شده‌اند. محاسبه تعادل و استفاده مجدد از آن برای انتخاب عمل و به‌روزرسانی توابع مقدار، از مهم‌ترین مراحل روش‌های مبتنی بر تعادل نیز محسوب

از جمله کارهای قبلی در حوزه یادگیری تقویتی می‌توان به الگوریتم یادگیری CQ<sup>۱</sup> [۱۷] اشاره کرد که در آن عامل‌ها می‌توانند به طور هم‌زمان از رشد نمایی فضای حالت جلوگیری کرده و همچنین سیاست‌های کارآمدی را در محیط‌های چندعاملی بیاموزند. ضعف این الگوریتم این است که جزو روش‌های مبتنی بر تعادل نیست و در آن عامل‌ها در یک حالت مشترک خاص به جای استراتژی‌های تعادلی، به صورت حریصانه عمل می‌کنند که در نهایت ممکن است، منجر به برخورد شوند. همان‌طور که در جدول (۱) نشان داده شده، روش‌های مبتنی بر تعادل NegoQ<sup>۲</sup> [۱۸]، یکی از مهم‌ترین روش‌های یادگیری تقویتی چندعاملی هستند که مفاهیم تعادلی را در تئوری بازی‌ها در نظر می‌گیرند، به طوری که عامل‌ها در هر حالت، استراتژی‌های تعادلی را انتخاب می‌کنند. این روش‌ها از مدل بازی مارکوف، به عنوان چارچوب یادگیری تقویتی و از مفاهیم تعادلی مختلف برای یادگیری سیاست بهینه استفاده می‌کنند. در روش‌های مبتنی بر تعادل، مفهوم تعادل هم به یافتن سیاست بهینه توسط عامل‌ها و هم به جلوگیری از برخورد بین عامل‌ها کمک می‌کند؛ با این حال، در بیشتر الگوریتم‌های یادگیری تقویتی چندعاملی مبتنی بر تعادل، عامل‌ها باید توابع مقدار عامل‌های دیگر را رونوشت کرده و از استراتژی پروفایل احتمالی<sup>۳</sup> استفاده کنند. با توجه به اینکه تابع مقدار، جزو اطلاعات خصوصی یک عامل است،

<sup>۴</sup> Pure strategy profile<sup>۵</sup> Negotiation-based multi-agent reinforcement learning with sparse interaction<sup>۶</sup> Negotiation-based Q-learning with value function transfer<sup>۱</sup> CQ-Learning<sup>۲</sup> Negotiation-based Q-learning<sup>۳</sup> Mixed strategy profile

### ۳- تعریف مسأله

در این بخش، مفاهیم و تعاریف اصلی در یادگیری تقویتی چندعاملی ارائه شده و در ادامه تعریف مسئله به طور کامل شرح داده می شود.

#### ۳-۱- مدل فرآیند تصمیم گیری مارکوف<sup>۱</sup> (MDP)

این مدل یکی از مدل های استاندارد تصمیم گیری در یادگیری تقویتی است. فرآیند تصمیم گیری مارکوف، یک مسئله تصمیم گیری متوالی را مطرح می کند که در آن عامل در هر گام زمانی، عملی را برای افزایش پاداش انتخاب می کند. MDP، یک چندتایی  $\langle S, A, R, T \rangle$  است که  $S$  فضای حالت ها و  $A$  فضای عمل های یک عامل،  $R: S \times A \rightarrow \mathbb{R}$  تابع نگاشت پاداش به زوج حالت-عمل و  $T: S \times A \times S \rightarrow [0,1]$  تابع انتقال را مشخص می کند. هدف مدل MDP این است که عامل ها بتوانند سیاست بهینه ای را پیدا کنند تا بعضی معیارهای بهینه سازی مبتنی بر پاداش را ماکزیمم کنند. از جمله این معیارها، مجموع پاداش های مورد انتظار به صورت رابطه (۱) است [۸]:

$$V^*(s) = \max_{\pi} E_{\pi} \left\{ \sum_{k=0}^{\infty} \gamma^k r^{t+k} \mid s^t = s \right\} \quad (1)$$

که  $V^*(s)$  مقدار حاصل در حالت  $s$  تحت سیاست بهینه را نشان می دهد.  $\pi: S \times A \rightarrow [0,1]$  سیاست هر عامل،  $t$  هر گام زمانی،  $k$  گام زمانی بعدی،  $r^{t+k}$  پاداش حاصل در گام زمانی  $(t+k)$  و  $\gamma \in [0,1]$  پارامتر تنزیل<sup>۲</sup> را مشخص می کنند.

هدف MDP را می توان با مقادیر  $Q$  برای هر زوج حالت-عمل به صورت رابطه (۲) بیان کرد [۸]:

$$Q^*(s, a) = r(s, a) + \gamma \sum_s T(s, a, s') \max_{a'} Q^*(s', a') \quad (2)$$

که  $Q^*(s, a)$  مقدار بهینه در زوج حالت-عمل  $(s, a)$  تحت سیاست بهینه،  $s'$  حالت بعدی،  $a'$  عمل بعدی،  $r(s, a)$  پاداش فوری عامل با انتخاب عمل  $a$  در حالت  $s$  و  $T(s, a, s')$  احتمال انتقال از حالت  $s$  به حالت  $s'$  با انجام عمل  $a$  را نشان می دهند.

از الگوریتم یادگیری  $Q$  می توان برای تخمین  $Q^*(s, a)$  استفاده کرد که قانون به روزرسانی یک مرحله ای آن به صورت رابطه (۳) است [۸]:

$$Q(s, a) \leftarrow (1 - \alpha) Q(s, a) + \alpha [r(s, a) + \gamma \max_{a'} Q(s', a')] \quad (3)$$

می شود. به طور کلی، مفهوم انتقال دانش، بدین معنی است که عامل پس از حل یک مسئله یادگیری تقویتی ساده، دانش آن را به یک مسئله پیچیده تر منتقل می کند [۲۱]. الگوریتم NegoQ with VFT، به جای قوانین یادگیری  $Q$ ، از تئوری بازی ها به عنوان رویکرد پایه یادگیری استفاده می کند. این الگوریتم، جزو روش های مبتنی بر انتقال دانش و روش های مبتنی بر تعاملات پراکنده است که از سه سازوکار انتقال دانش (انتقال تابع مقدار، انتقال تابع مقدار انتخابی و انتزاع بازی مبتنی بر انتقال مدل) برای یادگیری سیاست های بهتر در سیستم های چندعاملی با تعاملات پراکنده استفاده می کند.

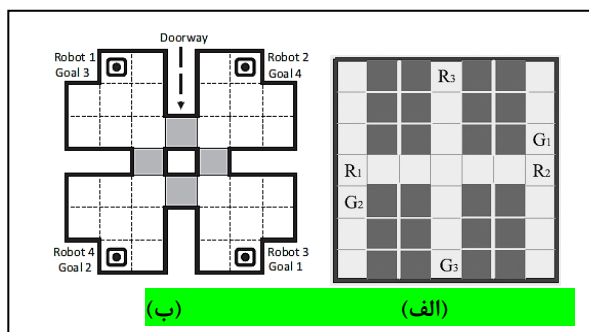
نوع سوم از روش های یادگیری تقویتی چندعاملی، روش های مبتنی بر تعاملات پراکنده [۱۱، ۲۰، ۲۲ و ۲۳] هستند. در این روش ها، به دلیل اینکه امکان دارد در همه حالت ها، عامل ها با یکدیگر تعامل نداشته باشند و همیشه این تعاملات، شامل همه عامل ها نباشد، تعاملات بین عامل ها پراکنده است. استفاده از تعاملات پراکنده در سیستم های چندعاملی، می تواند عملکرد یادگیری تقویتی را بهبود بیشتری ببخشد. در الگوریتم NegoQ with VFT، به دلیل پراکنده بودن تعاملات بین عامل ها، هر عامل می تواند با اندک دانشی که در اختیار دارد، خود را با کمترین تغییرات محیط تطبیق دهد.

در این مقاله، برای تمرکز بر حل مسائل یادگیری در سیستم های پیچیده، الگوریتم RKT-MARL به صورت ترکیبی از روش های مبتنی بر تعادل، روش های مبتنی بر تعاملات پراکنده و روش های مبتنی بر انتقال دانش پیشنهاد می شود. با توجه به اینکه الگوریتم های مبتنی بر تعادل با اتصالات محکم، در سیستم های پیچیده غیرعملی بوده و همچنین، الگوریتم های مبتنی بر تعاملات پراکنده، موجب برخوردهای زیادی می شوند، در این مقاله یک روش بهینه برای حل این مشکلات پیشنهاد شده است. بدین ترتیب در الگوریتم پیشنهادی، یک چارچوب یادگیری مبتنی بر تعاملات پراکنده در نظر گرفته شده، به طوری که هر عامل، عمل مشترک تعادلی را در مناطق هماهنگی انتخاب می کند. بنابراین در مقایسه با الگوریتم های مبتنی بر تعاملات پراکنده، در این الگوریتم، عامل ها هنگام انتخاب عمل مشترک، تعادل را نیز در نظرمی گیرند. همچنین برخلاف الگوریتم های مبتنی بر انتقال دانش قبلی، در این الگوریتم برای مقداردهی اولیه  $Q$  از ضرب تنظیم پذیری بین اطلاعات آموزش داده شده قبلی و اطلاعات محیطی جاری استفاده شده که بین دو دانش، تنظیم ایجاد می کند و همچنین در محیط های پیچیده تر برای بالابردن سرعت یادگیری، دانش مناسب تری به عامل ها انتقال داده می شود.

<sup>1</sup> Markov decision process

<sup>2</sup> Discount factor

هماهنگی برقرار کنند؛ اما این فرضیه همیشه در عمل برقرار نیست و حقیقت این است که در بیش‌تر سیستم‌های چندعاملی، عامل‌های یادگیرنده با توجه به تعاملات محدود در بعضی مناطق، هماهنگی ضعیف‌تری با یکدیگر دارند، زیرا ممکن است، تعاملات بین عامل‌ها در تمام حالت‌ها رخ ندهد و یا همیشه شامل همه عامل‌ها نباشد. به‌عنوان مثالی از سیستم‌های چندعاملی با تعاملات پراکنده، سیستم انبار هوشمند<sup>۲</sup> در شکل (۱) (الف) در نظر گرفته می‌شود که در آن، ربات‌های خودمختار تنها زمانی به ربات‌های دیگر توجه می‌کنند که به‌اندازه کافی به یکدیگر نزدیک باشند، در غیر این صورت به‌طور مستقل عمل می‌کنند.



(شکل-۱): مثال هایی از سیستم‌های چندعاملی با تعاملات پراکنده. (الف) بازی انبار هوشمند که در آن نماد R و G به ترتیب نشان‌دهنده ربات‌های خودمختار و هدف هر یک از عامل‌هاست. (ب) بازی چندعاملی مبتنی بر تعاملات پراکنده [۲۰].

(Figure-1): Examples of multi-agent systems with sparse interactions. a) Intelligent warehouse game, in which the R and G symbols represent the autonomous robots and the goal of each agent, respectively. b) Multi-agent game based on the sparse interactions [20].

شکل (۱) (ب) را نیز می‌توان به‌عنوان مثالی دیگر از تعاملات پراکنده در نظر گرفت که در آن هر چهار ربات باید راهی برای رسیدن به اهداف خود پیدا کنند. اکثر اوقات، ربات‌ها می‌توانند آزادانه حرکت کنند ولی در دروازه باریک، هر زمان فقط یک ربات می‌تواند عبور کند و این بدین معنی است که ربات‌ها باید در اطراف دروازه هماهنگی برقرار کنند.

#### ۴- روش پیشنهادی

در این بخش، الگوریتم RKT-MARL برای حل مسائل یادگیری در سیستم‌های پیچیده، با جزئیات بیشتر توضیح داده می‌شود. هدف روش پیشنهادی در این مقاله ارائه الگوریتمی است که با بهره‌گیری از راه‌حل‌های تعادلی، از

که  $Q(s, a)$  مقدار زوج حالت-عمل  $(s, a)$  بوده و  $\alpha \in [0, 1]$  پارامتری به نام نرخ یادگیری<sup>۱</sup> را نشان می‌دهد.

#### ۳-۲- مدل بازی مارکوف (MG)

مدل بازی مارکوف به‌عنوان MDP چندگانه محسوب می‌شود که در آن احتمال انتقال و پاداش به زوج حالت عمل مشترک همه عامل‌ها بستگی دارد. در یک حالت مشخص، مجموعه عمل‌های فردی عامل‌ها، یک بازی تکرارشونده را ایجاد می‌کند که با روش نظریه بازی‌ها حل می‌شود؛ بنابراین MG یک چارچوب قوی‌تری به شمار می‌آید که می‌تواند هم MDP و هم بازی تکرارشونده را عمومیت بخشد.

بازی مارکوف  $n$  عاملی، یک چندتایی  $\langle n, \{S_i\}_{i=1, \dots, n}, \{A_i\}_{i=1, \dots, n}, \{R_i\}_{i=1, \dots, n}, T \rangle$  است که در آن  $n$  تعداد عامل‌ها،  $S_i$  مجموعه فضای حالت‌های عامل  $i$  ام،  $A_i$  مجموعه فضای عمل‌های تمام عامل‌ها،  $A = \{A_i\}_{i=1, \dots, n}$  ام،  $i$  ام، مجموعه فضای عمل‌های تمام عامل‌ها و  $\pi = (\pi_1, \dots, \pi_n)$  سیاست مشترک تمام عامل‌ها را مشخص می‌کنند. مقادیر زوج حالت-عمل مشترک برای عامل  $i$  تحت سیاست مشترک  $\pi$  به‌صورت رابطه (۴) تعریف می‌شود [۱۱]:

$$Q_i^\pi(\vec{s}, \vec{a}) = E_\pi \left\{ \sum_{k=0}^{\infty} \gamma^k r_i^{t+k} \mid \vec{s}^t = \vec{s}, \vec{a}^t = \vec{a} \right\} \quad (4)$$

که،  $\vec{s} \in S$  حالت مشترک،  $\vec{a} \in A$  عمل مشترک و  $r_i^{t+k}$  پاداش دریافتی در گام زمانی  $(t+k)$  است. برخلاف هدف بهینه‌سازی MDP، هدف مدل MG پیدا کردن سیاست مشترک تعادلی  $\pi$  به جای سیاست مشترک بهینه برای همه عامل‌هاست. بنابراین مقدار  $Q$  در زوج حالت-عمل مشترک بر اساس تعادل به‌صورت رابطه (۵) به‌روز می‌شود که در آن  $\emptyset_i(\vec{s}')$  مقدار مورد انتظار تعادل در حالت مشترک بعدی برای عامل  $i$  است [۱۱]:

$$Q_i(\vec{s}, \vec{a}) \leftarrow (1 - \alpha) Q_i(\vec{s}, \vec{a}) + \alpha (r_i(\vec{s}, \vec{a}) + \gamma \emptyset_i(\vec{s}')) \quad (5)$$

#### ۳-۳- تعاملات پراکنده

بر اساس تعریفی که در مدل MG ذکر شد، ابتدا همه عامل‌ها باید سیاست خود را در تمام فضای حالت-عمل مشترک بیاموزند و سپس بتوانند با یکدیگر

<sup>۲</sup> Intelligent warehouse system

<sup>۱</sup> Learning rate

برخوردهای احتمالی بین عامل‌ها در سیستم‌های چندعاملی مبتنی بر تعاملات پراکنده جلوگیری کند. در همین راستا، در این الگوریتم، عامل‌ها در زمان انتخاب عمل مشترک، برای به‌دست‌آوردن مجموعه راه‌حل‌های تعادلی، از استراتژی پروفایل تعادلی غالب غیرقطعی<sup>۱</sup> (Non-strict EDSP) یا تعادل متا<sup>۲</sup> استفاده می‌کنند.

در الگوریتم RKT-MARL، ابتدا عامل‌ها در یک محیط ساده و بدون توجه به حضور عامل‌های دیگر، سیاست بهینه تک‌عاملی خود را به کمک مدل فرآیند تصمیم‌گیری مارکوف به‌دست می‌آورند. سپس زمانی که در حالت هماهنگی قرار بگیرند، برای به‌دست‌آوردن مجموعه تعادل با عامل‌های دیگر مذاکره می‌کنند و برای انتخاب بهترین عمل مشترک در مجموعه تعادل حاصل، از روش کمینه واریانس استفاده می‌کنند. پس از آن، عامل‌ها بعد از انتخاب بهترین عمل مشترک به حالت بعدی حرکت کرده و پاداش فوری خود را دریافت می‌کنند. در نهایت، مقدار  $Q$  با توجه به سود تعادلی و پاداش عامل‌ها به‌روز می‌شود.

#### ۴-۱- چارچوب کلی

زمانی که افراد در محیط‌هایی با برخوردهای احتمالی کار می‌کنند، ابتدا یاد می‌گیرند که چگونه وظایف فردی خود را به پایان برسانند؛ سپس در این محیط چگونه با دیگران هماهنگ شوند. با الهام از این فرضیه رایج، در الگوریتم RKT-MARL فرآیند یادگیری را در دو زیر فرآیند به‌صورت زیر تجزیه می‌کنیم: (۱) هر عامل بدون توجه به حضور عامل‌های دیگر، سیاست بهینه تک‌عاملی خود را در یک محیط ایستا یاد می‌گیرد، و (۲) در مواقعی که حضور عامل‌ها برهم‌تأثیرگذار باشد و برخوردی بین آن‌ها رخ دهد، پاداش فوری<sup>۳</sup> دریافتی توسط عامل‌ها تغییر می‌یابد؛ بنابراین در این شرایط لازم است که هر عامل بتواند بر اساس تغییرات پاداش فوری با عامل‌های دیگر هماهنگی برقرار کند. همچنین وجود سازوکار انتقال دانش در این الگوریتم، به یادگیری عامل‌ها سرعت بخشیده و آن‌ها را در برقراری هماهنگی با یکدیگر توانمندتر می‌کند. در ادامه، چهار مرحله اصلی الگوریتم RKT-MARL به‌طور کامل شرح داده می‌شود.

#### ۴-۱-۱- ایجاد یک چارچوب مبتنی بر مذاکره برای تعاملات پراکنده

همان‌طور که در قبل نیز ذکر شده، فرض اولیه بر این است که عامل‌ها سیاست بهینه تک‌عاملی و مدل پاداش

خود را قبل از یادگیری در محیط چندعاملی به دست می‌آورند؛ ولی زمانی که عامل‌ها در محیط چندعاملی قرار می‌گیرند، دو حالت ممکن است، رخ دهد: اگر در یک زوج حالت-عمل، پاداش فوری دریافتی با مقدار مورد انتظار از مدل پاداش یکسان باشند، عامل به‌صورت مستقل عمل می‌کند. در غیر این صورت، به‌منظور هماهنگی بهتر، عامل‌ها باید زوج حالت-عمل فردی خود را با افزودن زوج‌های حالت-عمل عامل‌های دیگر به نوع مشترک بسط دهند. مراحل ایجاد چارچوب مبتنی بر مذاکره برای تعاملات پراکنده به‌صورت زیر است:

**منتشر کردن حالت مشترک:** اگر عامل‌ها در یک حالت خاص با انتخاب یک عمل، تغییری در پاداش فوری تشخیص دهند، این زوج حالت-عمل به‌عنوان "خطرناک" و این عامل‌ها به‌عنوان "عامل‌های هماهنگ‌کننده"<sup>۴</sup> علامت‌گذاری می‌شوند؛ سپس این عامل‌ها اطلاعات حالت-عمل خود را به همه عامل‌های دیگر منتشر کرده و اطلاعات آن‌ها را دریافت می‌کنند. این زوج حالت-عمل به همراه تغییرات پاداش، یک زوج مشترک به نام "زوج هماهنگی"<sup>۵</sup> تشکیل می‌دهند و حالتی که در آن هماهنگی نیاز است به‌عنوان "حالت هماهنگی"<sup>۶</sup> علامت‌گذاری می‌شود که در فضای حالت همه عامل‌های هماهنگ‌کننده وجود دارد.

**انجام مذاکره برای به‌دست‌آوردن سیاست تعادلی:** عامل‌ها هنگام انتخاب یک زوج حالت-عمل خطرناک، با انتشار اطلاعات خود به یکدیگر، مشخص می‌کنند که آیا در زوج هماهنگی قرار دارند یا نه. اگر چنین باشد، برای جلوگیری از برخورد، عامل‌ها باید به جای سیاست حریصانه به دنبال پیدا کردن سیاست تعادلی باشند. بدین ترتیب، عامل‌ها می‌توانند از سازوکار مذاکره برای پیدا کردن سیاست تعادلی استفاده کنند. در چنین شرایطی اگر تعداد عامل‌ها بیش از دو باشد، ممکن است چندین بسط متفاوت برای حالت‌های هماهنگی مختلف در یک حالت وجود داشته باشد. همان‌طور که در شکل (۲) مشاهده می‌شود، عامل‌های ۱ تا ۴ به ترتیب دارای ۱۲، ۴، ۵ و ۳ حالت هستند. به‌عنوان مثالی از بسط حالت مشترک، عامل ۱ تمام حالت‌های عامل‌های دیگر را حتی اگر به بخش کوچکی از آن‌ها نیاز داشته باشد، به حالت مشترک خود اضافه می‌کند. حالت هشتم از عامل ۱ به دو حالت مشترک، یکی با عامل ۴ و دیگری با عامل ۳ و ۴، بسط داده می‌شود. به‌طور مشابه، حالت نهم به سه حالت مشترک بسط داده شده‌است.

<sup>4</sup> Coordinating agents

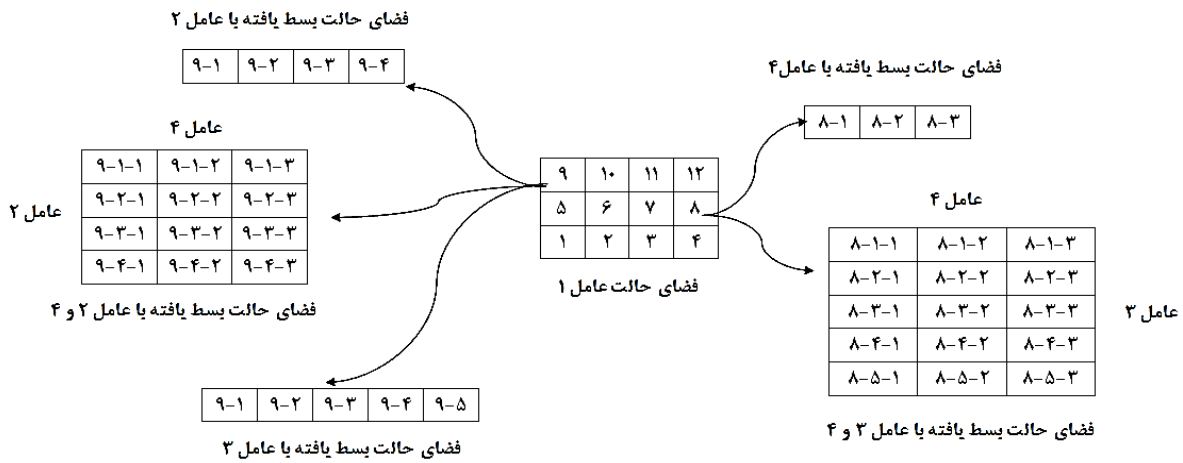
<sup>5</sup> Coordination pair

<sup>6</sup> Coordination state

<sup>1</sup> Non-strict equilibrium dominating strategy profile

<sup>2</sup> Meta equilibrium

<sup>3</sup> Immediate reward



(شکل-۲): مثالی از بسط فضای حالت عامل‌ها، زمانی که در محل برخورد باهم مذاکره می‌کنند. حالت هشتم از عامل ۱ به دو حالت مشترک، یکی با عامل ۴ و دیگری با عامل ۳ و ۴، بسط داده می‌شود. حالت نهم نیز به سه حالت مشترک با عامل ۲، عامل ۳ و عامل‌های ۲ و ۴ بسط داده می‌شود.

(Figure-2). Example of the expansion of an agent's state space, when they negotiate at the collision place. The 8th state of agent 1 is expanded to two joint states, the first one is with agent 4 and the second one is with agents 3 and 4. Similarly, the 9th state of agent 1 is expanded to three joint states, the first one is with agent 2, the second one is with agents 2 and 4, and the third one is with

منا به عنوان مفهوم تعادل محاسبه می‌شود که همیشه غیر تهی است.

(۲) شرایط لازم و کافی برای تعادل منا: در یک بازی  $n$  عاملی به شکل نرمال، عمل مشترک  $\vec{a}$  در بازی منا  $k_1 k_2 \dots k_r \Gamma$ ، به شرطی یک تعادل منا به شمار می‌آید که اگر و تنها اگر رابطه زیر برقرار باشد:

$$U_i(\vec{a}) \geq \min_{\vec{a}_{-i}} \max_{a_i} \min_{\vec{a}_{-i}} U_i(\vec{a}_{-i}, a_i, \vec{a}_{-i}) \quad (۷)$$

که  $p_i$  مجموعه عامل‌هایی است که در پیشوند  $k_1 k_2 \dots k_r$  قبل از علامت  $i$  و  $s_i$  عامل‌هایی که بعد از علامت  $i$  قرار دارند.

برای مثال در یک بازی منا سه عاملی  $213 \Gamma$  تعادل متای  $\vec{a}$  باید شرایط رابطه (۸) را داشته باشد:

$$\begin{aligned} U_1(\vec{a}) &\geq \min_{a_2} \max_{a_1} \min_{a_3} U_1(a_1, a_2, a_3) \\ U_2(\vec{a}) &\geq \max_{a_2} \min_{a_1} \min_{a_3} U_2(a_1, a_2, a_3) \\ U_3(\vec{a}) &\geq \min_{a_2} \min_{a_1} \max_{a_3} U_3(a_1, a_2, a_3) \end{aligned} \quad (۸)$$

#### ۴-۱-۳- روش کمینه واریانس برای انتخاب عمل مشترک

در مرحله قبل، فرآیندی برای همه عامل‌های هماهنگ‌کننده ارائه شده است تا برای به دست آوردن مجموعه عمل مشترک تعادلی با یکدیگر مذاکره کنند. با این حال، مجموعه حاصل به طور معمول شامل چندین استراتژی پروفایل بوده و عامل‌ها توانایی انتخاب بهترین استراتژی را ندارند. بنابراین در الگوریتم RKT-MARL، عامل‌های هماهنگ‌کننده برای انتخاب عمل مشترک با

#### ۴-۱-۲- مذاکره برای به دست آوردن مجموعه تعادل

هدف اصلی الگوریتم RKT-MARL، این است که عامل‌ها بتوانند راه حل‌های تعادلی را در حالت‌هایی که هماهنگی مورد نیاز است، پیدا کنند. بدین ترتیب، در این بخش دو استراتژی پروفایل غیر احتمالی، non-strict EDSP و تعادل منا به عنوان مفاهیم تعادلی به طور کامل شرح داده می‌شوند.

(۱) non-strict EDSP: در یک بازی  $n$  عاملی ( $n \geq 2$ ) به شکل نرمال، اگر  $\vec{e}_i \in A$  ( $i = 1, 2, \dots, m$ ) مجموعه تعادل نش‌های بازی باشد، عمل مشترک  $\vec{a} \in A$  به شرط زیر (رابطه (۶)) یک non-strict EDSP است:

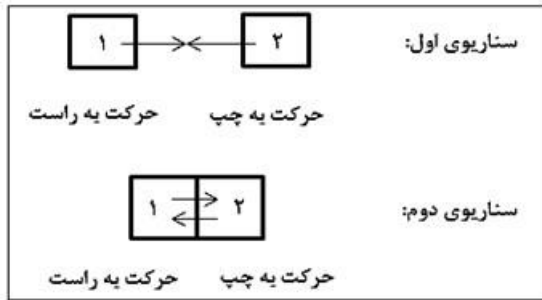
$$\forall j \leq n, U_j(\vec{a}) \geq \min_i U_j(\vec{e}_i) \quad (۶)$$

که  $U_j(\vec{e}_i)$  و  $U_j(\vec{a})$  به ترتیب مقادیر سودهایی هستند که عامل  $j$  در عمل مشترک  $\vec{a}$  و تعادل نش  $\vec{e}_i$  به دست می‌آورد. در این رابطه، هدف این است که سود عمل مشترک  $\vec{a}$  (سمت چپ نامساوی) از کمترین سود مجموعه تعادل نش‌ها، بیشتر باشد.

فرآیند مذاکره برای به دست آوردن مجموعه non-strict EDSP شامل دو گام زیر است: (۱) هر کدام از عامل‌ها با توجه به سود فردی خود، مجموعه استراتژی پروفایل‌هایی که non-strict EDSP هستند را پیدا می‌کنند، و (۲) سپس اشتراک مجموعه‌های به دست آمده در گام قبلی را پیدا کرده و یک مجموعه non-strict EDSP را انتخاب می‌کنند.

با توجه به اینکه در بعضی موارد به دلیل عدم وجود تعادل نش، مجموعه non-strict EDSP تهی است، تعادل

دانش آموزش داده شده قبلی می باشند. همچنین در این رابطه، عبارت  $Q_i^J((s_i, s_{-i}), (a_i, a_{-i}))$  با عبارت  $Q_i^J(\vec{s}, \vec{a})$  یکسان است. مشابه روش ورنکس، حالت مشترک  $\vec{s}$  در  $Q_i^{CT}(\vec{s}, \vec{a})$  با موقعیت نسبی  $(\Delta x, \Delta y)$  محاسبه می شود. همانطور که در شکل (۳) مشاهده می شود، زمانی که عامل ها به یک موقعیت مکانی مشترک یا مکان همدیگر حرکت کنند، امتیاز ۱۰- را به عنوان جریمه دریافت می کنند و درغیراینصورت پاداش صفر می گیرند.



(شکل-۳). دو سناریو از برخوردهای احتمالی بین دو عامل.  
سناریوی اول: حرکت دو عامل به یک موقعیت مکانی مشترک.  
سناریوی دوم: حرکت عامل ها به مکان همدیگر. [11]  
(Figure-3). Two scenarios of possible collisions between two agents. Scenario 1) moving two agents into the same grid location. Scenario 2) moving two agents into their previous locations [11].

|   | ۱ | ۲ | ۳ | ۴ |
|---|---|---|---|---|
| ۱ |   |   |   |   |
| ۲ |   | ۱ |   |   |
| ۳ |   |   | ۲ |   |
| ۴ |   |   |   |   |

(الف)

|  |       |   |   |       |
|--|-------|---|---|-------|
|  |       | ۱ |   |       |
|  |       |   | ۲ |       |
|  |       |   |   | هدف ۱ |
|  | هدف ۲ |   |   |       |

(ب)

(شکل-۴). انتقال مقادیر  $Q$  محلی در سیستم دو عاملی. (الف) محیط ساده  $4 \times 4$  با دو عامل، (ب): یک محیط پیچیده.

(Figure-4). Example of the local Q-value transfer in a two-agent system. a)  $4 \times 4$  simple environment with two agents. b) A complex environment.

بالاترین مجموع سودمندی و پایین ترین واریانس سودمندی، از روش کمینه واریانس [۱۱] استفاده می کنند. این روش ویژگی همکاری بین عامل ها و عادلانه بودن فرآیند یادگیری را تضمین می کند. عامل ها بعد از استفاده از روش کمینه واریانس برای به دست آوردن یک راه حل تعادلی و انتخاب عمل مشترک، حالت های خود را به روز کرده و پاداش فوری را دریافت می کنند که برای به روز کردن مقادیر  $Q$  محلی استفاده می شود.

#### ۴-۱-۴ انتقال دانش مقادیر $Q$ محلی

در ابتدای فرآیند یادگیری، الگوریتم RKT-MARL، دانشی درمورد سیاست بهینه تک عاملی به عامل ها برای سرعت بخشیدن به یادگیری انتقال داده می شود. علاوه بر این، با انتقال مقادیر  $Q$  محلی، عملکرد الگوریتم نیز بهبود می یابد. در الگوریتم یادگیری CQ [۱۷ و ۲۴]، مقدار اولیه  $Q$  در حالت های مشترک بسط داده شده، صفر است، ولی در روش یادگیری انتقالی که توسط ورنکس<sup>۱</sup> و همکارانش [۲۵] پیشنهاد شده،  $Q$  با مقادیری که از قبل روی یک منبع آموزش داده شده، مقداردهی می شود. زمانی که افراد در مسیر رسیدن به مقصد با یکدیگر ملاقات می کنند، به طور معمول دانش قبلی برای جلوگیری از برخورد دارند. بر اساس این فرضیه و داشتن اطلاعاتی درمورد به اتمام رساندن وظایف فردی، عامل ها یاد می گیرند که هنگام نزدیک شدن به همدیگر مذاکره کرده و استراتژی های مناسبی برای محیط های خاصی پیدا کنند. ورنکس تنها از دانش هماهنگی برای مقداردهی اولیه  $Q$  استفاده کرده و اطلاعات محیطی را نادیده گرفته است، ولی در الگوریتم RKT-MARL، برای مقداردهی اولیه  $Q$  هم دانش هماهنگی که از قبل آموزش داده شده و هم اطلاعات محیطی در نظر گرفته شده است. بدین صورت که  $Q$  در حالت هماهنگی بسط داده شده  $Q_i^J((s_i, s_{-i}), (a_i, a_{-i}))$  توسط ضرایب تنظیم (a و b) بین دانش آموزش داده شده قبلی  $(Q_i^{CT}(\vec{s}, \vec{a}))$  و اطلاعات محیطی جاری  $(Q_i(s_i, a_i))$ ، به صورت زیر مقداردهی اولیه می شود:

$$Q_i^J((s_i, s_{-i}), (a_i, a_{-i})) \leftarrow a * (Q_i(s_i, a_i)) + b * (Q_i^{CT}(\vec{s}, \vec{a})) \quad (9)$$

که  $s_{-i}$  مجموعه حالت های مشترک به جز حالت  $s_i$ ،  $a_{-i}$  مجموعه عمل های مشترک به جز عمل  $a_i$  و  $Q_i^{CT}(\vec{s}, \vec{a})$  مقدار  $Q$  انتقال یافته از یک کار منبع در حالت  $\vec{s}$  یا همان

<sup>1</sup> Vrancx

$$Q_1'(\vec{s}_{1..}) = Q_1(s_{1..})^T \times (1,1,1,1) + Q_1^{CT}(\vec{s}_{1..}) = \begin{pmatrix} -10 & -10 & -10 & -10 \\ -1 & -1 & -11 & -1 \\ -5 & -5 & -5 & -5 \\ -11 & -1 & -1 & -1 \end{pmatrix} \quad (12)$$

$$Q_2'(\vec{s}_{2..}) = (1,1,1,1)^T \times Q_2(s_{2..}) + Q_2^{CT}(\vec{s}_{2..}) = \begin{pmatrix} -5 & -1 & -1 & -5 \\ -5 & -1 & -11 & -5 \\ -5 & -1 & -1 & -5 \\ -15 & -1 & -1 & -5 \end{pmatrix}$$

در این رابطه،  $Q_1(s_{1..})$  و  $Q_2(s_{2..})$  به ترتیب اطلاعات محیطی جاری برای عامل‌های ۱ و ۲ هستند که از رابطه (۱۲) به دست آمده‌اند. همچنین  $Q_1^{CT}(\vec{s}_{1..})$  و  $Q_2^{CT}(\vec{s}_{2..})$  بیان‌گر دانش انتقالی از محیط ساده‌تر می‌باشند که براساس رابطه (۱۱) مقداردهی شده‌اند. بدین ترتیب، مجموعه تعادل نش در این حالت مشترک  $\{(چپ راست), (پایین راست) و (پایین پایین)\}$  است.

## ۵- تنظیمات اولیه محیط آزمایش

به منظور بررسی اهمیت الگوریتم ارائه شده، آزمایش‌های مختلفی برای مسائل یادگیری در سیستم‌های چندعاملی بر روی چندین نوع بازی جهانی انجام شده‌است. در ادامه جزئیات مجموعه داده‌های ارزیابی شده و الگوریتم‌های مورد مقایسه شرح داده می‌شود.

### ۵-۱- معرفی مجموعه داده‌ها

کارایی الگوریتم پیشنهادشده در این مقاله بر روی پنج نوع بازی جهانی شبکه شده<sup>۱</sup> [۲۴ و ۲۶] مختلف ارزیابی شده که در شکل (۵) نشان داده شده‌اند. همچنین این بازی‌ها در پژوهش‌های متعددی برای ارزیابی الگوریتم‌های مختلف [۲۰ و ۱۷، ۱۱] مورد استفاده قرار گرفته‌اند. برای چهار بازی ابتدایی (PENTAGON, SUNY, ISR) و MIT از جمله بازی‌های دو عاملی هستند که در هر کدام علامت ضربدر موقعیت اولیه هر عامل و هدف عامل دیگر را نشان می‌دهد. بازی پنجم (GW-nju) نیز شامل دو عامل است که موقعیت ابتدایی عامل  $i$  ام با  $i$  و هدف آن با  $G_i$  مشخص شده و از جمله بازی‌های با رقابت بالا به حساب می‌آید. در همه بازی‌ها، هر عامل در یک حالت می‌تواند در چهار جهت حرکت کرده و چهار عمل بالا، پایین، چپ و راست را انتخاب کند. همچنین در هر بازی عامل‌ها با طی تعدادی گام و جلوگیری از برخورد، سعی در رسیدن به هدف‌های خود دارند.

همچنین یادگیرنده‌های عمل مشترک در این محیط ساده، بدون توجه به اطلاعات محیطی، برقراری هماهنگی با عامل‌های دیگر را یاد می‌گیرند.

در این محیط ساده، مکان عامل اول  $s_1 = (2,2)$  و مکان عامل دوم  $s_2 = (3,3)$  است و دو عامل در خانه‌های (2,3) و (3,2) احتمال برخورد دارند؛ بنابراین  $\vec{s} = (x_1 - x_2, y_1 - y_2) = (1,1)$  یا  $(x_2 - x_1, y_2 - y_1) = (-1, -1)$  بیانگر موقعیت نسبی عامل‌ها نسبت به موقعیت فعلی خودشان است. عمل مشترک دو عامل نیز به صورت رابطه (۱۰) تعریف می‌شود:

$$\vec{a} = ((a_1, a_2)) = \begin{pmatrix} (\uparrow, \uparrow) & (\uparrow, \downarrow) & (\uparrow, \leftarrow) & (\uparrow, \rightarrow) \\ (\downarrow, \uparrow) & (\downarrow, \downarrow) & (\downarrow, \leftarrow) & (\downarrow, \rightarrow) \\ (\leftarrow, \uparrow) & (\leftarrow, \downarrow) & (\leftarrow, \leftarrow) & (\leftarrow, \rightarrow) \\ (\rightarrow, \uparrow) & (\rightarrow, \downarrow) & (\rightarrow, \leftarrow) & (\rightarrow, \rightarrow) \end{pmatrix} \quad (10)$$

در این ماتریس،  $\{\uparrow, \downarrow, \leftarrow, \rightarrow\}$  مجموعه عمل‌های هر عامل است که به ترتیب نشان‌دهنده حرکت به بالا، پایین، چپ و راست هستند. بعد از چندین گام یادگیری، عامل‌ها ماتریس مقادیر  $Q$  خود را در حالت مشترک  $\vec{s}$  طبق رابطه (۱۱) یاد می‌گیرند:

$$Q_1^{CT}(\vec{s}_{1..}) = Q_2^{CT}(\vec{s}_{2..}) = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & -10 & 0 \\ 0 & 0 & 0 & 0 \\ -10 & 0 & 0 & 0 \end{pmatrix} \quad (11)$$

در این ماتریس، جریمه ۱۰- برای هر عامل نشانگر این است که عامل‌های ۱ و ۲ با انتخاب عمل‌های  $(\rightarrow, \uparrow)$  یا  $(\downarrow, \leftarrow)$  منجر به برخورد می‌شوند و در بقیه موارد پاداش صفر می‌گیرند. حال فرض کنید عامل‌ها در محیط پیچیده‌تری قرار داشته باشند (شکل (۴) (ب)). زمانی که عامل‌ها به صورت مستقل در محیط عمل می‌کنند، بردارهای مقادیر  $Q$  بهینه تک‌عاملی در حالت‌های  $s_1$  و  $s_2$  به صورت رابطه (۱۲) می‌باشند:

$$Q_1(s_{1..}) = (-10, -1, -5, -1) \\ Q_2(s_{2..}) = (-5, -1, -1, -5) \quad (12)$$

در این بردارها، مقادیر ۵-، ۱۰- و ۱- به ترتیب نشان‌دهنده دور شدن از محیط، خروج از محیط شبکه شده و نزدیک شدن به هدف همراه با صرف انرژی هستند. این مقادیر به دست آمده همان اطلاعات محیطی است که عامل‌ها در محیط پیچیده به دست می‌آورند و سپس ماتریس مقادیر  $Q$  محلی در حالت هماهنگی برای هریک از عامل‌ها به صورت رابطه (۱۳) مقداردهی می‌شود:

<sup>1</sup> Grid world games

(جدول-۲). مقادیر بهینه پارامترهای  $a$  و  $b$  برای هر پنج بازی جهانی شبکه بندی شده.

(Table-2). The optimal values of the parameters  $a$  and  $b$  for each five grid world games.

| پارامترها | ISR | SUNY | MIT  | PENTAGON | GW_nju |
|-----------|-----|------|------|----------|--------|
| $a$       | 10  | 25   | 25   | 1        | 20     |
| $b$       | 1   | 25   | 0.01 | 10       | 25     |

## ۶- نتایج و بحث ها

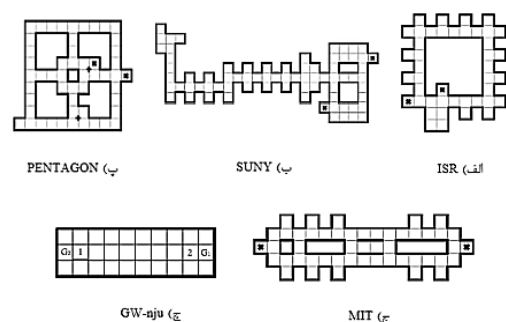
در این بخش، عملکرد الگوریتم RKT-MARL در مقایسه با الگوریتم های دیگر در حوزه یادگیری تقویتی چندعاملی مورد تحلیل و ارزیابی قرار می گیرد. نتایج به دست آمده برای یادگیری عامل ها در محیط های پیچیده تر روی بازی های معرفی شده نیز به صورت جداگانه ارزیابی می شود.

### ۶-۱- ارزیابی نتایج

ابتدا عملکرد الگوریتم های مورد مقایسه با توجه به میانگین گام های یادگیری آخرین اپیزود (معیار SFE) بر روی هر پنج بازی ارزیابی می شود. معیار SFE برای الگوریتم بر روی هر پنج بازی در جدول (۳) نشان داده شده است. سطر آخر جدول (۳) مقادیر بهینه های انتقال تابع مقدار در الگوریتم های ILVFT و NegoQVFT به هم گرایی سریع در تمام بازی ها به جزء SUNY می کند. سرعت همگرایی بالاتر نشان دهنده این است که عامل ها بهتر و سریع تر می توانند حالت های تعادل را یاد گرفته و مسئله یادگیری تقویتی را حل کرده و به هدف خود نزدیک شوند. در الگوریتم CQ-learning در بازی های PENTAGON و GW\_nju، سیاست بهینه ای که عامل ها از قبل آموخته اند، همیشه دارای برخورد است و CQ-learning نمی تواند عمل مشترک تعادلی را برای برقراری هماهنگی بین عامل ها پیدا کند. زمانی که برخوردهای زیادی رخ می دهد، عامل ها به صورت متوالی به حالت های قبلی خود برمی گردند و مجبورند گام های زیادی قبل از رسیدن به هدف طی کنند [27,28]. در همه بازی ها الگوریتم RKT-MARL نزدیک ترین مقدار را به مقدار بهینه نشان می دهد؛ زیرا در این الگوریتم عامل ها هم دانش قبلی در مورد برخوردها دارند و هم با انتقال این دانش و انجام مذاکره در حالت های هماهنگی، از برخوردهای احتمالی جلوگیری می کنند. در نتیجه عامل ها

## ۵-۲- ارزیابی الگوریتم ها

الگوریتم هایی که RKT-MARL با آن ها مقایسه شده است، عبارتند از: یادگیری CQ [۱۷]، NegoQ with VFT [۲۰]، ILVFT و NegoSI [۱۱]. همه الگوریتم ها در محیط و شرایط آزمایشی یکسانی مورد بررسی قرار گرفته اند. عملکرد الگوریتم پیشنهادی با بهترین نتایج گزارش شده از الگوریتم های مورد مقایسه، ارزیابی شده و نتایج حاصل از اعمال الگوریتم ها بر روی پنج بازی جهانی، نشان دهنده برتری قابل ملاحظه الگوریتم RKT-MARL نسبت به سایر الگوریتم ها است. در این مقاله از دو معیار پاداش های هر اپیزود<sup>۱</sup> (REE) و تعداد گام های آخرین اپیزود<sup>۲</sup> (SFE) برای مقایسه این الگوریتم ها استفاده می شود. قوانین پاداش دریافتی هر عامل نیز به صورت زیر است:



(شکل-۵): بازی های جهانی شبکه بندی شده.  
(Figure-5). Grid world games.

## ۵-۳- مفروضات پیاده سازی

روش پیشنهادی دارای پنج پارامتر شامل نرخ یادگیری  $\alpha$ ، پارامتر  $\epsilon$  یا همان پارامتر اکتشافی که در انتخاب یک عمل به صورت رندوم به عنوان یک احتمال ثابت در نظر گرفته می شود،  $a$  و  $b$  (در رابطه ۹) پارامترهای تنظیم بین دانش آموزش داده شده قبلی و اطلاعات محیطی که با مقادیر مختلف مورد ارزیابی قرار می گیرند و پارامتر تنزیل  $\gamma$  است. مقادیر پارامترهای  $\epsilon$ ،  $\alpha$  و  $\gamma$  در همه الگوریتم ها به ترتیب ۰/۰۱، ۰/۹ و ۰/۱ است. همچنین مقادیر بهینه دو پارامتر تنظیم  $a$  و  $b$  که در الگوریتم RKT-MARL برای تنظیم پذیری بین دانش آموزش داده شده قبلی و اطلاعات محیطی جاری استفاده می شوند، در سطر آخر جدول (۲) آورده شده است. تمام الگوریتم های مورد تست به تعداد دوهزار تکرار اجرا شده اند. پیاده سازی الگوریتم پیشنهادی RKT-MARL نرم افزار متلب انجام گرفته است.

Rewards of each episode

<sup>2</sup> Steps of final episode

برای رسیدن به هدف تعداد گام‌های کمتری را نسبت به زمانی که برخورد رخ دهد، طی می‌کنند.

(جدول-۳). میانگین گام‌های یادگیری آخرین اپیزود در هر

بازی تست شده برای هر الگوریتم مورد مقایسه.

(Table-3). Average learning steps of the final episode in each tested game for each algorithm.

| الگوریتم‌ها        | ISR   | SUN<br>Y | MIT   | PENTA<br>GON | GW_<br>nju |
|--------------------|-------|----------|-------|--------------|------------|
| CQ-learning        | 8.91  | 10.70    | 19.81 | 15.32        | 12.65      |
| ILVFT              | 13.11 | 10.38    | 18.67 | 14.18        | 10.94      |
| NegoQVFT           | 8.36  | 12.98    | 19.81 | 8.55         | 10.95      |
| NegoSI             | 7.48  | 10.92    | 21.29 | 10.30        | 12.11      |
| RKT-MARL           | 7.1   | 10       | 18.6  | 8.2          | 10.05      |
| The optimal policy | 6     | 10       | 18    | 8            | 10         |

سپس در شکل‌های (۶) تا (۱۰)، پاداش‌های هر اپیزود (معیار REE) در اپیزودهای مختلف بر روی الگوریتم‌های مورد مقایسه بررسی شده و نتایج حاصل حاکی از آن است که مجموع پاداش به دست آمده برای عامل‌ها در هر بازی و برحسب الگوریتم مورد نظر متفاوت است.

در بازی ISR شکل (۶)، در هر دو الگوریتم NegoSI و RKT-MARL به دلیل انتقال مناسب دانش اولیه به عامل‌ها، شروع مجموع پاداش هر دو عامل بالاتر از چهل امتیاز است. در الگوریتم CQ-learning پاداش منفی بیانگر این است که عامل‌ها توانایی یافتن عمل مشترک تعادلی برای برقراری هماهنگی با یکدیگر را نداشته و بین آن‌ها تعداد زیادی برخورد رخ داده است. در اثر هر یک از این برخوردها، عامل‌ها یک گام به عقب حرکت کرده، تعداد گام‌های رسیدن به هدف افزایش و مجموع پاداش‌ها عددی منفی می‌شود. الگوریتم NegoSI در این بازی در تمام فرآیند یادگیری نسبت به الگوریتم‌های دیگر بالاترین پاداش را دارد که نشان‌دهنده برخوردهای کمتر و هماهنگی بیشتر بین عامل‌هاست. همچنین الگوریتم NegoSI نسبت به الگوریتم‌های دیگر عادلانه‌تر است؛ زیرا مقادیر پاداش دو عامل در نمودار بهم نزدیک‌تر است.

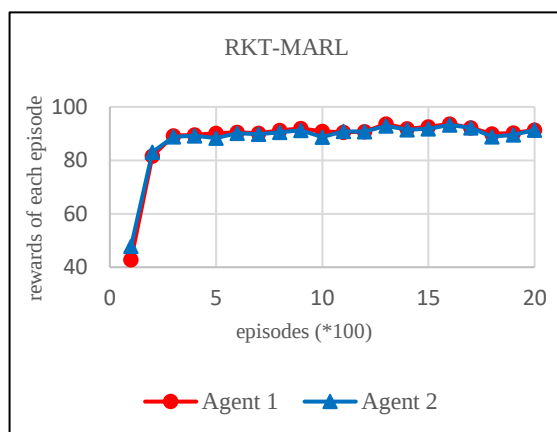
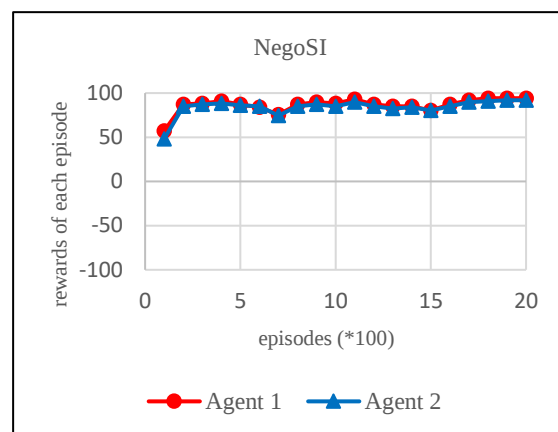
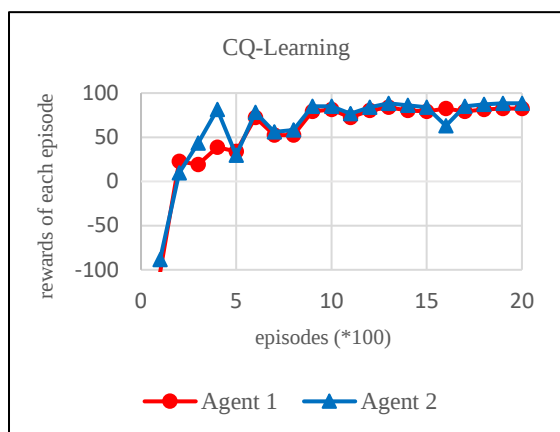
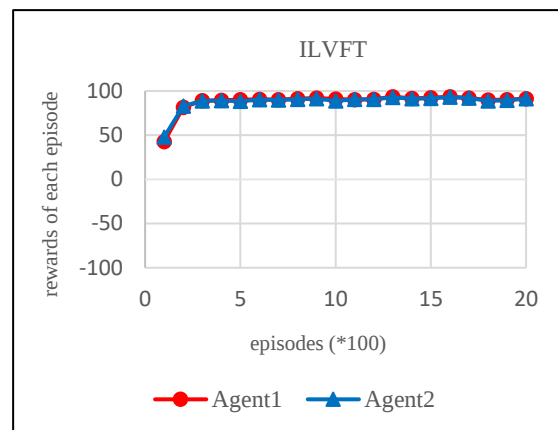
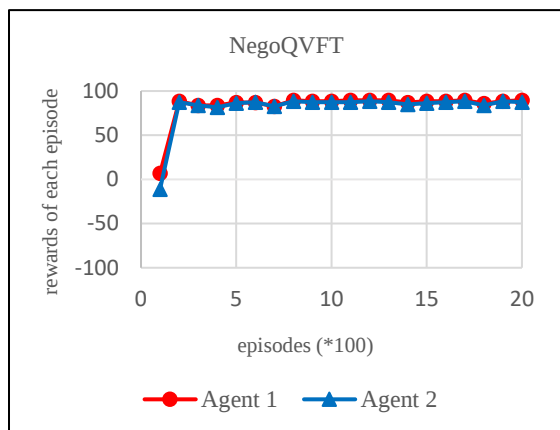
در بازی SUNY شکل (۷)، عامل‌ها مستقل از هم هستند و هرکدام سه انتخاب برای سیاست بهینه تک‌عاملی خود دارند. تفاوت بین مقادیر پاداش دو عامل در الگوریتم RKT-MARL کم است و این نشان‌دهنده عادلانه بودن فرآیند یادگیری این الگوریتم است. با اینکه در الگوریتم ILVFT یکی از عامل‌ها عملکرد بهتری دارد، ولی در کل نسبت به الگوریتم RKT-MARL عادلانه نیست. عامل‌ها در الگوریتم RKT-MARL دانش اولیه مناسب‌تری دریافت کرده و نسبت به سه الگوریتم

NegoQVFT، CQ-learning و NegoSI بالاترین مقدار پاداش را دارند. در الگوریتم RKT-MARL از اپیزود ۲۰۰ام به بعد منحنی یادگیری آن نسبت به سایر الگوریتم‌ها حالت پایداری به خود می‌گیرد.

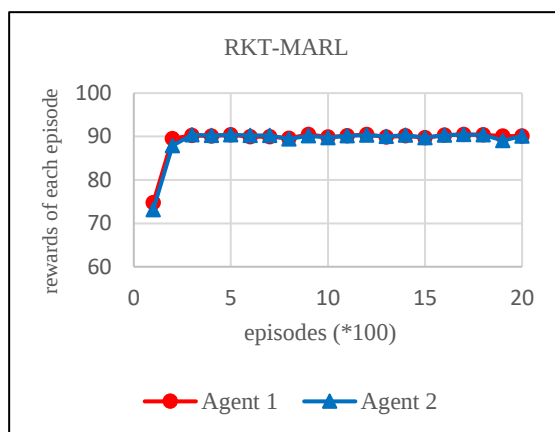
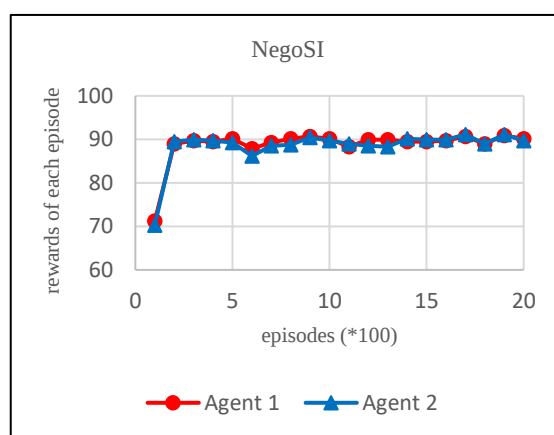
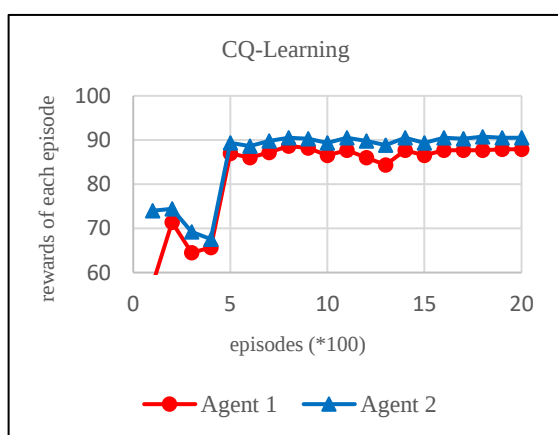
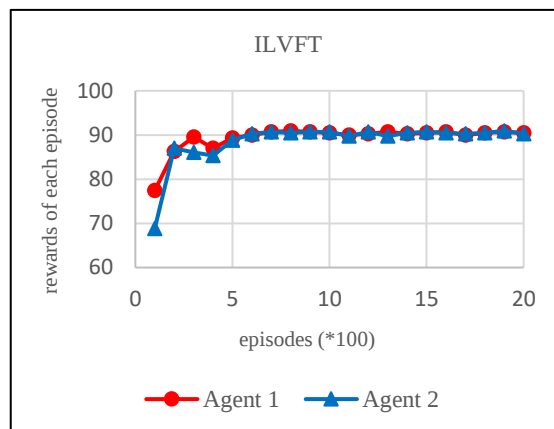
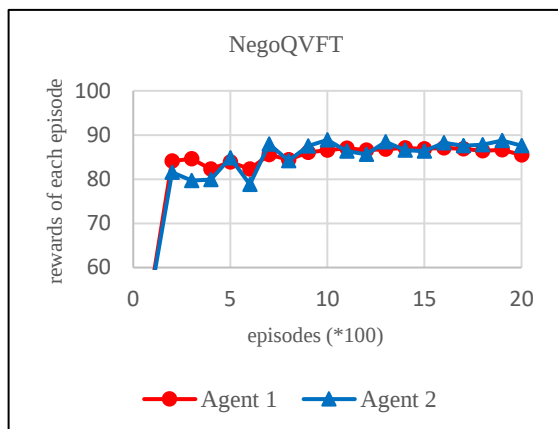
در بازی MIT شکل (۸) که عامل‌ها انتخاب بیشتری برای جلوگیری از برخورد دارند، همه الگوریتم‌های مورد تست، پاداش نهایی حدود ۸۰ امتیاز را کسب کرده‌اند. همان‌طور که در شکل (۸) مشاهده می‌شود، دانش اولیه الگوریتم RKT-MARL و CQ-learning نسبت به بقیه روش‌ها بیشتر است، ولی مشکل الگوریتم CQ-learning این است که منحنی یادگیری آن به‌نسبه ناپایدار هست.

در بازی PENTAGON شکل (۹)، در الگوریتم RKT-MARL ویژگی عادلانه بودن فرآیند یادگیری بین دو عامل به‌وضوح مشخص است و از اپیزود ۸۰۰ام به بعد منحنی یادگیری آن حالت پایداری به خود می‌گیرد. همچنین این الگوریتم به دلیل داشتن سازوکار انتقال دانش و دریافت دانش اولیه مناسب عملکرد بهتری دارد. الگوریتم CQ-learning نیز به دلیل برخوردهای زیاد بین عامل‌ها، منحنی یادگیری پایداری ندارد.

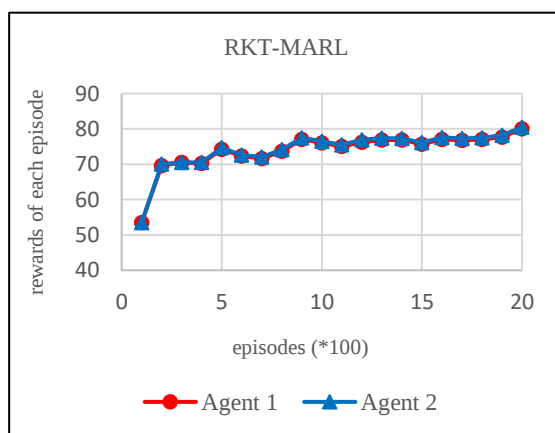
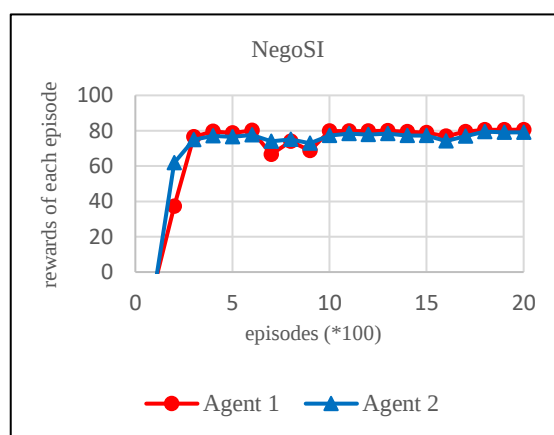
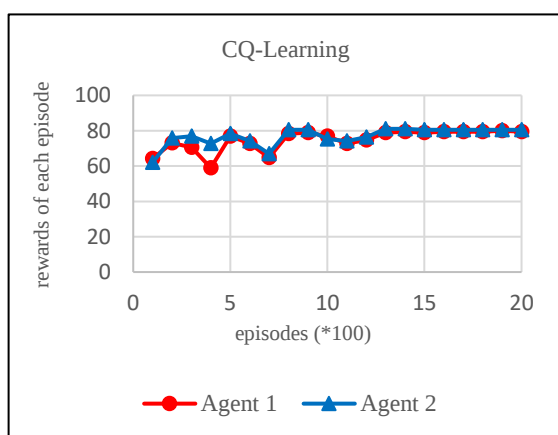
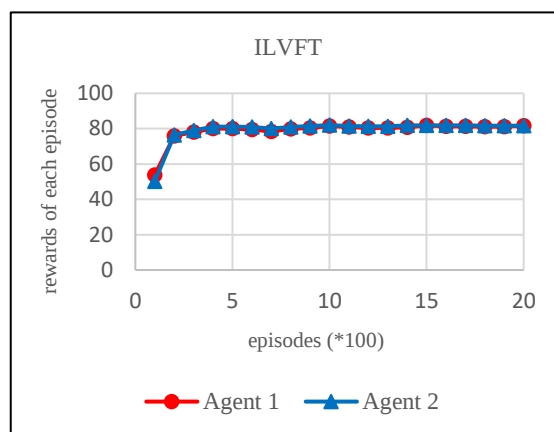
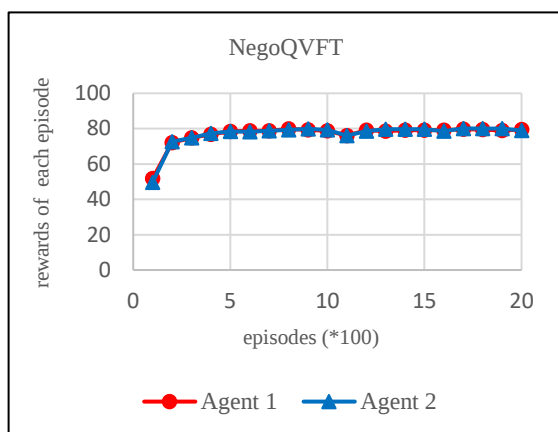
در بازی GW\_nju شکل (۱۰)، سیاست‌های بهینه تک‌عاملی عامل‌ها با هم برخورد دارند و نیاز به برقراری هماهنگی مناسب‌تری دارند. در این بازی مقادیر بهینه برای پاداش نهایی عامل‌های ۱ و ۲ به ترتیب ۹۳ و ۹۱ است. در الگوریتم‌های RKT-MARL، NegoSI و NegoQVFT منحنی یادگیری هر دو عامل در نهایت به مقادیری نزدیک مقدار بهینه همگرا می‌شود، با این تفاوت که برخلاف روش NegoQVFT، در روش‌های RKT-MARL و NegoSI، ابتدا عامل‌ها می‌توانند برای به دست آوردن سیاست امن و ثابت مذاکره کرده و سپس به صورت حریصانه حرکت کنند که این بیانگر توانایی برقراری هماهنگی بهتر در این روش‌هاست. در نهایت نکته قابل توجه این است که با وجود سازوکار مذاکره و انتقال دانش، عامل‌ها یاد می‌گیرند که هم سود بیشتری به خود و هم ضرر کمتری به عامل‌های دیگر برسانند و این شاهده بر وجود همکاری بین عامل‌هاست. همچنین در الگوریتم RKT-MARL، عامل‌ها توانایی جلوگیری از برخورد را برای کسب پاداش بیشتر دارند، درحالی‌که در الگوریتم‌های یادگیری تقویتی قبلی، تمایل کمتری در برقراری هماهنگی داشتند و این موجب از دست دادن پاداش توسط عامل‌ها می‌شد.



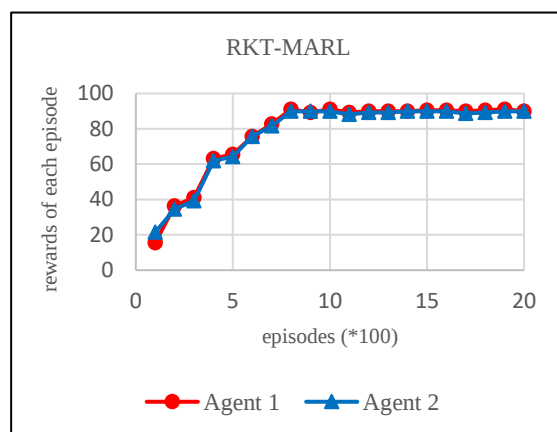
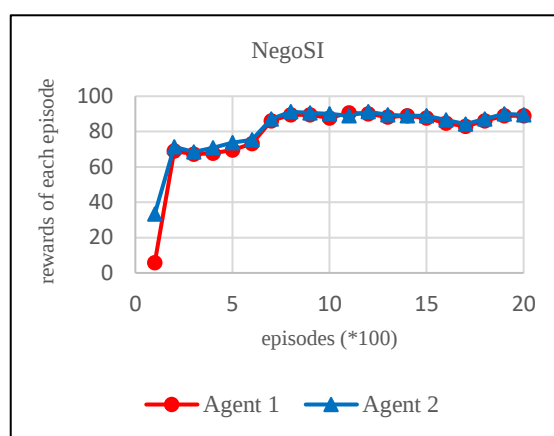
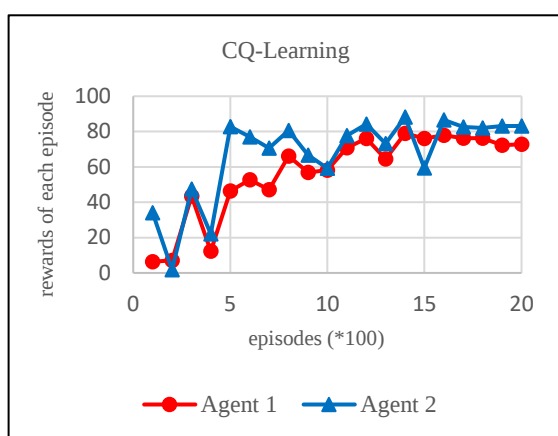
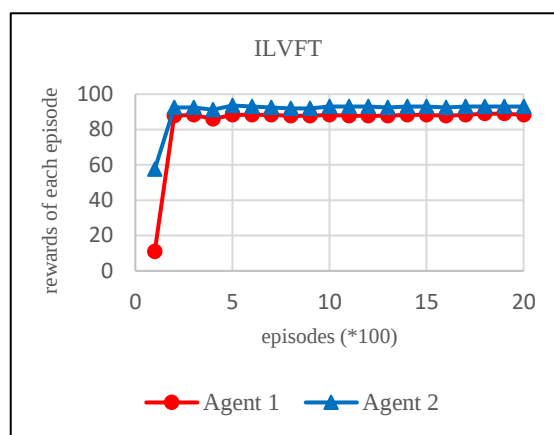
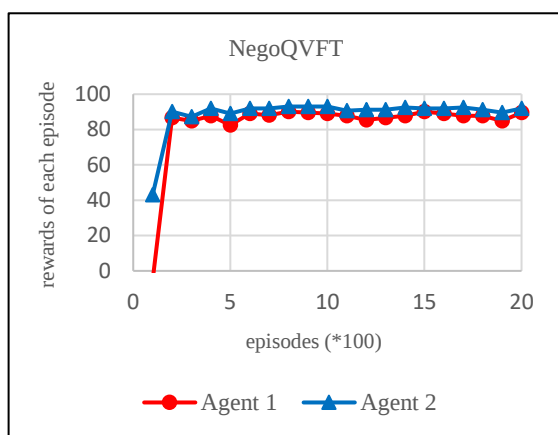
(شکل ۶-): معیار REE برای هر الگوریتم در بازی ISR.  
(Figure-6). REE criterion for each algorithm in ISR game.



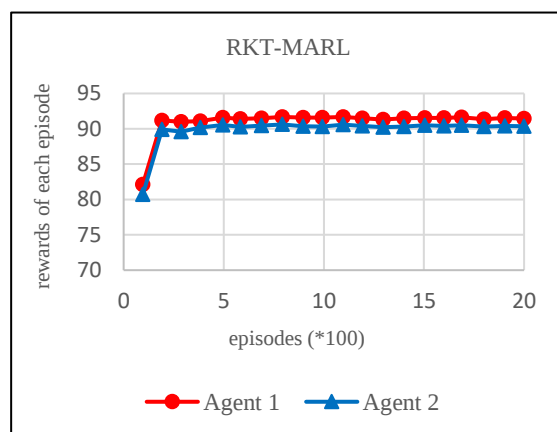
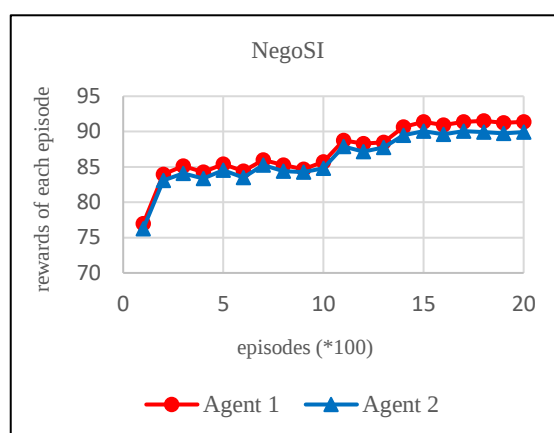
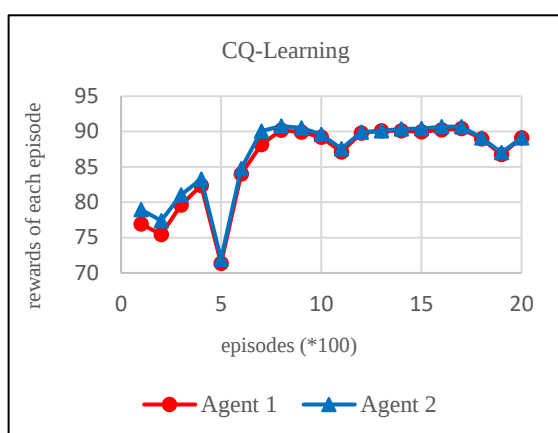
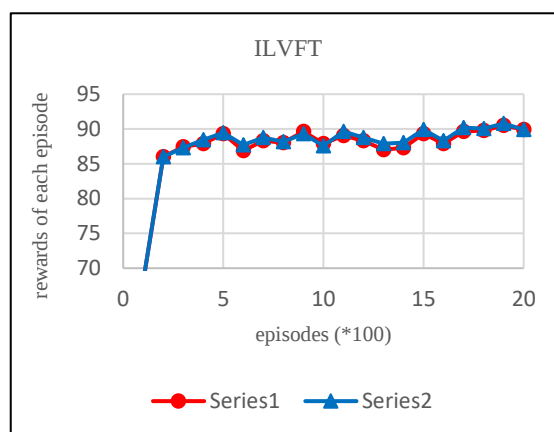
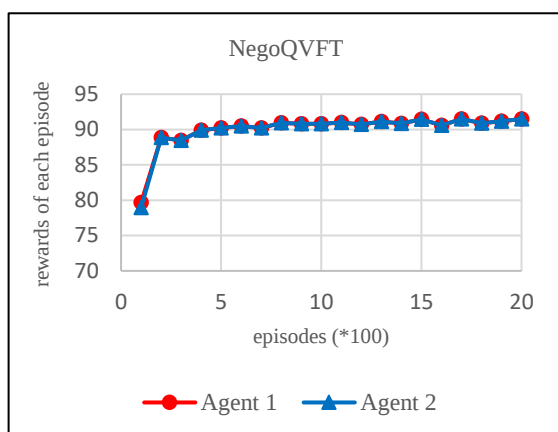
(شکل-۷). معیار REE برای هر الگوریتم در بازی SUNY.  
(Figure-7). REE criterion for each algorithm in SUNY game.



(شکل-۸). معیار REE برای هر الگوریتم در بازی MIT.  
(Figure-8). REE criterion for each algorithm in MIT game.



(شکل-۹). معیار REE برای هر الگوریتم در بازی PENTAGON.  
(Figure-9). REE criterion for each algorithm in PENTAGON game.



(شکل-۱۰). معیار REE برای هر الگوریتم در بازی GW\_nju.  
(Figure-10). REE criterion for each algorithm in GW\_nju game.

یادگیری تقویتی چندعاملی یکی از مهم‌ترین روش‌های یادگیری در حوزه سیستم‌های چندعاملی و یادگیری ماشین است که در آن عامل‌ها می‌آموزند که برای دریافت پاداش بیشتر از محیط، چه عمل‌هایی را باید انجام دهند. در این مقاله یک الگوریتم یادگیری تقویتی چندعاملی برای حل مسائل یادگیری و هماهنگی در سیستم‌های چندعاملی تحت عنوان RKT-MARL معرفی شده که به صورت ترکیبی از روش‌های مبتنی بر تعاملات پراکنده، راه‌حل‌های تعادلی و انتقال دانش است. در این الگوریتم مفاهیم تعادلی و مذاکره برای جلوگیری از برخورد بین عامل‌ها و سازوکار انتقال دانش برای بهبود سرعت یادگیری عامل‌ها به کار می‌رود. همچنین از ضرایب تنظیم برای برقراری تعادل بین دانش انتقالی و اطلاعات محیطی استفاده می‌شود. نتایج حاصل، اهمیت و برتری الگوریتم پیشنهادی را از دیدگاه همگرایی سریع و پیچیدگی محاسباتی کم نسبت به الگوریتم‌های دیگر در این حوزه نشان می‌دهد.

برای ادامه کار، در حال برنامه‌ریزی جهت توسعه روش RKT-MARL برای سیستم‌های سه عاملی هستیم. همچنین به دنبال استفاده از یادگیری تقویتی چند هدفی می‌باشیم تا به صورت گسترده‌ای اطلاعات محیطی و دانش هماهنگی را برای یادگیری محلی ترکیب کنیم.

## ۸- مراجع

- [1] C. Yu, M. Zhang, F. Ren and G. Tan, "Multiagent learning of coordination in loosely coupled multiagent systems," *IEEE Transactions on Cybernetics*, vol. 45, no.12, pp. 2853-2867, 2015.
- [2] J. Kober, J. A. Bagnell and J. Peters, "Reinforcement learning in robotics: A survey," *The International Journal of Robotics Research*, vol. 32, no. 11, pp. 1238-1274, 2013.
- [3] R. Babuška, L. Busoniu and B. D. Schutter, "Reinforcement learning for multi-agent systems," *IEEE International Conference on Emerging Technologies and Factory Automation*, 2006.
- [4] K. Arulkumaran, M. P. Deisenroth, M. Brundage and A. A. Bharath, "A brief survey of deep reinforcement learning," *arXiv preprint arXiv:1708.05866*, 2017.
- [5] Q. Zhang, P. Jiao, Q. Yin and L. Sun, "Coordinated Learning by Model Difference Identification in Multiagent Systems with Sparse Interactions," *Discrete Dynamics in Nature and Society*, 2016.
- [6] A. OroojlooyJadid and D. Hajinezhad, "A review of cooperative multi-agent deep reinforcement learning," *arXiv preprint arXiv: 1908.03963*, 2019.
- [7] A. Nowé, P. Vrancx and Y. M. D. Hauwere, "Game theory and multi-agent reinforcement learning," *In Reinforcement Learning*, Springer, Berlin, Heidelberg, pp. 441-470, 2012.

- [8] F. S. Melo and M. Velso, "Decentralized MDPs with sparse interactions," *Artificial Intelligence*, vol. 175, no.11, pp. 1757-1789, 2011.
- [9] D. S. Bernstein, R. Givan, N. Immerman and S. Zilberstein, "The complexity of decentralized control of Markov decision processes," *Mathematics of operations research*, vol. 27, no. 4, pp. 819-840, 2002.
- [10] A. M. Metelli, M. Mutti and M. Restelli, "Configurable Markov decision processes," *In International Conference on Machine Learning*, pp. 3491-3500, PMLR, 2018.
- [11] L. Zhou, P. Yang, C. Chen and Y. Gao, "Multiagent reinforcement learning with sparse interactions by negotiation and knowledge transfer," *IEEE transactions on cybernetics*, vol. 47, no. 5, pp. 1238-1250, 2017.
- [12] L. Canese, G.C. Cardarilli, L. Di Nunzio, R. Fazzolari, D. Giardino, M. Re and S. Spano, "Multi-Agent Reinforcement Learning: A Review of Challenges and Applications," *Applied Sciences*, p. 4948, 2021.
- [13] J. Tahmoresnezhad and S. Hashemi, "Exploiting kernel-based feature weighting and instance clustering to transfer knowledge across domains," *Turkish Journal of Electrical Engineering & Computer Sciences*, vol. 25, no. 1, pp. 292-307, 2017.
- [14] J. Tahmoresnezhad and S. Hashemi, "Visual domain adaptation via transfer feature learning," *Knowledge and Information Systems*, vol. 50, no. 2, pp. 585-605, 2017.
- [15] Y. Hu, Y. Gao and B. An, "Accelerating multiagent reinforcement learning by equilibrium transfer," *IEEE transactions on cybernetics*, vol. 45, no. 7, pp. 1289-1302, 2015.
- [16] C. J. C. H. Watkins, "Learning from delayed rewards," (Doctoral dissertation, King's College, Cambridge), 1989.
- [17] Y. M. D. Hauwere, P. Vrancx and A. Nowé, "Learning multi-agent state space representations," *In Proceedings of the 9th International Conference on Autonomous Agents and Multiagent Systems*, vol. 1, pp. 715-722, 2010.
- [18] Y. Hu, Y. Gao and B. An, "Multiagent reinforcement learning with unshared value functions," *IEEE transactions on cybernetics*, vol. 45, no. 4, pp. 647-662, 2015.
- [19] D. Abel, Y. Jinnai, S. Y. Guo, G. Konidaris and M. Littman, "Policy and Value Transfer in Lifelong Reinforcement Learning," *In International Conference on Machine Learning*, pp. 20-29, 2018.
- [20] Y. Hu, Y. Gao and B. An, "Learning in multi-agent systems with sparse interactions by knowledge transfer and game abstraction," *In Proceedings of the 2015 International Conference on Autonomous Agents and Multiagent Systems*, pp. 753-761, 2015.
- [21] P. Mannion, "Knowledge-based multi-objective multi-agent reinforcement learning (Doctoral dissertation), 2017.
- [22] T. Kujirai and T. Yokota, "Breaking Deadlocks in Multi-agent Reinforcement Learning with Sparse Interaction," *In Pacific Rim International Conference on Artificial Intelligence*, pp. 746-759, 2019.
- [23] Y. Liu, Y. Hu, Y. Gao, Y. Chen and C. Fan, "Value Function Transfer for Deep Multi-Agent Reinforcement Learning Based on N-Step Returns," *In IJCAI*, pp. 457-463, 2019.
- [24] F. S. Melo and M. Veloso, "Learning of coordination: Exploiting sparse interactions in multiagent systems," *In Proceedings of The 8th International Conference on Autonomous Agents and Multiagent Systems*, vol. 2, pp. 773-780, 2009.
- [25] P. Vrancx, Y. M. D. Hauwere and A. Nowé, "Transfer Learning for Multi-agent Coordination," *In ICAART (2)*, pp. 263-272, 2011.
- [26] Melo, Francisco S., and Manuela Veloso. "Decentralized MDPs with sparse interactions," *Artificial Intelligence* 175.11 (2011).
- [27] Sherafati F, Tahmoresnezhad J. Image Classification via Sparse Representation and Subspace Alignment. *JSDP* 2020; 17 (2) :58-47.
- [28] Zandifar M, Tahmoresnezhad J. Sample-oriented Domain Adaptation for Image Classification. *JSDP* 2019; 16 (3) :148-129.



**جعفر طهمورث نژاد** مدرک دکترای

مهندسی کامپیوتر-هوش مصنوعی  
خود را در سال ۱۳۹۴ از دانشگاه  
شیراز دریافت کردند. ایشان در  
راستای فعالیت‌های علمی خود، در

حال حاضر به‌عنوان دانشیار دانشگاه صنعتی ارومیه در  
حال فعالیت هستند. علایق پژوهشی ایشان شامل  
حوزه‌های یادگیری ماشین، یادگیری انتقالی، داده‌کاوی و  
سامانه‌های امنیتی است.

نشانی رایانامه ایشان عبارت است از:

**j.tahmores@it.uut.ac.ir**



**نیلوفر علوی** مدرک کارشناسی خود

را در رشته مهندسی فناوری اطلاعات  
در سال ۱۳۹۵ از دانشگاه شهید مدنی  
آذربایجان و مدرک کارشناسی ارشد  
رشته مهندسی فناوری اطلاعات را از

دانشگاه صنعتی ارومیه در سال ۱۳۹۹ دریافت کرده‌است.  
علایق پژوهشی ایشان شامل حوزه‌های یادگیری انتقالی و  
یادگیری تقویتی است.

نشانی رایانامه ایشان عبارت است از:

**niloolavi504@it.uut.ac.ir**

