



Solution for Reducing Complexity and Increasing Efficiency in AI-Based Collaborative Robots

Seyed Reza Mazlooman¹, Saeed Setayeshi^{2*}, Mohsen Jahanshahi³

^{1,3} Department of Computer Engineering, Central Tehran Branch, Islamic Azad University, Tehran, Iran.

² Department of Medical Radiation Engineering, Amirkabir University of Technology, Tehran, Iran.

ARTICLE INFO

Article Type:

Original Research

Received: 11.03.2025

Revised: 01.09.2025

Accepted: 15.10.2025

Keyword:

Job Shop Scheduling Problem
Cobot, Markov Decision Process
Multi-Agent Reinforcement
Learning
Intrinsic Motivation-Based
Approach

*Corresponding Author:

Saeed Setayeshi

Email: setayesh@aut.ac.ir

ABSTRACT

Significant advancements have been made in the development of processing tasks across various industries to enhance productivity, reduce costs, and improve flexibility. However, in the manufacturing industry, resource constraints remain a persistent challenge, often manifesting as precedence relationships among factors. Therefore, task scheduling requires greater precision and adaptability. To address this issue, this study presents a scheduling algorithm that utilizes a multi-agent system to solve the Flexible Job Shop Scheduling Problem (FJSP) in a dynamic input environment. The resource precedence environment is modeled as a decentralized partially observable Markov decision process (Dec-POMDP), where tasks act as independent agents and select an available collaborative robot (cobot) based on their current observations. To overcome the challenge of managing the high-dimensional operational space in concurrent multi-task scheduling, an architecture is designed within a multi-agent reinforcement learning (MARL) system with intrinsic motivation. This model simultaneously considers both coordinated behavior among agents and individual agent performance. The novelty of this research lies in the application of an intrinsic motivation-based multi-agent system for problem-solving and optimizing overall job execution time. Compared to well-known methods such as rule-based approaches, metaheuristics, and cooperative strategies, the proposed model demonstrates superior performance. The results from the case study highlight the flexibility and training stability of this model, suggesting that it can serve as an effective solution for addressing task scheduling challenges in various industries.



EXTENDED ABSTRACT

Introduction

In smart manufacturing, robots such as cell robots and collaborative robots (cobots) play a crucial role in logistics, supply chains, and flexible production environments. However, due to limited resources like robots, machines, and tools, tasks cannot always be executed simultaneously. Efficient scheduling is therefore essential for optimizing production processes, resource allocation, and job sequencing, particularly in modern manufacturing systems like cloud manufacturing, mass personalization, and high-mix-low-volume (HMLV) production. One major challenge in smart manufacturing is the Flexible Job Shop Scheduling Problem (FJSP), which is classified as NP-hard, making it computationally complex to solve optimally. Traditional approaches, including meta-heuristics and dispatching rules, perform well under static conditions but struggle with dynamic job arrivals and diverse requirements. Real-time decision-making is necessary to address these challenges. To tackle this issue, we propose a multi-agent scheduling framework for HMLV environments that adapts to dynamically arriving jobs. Various scheduling methods, including meta-heuristics, single-agent models, multi-agent systems, and deep reinforcement learning (DRL), have been explored. However, many rely on assumptions like fixed transportation times, system centralization, or simplified job conditions, which limit their effectiveness. Our approach seeks to overcome these limitations by developing a multi-agent reinforcement learning (MARL) system that minimizes makespan while optimizing cobot and resource utilization.

The proposed system models a flexible job shop environment with multiple workstations, each containing multiple machines capable of performing specific operations. Jobs arrive according to a stochastic distribution, and multiple robots transport them between workstations. The scheduling model incorporates preemptive job execution, allowing higher-priority jobs to interrupt and reschedule lower-priority ones. To manage this complexity, we integrate a preemptive decision-making policy within our MARL framework, enabling job agents to learn optimal scheduling strategies while adapting to unexpected changes. This problem is inherently decentralized, as jobs act as independent agents making scheduling decisions based on system conditions. Our approach employs MARL to optimize both individual job decision-making and overall system coordination. Specifically, we introduce a Multi-agent Scheduling System (MASS) with three key contributions:

- i) Intrinsic motivation mechanisms to handle job priority challenges in DRL.
- ii) A novel scheduling model via Centralized Training and Decentralized Execution (CTDE).
- iii) A dynamic FJSP incorporating stochastic job arrivals and adaptive scheduling policies.

The remainder of this paper covers related work, problem formulation, system architecture, scheduling algorithms, experimental results, and conclusions.

Methodology

We represent the resource preemption environment as a decentralized partially observable Markov Decision Process (MDP), where individual jobs act as agents that select an available cobot based on their current observations. To tackle the challenge of managing a high-dimensional action space in multi-task simultaneous scheduling, we develop an architecture within a multi-agent reinforcement learning system that is driven by intrinsic motivation. This model considers both agent coordination and individual agent behavior. To evaluate the effectiveness of our model, we simulate various FJSP scenarios. We use makespan as the primary performance indicator and compare the performance of our MASS approach with other scheduling algorithms such as PDR, HHA, and MADRL. We conduct comparative experiments on total makespan and training stability, and perform scalability experiments by varying the number of job agents. Detailed scheduling results are presented through Gantt charts. The development platform is based on Python 2.7.3 and PyTorch 1.4.1, running on a single NVIDIA GeForce RTX 2080Ti with 16GB of memory and an Intel Core i7-9700KF CPU. The environment is developed using the OpenAI gym package, which provides interfaces and specifications for RL environments. The unit of time consumption mentioned in this section is minutes.

Results and discussion

We compare the performance of our MASS against several established algorithms: QMIX, a simulated annealing-based hyper-heuristic (SA-HH) method, and two RBSP methods including SPT and earliest due date (EDD), which are widely used PDRs in real-world production settings. In reinforcement learning-based scheduling algorithms, training stability is crucial for ensuring reproducibility and reliability. QMIX is known for its effectiveness in multi-agent systems with a small number of agents. Conversely, MASS is tailored for decentralized MARL environments, where each agent has partial observability and learns from its own experiences while coordinating with other agents. Thus, MASS can be more efficient in settings with a larger number of agents and more complex tasks. To evaluate the robustness of the MASS model, we compare the training processes of MASS and QMIX using three experimental instances (D01-D10). Figure 2 illustrates the total makespan achieved by MASS within the first 8000 training steps, highlighting improved agent coordination and effective internal motivation, which helps avoid delays. The results suggest that MASS efficiently adapts to various situations by learning appropriate behaviors. The intrinsic reward mechanism enhances policy learning: agents receive higher intrinsic rewards for valid policies and lower rewards for faulty ones. The plots Figures 1 depict total rewards. Agents strive to make optimal decisions to maximize rewards, progressing towards both individual and collaborative goals. The results indicate that agents have learned to prioritize tasks based on arrival rates. Further analysis reveals how the attention mechanism aids intrinsic reward learning.

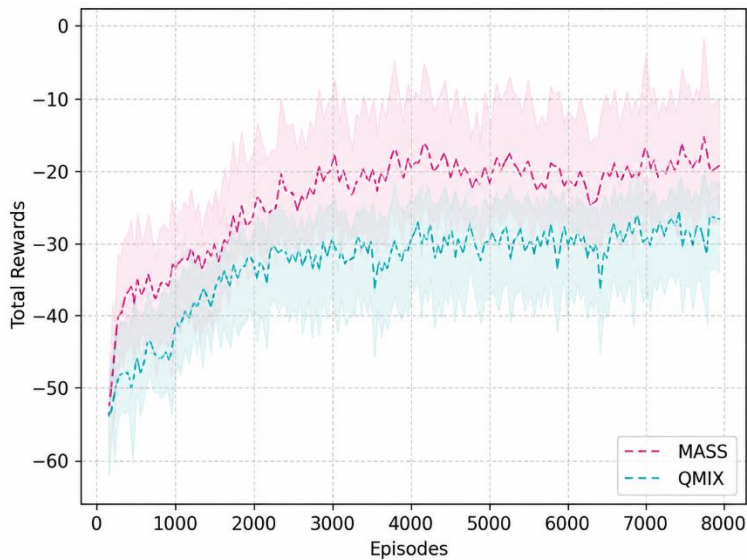


Figure 2. Curve training instance 11, training compare between MASS and QMIX.

Conclusion

This paper presents a Multi-Agent Reinforcement Learning solution based on intrinsic motivation for dynamic Flexible Job Shop Scheduling Problems. We propose a scheduling system that incorporates centralized training, facilitating decentralized scheduling execution. Our approach addresses job scheduling in a multi-agent environment through a decentralized MDP that balances agent coordination and individual behavior. Our intrinsic motivation-oriented multi-agent system outperforms well-known methods in optimizing makespan and can manage dynamic job arrivals under a MARL framework. The study demonstrates the flexibility and training stability of our model, suggesting its suitability for various industrial job scheduling challenges.



راهکاری برای کاهش پیچیدگی و افزایش بهره‌وری از روبات‌های همکار مبتنی بر هوش مصنوعی

سید رضا مظلومان^۱، سعید ستایشی^{۲*}، محسن جهانشاهی^۳

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

۲- گروه مهندسی انرژی و فیزیک، دانشگاه صنعتی امیرکبیر، تهران، ایران.

چکیده

اطلاعات مقاله

نوع مقاله: مقاله پژوهشی

دریافت مقاله: ۱۴۰۱/۱۲/۲۱

بازنگری مقاله: ۱۴۰۴/۰۶/۱۰

پذیرش مقاله: ۱۴۰۴/۰۷/۲۳

کلید واژگان:

مسئله زمان‌بندی کارگاهی
کوبوت
فرآیند تصمیم‌گیری مارکوف
یادگیری تقویتی چندعاملی
محوریت بر انگیزه درونی

*نویسنده مسئول: سعید ستایشی

پست الکترونیکی:

setayesh@aut.ac.ir

پیشرفت‌های قابل توجهی در توسعه وظایف پردازشی در صنایع مختلف به منظور افزایش بهره‌وری، کاهش هزینه‌ها و بهبود انعطاف‌پذیری صورت گرفته است. با این حال، در صنعت تولید، محدودیت منابع به طور معمول یک چالش پایدار محسوب می‌شود که به شکل روابط تقدم و تأخر میان عوامل نمود پیدا می‌کند. بنابراین، زمان‌بندی وظایف نیازمند دقت و انعطاف‌پذیری بیشتری است. برای رفع این مسئله، این مطالعه یک الگوریتم زمان‌بندی ارائه می‌دهد که از یک سیستم چندعاملی برای حل مسئله زمان‌بندی کارگاهی انعطاف‌پذیر در محیطی با ورودی پویا استفاده می‌کند. محیط تقدم منابع به صورت یک فرآیند تصمیم‌گیری مارکوف غیرمتمرکز و قابل مشاهده جزئی مدل‌سازی شده است، که در آن وظایف به عنوان عوامل جداگانه عمل کرده و بر اساس مشاهدات فعلی خود، یک همکار روباتیک در دسترس را انتخاب می‌کنند. برای غلبه بر چالش مدیریت فضای عملیاتی با ابعاد بالا در زمان‌بندی همزمان چندوظیفه‌ای، یک معماری در یک سیستم یادگیری تقویتی چندعاملی با انگیزه درونی طراحی شده است. این مدل هم‌زمان رفتار هماهنگ بین عوامل و عملکرد فردی هر عامل را در نظر می‌گیرد. نوآوری این پژوهش در استفاده از سیستم چندعاملی مبتنی بر انگیزه درونی برای حل مسئله و بهینه‌سازی زمان کل اجرای کارها است، که نسبت به روش‌های شناخته‌شده‌ای مانند روش‌های مبتنی بر قوانین، متاهیوریستیک و روش‌های تعاونی عملکرد بهتری دارد. نتایج به دست آمده از مطالعه موردی نشان‌دهنده انعطاف‌پذیری و پایداری آموزشی این مدل بوده و پیشنهاد می‌کند که این روش می‌تواند راه‌حل مناسبی برای مقابله با چالش‌های زمان‌بندی وظایف در صنایع مختلف باشد.



۱. مقدمه

در سال‌های اخیر، با پیشرفت چشمگیر فناوری‌های نوین، تولید هوشمند به یکی از ارکان اساسی صنایع پیشرفته تبدیل شده است. یکی از مهم‌ترین پیشرفت‌ها در این حوزه، توسعه و کاربرد گسترده ربات‌های سلولی و همکار^۱ است [۱] که نقش مهمی در بهبود بهره‌وری، کاهش هزینه‌ها و افزایش انعطاف‌پذیری فرآیندهای تولید ایفا می‌کنند [۲]. این روبات‌ها علاوه بر انجام وظایف پیچیده مانند جابجایی مواد و پردازش عملیات، توانایی تطبیق با محیط‌های پویا را دارند [۳]. با وجود این پیشرفت‌ها، مدیریت زمان‌بندی وظایف همچنان چالشی اساسی در صنایع تولیدی محسوب می‌شود. محدودیت منابع نظیر ماشین‌آلات و ابزارها باعث می‌شود انجام هم‌زمان تمامی وظایف ممکن نباشد [۴]. از این رو، زمان‌بندی مؤثر به‌عنوان عاملی کلیدی برای بهینه‌سازی فرآیندهای تولید و تخصیص منابع مطرح شده است. مسئله زمان‌بندی کارگاهی انعطاف‌پذیر (FJSP)^۲ یکی از چالش‌های رایج در این زمینه است که به دلیل پیچیدگی‌های ترکیبی خود به‌عنوان یک مسئله NP-سخت شناخته می‌شود [۵]. روش‌های سنتی نظیر الگوریتم‌های متاهوریستیک و قوانین اعزام، در شرایط ایستای نتایج نسبتاً مطلوبی ارائه می‌دهند، اما در مواجهه با محیط‌های پویا و وظایف با ورود تصادفی عملکرد ضعیفی دارند [۶]. در این مقاله، مدلی مبتنی بر یادگیری تقویتی چندعاملی^۳ ارائه می‌شود که با استفاده از انگیزه درونی، عملکردی مؤثر در حل مسئله FJSP دارد. این مدل به‌گونه‌ای طراحی شده است که عامل‌ها بتوانند به صورت غیرمتمرکز عمل کرده، از تجربیات خود یاد بگیرند و وظایف را بهینه‌سازی کنند. هدف اصلی مدل پیشنهادی، کاهش زمان کل اجرای وظایف^۴ و افزایش بهره‌وری سیستم است.

۲. پیشینه پژوهش

مسئله زمان‌بندی کارگاهی انعطاف‌پذیر یکی از موضوعات اساسی و چالش‌برانگیز در حوزه تولید هوشمند است که توجه بسیاری از محققان را به خود جلب کرده است [۷]. در طول سال‌های اخیر، روش‌های مختلفی برای حل این مسئله ارائه شده است که هرکدام مزایا و معایب خاص خود را دارند. روش‌های بهینه‌سازی عملیاتی نظیر برنامه‌ریزی خطی و پویا در مسائل کوچک عملکرد خوبی نشان داده‌اند، اما در مسائل بزرگ‌تر و شرایط پویا، کارایی خود را از دست می‌دهند [۸]. از سوی دیگر، روش‌های مبتنی بر قوانین^۵، مانند قوانین اعزام وظایف، به دلیل سادگی و سرعت، در بسیاری از محیط‌های صنعتی به کار گرفته می‌شوند. این روش‌ها شامل الگوریتم‌هایی نظیر کمترین زمان پردازش و اولویت‌بندی بر اساس مهلت انجام کار هستند. با این حال، این روش‌ها در مدیریت محیط‌های پویا و شرایطی که اختلالات مکرر رخ می‌دهند، محدودیت‌های جدی دارند [۹]. الگوریتم‌های ابر-اکتشافی^۶ مانند الگوریتم‌های ژنتیک [۴]، سیستم‌های ایمنی [۱۰] و بهینه‌سازی کلونی مورچگان [۱۱]، توانایی ارائه راه‌حل‌های نزدیک به بهینه را دارند. این روش‌ها با جست‌وجوی فضای راه‌حل‌ها، نتایج قابل قبولی در مسائل ایستای ارائه می‌دهند. با این وجود، به دلیل عدم انعطاف‌پذیری و هوشمندی کافی، در شرایط پویا عملکرد مطلوبی ندارند. برای رفع این محدودیت‌ها، مدل‌های ترکیبی ابر-اکتشافی و مبتنی بر قوانین ارائه شده‌اند که سعی در بهبود عملکرد این روش‌ها دارند. اگرچه این مدل‌ها تا حدی نتایج بهتری ارائه داده‌اند، اما همچنان در مواجهه با پیچیدگی‌های مسائل زمان‌بندی در محیط‌های متغیر و پویا، کارایی

¹ Cobots² Flexible Job-Shop Scheduling Problem (FJSP)³ Multi-agent Reinforcement Learning (MARL)⁴ Makespan⁵ Rule-based Scheduling Methods (RBSM)⁶ Hyper-Heuristic Algorithms (HHA)

کافی ندارند [۸]. در سال‌های اخیر، یادگیری تقویتی^۱ به‌عنوان یکی از روش‌های نوین برای حل مسائل پیچیده زمان‌بندی مطرح شده است. این روش با استفاده از تجربیات عامل‌ها، سیاست‌هایی نزدیک به بهینه را یاد می‌گیرد و در مسائل متنوعی مانند کنترل روبات‌ها، مدیریت و سایل نقلیه خودکار و زمان‌بندی تولید به‌کار گرفته شده است. ترکیب یادگیری تقویتی با یادگیری عمیق [۱۲]، پیشرفت‌های قابل توجهی را در حل مسائل بزرگ و پیچیده فراهم کرده است. برخی مطالعات نیز روش‌های یادگیری تقویتی را با قوانین اعزام وظایف و ابر-اکتشافی ترکیب کرده‌اند تا راه‌حل‌های بهتری ارائه دهند [۱۳]. سیستم‌های یادگیری تقویتی چندعاملی (MARL) نیز به دلیل توانایی آن‌ها در توزیع وظایف و خودمختاری عامل‌ها، به یکی از ابزارهای مهم در مدیریت محیط‌های صنعتی پویا تبدیل شده‌اند. این سیستم‌ها به عامل‌ها امکان می‌دهند که به‌صورت غیرمتمرکز عمل کنند و با همکاری یکدیگر به بهینه‌سازی زمان‌بندی بپردازند. مطالعات اخیر نشان داده‌اند که استفاده از مکانیزم‌های اختصاص اعتبار چندعاملی^۲ (MCA) می‌تواند تعامل میان عامل‌ها و دینامیک محیط را بهبود بخشد. علاوه بر این، پژوهش‌های متعددی از چارچوب‌های مبتنی بر یادگیری تقویتی عمیق برای حل مسائل زمان‌بندی استفاده کرده‌اند [۱۴]. در این میان، ترکیب یادگیری تقویتی چندعاملی با محرک‌های انگیزه درونی، گامی نوین در بهبود عملکرد سیستم‌های زمان‌بندی محسوب می‌شود [۱۵]. این رویکرد نه تنها بهره‌وری سیستم را افزایش می‌دهد، بلکه به عامل‌ها امکان می‌دهد تا با استفاده از پاداش‌های درونی، وظایف خود را به‌صورت بهینه‌تری اولویت‌بندی کنند.

مطالعات نشان داده‌اند که استفاده از این رویکرد می‌تواند باعث بهبود چشم‌گیر در زمان‌بندی وظایف، کاهش زمان اجرای کل و افزایش بهره‌وری منابع شود [۱۶]. پژوهش ژانگ و همکاران [۱۷] مدلی مبتنی بر یادگیری تقویتی عمیق^۳ و گراف‌های چندعاملی برای FJSP ارائه کرده است که در آن عامل‌ها وظایف را با توجه به توالی و مسیر بهینه انتخاب می‌کنند. مطالعات اخیر، مدل‌های مختلفی برای حل مسئله FJSP پیشنهاد کرده‌اند [۱۸]. یوهانسون و همکاران [۶] الگوریتمی مبتنی بر Double DQN برای زمان‌بندی پویا در یک سلول مونتاژ روباتیک توسعه داد. سان و همکاران [14] از یادگیری تقویتی چندعاملی برای تصمیم‌گیری در زمان واقعی در محیط‌های تولید با ظرفیت محدود استفاده کرد. کوین و همکاران [8] مدل DRL چندعاملی برای زمان‌بندی پویا ارائه کرد که به ترکیب آموزش استاتیک با اجرای پویا پرداخت. مدل پیشنهادی این پژوهش، با بهره‌گیری از محرک‌های انگیزه درونی و رویکرد CTDE، گامی نو در حل مسئله زمان‌بندی کارگاهی انعطاف‌پذیر محسوب می‌شود. این مدل با ترکیب سیاست‌های غیرمتمرکز و یادگیری عمیق، توانایی خود را در مدیریت وظایف پویا و بهینه‌سازی منابع به اثبات رسانده است.

۳. روش‌شناسی

۳.۱. بیان مسئله و ساختار سیستم

در این پژوهش، هدف اصلی ارائه یک مدل پویا برای مسئله زمان‌بندی کارگاهی انعطاف‌پذیر در یک محیط تولید چندعاملی^۱ است. این مدل، وظایف تولیدی مختلف را با استفاده از روبات‌های همکار (کوبوت‌ها) در ایستگاه‌های کاری پردازش می‌کند. در این محیط، وظایف تولیدی به‌صورت تصادفی و با نرخ ورود مشخصی (λ) بر اساس توزیع پواسون وارد سیستم می‌شوند [۱۹]. وظایف به‌صورت مجموعه‌ای از عملیات تعریف می‌شوند که باید در ایستگاه‌های کاری مشخص و با استفاده از منابع موجود انجام شوند. هر وظیفه $mj_i \in Mj$ دارای یک تاریخ سررسید $T_{mj_i}^{DD}$ است که باید

^۱ Reinforcement Learning (RL)

^۲ Multi-agent Credit Assignment (MCA)

^۳ Deep Reinforcement Learning (DRL)

تا آن زمان تکمیل شود. وظایف بین ایستگاه‌های کاری توسط کوپوت‌ها منتقل می‌شوند، و هر کوپوت وظایف خاصی را مدیریت می‌کند. نقش کلیدی کوپوت‌ها، کاهش تأخیرها و بهینه‌سازی فرآیند انتقال وظایف بین ایستگاه‌ها است.

۳.۲. محدودیت‌ها و مفروضات مدل

برای طراحی یک سیستم مؤثر، مفروضات در نظر گرفته شده شامل: (i) پردازش وظایف: هر وظیفه از مجموعه‌ای از عملیات تشکیل شده است که باید به ترتیب مشخصی پردازش شوند، (ii) قابلیت کوپوت‌ها: چندین کوپوت مشابه در سیستم وجود دارد، اما هر کوپوت تنها می‌تواند یک عملیات را در هر لحظه انجام دهد، (iii) ظرفیت ایستگاه‌های کاری: هر ایستگاه کاری فقط یک عملیات را در هر لحظه پردازش می‌کند، (iv) محدودیت منابع: منابع موجود در هر ایستگاه کاری ممکن است در هر زمان خاص محدود باشند و (v) زمان لجستیک: زمان انتقال وظایف بین ایستگاه‌های کاری متناسب با فاصله آن‌ها تعیین می‌شود، اگرچه عوامل دیگری نیز ممکن است این زمان را تحت تأثیر قرار دهند. توضیحات نمادها در جدول ۱ ارائه شده‌اند.

۳.۳. تابع هدف: کاهش زمان کل اجرای وظایف

تابع هدف اصلی، حداقل سازی زمان کل مورد نیاز برای تکمیل تمامی وظایف است [۲۰]. این زمان شامل سه بخش است: (i) زمان انتقال وظایف توسط کوپوت‌ها، (ii) زمان توقف عملیات در صورت بروز اختلالات و (iii) زمان پردازش عملیات در ایستگاه‌های کاری. فرض کنیم $mw_{i,k}$ و $mr_{i,l}$ ایستگاه کاری و منبع انتخابی برای وظیفه i در مرحله t باشند. بنابراین، تابع زمان کل اجرای وظایف به صورت فرمول ۱ تعریف می‌شود که مجموع زمان مصرف شده برای هر عملیات در وظایف تولیدی است:

$$\min Obj_{makespan} = \max_{mj \in MJ} \left(\sum_{t \in T} T_{mw_{t,l}}^{lgs} + T_{mw_{t,k},mr_{t,l}}^P + T_{mojt}^{sus} \right) \quad (1)$$

با این شرط که:

$$\forall mj_i \in MJ, \sum_{moj \in mj_i} T_{moj}^P < T_{mj_i}^{DD}$$

جدول ۱. نمادهای به کار رفته در توصیف محیط مسئله.

شاخص‌ها	ثابت‌ها
شاخص وظایف: i	تعداد وظایف تولیدی: n^{mj}
شاخص عملیات: j	تعداد ایستگاه‌های کاری: n^{mw}
شاخص کوپوت: q	تعداد کوپوت‌ها: n^c
شاخص ایستگاه‌های کاری: k	نرخ ورود وظایف: λ
شاخص منابع: l	
مجموعه‌ها	
مجموعه وظایف تولیدی: MJ	
$MJ = \{mj_1, mj_2, \dots, mj_{n^{mj}}\}$	

MW : مجموعه ایستگاه‌های کاری $MW = \{mw_1, mw_2, \dots, mw_n^{mw}\}$ مجموعه کوپوت ها : C $C = \{C_1, C_2, \dots, C_n^c\}$ مجموعه عملیات : MOJ $MOJ = \{moj_j j = 1, 2, \dots\}$ توزیع منابع در ایستگاه‌های کاری : D_R ماتریس زمان انتقال کوپوت ها : lgs^c توالی عملیات وظیفه i م : Seq_i مجموعه منابع : MR $MR = \{mr_l l = 1, 2, \dots\}$	متغیرها
انتقال وظایف توسط کوپوت q بین ایستگاه کاری k و k' : $lgs_{mw_k, mw_{k'}}^q$ عملیات j از وظیفه i : $moj_{i,j}$ زمان لجستیک ایستگاه کاری l در زمان t : $T_{mw_t, l}^{lgs}$ زمان پردازش ایستگاه کاری k در زمان t : $T_{mw_t, k, mr_t, l}^P$ زمان کل تعلیق عملیات در زمان t : $T_{moj_t}^{sus}$ تاریخ سر رسید وظیفه i : $T_{mj_i}^{DD}$	

۳.۴. معماری سیستم یادگیری تقویتی چندعاملی

سیستم یادگیری تقویتی چندعاملی پیشنهادی در این پژوهش، بر همکاری عامل‌ها برای حل مسئله زمان‌بندی متمرکز است. در این معماری، هر وظیفه به‌عنوان یک عامل مستقل در نظر گرفته می‌شود که به‌طور غیرمتمرکز عمل کرده و تصمیم‌گیری‌های خود را بر اساس وضعیت محیط و سیاست‌های یادگیری اتخاذ می‌کند [۲۱]. این سیستم، با ترکیب مؤلفه‌های یادگیری انگیزه درونی و یادگیری تقویتی، بهینه‌سازی فرآیند زمان‌بندی را امکان‌پذیر می‌کند. معماری سیستم شامل سه بخش اصلی است که در تعامل با یکدیگر، عملکرد کلی سیستم را بهبود می‌بخشند. نخستین بخش، سیاست یادگیری غیرمتمرکز است که به هر عامل اجازه می‌دهد به‌طور مستقل از تجربیات خود برای یادگیری و بهبود تصمیم‌گیری استفاده کند. عامل‌ها بدون وابستگی به یک کنترل‌کننده مرکزی، از اطلاعات محلی مانند وضعیت ایستگاه‌های کاری، پیشرفت وظایف، و اولویت‌بندی عملیات بهره می‌برند تا بهترین اقدام ممکن را انتخاب کنند. بخش دوم، محرک انگیزه درونی است که برای هر عامل، یک مکانیزم پاداش درونی فراهم می‌کند. این مکانیزم، عامل‌ها را تشویق می‌کند تا با اولویت‌بندی بهینه وظایف، از تأخیرهای احتمالی جلوگیری کنند. پاداش درونی مبتنی بر وضعیت و اقدام عامل‌ها تعریف می‌شود و مکمل پاداش‌های بیرونی محیط است. این ویژگی به عامل‌ها کمک می‌کند تا در کنار بهبود عملکرد فردی، هماهنگی بهتری با سایر عامل‌ها داشته باشند. در نهایت، بخش سوم شامل منتقد متمرکز است که به‌عنوان یک ناظر عمل کرده و عملکرد عامل‌ها را ارزیابی می‌کند. منتقد متمرکز، با استفاده از اطلاعات جمع‌آوری‌شده از محیط و تعاملات عامل‌ها، سیاست‌های یادگیری را به‌روزرسانی می‌کند. این به‌روزرسانی‌ها به عامل‌ها کمک می‌کند تا رفتار خود را در جهت بهبود عملکرد کلی سیستم تنظیم کنند. به‌طور کلی، این معماری با ترکیب یادگیری غیرمتمرکز، انگیزه درونی، و بازخورد منتقد متمرکز، یک رویکرد نوآورانه برای بهینه‌سازی زمان‌بندی در

محیط‌های تولید پویا ارائه می‌دهد. عامل‌ها علاوه بر تمرکز بر وظایف فردی خود، با سایر عامل‌ها نیز تعامل کرده و در جهت کاهش زمان کل اجرای وظایف و بهبود بهره‌وری منابع تلاش می‌کنند.

۳.۵. فرآیند تصمیم‌گیری مارکوف

به‌جای بهینه‌سازی مستقیم هدف از طریق مدل‌های ریاضی یا الگوریتم‌های ابتکاری، MARL مسائل تصمیم‌گیری متوالی را با مدل‌سازی آن‌ها به‌عنوان یک MDP حل می‌کند [۱۳]. در این پژوهش، ما از یک سیستم کاملاً مشارکتی چندعاملی استفاده می‌کنیم که از الگوریتم DRL غیرمتمرکز بهره می‌برد. در این سیستم، عامل‌ها باید به‌صورت مستقل عمل کنند و بدون کنترل مرکزی، سیاست‌های بهینه تصمیم‌گیری را برای هر وظیفه یاد بگیرند. مدل پیشنهادی بر اساس فرآیند تصمیم‌گیری مارکوف غیرمتمرکز (Dec-POMDP) طراحی شده است. یک وظیفه تصمیم‌گیری در این زمینه به‌صورت یک Dec-POMDP تعریف می‌شود که شامل مجموعه زیر است:

$\langle \mathcal{S}, \mathcal{A}, \mathcal{U}, \mathcal{T}, \mathcal{O}, P, r^{ext}, \gamma \rangle$ که $\mathcal{S}$ مجموعه‌ای از مشاهدات جزئی محیط توسط عامل‌ها، $\mathcal{A}$ مجموعه‌ای از عامل‌ها، $\mathcal{U}$ مجموعه اقدامات ممکن، $\mathcal{T}$ تابع انتقال وضعیت، r^{ext} تابع پاداش بیرونی و γ عامل تخفیف برای پاداش‌های آینده است. این مدل، $\mathcal{O}$ به‌عنوان تابع مشاهده تعریف می‌شود که اطلاعات جزئی هر عامل را از وضعیت کلی محیط استخراج می‌کند، مانند وضعیت ایستگاه کاری مقصد، تعداد وظایف در صف، و زمان‌بندی‌های قبلی. پارامتر P نیز نشان‌دهنده احتمال مشاهده وضعیت t^S توسط عامل i با استفاده از مشاهده محلی $\mathcal{O}_i^t$ است. این تابع نقش کلیدی در تعریف عدم قطعیت مشاهده‌ای عامل‌ها ایفا می‌کند، به‌ویژه در محیط‌هایی با دید ناقص. در هر لحظه، عامل‌ها مجموعه مشاهدات محلی خود را بررسی کرده و اقداماتی انجام می‌دهند که منجر به تغییر وضعیت محیط می‌شود. هدف هر عامل، حداکثرسازی مجموع پاداش‌های آینده است.

هدف هر عامل، حداکثرسازی مجموع پاداش‌های آینده است.

۳.۶. تابع پاداش

تابع پاداش بررسی می‌کند که آیا تصمیمات گرفته‌شده توسط شبکه صحیح هستند یا خیر. برای ارزیابی عملکرد عامل‌ها، تابع پاداش زیر تعریف شده است: (i) پاداش مثبت برای تکمیل وظایف بدون تأخیر، (ii) پاداش منفی برای وظایف به تأخیر افتاده یا متوقف شده. تابع پاداش به‌صورت فرمول ۲ بیان می‌شود:

$$r_t = \begin{cases} N_t^{fin} + N_t^{halt} - N_t^{DD} - \frac{t}{k_1}, & t \neq t_{done} \\ \frac{k_2}{T^{total}}, & t = t_{done} \end{cases} \quad (2)$$

که در آن در اینجا N_t^{fin} نشان‌دهنده تعداد وظایفی است که در گام زمانی t تکمیل شده‌اند. این مقدار به‌صورت جمع تمامی اقداماتی است که $u_t = act_t(finish)$ باشند و $u_t \in \mathcal{U}$ از طرفی N_t^{halt} تعداد وظایفی است که در گام زمانی t متوقف شده‌اند و این مقدار نیز به‌صورت جمع تمامی اقداماتی است که $u_t = act_t(halt)$ باشند و $u_t \in \mathcal{U}$ همچنین N_t^{DD} تعداد وظایفی است که از تاریخ سررسید خود فراتر رفته و تکمیل نشده‌اند. علاوه بر این k_1, k_2 پارامترهای مثبت صحیح ($\mathbb{Z}^+$) هستند که وزن اهمیت دو مؤلفه پاداش (پاداش مثبت برای اتمام موفق وظایف و پاداش نهایی پایان اپیزود) را تعیین می‌کنند. مقداردهی این ضرایب به‌صورت تجربی و با توجه به هدف کنترل شدت اثرگذاری هر بخش انجام شده است. بر اساس تحلیل حساسیت، مقادیر ۱ و ۵ به ترتیب برای k_2, k_1 پایداری بهتری در فرآیند آموزش ایجاد کردند. T^{total} نشان‌دهنده زمان کل اجرای تمامی وظایف است، یعنی مجموع

زمان‌های لازم برای تکمیل تمام عملیات در سیستم. این مقادیر برای ارزیابی عملکرد سیستم در هر گام زمانی و بهینه‌سازی فرآیند زمان‌بندی استفاده می‌شوند.

۳.۷. بازنویسی معماری الگوریتم با ساختار جدید

در این پژوهش، برای حل مسئله زمان‌بندی کارگاهی انعطاف‌پذیر در محیط‌های تولید پویا، یک معماری یادگیری تقویتی چندعاملی معرفی شده است که از مدل یادگیری پاداش درونی فردی (LIIR) الهام گرفته شده است [۴]. این رویکرد ترکیبی از یادگیری غیرمتمرکز، پاداش‌های انگیزه درونی و سیستم‌های منتقد بیرونی را برای بهینه‌سازی تصمیم‌گیری عامل‌ها و بهبود کارایی سیستم ارائه می‌دهد. مزیت اصلی این معماری، توانایی آن در مدل‌سازی رفتار عامل‌ها در محیط‌های پویا و پیچیده با کمترین نیاز به ساده‌سازی‌های مصنوعی است. معماری پیشنهادی با تمرکز بر سه اصل طراحی شده است: (i) عامل‌ها بدون وابستگی به یک کنترل‌کننده مرکزی، بر اساس سیاست‌های مستقل تصمیم‌گیری می‌کنند. این سیاست‌ها از طریق تعاملات محلی با محیط و سایر عامل‌ها بهبود می‌یابد، (ii) استفاده از پاداش‌های انگیزه درونی به عامل‌ها کمک می‌کند تا علاوه بر تمرکز بر اهداف کوتاه‌مدت، بر اهداف بلندمدت سیستم نیز تمرکز داشته باشند، (iii) ارزیابی و بهینه‌سازی دو سطحی به دین معنی که فرآیند یادگیری شامل دو سطح است: حداکثرسازی پاداش‌های فردی عامل‌ها و بهینه‌سازی پاداش‌های کلی سیستم. ر این معماری، پاداش کلی برای عامل‌ها به دو بخش اصلی تقسیم می‌شود: پاداش بیرونی و پاداش درونی. پاداش بیرونی (r_t^{ext}) از محیط دریافت می‌شود و مرتبط با عملکرد کلی سیستم است. پاداش درونی ($r_{\omega,t}^{in}$) بر اساس اولویت‌های عامل و شرایط محلی تعریف می‌شود. برای ترکیب این دو نوع پاداش، یک تابع پاداش پروکسی ($r_{i,t}^{proxy}$) تعریف شده است که به صورت زیر در فرمول ۳ بیان می‌شود:

$$r_{i,t}^{proxy} = r_t^{ext}(\alpha - 1) + r_{\omega,t}^{in}(s_{i,t}, u_{i,t}) * \alpha \quad (3)$$

در این فرمول، α یک ابرپارامتر برای تنظیم تأثیر پاداش درونی نسبت به پاداش بیرونی است. این رویکرد به عامل‌ها اجازه می‌دهد که ضمن بهینه‌سازی اهداف کوتاه‌مدت خود، به اهداف بلندمدت سیستم نیز توجه کنند. مقدار پارامتر α که وزن ترکیب پاداش درونی و بیرونی را تعیین می‌کند، در این پژوهش به صورت تجربی تنظیم شده است. برای تعیین مقدار بهینه، تحلیل حساسیت در بازه $[0.2, 0.8]$ با گام ۰.۱ انجام شد و نتایج نشان دادند که بازه ۰.۴ تا ۰.۶ عملکرد پایدارتری برای مدل ایجاد می‌کند. در نهایت، مقدار $\alpha = 0.5$ انتخاب و برای تمامی آزمایش‌ها استفاده گردید.

۳.۸. ساختار سیستم یادگیری تقویتی

معماری سیستم شامل سه بخش اصلی است، (i) سیاست عامل‌ها و یادگیری مستقل: هر عامل در سیستم، به‌عنوان یک شبکه عصبی سه‌لایه مدل‌سازی می‌شود که شامل لایه‌های ورودی، بازگشتی و خروجی است. داده‌های ورودی شامل مشاهدات محیطی، کدگذاری تک‌حالت‌ه نمایه عامل و اقدام قبلی آن است. این اطلاعات پس از پردازش در لایه بازگشتی دروازه‌دار (GRU)، به‌عنوان یک تابع سیاست به عامل کمک می‌کند که اقدام بهینه را انتخاب کند. لایه خروجی شبکه با استفاده از یک تابع فعال‌سازی Leaky ReLU طراحی شده است تا از مشکلات مرتبط با انفجار گرادینت جلوگیری کند، (ii) پاداش درونی و نقش آن در یادگیری: پاداش درونی یکی از اجزای کلیدی این معماری است که برای هر عامل به‌صورت مستقل تعریف می‌شود. این پاداش‌ها به عامل‌ها کمک می‌کنند تا با اولویت‌بندی بهتر وظایف خود، تأخیرها را کاهش دهند. از آنجا که تعریف دقیق پاداش درونی ممکن است چالش‌برانگیز باشد، یک تابع پاداش پروکسی طراحی شده است که پیچیدگی محیط را بدون نیاز به تعریف جفت حالت-اقدام منعکس می‌کند. این

تابع، پاداش درونی و بیرونی را ترکیب کرده و به عامل‌ها انگیزه می‌دهد تا رفتار بهینه‌ای داشته باشند، iii) منتقد بیرونی و ارزیابی سیستم: برای ارزیابی عملکرد عامل‌ها، یک شبکه منتقد بیرونی طراحی شده است که به تخمین ارزش حالت‌ها و اقدامات در محیط می‌پردازد. این شبکه از تابع مزیت (A_π) استفاده می‌کند که به هر عامل اجازه می‌دهد عملکرد خود را در مقایسه با اقدامات میانگین ارزیابی کند. این تابع به صورت فرمول ۴ تعریف می‌شود:

$$A_\pi(s, u) = r^{ex}(s, u) + V_\theta^{ex}(s') - V_\theta^{ex}(s) \quad (4)$$

در این فرمول، s حالت فعلی، s' حالت بعدی پس از انجام اقدام u و V_θ^{ex} ارزش حالت بر اساس شبکه منتقد بیرونی است. شبکه منتقد بیرونی، ارزش حالت‌ها را تخمین زده و با استفاده از تابع مزیت، اقدامات عامل‌ها را ارزیابی می‌کند. این ارزیابی به عامل‌ها کمک می‌کند که سیاست‌های خود را بهینه کنند. علاوه بر این، یک شبکه منتقد پروکسی نیز طراحی شده است که ارزش حالت‌ها را از دیدگاه عامل محاسبه می‌کند. تابع مزیت پروکسی به صورت فرمول شماره ۵ تعریف می‌شود:

$$A_i^{proxy}(s_i, u_i) = r_i^{proxy}(s_i, u_i) + V_{\theta_i}^{proxy}(s') - V_{\theta_i}^{proxy}(s) \quad (5)$$

در اینجا، $V_{\theta_i}^{proxy}$ نشان‌دهنده پارامتر ارزش پروکسی برای عامل i است و s' حالت بعدی عامل i در مسیر حرکت آن را نشان می‌دهد. ما مدل سیاست و سیستم MASS را به‌طور همزمان با استفاده از پس‌انتشار گرادیان و نزول گرادیان در الگوریتم ۱ آموزش می‌دهیم. برای هر عامل، پاداش درونی به صورت مستقل تعریف شده و با هدف هدایت رفتار عامل در جهت کاهش تأخیر، مدیریت بار ایستگاه کاری و رعایت اولویت وظایف طراحی شده است. پارامترهایی مانند میزان تأخیر وظیفه، تراکم ایستگاه مقصد، اهمیت نسبی وظیفه و نرخ ورود فعلی وظایف در تعریف این پاداش مؤثر هستند. این مقادیر ابتدا به‌عنوان ورودی به یک تابع غیرخطی مبتنی بر شبکه عصبی سبک وارد می‌شوند و خروجی آن به‌عنوان پاداش درونی مورد استفاده قرار می‌گیرد. این پاداش سپس با پاداش بیرونی محیط ترکیب شده و تابع نهایی پاداش پروکسی r_i^{proxy} را شکل می‌دهد، که مبنای به‌روزرسانی سیاست‌های عامل‌ها است. لازم به ذکر است که در این مدل، شبکه منتقد بیرونی صرفاً در مرحله آموزش به‌کار گرفته می‌شود و در زمان اجرا، تمامی عامل‌ها به صورت مستقل و غیرمتمرکز، بر اساس مشاهدات محلی خود اقدام می‌کنند. این طراحی منطبق با رویکرد CTDE بوده و از نظر عملیاتی با الزامات محیط‌های واقعی صنعتی سازگار است.

۳.۹. فرآیند آموزش معماری

الگوریتم آموزشی پیشنهادی به‌صورت کلی به سه بخش تقسیم می‌شود: i) مقداردهی اولیه: پارامترهای شبکه‌های سیاست، پاداش درونی و حافظه بازپخش مقداردهی می‌شوند، ii) تعامل عامل‌ها با محیط: عامل‌ها بر اساس سیاست‌های فعلی خود، اقداماتی انجام می‌دهند که منجر به تغییر وضعیت محیط و ثبت تجربیات جدید می‌شود، iii) به‌روزرسانی سیاست‌ها: با استفاده از داده‌های ثبت شده در حافظه بازپخش، پارامترهای شبکه‌های سیاست و پاداش درونی به‌روزرسانی می‌شوند. این به‌روزرسانی‌ها بر اساس الگوریتم نزول گرادیان و توابع مزیت انجام می‌شود.

۳.۱۰. مزایای معماری پیشنهادی

این معماری قادر است در محیط‌های تولید پویا و با تعداد زیادی از عامل‌ها عملکرد مطلوبی داشته باشد. با استفاده از ترکیب پاداش‌های درونی و بیرونی، عامل‌ها به‌طور همزمان بر اهداف فردی و سیستم تمرکز دارند و استفاده از شبکه‌های بازگشتی و توابع مزیت به بهینه‌سازی فرآیند تصمیم‌گیری کمک می‌کند. معماری پیشنهادی نشان‌دهنده یک گام رو به جلو در استفاده از یادگیری تقویتی برای حل مسائل پیچیده زمان‌بندی در محیط‌های صنعتی است. این روش نه تنها عملکرد فردی عامل‌ها را بهبود می‌بخشد، بلکه بهره‌وری کلی سیستم را نیز ارتقا می‌دهد.

۴. مطالعه تجربی: ارزیابی عملکرد مدل MASS

برای ارزیابی مدل پیشنهادی MASS، سناریوهای مختلفی از مسئله زمان بندی کارگاهی انعطاف پذیر شبیه سازی شده است. در این مطالعه، شاخص اصلی برای اندازه گیری عملکرد، زمان کل اجرای وظایف است. ما عملکرد روش MASS را با الگوریتم های شناخته شده ای مانند PDR، HHA و MADRL مقایسه کرده و آزمایش های مختلفی برای بررسی پایداری آموزش، عملکرد، و مقیاس پذیری آن انجام داده ایم. نتایج زمان بندی از طریق نمودارهای گانت برای نمایش دقیق ارائه شده اند.

۴.۱. پلتفرم توسعه

شبیه سازی ها بر روی پلتفرمی مبتنی بر Python 2.7.3 و PyTorch 1.4.1 اجرا شده اند. سخت افزار استفاده شده شامل یک کارت گرافیک NVIDIA GeForce RTX 2080Ti با ۱۶ گیگابایت حافظه و یک پردازنده Intel Core i7-9700KF است. محیط آزمایش با استفاده از بسته OpenAI Gym طراحی شده که ابزارهای لازم برای توسعه محیط های یادگیری تقویتی را فراهم می کند. زمان مصرف شده در این آزمایش ها به دقیقه محاسبه شده است.

الگوریتم ۱: الگوریتم آموزش MASS

۱. تا زمانی که شرط خاتمه برقرار نشده است، انجام دهید:
 ۱. بازنشانی محیط MASS.
 ۲. مقداردهی اولیه:
 - پارامترهای سیاست بازیگرهای زمان بندی θ
 - پارامتر پاداش درونی ω
 - حافظه بازپخش $\mathcal{D}$
۲. برای هر $t=1$ تا $\max_episode$ ، انجام دهید.
 ۱. مشاهدات عامل o_t از محیط را دریافت کنید.
 ۲. برای هر بازیگر (actor) در گام زمانی t
 - از شبکه سیاست برای انتخاب اقدامات موجود $u_i|o_t$ استفاده کنید.
 - اقدام u_i را در محیط اجرا کنید.
 - وضعیت جدید o'_t و مشاهدات فعلی r_t^{ex} را مشاهده کنید.
 ۳. پاداش درونی r_t^{in} را محاسبه کنید.
 ۴. تجربه $(o_t, r_t^{in}, r_t^{ext}, o'_t, s_{t+1})$ را در حافظه $\mathcal{D}$ ذخیره کنید.
۳. نمونه برداری از یک مینی بچ (Minibatch) از حافظه $\mathcal{D}$:
 ۱. به روزرسانی ω با استفاده از تابع مزیت و نرخ یادگیری β طبق رابطه (8):

$$\omega' = \omega + \beta \nabla \log(\pi|s) A_{\pi}(s, u)$$
 ۲. به روزرسانی θ با استفاده از تابع مزیت پروکسی و نرخ یادگیری α طبق رابطه:

$$\theta' = \theta + \alpha \nabla \log(\pi|s) A^{prox}(s, u) \quad (9)$$

پایان الگوریتم.

۴.۲. داده های شبیه سازی شده

وظایف تولیدی در طول شبیه‌سازی‌ها بر اساس توزیع پواسون با نرخ ورود λ وارد سیستم می‌شوند. برای هر وظیفه، احتمال انتخاب یکسانی برای محصولات مختلف وجود دارد. فرض شده است که تعداد کوبوت‌ها برابر با تعداد ایستگاه‌های کاری است و زمان لجستیک در بازه‌ای از ۱ تا ۷ قرار دارد. داده‌های شبیه‌سازی شده به صورت تصادفی تولید شده و رویدادهای احتمالی بر اساس فرآیند پواسون تعریف شده‌اند. نرخ ورود λ_t در هر زمان t بر اساس فرمول شماره ۶ محاسبه می‌شود:

$$\lambda_t = \frac{\sum_{i=1}^{n^{mj}} \sum_{j=1}^{n^{moj}} T_{i,j}^P + \sum_{k=1}^{n^{mw}} \sum_{kt=1}^{n^{mw}} T_{mw_{t-1,k}, mw_{t-1,kt}}^{lgs}}{n^{mj} n^{mw} \Omega + n^c v} \quad (6)$$

در اینجا، Ω نرخ استفاده از ایستگاه کاری، v نرخ استفاده از کوبوت‌ها است و به صورت ثابت برابر با ۹۵٪ در نظر گرفته شده است. توضیحات سایر نمادها در جدول ۱ ارائه شده‌اند. لازم به ذکر است که در این شبیه‌سازی‌ها، تنها ورود وظایف به صورت تصادفی و پویا در نظر گرفته شده و اختلالاتی نظیر خرابی منابع یا توقف ناگهانی ایستگاه‌ها مدل‌سازی نشده‌اند. هدف اصلی در این مرحله، ارزیابی عملکرد مدل در برابر نوسانات ورودی و بررسی پایداری آموزش عامل‌ها بوده است. با این حال، معماری پیشنهادی قابلیت گسترش به سناریوهای شامل اختلالات عملیاتی را نیز دارا می‌باشد. انتخاب توزیع پواسون برای مدل‌سازی ورود وظایف، مبتنی بر ویژگی‌های آماری شناخته شده این توزیع در مدل‌سازی ورود گسسته سیستم‌های تولیدی است. این انتخاب با مطالعات پیشین در حوزه زمان‌بندی تولید هم‌راستا بوده و امکان تنظیم دقیق نرخ ورود و بررسی پایداری عملکرد سیستم در شرایط مختلف را فراهم کرده است.

۴.۳. محدودیت زمان سررسید (DDT)

در شبیه‌سازی‌ها، محدودیت زمان سررسید (DDT) به عنوان میزان زمانی تعریف می‌شود که وظیفه می‌تواند قبل از رسیدن به تاریخ سررسید تکمیل شود. تاریخ سررسید T_{mji}^{DD} هر وظیفه با استفاده از فرمول شماره ۷ تعیین می‌شود:

$$T_{mji}^{DD} = A_{mji} + \left(\sum_{j=1}^{moji} T_{mji,j}^P \cdot DDT \right) \quad (7)$$

که در آن: A_{mji} زمان ورود وظیفه، $T_{mji,j}^P$ زمان پردازش عملیات، DDT پارامتری که نشان‌دهنده سختی یا انعطاف‌پذیری زمان سررسید است. نتایج شبیه‌سازی برای مقادیر DDT برابر با ۱، ۱.۵ و ۲ مورد بررسی قرار گرفته‌اند.

۴.۴. نمونه‌های مسئله

جدول ۲ تنظیمات اصلی نمونه‌های شبیه‌سازی شده را نشان می‌دهد. این نمونه‌ها شامل تعداد وظایف، تعداد ایستگاه‌های کاری، و دیگر ویژگی‌های مرتبط با عملیات هستند. لازم به ذکر است که با وجود انعطاف‌پذیری مدل MASS در برابر تغییر مقیاس سیستم، برای انطباق کامل با پیکربندی جدید ایستگاه‌های کاری یا کوبوت‌ها، نیاز به آموزش مجدد یا تنظیم مجدد وزن‌های شبکه‌ها وجود دارد. این ویژگی در چارچوب مدل پیش‌بینی شده و با شرایط جدید سازگار می‌گردد.

۴.۵. الگوریتم‌های مقایسه‌ای، آموزش و آزمون

برای بررسی اثربخشی و توانایی مدل MASS در حل مسائل زمان‌بندی کارگاهی انعطاف‌پذیر، مقایسه‌ای جامع میان این مدل و چند الگوریتم معتبر دیگر انجام شده است. این الگوریتم‌ها به دلیل کاربرد گسترده و نتایج برجسته در

¹ Due Date Tightness (DDT)

مسائل مشابه انتخاب شده‌اند. یکی از این الگوریتم‌ها، QMIX است که به دلیل قابلیت‌های بالا در مدیریت سیستم‌های چندعاملی کوچک و استفاده از یک شبکه ترکیبی برای تصمیم‌گیری متمرکز شناخته شده است. این الگوریتم در شرایطی که تعداد عامل‌ها محدود بوده و پیچیدگی محیط کمتر است، عملکرد بسیار خوبی از خود نشان می‌دهد. در مقابل، مدل MASS برای محیط‌های یادگیری تقویتی چندعاملی غیرمتمرکز طراحی شده است که در آن، هر عامل به صورت مستقل عمل کرده و از تعامل با سایر عامل‌ها یاد می‌گیرد. علاوه بر QMIX، یک روش ابر-اکتشافی مبتنی بر بازپخت شبیه‌سازی شده (SA-HH) نیز برای مقایسه انتخاب شده است. این روش به دلیل توانایی در جستجوی فضای راه‌حل و یافتن نتایج نزدیک به بهینه در مسائل پیچیده زمان‌بندی مورد توجه قرار گرفته است. همچنین، دو روش مبتنی بر قوانین اولویت‌بندی (PDR) شامل SPT (کوتاه‌ترین زمان پردازش) و EDD (اولین تاریخ سررسید) به عنوان الگوریتم‌های پایه‌ای و کاربرد در محیط‌های تولید واقعی در این مقایسه گنجانده شده‌اند. این دو روش به طور گسترده در صنعت استفاده می‌شوند و به دلیل سادگی و سرعت اجرا، در شرایطی که نیاز به تصمیم‌گیری سریع است، کاربرد دارند. برای بررسی قابلیت‌ها و پایداری الگوریتم MASS در مقایسه با QMIX، از فرآیندهای آموزشی استاندارد استفاده شد. یکی از معیارهای مهم در این آزمایش‌ها، پایداری آموزش است.

جدول ۲. مشخصات ورودی نمونه‌های زمان‌بندی برای شبیه‌سازی.

نمونه‌ها	تعداد وظایف	تعداد ایستگاه کاری	تعداد عملیات	حداکثر منابع	زمان پردازش
		(n^{mw})	(n^{mol})	$(\max(n^{MR}))$	(T_{ζ}^P)
D01	۱۰	۶	بازه تصادفی [5-7]	۳	بازه تصادفی [1-3]
D02	۱۰	۶	بازه تصادفی [5-7]	۶	بازه تصادفی [1-3]
D03	۱۵	۸	ثابت [10]	۵	بازه تصادفی [1-3]
D04	۱۵	۸	بازه تصادفی [3-10]	۳	بازه تصادفی [1-3]
D05	۱۵	۴	بازه تصادفی [5-10]	۲	بازه تصادفی [3-5]
D06	۱۰	۱۵	ثابت [15]	۵	بازه تصادفی [1-3]
D07	۲۰	۵	بازه تصادفی [5-10]	۵	بازه تصادفی [1-3]
D08	۲۰	۱۰	بازه تصادفی [5-10]	۲	بازه تصادفی [1-3]
D09	۲۰	۱۰	ثابت [10]	۵	بازه تصادفی [1-3]
D10	۲۰	۱۵	ثابت [10]	۲	بازه تصادفی [1-3]
D11	۳	۲	بازه تصادفی [5-7]	۲	بازه تصادفی [1-3]

پایداری در آموزش نه تنها به عنوان شاخصی از تکرارپذیری مدل بلکه به عنوان معیاری برای ارزیابی قابلیت اعتماد به آن نیز اهمیت دارد. MASS به دلیل استفاده از یک چارچوب غیرمتمرکز، توانایی بیشتری در مدیریت محیط‌هایی با تعداد زیاد عامل‌ها و وظایف پیچیده دارد. این امر ناشی از طراحی خاص مدل برای سازگاری با محیط‌های پویا و غیرمتمرکز است که در آن عامل‌ها از مشاهدات محلی خود یاد می‌گیرند. برای ارزیابی دقیق عملکرد، فرآیندهای آموزشی و آزمون در محیط‌های شبیه‌سازی شده انجام شد. محیط‌های آزمایشی انتخاب شده شامل سه نمونه شامل D01، D02، و D06 از جدول داده‌های شبیه‌سازی شده بودند. این نمونه‌ها به گونه‌ای انتخاب شده‌اند که تنوع در تعداد عامل‌ها، ایستگاه‌های کاری، و شرایط منابع را پوشش دهند. هدف این بود که تأثیر تغییرات در تعداد ایستگاه‌های کاری و منابع را بر عملکرد الگوریتم‌ها ارزیابی کنیم. برای هر الگوریتم، پاداش اپیزودیک (مجموع پاداش‌های دریافتی در یک اپیزود) به عنوان معیار اصلی عملکرد در نظر گرفته شد. پاداش اپیزودیک بالاتر نشان‌دهنده عملکرد بهتر مدل است.

نتایج نشان داد که در محیط‌هایی با تعداد محدود منابع یا ایستگاه‌های کاری، عملکرد MASS و QMIX در مراحل اولیه آموزش مشابه است. با این حال، در محیط‌های پیچیده‌تر و با تعداد عامل‌های بیشتر، MASS به‌طور قابل توجهی QMIX را پشت سر گذاشت. این عملکرد بهتر ناشی از مکانیزم انگیزه درونی و معماری یادگیری پیشرفته MASS است که به عامل‌ها کمک می‌کند تا رفتارهای خود را بهینه کرده و از تأخیر جلوگیری کنند. به‌طور کلی، مقایسه‌ها نشان می‌دهد که MASS علاوه بر ارائه نتایج پایدارتر، توانایی مقیاس‌پذیری بهتری نسبت به QMIX و سایر روش‌ها دارد. این یافته‌ها نشان‌دهنده پتانسیل بالای مدل پیشنهادی برای مدیریت مسائل پیچیده زمان‌بندی در محیط‌های تولیدی مدرن است. جدول ۳ مقایسه‌ای از عملکرد روش‌های مختلف زمان‌بندی شامل PDR، الگوریتم ابر-اکتشافی (SA-HH)، QMIX و MASS را بر اساس شاخص زمان کل اجرای وظایف ارائه می‌دهد. داده‌ها نشان می‌دهند که روش MASS در اکثر موارد زمان کمتری برای تکمیل وظایف نسبت به سایر روش‌ها به خود اختصاص داده است. در مقایسه با QMIX، مدل MASS نه تنها دارای ساختار غیرمتمرکز و مستقل در سیستم‌های عامل‌ها است، بلکه از پاداش درونی، شبکه‌های مستقل و مکانیزم توجه بهره می‌برد که آن را برای محیط‌های پویا با تعداد زیاد عامل، مقیاس‌پذیرتر می‌سازد.

۴.۶. نتایج آموزشی، مقایسه عملکرد و بررسی زمان کل اجرای وظایف

فرآیند آموزشی مدل‌های MASS و QMIX نشان می‌دهد که این دو رویکرد در بسیاری از موارد عملکرد مشابهی در مراحل اولیه آموزش دارند. این شباهت به‌ویژه در محیط‌هایی که منابع به مقدار کافی موجود است و پیچیدگی کمتری دارد، مشهود است. با این حال، هنگامی که تعداد منابع یا ایستگاه‌های کاری محدود می‌شود، MASS به دلیل طراحی مبتنی بر مکانیزم انگیزه درونی، برتری قابل توجهی نسبت به QMIX پیدا می‌کند اما در شرایطی که منابع محدود است، پایداری و کارایی بیشتری ارائه می‌دهد. این تفاوت ناشی از توانایی MASS در سازگاری سریع با شرایط متغیر و افزایش هماهنگی بین عامل‌ها است. پایداری در آموزش یک ویژگی حیاتی برای الگوریتم‌های یادگیری تقویتی چندعاملی است. این پایداری به معنای توانایی مدل در یادگیری مؤثر و جلوگیری از نوسانات شدید در فرآیند یادگیری است. شکل ۱، نتایج نشان‌دهنده این است که MASS با استفاده از مکانیزم‌های پیشرفته خود، پایداری بیشتری نسبت به روش‌های DPR-SPT، DPR-EDD و SA-HH دارد. این الگوریتم‌ها گرچه در شرایط ایستا قابل اتکا هستند، اما در شرایط پویا و با ورود وظایف به‌صورت تصادفی دچار افت عملکرد می‌شوند. این مدل با استفاده از مکانیزم‌های انگیزه درونی و هماهنگی بین عامل‌ها توانسته است زمان اجرای وظایف را به حداقل برساند. کاهش زمان اجرای وظایف به‌ویژه در محیط‌های تولیدی با ورود پویا اهمیت بالایی دارد، زیرا تأخیر در تکمیل وظایف می‌تواند منجر به افزایش هزینه‌ها و کاهش بهره‌وری شود. تحلیل نتایج حاکی از آن است که MASS نه تنها زمان اجرای وظایف را کاهش می‌دهد، بلکه به دلیل استفاده بهینه از منابع، توانسته است بهره‌وری کلی سیستم را نیز بهبود بخشد. شایان ذکر است که تمام مقادیر زمانی ارائه‌شده در این بخش، از جمله نتایج زمان کل اجرای وظایف، بر حسب دقیقه محاسبه شده‌اند. یکی از ویژگی‌های منحصر به فرد MASS، مکانیزم پاداش درونی است که به عامل‌ها کمک می‌کند تا تصمیمات بهینه‌تری بگیرند. این مکانیزم از طریق ارائه پاداش‌های مثبت به عامل‌هایی که سیاست‌های کارآمدتری اتخاذ می‌کنند و کاهش پاداش برای سیاست‌های ناکارآمد، فرآیند یادگیری را بهبود می‌بخشد بدین صورت عامل‌ها با استفاده از پاداش‌های درونی یاد گرفته‌اند که وظایف خود را به صورت اولویت‌بندی شده و بر اساس نرخ ورود مدیریت کنند.

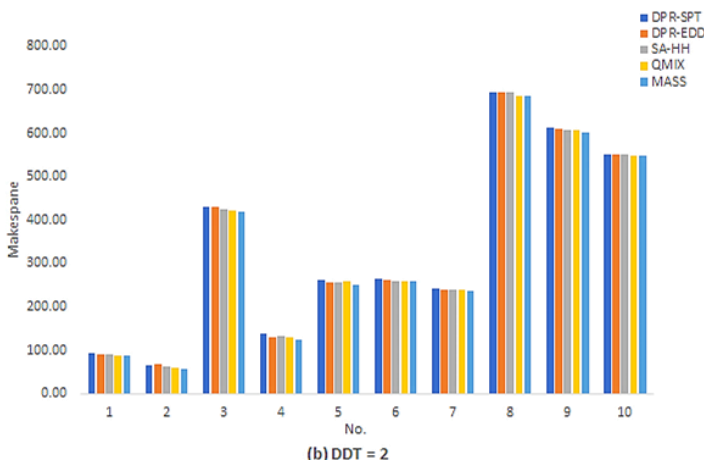
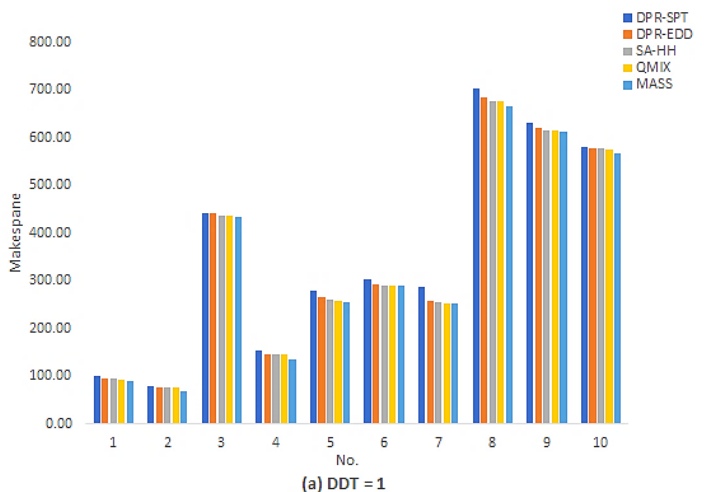
جدول ۳. مقایسه زمان کل اجرای وظایف در مدل های مختلف (برحسب دقیقه).

زمان کل اجرای وظایف											راه حل
D10	D09	D08	D07	D06	D05	D04	D03	D02	D01		
580.83	631.16	701.79	286.59	301.94	278.82	152.40	442.23	78.80	99.91	DDT = 1	PDR
550.19	611.53	694.06	242.26	263.39	261.50	139.00	431.44	66.14	92.27	DDT = 1.5	
537.05	592.34	683.19	201.84	233.45	246.51	127.78	403.43	51.77	89.63	DDT = 2	
578.56	619.45	685.51	256.52	290.94	264.36	146.04	441.93	76.74	95.50	DDT = 1	Hyper-Heuristic
550.42	608.91	693.15	239.29	262.03	257.52	131.10	430.64	66.85	91.48	DDT = 1.5	
532.48	589.28	681.67	200.47	231.28	246.41	126.11	402.10	47.63	89.90	DDT = 2	
578.27	616.10	675.49	254.94	290.65	260.68	146.04	436.87	76.88	93.94	DDT = 1	MADL
550.45	607.56	693.94	239.84	259.97	257.08	131.67	424.77	63.92	91.34	DDT = 1.5	
532.72	588.48	680.69	200.26	231.11	246.38	125.66	401.07	47.85	87.91	DDT = 2	
575.23	615.21	675.69	252.47	289.22	257.16	144.44	436.98	75.20	93.04	DDT = 1	QMIK
548.24	606.55	687.11	238.70	259.53	257.82	130.86	422.17	59.06	87.78	DDT = 1.5	
529.10	588.03	680.09	197.68	230.17	244.23	124.69	398.12	43.51	83.69	DDT = 2	
568.05	611.78	665.53	252.25	288.90	255.97	135.64	434.66	68.74	90.58	DDT = 1	MASS
547.16	602.34	686.90	236.97	257.73	250.39	124.47	420.16	58.17	87.63	DDT = 1.5	
524.10	585.09	676.87	194.94	230.64	243.64	122.14	397.64	43.17	83.47	DDT = 2	

مکانیسم توجه در مدل MASS نقش مهمی در بهبود فرآیند یادگیری پاداش درونی دارد. این مکانیسم به عامل‌ها کمک می‌کند تا در هر گام زمانی، وظایف و منابع مرتبط با اولویت بیشتری را شناسایی کنند. این موضوع باعث شده است که عامل‌ها بتوانند با تمرکز بیشتر بر وظایف حیاتی، تصمیمات مؤثرتری بگیرند. استفاده از این مکانیسم کارایی مدل را افزایش داده و توانایی آن را در مدیریت وظایف پیچیده بهبود بخشیده است.

۵. نتیجه‌گیری

در این پژوهش، یک راهکار یادگیری تقویتی چندعامله با بهره‌گیری از انگیزش درونی برای حل مسائل زمان‌بندی پویا در کارگاه‌های تولیدی منعطف ارائه شده است. سیستم پیشنهادی از یک فرآیند آموزش متمرکز بهره می‌برد که امکان اجرای غیرمتمرکز زمان‌بندی را فراهم می‌سازد. در این رویکرد، زمان‌بندی وظایف در یک محیط چندعامله با استفاده از یک فرآیند تصمیم‌گیری مارکوف غیرمتمرکز انجام می‌شود که هدف آن ایجاد تعادل میان هماهنگی عامل‌ها و استقلال عملکرد هر عامل است. نتایج حاصل از این مطالعه نشان می‌دهد که سیستم چندعامله مبتنی بر انگیزش درونی ما، در مقایسه با روش‌های متداول، کارایی بهتری در بهینه‌سازی زمان تکمیل کارها دارد. MASS در مراحل پیشرفته آموزش عملکرد پایدارتر و بهتری از خود نشان می‌دهد. این عملکرد به‌ویژه در نمونه‌های آزمایشی D1-D10 قابل مشاهده است. این روش قادر است در یک چارچوب یادگیری تقویتی چندعامله، ورود پویا و غیرقابل پیش‌بینی کارها را مدیریت کند، بدون اینکه پایداری فرآیند یادگیری مختل شود. علاوه بر بهبود کارایی و کاهش زمان تکمیل، مدل پیشنهادی انعطاف‌پذیری بالایی را در شرایط پویا نشان داده و پایداری فرایند آموزش آن تأیید شده است. همچنین، ساختار پیشنهادی این مدل با توجه به طراحی ماژولار و معماری غیرمتمرکز عامل‌ها، قابلیت تعمیم به سایر مسائل زمان‌بندی نظیر Flow Shop و Open Shop را نیز داراست، به شرط بازتعریف مناسب توالی عملیات و محدودیت‌های منابع در محیط آموزش.



شکل ۱. مقایسه روند آموزش مدل‌ها در مقادیر مختلف DDT

این ویژگی‌ها نشان می‌دهند که این مدل می‌تواند برای کاربردهای صنعتی گوناگون در حوزه زمان‌بندی تولید، از جمله صنایع پیچیده و محیط‌های تولیدی با تغییرات مداوم، مورد استفاده قرار گیرد. هرچند در مسیر پیاده‌سازی صنعتی مدل MASS، چالش‌هایی از جمله نیاز به داده‌های دقیق از محیط واقعی، تطبیق با سیستم‌های کنترلی موجود، مقاومت در برابر نویز محیطی، و آموزش مجدد در صورت تغییرات ساختاری وجود خواهد داشت. با این حال، با توجه به طراحی غیرمتمرکز و قابلیت انعطاف‌پذیری مدل، این چالش‌ها قابل مدیریت و حل هستند.

References

- [۱] Aubret, A., Matignon, L., & Hassas, S. (2023). *An information-theoretic perspective on intrinsic motivation in reinforcement learning: A survey*. Entropy, 25(2), 327. <https://www.mdpi.com/1099-4300/25/2/327>.
- [۲] Borboni, A., Reddy, K. V. V., Elamvazuthi, I., AL-Quraishi, M. S., Natarajan, E. & Azhar Ali, S. S. (2023). *The expanding role of artificial intelligence in collaborative robots for industrial applications: a systematic review of recent works*. Machines, 11(1), 111. <https://www.mdpi.com/2075-1702/11/1/111>.
- [۳] Brandimarte, P. (1993). *Routing and scheduling in a flexible job shop by tabu search*. Annals of Operations research, 41(3), 157-183. <https://doi.org/10.1007/BF02023073>.
- [۴] Du, Y., Han, L., Fang, M., Liu, J., Dai, T., & Tao, D. (2019). *Liir: Learning individual intrinsic reward in multi-agent reinforcement learning*. Advances in neural information processing systems, 32. https://proceedings.neurips.cc/paper_files/paper/2019/file/07a9d3fed4c5ea6b17e80258dee231fa-Paper.pdf.
- [۵] Hao, J., Yang, T., Tang, H., Bai, C., Liu, J., Meng, Z., Liu, P., & Wang, Z. (2023). *Exploration in deep reinforcement learning: From single-agent to multiagent domain*. IEEE Transactions on Neural Networks and Learning Systems. <https://ieeexplore.ieee.org/document/10021988>.
- [۶] Johnson, D., Chen, G., & Lu, Y. (2022). *Multi-agent reinforcement learning for real-time dynamic production scheduling in a robot assembly cell*. IEEE Robotics and Automation Letters, 7(3), 7684-7691. <https://ieeexplore.ieee.org/document/9801608>.
- [۷] Hoseini, F., Sepehrzadeh, H., & Kheyri, M. (202۴). *Presenting an "Adaption Ahead" Optimization Algorithm for Training Models Used in Deep Learning*. Karafan Journal. https://karafan.nus.ac.ir/article_210487.html.
- [۸] Keung, K. L., Lee, C. K., Liqiao, X., Chao, L., Bufan, L., & Ji, P. (2023). *A cyber-physical robotic mobile fulfillment system in smart manufacturing: The simulation aspect*. Robotics and Computer-Integrated Manufacturing, 83, 102578. <https://doi.org/10.1016/j.rcim.2023.102578>.
- [۹] Ladosz, P., Weng, L., Kim, M., & Oh, H. (2022). *Exploration in deep reinforcement learning: A survey*. Information Fusion, 85, 1-22. <https://doi.org/10.1016/j.inffus.2022.03.003>.
- [۱۰] Lim, K. C. W., Wong, L.-P., & Chin, J. F. (2023). *Hyper-heuristic for flexible job shop scheduling problem with stochastic job arrivals*. Manufacturing letters, 36, 5-8. <https://doi.org/10.1016/j.mfglet.2022.12.009>.
- [۱۱] Liu, C.-L., Chang, C.-C., & Tseng, C.-J. (2020). *Actor-critic deep reinforcement learning for solving job shop scheduling problems*. Ieee Access, 8, 71752-71762. <https://ieeexplore.ieee.org/document/9066984>.
- [۱۲] Rafiei, S. H., Ojaghi, M., & Sabouri, M. (2023). *The Effective Use of SVM to Detect Faults in Electric Machines by Finite Element Analysis*. Karafan Journal, 20(3), 367-392. https://karafan.tvu.ac.ir/article_181325.html?lang=en.

- [۱۳] Oroojlooy, A., & Hajinezhad, D. (2023). *A review of cooperative multi-agent deep reinforcement learning*. Applied Intelligence, 53(11), 13677-13722. <https://link.springer.com/article/10.1007/s10489-022-04105-y>.
- [۱۴] Sun, Y., Chung, S.-H., Wen, X., & Ma, H.-L. (2021). *Novel robotic job-shop scheduling models with deadlock and robot movement considerations*. Transportation Research Part E: Logistics and Transportation Review, 149, 102273. <https://doi.org/10.1016/j.tre.2021.102273>.
- [۱۵] Benhari, M., & Hosseini, R. (2024). *An Intelligent Ensemble Model of Uncertainty Management in Belief Network for the Classification of Microscopic Images to Detect Cervical Cancer*. Karafan Journal, 2۱(۱), 91-110. https://karafan.nus.ac.ir/article_181335.html?lang=en.
- [۱۶] Wang, M., Zhang, J., Zhang, P., Cui, L., & Zhang, G. (2022). *Independent double DQN-based multi-agent reinforcement learning approach for online two-stage hybrid flow shop scheduling with batch machines*. Journal of Manufacturing Systems, 65, 694-708. <https://doi.org/10.1016/j.jmsy.2022.11.001>.
- [۱۷] Xiong, W., & Fu, D. (2018). *A new immune multi-agent system for the flexible job shop scheduling problem*. Journal of Intelligent Manufacturing, 29, 857-873. <https://link.springer.com/article/10.1007/s10845-015-1137-2>.
- [۱۸] Xu, M., Zhang, F., Mei, Y., & Zhang, M. (2022). *Genetic programming with multi-case fitness for dynamic flexible job shop scheduling*. 2022 IEEE Congress on Evolutionary Computation (CEC). <https://ieeexplore.ieee.org/document/9870340>.
- [۱۹] Yang, S., Wang, J., & Xu, Z. (2022). *Real-time scheduling for distributed permutation flowshops with dynamic job arrivals using deep reinforcement learning*. Advanced Engineering Informatics, 54, 101776. <https://doi.org/10.1016/j.aei.2022.101776>.
- [۲۰] Zhang, J.-D., He, Z., Chan, W.-H., & Chow, C.-Y. (2023). *DeepMAG: Deep reinforcement learning with multi-agent graphs for flexible job shop scheduling*. Knowledge-Based Systems, 259, 110083. <https://doi.org/10.1016/j.knosys.2022.110083>.
- [۲۱] Zhang, Y., Zhu, H., Tang, D., Zhou, T., & Gui, Y. (2022). *Dynamic job shop scheduling based on deep reinforcement learning for multi-agent manufacturing systems*. Robotics and Computer-Integrated Manufacturing, 78, 102412. <https://doi.org/10.1016/j.rcim.2022.102412>.