

Fuzzy Logic-Driven Reward Optimization for Deep Reinforcement Learning in Carbon Emission Trading

Seyed Ali Hosseini
Politecnico di Milano
Department of Energy
Milan, Italy
seyedal1.hosseini@polimi.it

Francesco Grimaccia
Politecnico di Milano
Department of Energy
Milan, Italy
francesco.grimaccia@polimi.it

Alessandro Niccolai
Politecnico di Milano
Department of Energy
Milan, Italy
alessandro.niccolai@polimi.it

Abstract—This study presents a novel algorithmic trading framework for carbon emission markets, integrating fuzzy logic with deep reinforcement learning (DRL) to optimize trading strategies under uncertainty. We employ an Advantage Actor-Critic (A2C) algorithm with a fuzzy logic-enhanced reward function, leveraging triangular membership functions and a pre-defined rule base to model volatile market dynamics, including price fluctuations and technical indicators. Using five years of tick-by-tick carbon emission data, the framework is validated through a rigorous walk-forward approach, demonstrating robust generalization. Results show significant improvements in cumulative return, Sharpe ratio, and Sortino ratio, with the fuzzy reward function enabling stable and adaptive decision-making in complex market conditions. By synergizing fuzzy systems with DRL, this work advances trading strategy development, offering a scalable approach for dynamic financial environments.

Index Terms—Financial Data Analysis, Reinforcement Learning, Deep Neural Networks, Fuzzy Logic, Algorithmic Trading

I. INTRODUCTION

The global push for sustainability and decarbonization has elevated carbon emission allowance markets as critical mechanisms for reducing greenhouse gas emissions while supporting economic performance [1]. These cap-and-trade systems incentivize companies to lower emissions, directly impacting their financial and operational strategies by linking compliance to market-driven trading opportunities [2]. However, the carbon market's complexity, marked by extreme price volatility, supply-demand imbalances, regulatory uncertainties, and nonlinear dynamics, poses significant challenges for algorithmic trading.

Traditional trading methods, such as rule-based or statistical approaches, are often inadequate in this data-rich, stochastic environment. These methods rely on static patterns or assumptions of market stability, limiting their ability to process vast datasets or balance multiple criteria, such as profitability and environmental compliance [3], [4]. The high-dimensional, dynamic nature of carbon markets demands adaptive strategies capable of real-time decision-making under uncertainty.

Deep reinforcement learning (DRL) offers a powerful solution, acting as an intelligent trader that learns optimal strategies through trial and error [5]. Unlike traditional methods,

DRL excels in stochastic environments, adapting to volatile market conditions and optimizing complex objectives, as demonstrated in energy pricing and bidding applications [6], [7]. However, applying DRL to carbon emission markets remains underexplored, particularly due to the critical challenge of designing effective reward functions. Conventional reward functions often fail to capture the noise, ambiguity, and multi-objective nature of carbon trading, leading to suboptimal performance.

To address this, we propose a novel DRL-based trading framework enhanced by a fuzzy logic-driven reward function, tailored for carbon emission markets [8], [9]. Fuzzy logic enables the incorporation of expert knowledge and human-like reasoning, using fuzzy sets and rules to model imprecise market inputs, such as price volatility and regulatory constraints [10]. This fuzzy reward mechanism overcomes the limitations of rigid reward designs, enabling adaptive and robust trading strategies. By synergizing fuzzy systems with DRL, our approach enhances trading performance in volatile, uncertain environments, advancing the application of fuzzy logic in financial decision-making.

This paper is structured as follows: Section II outlines the proposed methodology and related approaches. Section III details the experimental results and their evaluation. Finally, Section IV provides the conclusions of the study and discusses directions for future research.

II. METHODOLOGY

This section outlines the key components of the proposed framework, including data processing, reinforcement learning architecture, and our novel fuzzy-based reward mechanism. Figure 1 presents an overview of the complete workflow.

A. Data Preparation

The dataset consists of tick-level carbon emission allowance trading data from 2018 to mid-2023. To align with algorithmic requirements and improve efficiency, the data was aggregated into daily intervals. Logarithmic normalization (Equation 1) was applied to reduce the impact of scale and improve learning stability:

$$\text{Normalized Price}_t = \frac{\ln(\text{Price}_t)}{\ln(\text{Price}_{t-1})} \quad (1)$$

To mimic real-world conditions, the walk-forward method was used for train-test splits, as illustrated in Figure 2.

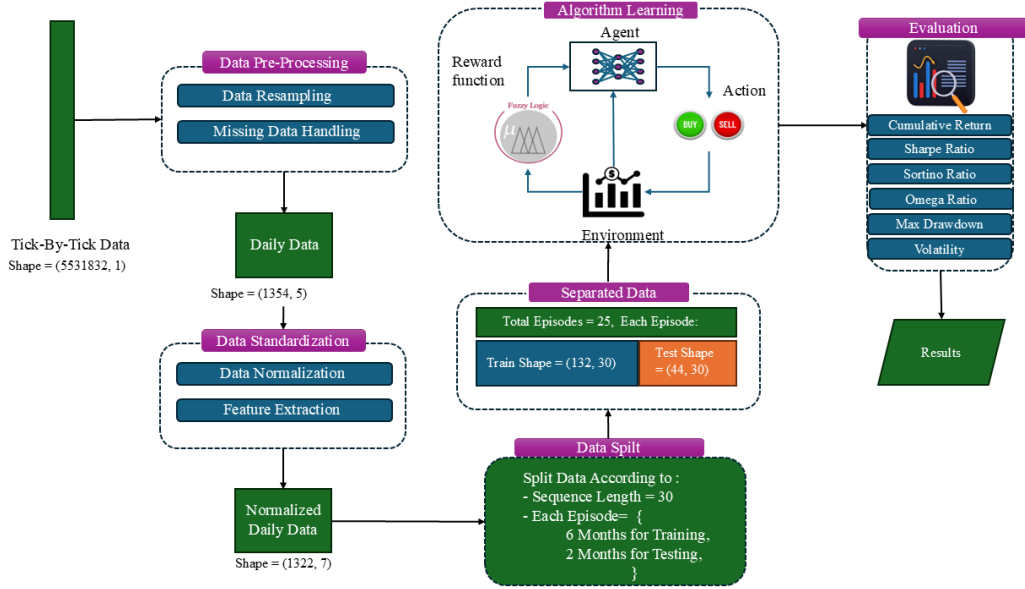


Fig. 1: Main flow chart of implemented methodology

This setup ensures chronological integrity and avoids data leakage by strictly training on historical data and testing only on future data. Contrary to conventional walk-forward strategies that retrain the model on each window, risking overfitting due to repeated exposure to similar patterns, we train our model only once over the initial training period and evaluate it sequentially across the remaining time windows. This single-training-pass design prevents the test data from participating in the learning process, thereby preserving the objectivity of the evaluation and mitigating overfitting risks. Moreover, by training and testing according to this method, we are able to evaluate the model's performance across different segments of the dataset, each potentially exhibiting distinct market conditions. This enhances the robustness and capability to generalize of the proposed strategy in dynamic trading environments [11].

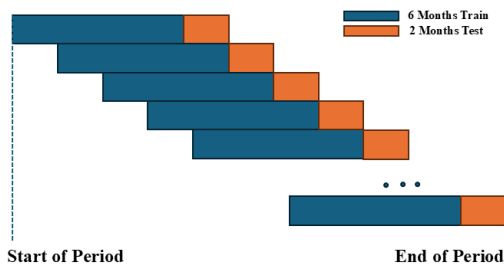


Fig. 2: Walk-forward train/test separation using a sliding time window

Key and famous financial indicators among traders, RSI, MACD, Bollinger Bands, OBV, and ATR, were computed, normalized, and appended to each timestep [12]. These indicators were selected for their relevance to market dynamics and later serve as inputs to the custom reward mechanism.

B. Reinforcement Learning Architecture

The trading environment is modeled as a Markov Decision Process (S, A, R, π, γ) , where R is the reward function

providing immediate scalar feedback and $\gamma \in [0, 1)$ is the discount factor that balances short-term versus long-term rewards. The state S includes a sequence of 30 past normalized close prices. The action space A comprises discrete trading decisions (buy, sell). We employ the Advantage Actor-Critic (A2C) algorithm [13], a synchronous, policy-gradient method suitable for continuous environments. The policy π is parameterized by a neural network (Figure 3), and the value function is updated concurrently to reduce variance. Though A2C is a standard RL method, its integration with our custom state representation and reward structure tailors it to carbon market dynamics.

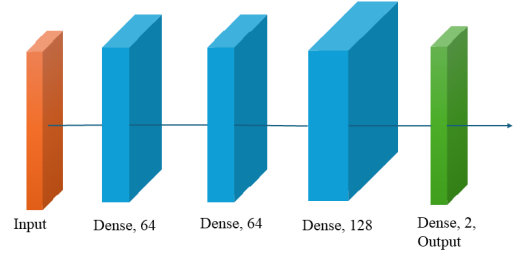


Fig. 3: Deep Neural Network Structure

C. Fuzzy Logic-Based Reward Shaping with Expert Knowledge (Main Contribution)

To overcome the limitations of sparse and delayed rewards in carbon emission allowance trading, we propose a novel fuzzy logic-driven reward function that incorporates domain-specific expert knowledge while preserving interpretability. Unlike traditional profit-based or binary rewards, our approach generates dense, context-aware reward signals that guide the deep reinforcement learning agent more effectively in the highly volatile and policy-sensitive carbon market.

The fuzzy reward system follows the standard Mamdani architecture (Figure 4) and consists of four stages: fuzzification, rule inference, aggregation, and defuzzification.

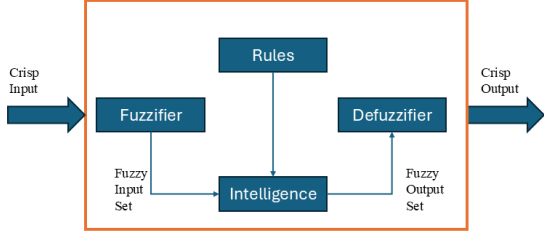


Fig. 4: Mamdani fuzzy inference process used in our reward function

1) *Fuzzification with Triangular Membership Functions*: five key and technical indicators, plus the agent’s instantaneous return, are used as inputs: RSI, MACD histogram, Bollinger Band position, On-Balance Volume (OBV) trend, and Average True Range (ATR). Each input is fuzzified into three linguistic terms: Bearish, Neutral, Bullish.

We deliberately selected triangular membership functions (Figure 5) for the following reasons:

- **Computational efficiency**: triangular functions are piecewise linear and require minimal computation, critical when the fuzzy system is called at every trading step inside the RL training loop.
- **Smooth overlap**: exactly two adjacent functions overlap at any point with a total membership degree of 1.0, ensuring smooth transitions and avoiding abrupt reward changes.
- **Interpretability**: the three parameters (Bearish, Neutral, Bullish) directly correspond to intuitive market thresholds set by carbon trading experts (e.g., $RSI < 30$ = oversold, $RSI > 70$ = overbought).
- **Proven effectiveness**: triangular MFs are the most widely used in financial applications [14] due to their balance between simplicity and expressiveness.

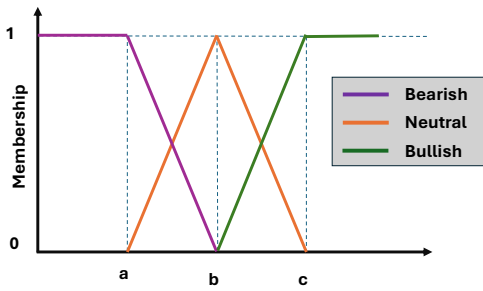


Fig. 5: Employed Triangular membership function

The membership degree $\mu_A(x)$ for a triangular function is defined as:

$$\mu_A(x) = \max \left(\min \left(\frac{x-a}{b-a}, \frac{c-x}{c-b} \right), 0 \right)$$

where a , b , and c denote the left, peak, and right vertices of the triangular membership function, calibrated using historical financial features derived from EUA (European Union Allowance) price data spanning 2018–2023, as described in the Methodology section.

2) *Custom Expert Rule Base for Carbon Trading*: The rule base (Table I) was constructed through collaboration with carbon market analysts and encodes well-established trading heuristics adapted to the carbon allowance market, which exhibits mean-reversion around policy announcement periods, sudden volatility spikes after compliance events, and strong volume-price divergence before major auctions.

Examples of key rules:

- **Rule 1**: if the RSI is *Bearish*, the Bollinger Bands position is *Bearish*, and the OBV is *Bullish*, then a **high positive reward** is assigned, reflecting a strong mean-reversion regime that often precedes policy-driven price recoveries.
- **Rule 2**: if the ATR is *Bullish* and the returns are *Bearish*, then a **low (negative) reward** is assigned, penalizing exposure during extreme volatility episodes typically associated with abrupt regulatory or macroeconomic shocks.
- **Rule 3**: if the MACD is *Bullish* and the RSI is *Neutral*, then a **moderate positive reward** is assigned, encouraging trend-following behavior when bullish momentum emerges without overbought market conditions.

In total, 12 core rules (plus automatic expansion via neutrality) cover the most informative market regimes while keeping the system computationally tractable.

3) *Inference and Defuzzification*: during inference:

- 1) For each active rule, the firing strength is calculated as the **minimum** (Mamdani implication) of the antecedent membership degrees.
- 2) The output fuzzy set (reward) is truncated at the firing strength.
- 3) All truncated output sets are aggregated using **maximum**.
- 4) Final crisp reward r_t is obtained via **centroid defuzzification**:

$$r_t = \frac{\int z \cdot \mu_{\text{agg}}(z) dz}{\int \mu_{\text{agg}}(z) dz}$$

where the output universe for reward is defined on $[-2, +2]$ (later scaled to match the magnitude of financial returns).

This produces a smooth, dense reward signal ranging continuously from strong punishment (-2) to strong encouragement (+2), rather than sparse terminal rewards.

4) *Advantages for Deep RL in Carbon Trading*: the proposed fuzzy reward offers several critical advantages:

- **Dense guidance**: instead of receiving feedback only at episode end, the agent gets meaningful signals at every step.
- **Expert knowledge injection** without losing end-to-end differentiability (the fuzzy module is non-learnable but fully deterministic).
- **Robustness to market regime shifts** typical in carbon markets (policy announcements, compliance deadlines).
- **Full interpretability**: every reward value can be explained by tracing activated rules and membership degrees.

Thus, the fuzzy logic component is not an ad-hoc addition but a principled, domain-aware reward shaping mechanism that directly addresses the credit assignment problem in carbon emission trading environments.

TABLE I: Fuzzy Rule Base for Carbon Emission Allowance Trading

Condition 1	Condition 2	Condition 3	Consequent (Reward)
RSI: <i>Bearish</i>	MACD: <i>Bullish</i>	N/A	<i>High</i>
RSI: <i>Bullish</i>	MACD: <i>Bearish</i>	N/A	<i>Low</i>
Bollinger Bands: <i>Bearish</i>	OBV: <i>Bullish</i>	N/A	<i>High</i>
Bollinger Bands: <i>Bullish</i>	OBV: <i>Bearish</i>	N/A	<i>Low</i>
ATR: <i>Bullish</i>	Returns: <i>Bearish</i>	N/A	<i>Low</i>
ATR: <i>Bearish</i>	Returns: <i>Bullish</i>	N/A	<i>High</i>
RSI: <i>Bearish</i>	Bollinger Bands: <i>Bearish</i>	OBV: <i>Bullish</i>	<i>Very High</i>
RSI: <i>Bullish</i>	Bollinger Bands: <i>Bullish</i>	OBV: <i>Bearish</i>	<i>Very Low</i>
Any <i>Neutral</i> configuration	N/A	N/A	<i>Neutral</i>

D. Performance Metrics

To evaluate the performance of the proposed deep reinforcement learning algorithm, several key performance metrics were used. These metrics provide a comprehensive assessment of both the profitability and risk management capabilities of the implemented trading strategy. Below, each metric is described in detail, along with its mathematical formulation.

1) *Cumulative Return*: Cumulative return measures the total percentage change in the value of a portfolio over a specified period [15]. It is calculated as:

$$R_c = \left(\frac{P_{end} - P_{start}}{P_{start}} \right) \times 100 \quad (2)$$

where P_{start} is the initial portfolio value and P_{end} is the final portfolio value. A higher cumulative return indicates greater overall profitability.

2) *Sharpe Ratio*: The Sharpe ratio evaluates the risk-adjusted return of a portfolio by comparing its excess return to the level of risk (measured by standard deviation) [16]. It is given by:

$$S = \frac{\bar{R}_p - R_f}{\sigma_p} \quad (3)$$

where $\bar{R}_p$ is the mean portfolio return, R_f is the risk-free rate, and σ_p is the standard deviation of portfolio returns. A higher Sharpe ratio indicates better performance per unit of risk.

3) *Sortino Ratio*: The Sortino ratio, a variation of the Sharpe ratio, focuses on downside risk by considering only negative deviations from the mean [17]. It is calculated as:

$$S_{Sortino} = \frac{\bar{R}_p - R_f}{\sigma_d} \quad (4)$$

where σ_d is the standard deviation of negative returns (downside risk). The Sortino ratio provides a more refined measure of risk-adjusted return in scenarios where downside risk is more relevant.

4) *Omega Ratio*: The Omega ratio evaluates the probability-weighted ratio of gains to losses above a specified threshold return, R_T [18]. It is defined as:

$$\Omega = \frac{\int_{R_T}^{\infty} (1 - F(R)) dR}{\int_{-\infty}^{R_T} F(R) dR} \quad (5)$$

where $F(R)$ is the cumulative distribution function of portfolio returns. A higher Omega ratio indicates better performance with respect to the chosen threshold.

5) *Maximum Drawdown*: Maximum drawdown (MDD) represents the largest peak-to-trough decline in portfolio value during the test period, expressed as a percentage [19]. It is given by:

$$MDD = \max_{t \in [0, T]} \left(\frac{\text{Peak}_t - \text{Trough}_t}{\text{Peak}_t} \right) \times 100 \quad (6)$$

where Peak_t is the highest portfolio value up to time t , and Trough_t is the lowest portfolio value following the peak. A lower MDD indicates a more robust trading strategy.

6) *Volatility*: Volatility measures the variability of returns over time, indicating the level of risk [20]. It is calculated as the standard deviation of portfolio returns:

$$\sigma = \sqrt{\frac{1}{N-1} \sum_{i=1}^N (R_i - \bar{R})^2} \quad (7)$$

where R_i is the return at time i , $\bar{R}$ is the mean return, and N is the number of observations. Lower volatility reflects more stable performance.

By combining these metrics, a comprehensive evaluation of the proposed algorithm's performance is achieved, enabling the assessment of both profitability and risk management.

III. RESULTS AND DISCUSSION

To evaluate the performance of the deep reinforcement learning (DRL) model under different reward functions, three distinct reward mechanisms were tested: (1) return-based reward, (2) binary reward (+1 for profitable trades, -1 for unprofitable trades), and (3) a fuzzy logic-based reward. Each reward function was integrated into the DRL framework and tested on the same dataset to ensure comparability.

During each trading step, the reward was calculated based on the respective reward function, reflecting the agent's performance under different criteria. The fuzzy logic-based reward function uniquely incorporated market conditions and agent performance metrics, aiming to provide a more nuanced and adaptive measure of success.

To ensure reproducibility, all experiments were conducted with a fixed random seed for data splitting and model initialization. This ensured consistent results across multiple runs, demonstrating the robustness and repeatability of the proposed methods.

A. Performance Metrics

The effectiveness of the DRL model under each reward function was assessed using several key performance metrics during the test period. Table II summarizes the results for cumulative return, Sharpe ratio, Sortino ratio, Omega ratio, maximum drawdown, and volatility.

TABLE II: Performance Metrics During the Test Period

Metric	Return-Based Reward	Binary Reward (+1/-1)	Fuzzy-Based Reward
Cumulative Return (%)	88.82	69.95	111.14
Sharpe Ratio	1.21	1.10	1.51
Sortino Ratio	1.23	1.19	1.57
Omega Ratio	1.24	1.22	1.31
Max Drawdown (%)	-22.50	-20.52	-21.43
Volatility (%)	20.17	18.41	18.56

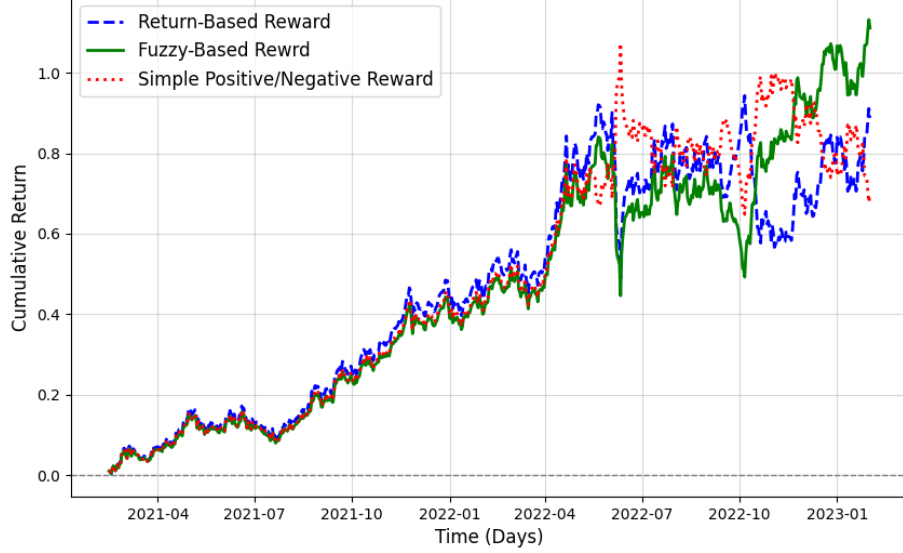


Fig. 6: Cumulative return trends for different reward functions

The fuzzy-based reward function outperformed the other two reward mechanisms across most metrics. It achieved the highest cumulative return (111.14%), the highest Sharpe ratio (1.51), and the highest Sortino ratio (1.57), indicating a strong balance between risk and return. Furthermore, its maximum drawdown highlights the robustness of the trading strategy under adverse market conditions.

B. Return Trend Analysis

The cumulative return trends for the DRL model under different reward functions are shown in Figure 6. The plot illustrates the performance of the model over time, highlighting the superior trajectory achieved with the fuzzy-based reward function. While the return-based reward and binary reward also showed positive trends, the fuzzy-based reward consistently led to higher returns and smoother performance, validating its ability to adapt to varying market conditions.

C. Discussion

The results clearly demonstrate that the fuzzy logic-based reward function significantly enhances the performance of the DRL algorithm for algorithmic trading. By incorporating both financial indicators and market conditions, the fuzzy-based reward guided the agent to make informed and balanced trading decisions, maximizing profitability while minimizing risk.

In comparison, the return-based and binary rewards, while simpler in design, failed to achieve the same level of performance. The binary reward, in particular, demonstrated a more limited ability to account for the complexity of financial

markets, as it provided no distinction between large and small profit/loss trades. Similarly, the return-based reward was effective but lacked the adaptability provided by the fuzzy logic system.

The superior performance of the fuzzy-based reward is evident in its higher Sharpe and Sortino ratios, which indicate better risk-adjusted returns. The low volatility and drawdown further highlight the robustness of this approach, making it well-suited for real-world trading applications.

Overall, the integration of fuzzy logic into the reward mechanism proved to be a significant advancement for DRL in algorithmic trading. These findings suggest that this approach can serve as a foundation for future research in developing sustainable and adaptive trading strategies in complex markets.

IV. CONCLUSION

This study presented a novel deep reinforcement learning framework for algorithmic trading in the carbon emission allowance market, enhanced by a fuzzy logic-based reward function. The integration of fuzzy logic allowed the agent to incorporate market conditions and financial indicators, such as RSI, MACD, Bollinger Bands, OBV, and ATR, into the reward calculation, providing a more dynamic and adaptive approach to trading decision-making.

The model employed the A2C algorithm as the primary reinforcement learning agent, with the state space defined as a sequence of the last 30 close prices. The training and testing process leveraged walk-forward validation, ensuring robust performance evaluation across changing market conditions.

Results demonstrated that the proposed framework achieved strong risk-adjusted returns, with metrics such as the Sharpe ratio, Sortino ratio, and Omega ratio highlighting the model's ability to balance profitability and risk. Additionally, the fuzzy reward function effectively guided the agent, contributing to stable performance and reduced volatility.

Reproducibility was a key focus, with fixed random seeds used throughout the experimental setup to ensure consistent results. These findings underscore the potential of combining deep reinforcement learning with fuzzy logic to address the challenges of dynamic and uncertain markets like the carbon emission allowance market.

A. Future Work

There are several avenues for future research to enhance the proposed framework further. One potential direction is the expansion and refinement of the fuzzy rule base. The current rule set, while effective, could be enriched by incorporating additional financial indicators or expert domain knowledge to better capture complex market conditions. Automated techniques, such as genetic algorithms or neural network-based rule extraction, could also be explored to optimize and dynamically adjust the rule base in response to evolving market trends.

Another promising direction is the investigation of alternative deep reinforcement learning algorithms. While this study employed the so called A2C algorithm, other methods such as Proximal Policy Optimization (PPO), Deep Deterministic Policy Gradient (DDPG), or Soft Actor-Critic (SAC) could be tested to evaluate their suitability and performance in the carbon emission allowance market. Comparative analyses of different algorithms could yield insights into their strengths and weaknesses under various market scenarios.

Additionally, applying the framework to other financial or environmental markets could validate its generalizability and adaptability. Future studies may also explore incorporating additional external factors, such as macroeconomic indicators or policy changes, into the model to further enhance decision-making in complex and dynamic environments.

This work demonstrates a significant step toward more efficient and adaptive algorithmic trading strategies in the carbon emission allowance market, contributing to the ongoing development of intelligent AI-based financial systems.

REFERENCES

- [1] A. Marzi, E. Marzi, and H. Marzi, "Achieving co2 emission targets for energy consumption at canadian manufacturing and beyond; using hybrid optimization model," *Proceedings of the International Joint Conference on Neural Networks*, 2013.
- [2] D. Castilho, M. R. Santos, R. Tinós, A. C. Carvalho, M. B. Paula, L. Ladeira, E. Guarnier, D. S. Filho, D. Y. Suiama, A. M. Edmur, and L. P. Alipio, "Feature selection using complex networks to support price trend forecast in energy markets," *Proceedings of the International Joint Conference on Neural Networks*, vol. 2023-June, 2023.
- [3] P. Hajek and V. Olej, "Hierarchical intuitionistic fuzzy system for bitcoin price forecasting," *IEEE International Conference on Fuzzy Systems*, 2023.
- [4] J. Wu, C. Wang, L. Xiong, and H. Sun, "Quantitative trading on stock market based on deep reinforcement learning," *Proceedings of the International Joint Conference on Neural Networks*, vol. 2019-July, 7 2019.
- [5] Y. Deng, F. Bao, Y. Kong, Z. Ren, and Q. Dai, "Deep direct reinforcement learning for financial signal representation and trading," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 28, pp. 653–664, 3 2017.
- [6] J. Wang, L. Li, and jiangfeng zhang, "Deep reinforcement learning for energy trading and load scheduling in residential peer-to-peer energy trading market," *SSRN Electronic Journal*, 5 2022.
- [7] T. Chen and W. Su, "Local energy trading behavior modeling with deep reinforcement learning," *IEEE Access*, vol. 6, pp. 62806–62814, 2018.
- [8] S. A. Mahmud, A. M. Ibrahim, M. R. Khan, O. S. Baqalaql, and P. D. Wilde, "Improving the learning efficiency of robot path planning via fuzzified deep reinforcement learning," *2025 IEEE International Conference on Fuzzy Systems (FUZZ)*, pp. 1–8, 10 2025.
- [9] A. Thavaneswaran, Y. Liang, Z. Zhu, and R. K. Thulasiram, "Novel data-driven fuzzy algorithmic volatility forecasting models with applications to algorithmic trading," *IEEE International Conference on Fuzzy Systems*, vol. 2020-July, 7 2020.
- [10] N. Naik, P. Jenkins, N. Savage, L. Yang, K. Naik, and J. Song, "Embedding fuzzy rules with yara rules for performance optimisation of malware analysis," *IEEE International Conference on Fuzzy Systems*, vol. 2020-July, 7 2020.
- [11] S. A. Hosseini, F. Grimaccia, A. Niccolai, and S. Trimarchi, "Pattern-based feature extraction for improved deep learning in financial time series classification," *IEEE Access*, vol. 13, pp. 113343–113355, 2025.
- [12] A. S. Saud and S. Shakya, "Technical indicator empowered intelligent strategies to predict stock trading signals," *Journal of Open Innovation: Technology, Market, and Complexity*, vol. 10, p. 100398, 12 2024.
- [13] J. Lu, J. Yang, S. Li, Y. Li, W. Jiang, J. Dai, and J. Hu, "A2c-drl: Dynamic scheduling for stochastic edge-cloud environments using a2c and deep reinforcement learning," *IEEE Internet of Things Journal*, vol. 11, pp. 16915–16927, 5 2024.
- [14] O. Kosheleva, V. Kreinovich, and T. N. Nguyen, "Why triangular membership functions are successfully used in f-transform applications: A global explanation to supplement the existing local ones," *Axioms* 2019, Vol. 8, Page 95, vol. 8, p. 95, 8 2019.
- [15] R. M. Anderson, S. W. Bianchi, and L. R. Goldberg, "Determinants of levered portfolio performance," *SSRN Electronic Journal*, 4 2014.
- [16] W. F. Sharpe, "The sharpe ratio," *The Journal of Portfolio Management*, vol. 21, pp. 49–58, 10 1994.
- [17] V. Mohan, J. G. Singh, and W. Ongsakul, "Sortino ratio based portfolio optimization considering evs and renewable energy in microgrid power market," *IEEE Transactions on Sustainable Energy*, vol. 8, pp. 219–229, 1 2017.
- [18] M. Kapsos, N. Christofides, and B. Rustem, "Worst-case robust omega ratio," *European Journal of Operational Research*, vol. 234, pp. 499–507, 4 2014.
- [19] S. Malekpour and B. R. Barmish, "A drawdown formula for stock trading via linear feedback in a market governed by brownian motion," *2013 European Control Conference, ECC 2013*, pp. 87–92, 2013.
- [20] Y. E. Arisoy, A. Salih, and L. Akdeniz, "Is volatility risk priced in the securities market? evidence from sp 500 index options," *Journal of Futures Markets*, vol. 27, pp. 617–642, 7 2007.