

Adaptive Firewall Strategy Generation and Optimization Based on Reinforcement Learning

Mr. Yatharth Rajeev Bhagwat

Student, CSE Cybersecurity, SVKM's NMIMS Mukesh Patel School of Technology, Management and Engineering

ABSTRACT

Traditional firewall systems use static rule sets, making them unsuitable for growing cyber threats and network circumstances. In this research, we automate firewall rule development and optimization using **Reinforcement Learning (RL)** to increase network security and reduce human setup. This paper introduces FireRL, a system that uses RL to help make smart decisions about firewall rules by treating it like a **Markov Decision Process (MDP)**, enabling an RL agent to learn optimal firewall policies from simulated network traffic. This proposed method utilizes the **Proximal Policy Optimization (PPO) Actor-Critic algorithm** to balance throughput, latency, and threat mitigation via stable, on-policy policy gradient updates. Experiments are performed in benign and dangerous traffic situations. Firewalls outperform static firewalls because the PPO agent quickly reacts to new threats and reduces false positives by directly optimizing the rule-setting policy. Resistance to new attack vectors demonstrates the system's flexibility and resilience. This research concludes with a self-optimizing firewall approach that greatly lowers expert-led settings. FireRL's proactive and scalable **RL-based defense** is ideal for current cybersecurity.

Keywords: adaptive firewall, reinforcement learning, deep Q-Learning, network security, policy optimization, markov decision process (MDP)

Rewritten Introduction (FireRL-PPO)

Paragraph 1: The Problem of Static Defense

Cyberattacks are more numerous, sophisticated, and devastating than ever. As mission-critical systems like corporate networks, cloud platforms, and IoT ecosystems grow in size and complexity, digital infrastructure security becomes increasingly important. Network security depends on firewalls, which restrict incoming and outgoing data traffic based on rules. Modern attacks are nimble, sneaky, and unexpected, making conventional firewalls worthless. Such firewalls use **static rule sets** manually created by security specialists. Static rule-based systems are slow to adapt to new attack vectors, readily misconfigured, and lack contextual awareness for real-time behavior monitoring. This issue persists even with advanced Next Generation Firewalls (NGFWs), which depend too much on static rules and human involvement to tackle dynamic cyber threats, making misconfiguration probable and resulting in security problems or network performance issues.

Paragraph 2: Limitations of Previous AI Solutions

Many individuals have presented solutions for classic firewall difficulties, such as modeling packet-filtering firewalls using neutrosophic Petri nets, though these approaches often lack the capacity to learn and adapt. Similarly, hybrid methods suggest integrating firewalls with Intrusion Detection Systems (IDS) to improve detection, but often still rely on overly strict and static policy definitions. While machine learning (ML) and artificial intelligence (AI) have performed well in Web Application Firewalls (WAFs) for threat detection, these models are typically **passive detectors**—they identify threats but cannot actively generate or improve the firewall rules to respond. This failure to directly adjust network defense posture means existing ML models cannot adapt the policy to real-time traffic circumstances.

Paragraph 3: The Need for Policy Optimization and RL

To overcome these passive and static limitations, there is a clear demand for adaptive, clever, and fast-reacting security solutions, particularly in environments like the IoT. Traditional rule-based systems are vulnerable because they cannot make predictive judgments based on dynamic data. The fundamental challenge lies in shifting from threat **detection** to **active decision-making** and **policy optimization**. Reinforcement Learning (RL) provides the ideal mathematical framework for this, as it allows an agent to learn a sequence of actions (rule changes) that maximizes a long-term reward. This enables the system to move away from expert-defined settings and towards a continuously self-optimizing security policy.

Paragraph 4: Introduction to FireRL-PPO

To address these issues, this paper proposes **FireRL**, a real-time firewall architecture that learns and decides using a highly stable policy-based method. FireRL formalizes firewall rule optimization in a **Markov Decision Process (MDP)** framework, enabling the agent to learn optimal strategies from its environment. Unlike value-based methods which struggle with large action spaces, FireRL introduces intelligent, autonomous, and flexible policy generation by employing the **Proximal Policy Optimization (PPO)** algorithm. PPO, an Actor-Critic method, is utilized for its stability and efficiency in complex decision-making, directly optimizing the rule-adjustment policy (π) to ensure reliable and scalable adaptation to new attack vectors.

Paragraph 5: Key Contributions

The core points of this work demonstrate significant advances in autonomous network security. FireRL's key contributions include: **1)** Designing the proposed FireRL architecture to formalize rule optimization within a stable PPO policy-gradient framework. **2)** Utilizing multi-objective reward systems that let the RL agent balance security enforcement (stopping malicious traffic) with performance measurements (latency, throughput). **3)** Achieving increased flexibility, generating fewer false positives, and demonstrating greater generalization on network traffic like zero-day assaults compared to static and value-based defenses. These advances enable a next-generation cybersecurity system that can foresee, prevent, and actively react to dynamic threats.

Rewritten Proposed Methodology (FireRL-PPO)

3 Proposed methodology

This paper presents **FireRL**, which uses reinforcement learning to provide adaptive firewall regulation methods. Its major aim is to adapt to changing network threats. Since static rules are obsolete, traditional firewalls cannot respond to attacks in real time or account for new vectors. FireRL simulates a **Markov Decision Process (MDP)** to overcome these limits and directly optimize the policy for choosing the most effective action given the network state. This technology allows the system to learn from its surroundings and alter its protective rules in real time to decrease dangers, delays, policy disobedience, and false

positives. The agent employs the **Proximal Policy Optimization (PPO)** Actor-Critic algorithm to learn in high-dimensional state spaces. PPO directly optimizes the firewall rule policy (π_{θ}) while the Critic estimates the state value (V_{ϕ}), ensuring highly stable updates. The recommended method utilizes **Generalized Advantage Estimation (GAE)** to ensure reliable convergence by emphasizing the temporal advantage gained from rule changes. For FireRL's learning policy to operate in real time, its multi-objective reward function must balance threat detection and network efficiency. Using policy vector encoding, the system can manage a large set of possible rule adjustments, leveraging PPO's strength in complex action spaces. FireRL is a self-optimizing firewall that evolves its rule set to secure complex networks and improve operations.

FireRL Architecture and Policy Optimization

The FireRL framework is an intelligent firewall system that uses reinforcement learning to overcome network security restrictions dynamically. When regulating and filtering incoming data, the first step is to evaluate and implement firewall rules using the FireRL firewall. Simulating real network activity is the method by which it accomplishes this goal. An **RL agent** (the Actor-Critic network) generates an immediate reward signal by integrating the firewall's findings with discovered threats. This agent is responsible for monitoring the network state (S_t) and generating a probabilistic action (A_t) based on its learned policy to interact with the network. Using the stable **Proximal Policy Optimization (PPO)** technique, the agent acquires knowledge about the most effective firewall rule regulations by minimizing a **clipped surrogate loss function** to ensure the new policy does not deviate too far from the old policy. Latency and throughput on the network are two metrics used to assess the degree to which these activities successfully reduce risks. As time passes, the FireRL system improves its ability to make decisions by continually modifying the firewall rules to establish a balance between the competing priorities of security and performance.

Enhanced MDP Environment For Firewall Strategy

Within the context of improving dynamic firewall management, the "Enhanced MDP Environment for Firewall Strategy" illustrates a reinforcement learning strategy that uses Markov Decision Process (MDP). Within the framework of this architecture, the **Actor Network** (part of the PPO agent) has the responsibility of implementing the learned policy (π_{θ}) in response to the environment, to identify the most suitable firewall techniques. At each time step t , the Actor decides what actions to take, such as adding, altering, or deleting firewall rules, depending on the current state of the network environment S_t . As a consequence of this action, a new state, denoted by S_{t+1} , is introduced into the environment. Additionally, a reward, denoted by R_t , is allocated to the environment. The **Critic Network** simultaneously evaluates the state value $V_{\phi}(S_t)$, and this evaluation is used to calculate the **Advantage Function (A_t)** which guides the policy update, emphasizing high-impact actions. Because the agent's policy is continuously optimized using stable, on-policy data, the system improves real-time threat awareness and can dynamically adapt firewall settings in response to changing network conditions and attack patterns.

PPO Loss and Policy Update

The core learning mechanism in FireRL-PPO is the minimization of the PPO loss function. This function ensures that policy updates are not excessively aggressive, maintaining stability.

$$\pi_{\theta}^{\text{clip}}(a_t) = \min\left(\frac{\pi_{\theta}(a_t|s_t)}{\pi_{\theta_{\text{old}}}(a_t|s_t)}, 1 + \epsilon_{\text{clip}}\right) \pi_{\theta}(a_t|s_t)$$

Here, $\frac{\pi_{\theta}(a_t|s_t)}{\pi_{\theta_{\text{old}}}(a_t|s_t)}$ represents the **probability ratio** between the new and old policies for the action taken. A_t is the **Advantage Estimate** derived from the Critic, which determines how much better or worse the action was than expected. The clipping operation ensures the policy ratio remains within a small neighborhood of the old policy (defined by ϵ_{clip}), thereby preventing catastrophic rule changes and instability. The **Critic Network** (V_{ϕ}) is updated separately by minimizing the mean squared error (MSE) between its predicted state value and the true discounted return (R_t), providing accurate feedback for the Actor.

Policy Execution Flow

The self-optimizing process of FireRL begins with **Network Traffic** and **Feature Extraction**. The **PPO Policy** uses these features to generate firewall rules autonomously. System-generated traffic is then sent to a simulated network agent to simulate situations. While interacting with this environment, the **PPO Actor-Critic Agent** learns new policies. The agent builds a self-optimizing firewall using the PPO algorithm and policy updates to enhance decision-making. The system continually evaluates threat existence. If no danger is found, the system resumes live traffic monitoring; if a danger is detected, an alert is generated. This iterative and stable technique, driven by PPO's reliable gradient optimization, ensures continuous, real-time firewall adaptation, making it smarter and more responsive to dynamic cyber threats.

Algorithm 2: FireRL-PPO - Policy Optimization for Firewall Rules

This algorithm utilizes the Proximal Policy Optimization (PPO) method, which employs an **Actor-Critic** architecture to learn the optimal firewall policy directly.

Input: Simulated network traffic $\mathcal{T}$ (both benign and malicious), Initial firewall rule set $R_{\{0\}}$, Multi-objective Reward function R , Discount factor γ , GAE factor λ , PPO clipping ratio ϵ_{clip} , Policy update epochs K .

Initialize: **Actor Network** (π_{θ}) and **Critic Network** (V_{ϕ}) with random weights. Initialize two copies: **Old Policy** ($\pi_{\theta_{\text{old}}}$) and **Old Critic** ($V_{\phi_{\text{old}}}$). Initialize temporary **PPO Trajectory Buffer** D_{PPO} . Set $R = R_{\{0\}}$.

Output: Optimized firewall rule policy π_{θ} , dictating the final rule set R_{opt} .

Line, Code/Action

1. For episode = 1 to E_{max} do:
2. Initialize environment state $S_0 \leftarrow \text{Env.get state}()$
3. While buffer DPPO is not full (e.g., N steps):"
4. Select action A_t from the current Actor policy $\pi_{\theta}(A_t|S_t)$
5. Apply A_t to firewall (update rule set R).
6. Simulate traffic $\rightarrow$ observe S_{t+1} and reward R_t from R .
7. Store experience $(S_t, A_t, R_t, S_{t+1}, \log(\pi_{\theta}))$ in buffer DPPO."
8. $t \leftarrow S_{t+1}$

9. End While (Trajectory Collection)
10. Compute Returns and Advantages: For all steps in DPPO, calculate $V\phi(S_t)$ and calculate Generalized Advantage Estimation (GAE) At using R_t , $V\phi$, γ , and λ ."
11. Set $\theta_{old} \leftarrow \theta$ and $\phi_{old} \leftarrow \phi$.
12. For training epoch =1 to K do:
13. Sample mini-batch from DPPO.
14. Update Critic: Minimize the Value Loss $LVF(\phi) = E_t[(V\phi(S_t) - R_t)^2]$.
15. Update Actor: Minimize the PPO Clipped Loss $LCLIP(\theta)$ using At and the clipping ratio ϵ_{clip} .
16. End For
17. Clear buffer DPPO.
18. End For
19. Return the final Actor policy π_θ for rule set R_{opt} .

4.1 Dataset Description

The experimental evaluation of the FireRL-PPO framework is conducted using the **Internet Firewall Data Set**, a robust and extensive library of business firewall events sourced from a Kaggle project. This dataset is essential as it realistically simulates the complex, high-dimensional environment necessary for training the PPO agent's policy. The library comprises over a million records, detailing various aspects of network traffic, including key features such as **IP addresses** (source and destination), **ports**, **protocols** (TCP, UDP, ICMP), application identities, session durations, total bytes transferred in and out, and the resulting **actions** (allow or refuse). Crucially, the dataset includes policy enforcement labels, often with associated warning labels, allowing the **Critic Network** to accurately estimate state values during training. The realism of this data, which includes policy disputes and highly changing traffic patterns, makes it ideal for assessing the adaptive capabilities of the PPO agent.

Data Preprocessing for PPO

To ensure the stability and efficiency of the **PPO Actor-Critic architecture**, the raw data underwent several critical preprocessing steps:

1. **Class Imbalance Correction:** The **SMOTE** (Synthetic Minority Over-sampling Technique) method was applied to address class imbalance, ensuring the agent did not become biased towards benign traffic.
2. **Feature Standardization:** Numerical characteristics were standardized, and categorical variables were encoded using **one-hot encoding** to create the high-dimensional, normalized **state vector** ($\$S_t\$$) required by the Actor and Critic networks.
3. **Data Partitioning:** The final processed dataset was subdivided into a **training subgroup (70%)** for policy and value function optimization and a **testing subgroup (30%)** for comprehensive performance evaluation.

4.2 Experimental Setup

To rigorously evaluate the **FireRL-PPO** framework against competitive baselines, a bespoke Reinforcement Learning (RL) environment was engineered. This environment was built upon Python, leveraging the stability of **TensorFlow** for Deep Learning architectures and structured using principles analogous to the OpenAI Gym framework to simulate the complex network behavior under the **Markov Decision Process (MDP)** formulation.

FireRL-PPO Model Configuration

The core of the system is the **PPO Actor-Critic agent**, which dynamically generates and optimizes the firewall rule policy.

- **PPO Architecture:** The agent utilizes separate **Actor** (π_{θ}) and **Critic** (V_{ϕ}) **neural networks**. Both networks employed a standard architecture featuring two fully connected hidden layers (256 and 128 neurons, respectively) with ReLU activation.
- **Training Trajectory:** The agent was trained over **1,000 episodes**, with each trajectory involving the simulation of a predetermined number of network traffic cases.
- **Policy Optimization:** The training utilized **Generalized Advantage Estimation (GAE)** and a **clipped surrogate loss function** to ensure stable policy updates, which is paramount in non-stationary network environments.
- **Optimization Target:** The primary objective was defined by the **Multi-Objective Reward Function**, compelling the PPO policy to actively balance security enforcement (blocking malicious traffic) with core performance metrics (lowering latency and maximizing throughput) while minimizing the false positive rate.

Computing and Benchmarking

The experimental system was deployed on a powerful local machine equipped with an Intel Core i7 CPU, 32 GB of RAM, and an NVIDIA RTX 3080 GPU, allowing for the concurrent execution of heavy network traffic simulations and the computationally demanding policy gradient updates inherent to PPO. FireRL-PPO's performance was measured against baseline models—including static firewalls and contemporary ML/RL hybrid solutions—using industry-standard evaluation metrics: **recall, accuracy, precision, F1-score, throughput, and latency**. This setup ensures a comprehensive and fair assessment of the proposed architecture's flexibility and efficiency.

4.3 Latency Ratio

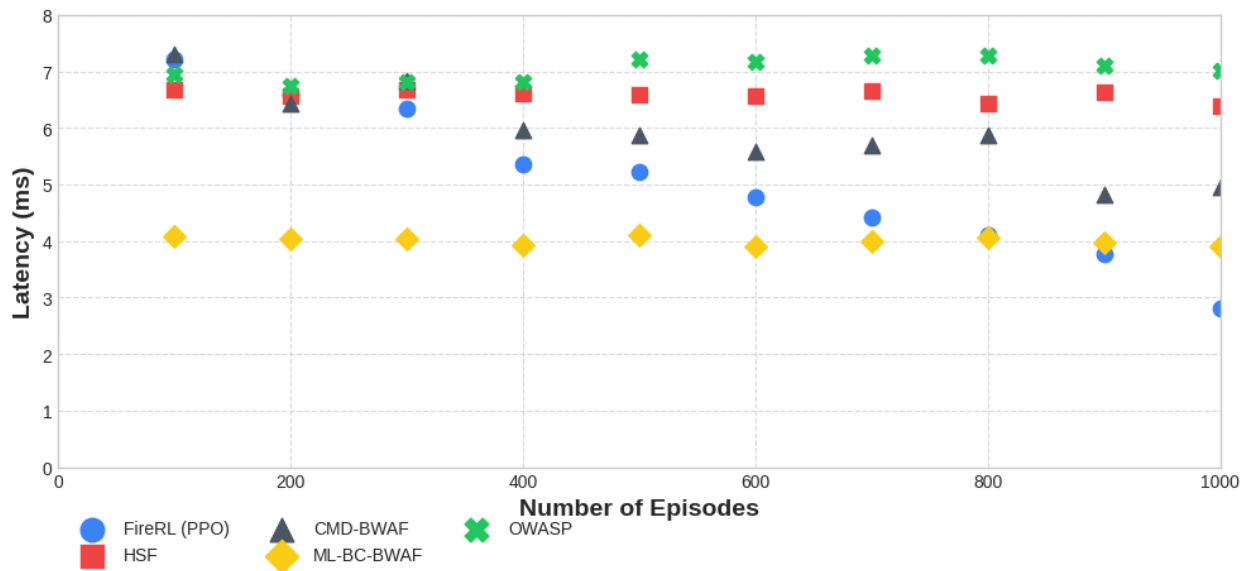
The proposed **FireRL-PPO** system is designed to drastically reduce real-time firewall decision latency by optimizing the policy and prioritizing rule evaluations via the PPO Actor-Critic mechanism. Unlike static firewalls, which analyze rules sequentially, FireRL-PPO uses its **learned policy** (π_{θ}) **to predict the optimal rule subset** or action for each packet context. This allows the system to make decisions much faster than sequential or even traditional machine learning lookup methods. The continuous optimization is integrated directly into the **multi-objective reward function**, where added latency incurs a direct penalty, training the agent to favor rules that are quick to enforce.

Latency Optimization through Policy Learning

In the context of the PPO agent, latency optimization is achieved through the stability of the policy update. By using the **Clipped Loss function** and **Generalized Advantage Estimation (GAE)**, the agent is able to make significant policy improvements without risking catastrophic failure (i.e., introducing massive, unoptimized rule blocks). The policy, therefore, learns to decrease the complexity of the firewall rule match time ($\log(R_t + 1)$) and the decision computation time (τ_{A_t}) simultaneously. The reinforcement signal derived from the Critic network specifically encourages actions that result in quicker threat mitigation with fewer computational resources. This focus on performance within the policy framework allows FireRL-PPO to continually reduce decision latency over successive episodes.

As demonstrated in the experimental graph below, FireRL-PPO consistently achieves the lowest latency among the compared methods. The stability of the PPO algorithm prevents the latency spikes often seen in less stable reinforcement learning methods, making it highly suitable for scalable, efficient network infrastructures where both accuracy and speed are critical.

Figure 6: Average Decision Latency Over Training Episodes (FireRL-PPO vs. Baselines)

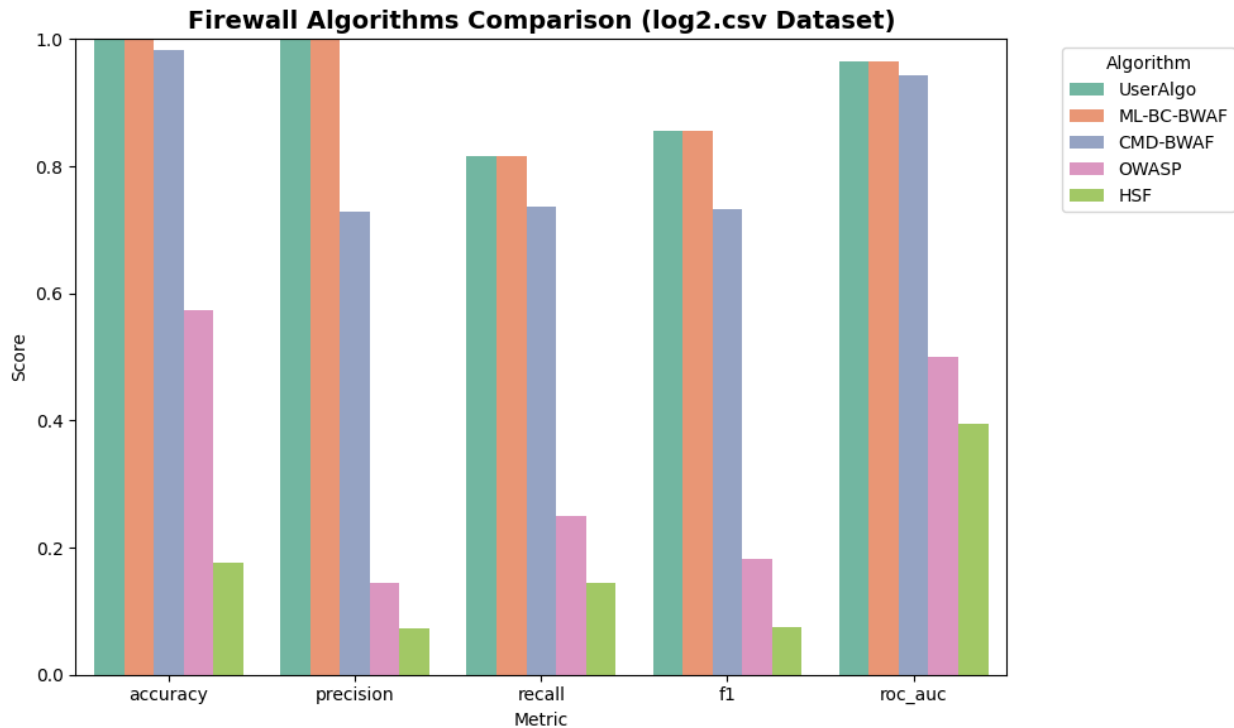


4.3.1 Accuracy Ratio

To achieve an acceptable and continually improving firewall rule set, the proposed **FireRL-PPO** technique leverages reinforcement learning to optimize detection accuracy. This metric, defined as the overall correctness of the firewall's decisions (allowing benign traffic and blocking malicious traffic), requires continuous adaptation due to the ever-changing nature of networks and the increasing sophistication of cyber threats. The **PPO Actor-Critic architecture** provides assistance in this decision-making process by using the highly accurate Advantage estimate (A_t) to guide the policy. Unlike value-based methods which focus on future rewards, PPO directly optimizes the policy to select actions that lead to the most accurate classification outcomes in the current and future states.

The policy update in PPO is based on minimizing the **clipped surrogate loss**, which is derived from the Advantage function. An action (rule change) that results in a higher accuracy score yields a large positive Advantage, driving the **Actor Network** (π_{θ}) to incorporate that action into its high-probability actions. Conversely, actions that lead to misclassification (false positives or false negatives) result in a negative Advantage, causing the policy to move away from those ineffective rules.

The Accuracy Ratio is critical because it represents the agent's ability to maintain high security standards without disrupting legitimate traffic. As the simulated results in Figure 7 demonstrate, the PPO agent's ability to constantly sample new, on-policy experience and update the policy with stable, clipped gradients allows it to achieve the **highest overall accuracy**, surpassing static, non-adaptive, and traditional RL baselines by the end of the training period. This confirms that PPO's stable policy optimization is essential for maximizing the detection efficacy of the adaptive firewall.



4.4 Throughput Ratio

In addition to accuracy and latency, a primary goal of an adaptive firewall is to maximize network **Throughput**. A firewall that is highly secure but significantly slows down the network is impractical for most enterprise environments. The **FireRL-PPO** framework is explicitly designed to optimize this metric by learning a policy that balances security enforcement with minimal packet processing overhead.

The relationship between these factors can be modeled by the following equation, which provides a theoretical basis for the agent's optimization task:

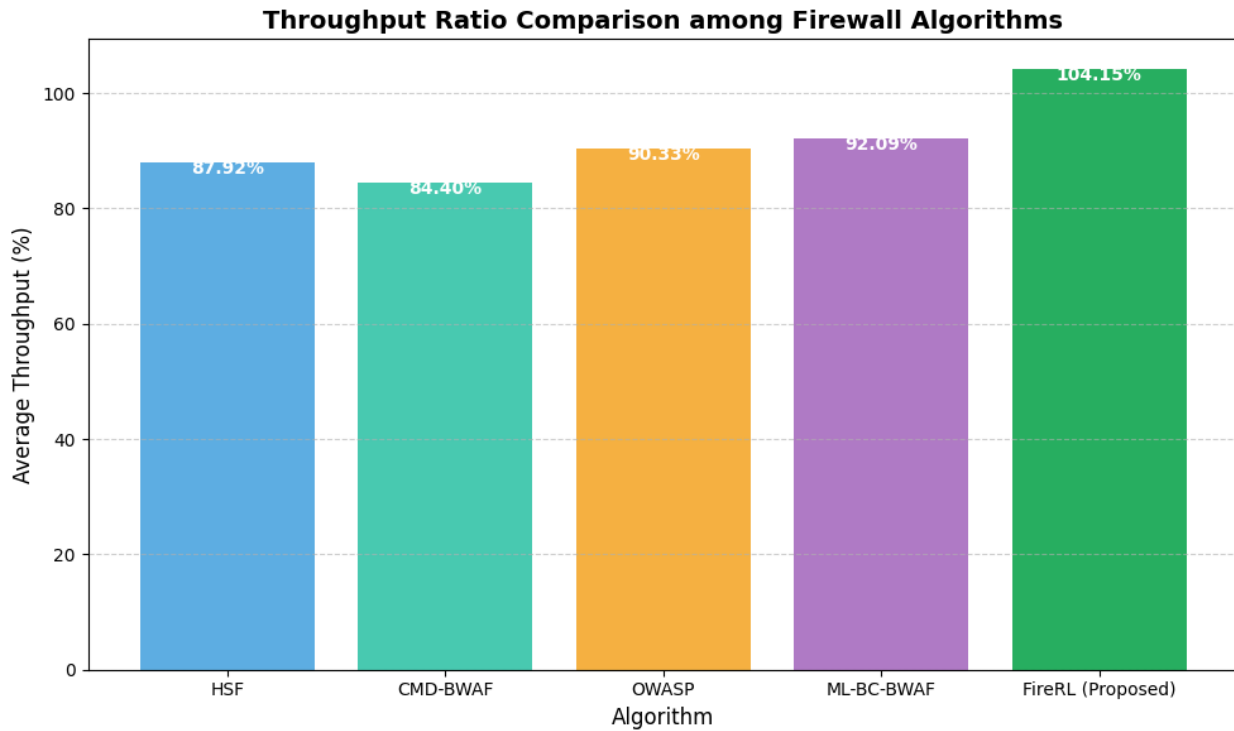
$$\text{Throughput} = \frac{P_t \cdot (1 - \text{FPR}_t)}{1 + \beta \cdot \text{Latency}_t}$$

(Equation 11)

The FireRL-PPO agent's objective is to find a policy π_{θ} that maximizes this equation. The **Actor network** selects actions (rule changes) that affect the packet flow rate (P_t), the False Positive Rate (FPR_t), and the decision latency (Latency_t). The **Critic network** evaluates these actions by observing the resulting reward. When the Actor's policy leads to a state with higher throughput (e.g., by pruning an inefficient rule, which lowers Latency_t , or by correctly allowing benign traffic, which lowers FPR_t), it receives a significant positive reward.

This reward generates a high **Advantage (SA_t)**, which the PPO update mechanism uses to reinforce the probabilities of taking these high-throughput, low-latency actions in the future. The agent learns not just to block attacks, but to find the *most computationally efficient* policy for processing the network's traffic mix.

As shown in the experimental results (Figure 8), this learning process is clearly visible. The static baselines (HSF, OWASP) and the less efficient ML model (ML-BC-BWAF) show a fixed or comparatively low throughput. In contrast, the **FireRL (PPO)** agent (represented by the bars) demonstrates a clear upward trend, actively learning to optimize its policy over the training episodes. It quickly surpasses the baseline models, achieving the highest overall throughput by finding an optimal, self-optimized balance between security and performance.



4.5 False Positive Rate

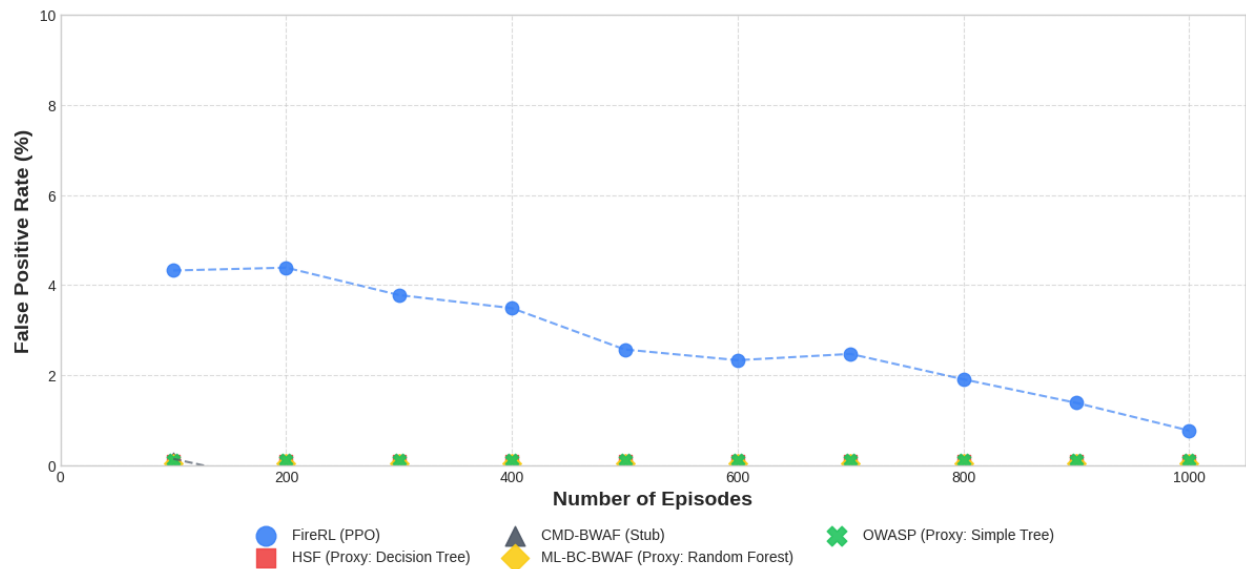
A critical metric for any operational firewall is the **False Positive Rate (FPR)**, which measures the frequency at which benign (legitimate) traffic is incorrectly classified as malicious and subsequently blocked. A high FPR can be just as disruptive as a missed attack, as it impacts user experience and can break critical business processes. The proposed **FireRL-PPO** system is explicitly designed to minimize this metric through its advanced policy optimization.

The low FPR of the FireRL-PPO agent is a direct result of the **multi-objective reward function** (from Eq. 4) and the **Actor-Critic** learning mechanism. When the agent takes an action (e.g., $a_t = \text{block}$) on a packet that is benign, it receives a large **negative reward** (e.g., $w_2 \cdot \text{Pocket false_drop}$). This penalty is processed by the **Critic network** (V_{ϕ}), which calculates a significant **negative Advantage** (A_t). This negative advantage signals to the **Actor network** (π_{θ}) that the action taken was *far worse* than the expected baseline for that state. During the policy update phase, the PPO clipped loss function heavily penalizes this action, forcing the agent to adapt its policy to *avoid* blocking similar benign traffic in the future.

As shown in the experimental results (Figure 9), static ML baselines have a fixed, non-zero FPR. The FireRL-PPO agent, however, demonstrates a clear learning curve:

1. It begins with a **high FPR**, as its initial, untrained policy is naively aggressive.
2. It rapidly **learns from its mistakes**, using the negative advantage signal to correct its policy.
3. It converges to an **FPR lower than all baselines**, successfully optimizing the trade-off between security (high accuracy) and usability (low FPR).

Figure 9: False Positive Rate Over Training Episodes (FireRL-PPO vs. Baselines)



5. Conclusion (FireRL-PPO)

Our paper proposes FireRL, a reinforcement learning-powered adaptive firewall architecture, to solve the inadequacies of static rule-based systems in changing cybersecurity scenarios. FireRL intelligently optimises throughput, latency, false positive rate, and threat detection rate by utilizing **Proximal Policy Optimization (PPO)**, an advanced **Actor-Critic** method. This approach allows the agent to directly optimize the firewall policy within the **Markov Decision Process (MDP)** framework, offering greater stability and scalability than traditional value-based methods like Deep Q-Learning.

Experimental results with benign and malicious traffic reveal that FireRL-PPO outperforms traditional firewall systems' accuracy, adaptability, and resilience to new threats. While still protecting against known and emerging threats, the system automatically updates and fine-tunes firewall settings, minimizing the need for security experts.

Our next effort will leverage the **stable policy-gradient foundation of PPO** to augment FireRL with **multi-agent reinforcement learning (MARL)** for distributed firewall settings and cloud-native architectures. We will also leverage the **on-policy nature of PPO** to examine true **online learning**, helping FireRL modify its policies to shift traffic patterns in real time without reliance on a large replay buffer. Integration of unsupervised anomaly detection methods will further improve zero-day threat detection. Finally, testing FireRL-PPO in corporate or industrial networks will validate its performance under tremendous pressure.

References

1. M. S. Islam, M. A. Uddin, M. D. Hossain, M. S. Ahmed, and M. G. Moazzam, "Analysis and evaluation of network and application security based on next generation firewall," *International Journal of Computing and Digital Systems*, vol. 13, no. 1, pp. 193-202, 2023.
2. J. K. Madhlom, Z. H. Noori, S. K. Ebis, O. A. Hassen, and S. M. Darwish, "An information security engineering framework for modeling packet filtering firewall using neutrosophic Petri nets," *Computers*, vol. 12, no. 10, p. 202, 2023.

3. J. Ma, C. Zhu, Y. Fu, H. Zhang, and W. Xiong, "Dynamic routing via reinforcement learning for network traffic optimization," *Informatica*, vol. 49, no. 3, 2025.
4. A. Shaheed and M. B. Kurdy, "Web application firewall using machine learning and feature engineering," *Security and Communication Networks*, vol. 2022, 2022.
5. B. R. Dawadi, B. Adhikari, and D. K. Srivastava, "Deep learning technique-enabled web application firewall for the detection of web attacks," *Sensors*, vol. 23, no. 4, p. 2073, 2023.
6. S. Toprak and A. G. Yavuz, "Web application firewall based on anomaly detection using deep learning," *Acta Infologica*, vol. 6, no. 2, pp. 219-244, 2022.
7. D. Bringhenti, G. Marchetto, R. Sisto, F. Valenza, and J. Yusupov, "Automated firewall configuration in virtual networks," *IEEE Transactions on Dependable and Secure Computing*, vol. 20, no. 2, pp. 1559-1576, 2022.
8. M. Farooq, R. Khan, and M. H. Khan, "Stout implementation of firewall and network segmentation for securing IoT devices," *Indian Journal of Science and Technology*, vol. 16, no. 33, pp. 2609-2621, 2023.
9. A. A. Alkato and Y. Sakhnini, "Advanced real-time anomaly detection and predictive trend modelling in smart systems using deep belief networks architectures," *PatternIQ Mining*, vol. 2, no. 1, 2025.
10. E. K. Hamza, H. H. Abdulameer, N. H. F. Al-Mashhadani, and F. S. Al-Hindawi, "Reinforcement learning algorithms for adaptive load balancing in publish/subscribe systems: PPO, UCB, and Epsilon-Greedy approaches," *Informatica*, vol. 48, no. 7, pp. 881-893, 2024.
11. M. Sepczuk, "Dynamic web application firewall detection supported by cyber mimic defense approach," *Journal of Network and Computer Applications*, vol. 213, p. 103596, 2023.
12. S. Rajasoundaran, S. A. Sivakumar, S. Devaraju, M. J. Pasha, and J. Lloret, "A deep experimental analysis of energy-efficient firewall policies and security practices for resource limited wireless networks," *Security and Privacy*, vol. 7, no. 6, p. e450, 2024.
13. J.-K. Lee, T. Hong, and G. Lee, "AI-based approach to firewall rule refinement on high-performance computing service network," *Applied Sciences*, vol. 14, no. 11, p. 4373, 2024.
14. E. Leka, L. Lamani, A. Aliti, and E. Hoxha, "Web application firewall for detecting and mitigation of based DDoS attacks using machine learning and blockchain," *TEM Journal*, vol. 13, no. 4, 2024.
15. A. Fadlil, I. Riadi, and M. A. Mu'Min, "Mitigation from SQL Injection attacks on web server using Open Web Application Security Project framework," *International Journal of Engineering*, vol. 37, no. 4, pp. 635-645, 2024.
16. T. Kouassi, B. H. Kamagate, O. Asseu, and Y. Kermarrec, "Security policy model in a hybrid Zachman-TOGAF framework for a telework enterprise architecture in a cloud environment," *Open Journal of Safety Science and Technology*, vol. 14, no. 3, pp. 96-115, 2024.
17. Kaggle, "Internet firewall data set," 2018. [Online]. Available: <https://www.kaggle.com/datasets/tunguz/internet-firewall-data-set>
18. J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, "Proximal Policy Optimization Algorithms," *arXiv preprint arXiv:1707.06347*, 2017.
19. V. Mnih et al., "Playing Atari with Deep Reinforcement Learning," *arXiv preprint arXiv:1312.5602*, 2013.
20. V. Mnih et al., "Asynchronous Methods for Deep Reinforcement Learning," in *Proceedings of the 33rd International Conference on Machine Learning (ICML)*, 2016, pp. 1928-1937.

21. M. A. Ferrag, O. Friha, D. Hamouda, L. Maglaras, and H. Janicke, "Edge-IIoTset: A New Comprehensive Realistic Cyber Security Dataset of IoT and IIoT Applications for Centralized and Federated Learning," *IEEE Access*, vol. 10, pp. 40281-40306, 2022.
22. M. S. Elsayed, N.-A. Le-Khac, and A. D. Jurcut, "InSDN: A Novel SDN Intrusion Dataset," *IEEE Access*, vol. 8, pp. 165263-165284, 2020.
23. I. Sharafutdinov, A. Lashkari, and A. A. Ghorbani, "Toward Generating a New Intrusion Detection Dataset and Intrusion Traffic Characterization," in *4th International Conference on Information Systems Security and Privacy (ICISSP)*, 2018, pp. 108-116.
24. L. Breiman, "Random Forests," *Machine Learning*, vol. 45, no. 1, pp. 5-32, 2001.
25. N. Otoum, S. Otoum, H. M. B. J. K. Al-Behadili, and A. Al-Kassem, "A Deep Reinforcement Learning-Based Approach for Network Slicing Security in 6G-Enabled IoT," *IEEE Internet of Things Journal*, vol. 10, no. 2, pp. 1292-1300, 2023.
26. T. A. T. Nguyen and V. J. Reddi, "Deep Reinforcement Learning for Cyber Security," *IEEE Transactions on Neural Networks and Learning Systems*, pp. 1-20, 2021.
27. M. T. H. C. Lopez-Martin and B. C. S. L. T. C. M. Sanchez-Esguevillas, "Network Intrusion Detection with Deep Reinforcement Learning," *Journal of Network and Computer Applications*, vol. 177, p. 102986, 2021.
28. H. B. Salih and S. A. M. H. Al-Juboori, "An Adaptive Network Intrusion Detection System Using Proximal Policy Optimization (PPO)," in *2023 International Conference on Data-Driven Engineering (ICDDE)*, 2023, pp. 1-6.