

Aplicación de modelos de inteligencia artificial para la detección de malware en infraestructuras críticas de telecomunicaciones

Implementing artificial intelligence models for malware detection in critical telecommunications infrastructures

Ing. Elizabeth Molina Mena^{*}, Ing. Heidy Rodríguez Malvares²,
M.Sc. Henry Raúl González Brito³

Recibido: 11/2025 | Aceptado: 11/2025 | Publicado: 12/2025

Resumen

El avance acelerado de la digitalización y la interconexión global ha incrementado la exposición de las infraestructuras críticas de telecomunicaciones a diversas amenazas cibernéticas, entre las cuales el *malware* representa una de las más persistentes y dañinas. En este contexto, la aplicación de modelos de inteligencia artificial (IA) para la detección temprana de *software* malicioso se ha convertido en una estrategia clave para garantizar la resistencia, la disponibilidad y la continuidad de los servicios esenciales. Este artículo analiza los principales enfoques basados en aprendizaje automático y aprendizaje profundo empleados en la identificación de *malware*,

^{*} Dirección de Seguridad Informática, Universidad de las Ciencias Informáticas.
elizabethmm@uci.cu

² Departamento de Infraestructuras Tecnológicas, Facultad de Ciberseguridad, Universidad de las Ciencias Informáticas. heidyrm@uci.cu

³ Dirección de Seguridad Informática, Universidad de las Ciencias Informáticas.
henryraul@uci.cu

con énfasis en su implementación dentro de entornos de telecomunicaciones críticos.

Se revisan técnicas clásicas como Support Vector Machines (SVM) y Random Forest, así como arquitecturas avanzadas de redes neuronales, incluyendo Convolutional Neural Networks (CNN), Long Short-Term Memory (LSTM) y Graph Neural Networks (GNN). La metodología se centra en una revisión sistemática de la literatura reciente y un análisis comparativo de modelos según métricas de rendimiento, capacidad de generalización y consumo de recursos computacionales. Los resultados muestran que los modelos híbridos basados en aprendizaje profundo y técnicas de análisis de flujo de red logran tasas de detección superiores al 98 % con una reducción significativa de falsos positivos. Por último, se discuten los desafíos actuales y las perspectivas futuras para la integración de IA en la ciberdefensa de infraestructuras críticas, y se destaca la importancia de la explicabilidad, la detección de amenazas desconocidas y la colaboración entre operadores e instituciones de investigación.

Palabras clave: inteligencia artificial, detección de *malware*, infraestructuras críticas, telecomunicaciones, ciberseguridad

Abstract

The rapid advance of digitization and global interconnection has made critical telecommunications infrastructure more vulnerable to various cyber threats, among which, malware is one of the most persistent and damaging. In this context, implementing Artificial Intelligence (AI) models to detect malicious software early on has become a key strategy for ensuring the resilience, availability, and continuity of essential services. This article analyzes the primary machine learning and deep learning approaches used for malware identification, focusing on their implementation in critical telecommunications environments. Classical techniques such as Support Vector Machines (SVM) and Random Forest are reviewed, as well as advanced neural network architectures, including Convolutional Neural Networks (CNN), Long Short-Term Memory (LSTM), and Graph Neural Networks (GNN). The methodology focuses on a systematic review of recent literature and a comparative analysis

of models based on performance metrics, generalization capability, and computational resource consumption. Results indicate that hybrid models combining deep learning and network flow analysis techniques can achieve 98% rates detection with a significantly fewer false positive. Finally, current challenges and future perspectives for the integration of AI into the cyber defense of critical infrastructures are discussed, highlighting the importance of explainability, unknown threat detection, and collaboration between Telcos carriers and research institutions.

Keywords: *Artificial Intelligence, malware detection, critical infrastructures, telecommunications, cybersecurity*

Introducción

Las infraestructuras críticas de telecomunicaciones constituyen el núcleo de los sistemas de comunicación modernos, que facilitan los servicios esenciales para el funcionamiento de gobiernos, empresas y ciudadanos. Estas infraestructuras abarcan redes de transmisión, centros de datos, sistemas de conmutación y plataformas de gestión que garantizan la conectividad nacional e internacional. Sin embargo, su creciente complejidad y dependencia de *software* las convierten en un objetivo prioritario para los ciberatacantes (García-Teodoro et al., 2023).

El *malware* (término que engloba a virus, troyanos, *ransomware*, *spyware* y *rootkits*) ha evolucionado en sofisticación, empleando técnicas de evasión como ofuscación de código, polimorfismo y *living-off-the-land* (LOTL), esto dificulta su detección por mecanismos tradicionales basados en firmas (Alasmary et al., 2022). En respuesta a esta problemática, la inteligencia artificial (IA) se ha posicionado como una herramienta fundamental para la identificación automatizada de patrones anómalos y la detección de amenazas emergentes en tiempo real (Sgandurra & Lupu, 2020).

El aprendizaje automático (*Machine Learning*, ML) y el aprendizaje profundo (*Deep Learning*, DL) han demostrado un potencial notable en la clasificación de *malware* a partir de grandes volúmenes de datos, extrayendo características tanto estáticas (del código binario) como dinámicas (del comportamiento en ejecución). Estos modelos

aprenden representaciones complejas que identifican variantes nuevas o desconocidas, incluso en escenarios de ataques dirigidos contra infraestructuras críticas (Ucci, Aniello & Baldoni, 2019).

El presente artículo tiene como objetivo analizar la aplicación de modelos de inteligencia artificial en la detección de *malware* dentro de sectores estratégicos de telecomunicaciones. Se aborda el problema desde una perspectiva técnico-científica, y combina el estudio de algoritmos, técnicas de extracción de características, y el contexto operativo de las redes de telecomunicaciones. Además, se examinan los desafíos asociados a la interpretabilidad de los modelos, la gestión de grandes volúmenes de datos y la integración con sistemas de seguridad existentes como los *Security Information and Event Management* (SIEM) y los *Intrusion Detection Systems* (IDS).

Este trabajo se estructura de la siguiente manera: en la sección de Materiales y Métodos se describe la metodología empleada para la revisión sistemática y la clasificación de modelos de IA; en Resultados y Discusión se analizan los hallazgos principales y su aplicabilidad en entornos reales; y, en Conclusiones se resumen los aportes expuestos y se plantean futuras líneas de investigación orientadas al fortalecimiento de la ciberresiliencia en el sector de las telecomunicaciones.

Materiales y métodos

Para la elaboración de este estudio sobre la aplicación de modelos de inteligencia artificial en la detección de *malware* en infraestructuras críticas de telecomunicaciones, se adoptó una metodología de revisión sistemática y análisis comparativo de enfoques existentes, complementada con una caracterización de los modelos más representativos implementados en entornos reales de telecomunicaciones.

Fuentes de información y criterios de selección

La búsqueda bibliográfica se realizó en bases de datos científicas de alto impacto como IEEE Xplore, SpringerLink, ScienceDirect, ACM Digital Library y Scopus, utilizando combinaciones de palabras clave como «AI-based malware detection», «deep learning for cyber defense», «critical infrastructure protection», y «telecommu-

nication cybersecurity».

Se establecieron los siguientes criterios de inclusión:

- Publicaciones revisadas por pares entre 2018 y 2025.
- Estudios enfocados en detección o clasificación de *malware* mediante IA.
- Investigaciones aplicadas a infraestructuras críticas o redes de telecomunicaciones.
- Trabajos que reportaran métricas de rendimiento cuantitativas (precisión, *recall*, F1-score o tasa de falsos positivos).

Asimismo, se excluyeron estudios con conjuntos de datos no verificables o con modelos sin reproducibilidad documentada. En total, se revisaron 96 publicaciones científicas, de las cuales 38 cumplieron los criterios establecidos y fueron incluidas en el análisis principal.

Metodología de análisis

La metodología adoptada se estructuró en tres fases:

1. Revisión sistemática de literatura: se empleó el protocolo PRISMA (*Preferred Reporting Items for Systematic Reviews and Meta-Analyses*) para identificar y filtrar los estudios relevantes. La búsqueda inicial generó 327 resultados, que se redujeron tras la aplicación de los filtros de calidad, año y ámbito de aplicación.

2. Clasificación de modelos de IA: los modelos se agruparon según su naturaleza algorítmica en cuatro categorías principales:

- Modelos clásicos de ML: SVM, Decision Trees, Random Forest, Naïve Bayes, KNN.
- Modelos de DL: CNN, RNN, LSTM, Autoencoders, GNN.
- Modelos híbridos: combinaciones de ML y DL, como CNN+SVM o Autoencoder+Random Forest.
- Modelos de detección basados en flujo de red y análisis de comportamiento, especialmente en entornos 5G y SDN (*Software Defined Networking*).

3. Evaluación comparativa: se analizaron las métricas de rendimiento de los modelos, comparando precisión, sensibilidad y tiempo de inferencia. Se utilizó un esquema de ponderación basado en la media geométrica de las métricas de precisión y *recall*, para determinar la eficiencia relativa de cada modelo (Zhang et al., 2022).

Infraestructuras críticas y contexto de aplicación

En el ámbito de las telecomunicaciones, las infraestructuras críticas comprenden redes troncales de fibra óptica, sistemas de señalización SS7 y SIP, nodos de conmutación, estaciones base, *routers* de *backbone* y plataformas de virtualización de funciones de red (NFV). Estas infraestructuras soportan servicios esenciales como telefonía, datos, radiocomunicación y redes satelitales (ENISA, 2023).

Debido a su carácter estratégico y dependencia intersectorial, un ataque de *malware* puede generar efectos en cascada sobre servicios energéticos, financieros y gubernamentales (CISA, 2022). Por ello, la detección temprana de anomalías mediante IA se considera un componente esencial de la ciberresiliencia operacional (Shafiq et al., 2020).

Modelos de aprendizaje automático empleados

Los algoritmos clásicos de aprendizaje supervisado como *Support Vector Machine* (SVM) y *Random Forest* (RF) han mostrado un buen rendimiento en la detección de *malware* mediante la extracción de características estáticas (Permana et al., 2020).

SVM se destaca por su capacidad para separar clases con márgenes óptimos, mientras que RF ofrece robustez ante sobreajuste al combinar múltiples árboles de decisión.

Sin embargo, los modelos de aprendizaje profundo han superado estos resultados en escenarios de detección de comportamiento dinámico. Las *Convolutional Neural Networks* (CNN) han sido aplicadas con éxito al análisis de secuencias de *bytes* y a la identificación de patrones binarios de *malware* (Kolosnjaji et al., 2018), mientras que las *Long Short-Term Memory* (LSTM) destacan en la detección de comportamientos maliciosos en tiempo real en flujos de red (Pektaş & Acarman, 2020).

Por su parte, los *Autoencoders* y las *Graph Neural Networks* (GNN) son empleados para la detección de anomalías sin supervisión, aprenden representaciones de alta dimensión de tráfico normal y detectan desviaciones significativas (Wang et al., 2022).

Datos y entornos de experimentación

Las bases de datos más utilizadas en los estudios revisados fueron:

- EMBER Dataset (*Endgame Malware Benchmark for Research*).
- CICIDS2017 y CSE-CIC-IDS2018 (*datasets* de tráfico de red).
- VirusShare y Microsoft Malware Classification Challenge (BIG 2015) para análisis estático.
- NSL-KDD y TON_IoT *datasets*, relevantes en redes IoT y 5G.

Implementación y herramientas utilizadas

La revisión evidenció que los entornos más empleados para la experimentación son Python con TensorFlow, Keras, PyTorch y Scikit-learn, junto con plataformas de orquestación en la nube como Google Colab, AWS SageMaker y Azure ML. En investigaciones centradas en entornos críticos, los modelos fueron desplegados sobre redes SDN mediante mininet y OpenDaylight, con soporte para análisis en tiempo real (Al-Yaseen et al., 2020).

Limitaciones del estudio

Aunque el análisis ofrece una visión global, se identifican algunas limitaciones:

- Escasez de estudios centrados específicamente en redes troncales y sistemas de señalización.
- Falta de disponibilidad de *datasets* específicos de telecomunicaciones críticas.
- Dificultad para reproducir algunos resultados debido a la confidencialidad de los datos.
- Carencia de modelos explicables que permitan interpretar las decisiones de los algoritmos ante incidentes complejos.

Resultados y discusión

Los resultados de la revisión sistemática y del análisis comparativo de modelos de IA aplicados a la detección de *malware* en infraestructuras críticas de telecomunicaciones, evidencian la evolución sostenida de los métodos de detección hacia enfoques basados en aprendizaje profundo y detección contextual adaptativa. En esta sección se presentan los hallazgos principales, organizados según el tipo de modelo, su rendimiento, las métricas utilizadas y las implicaciones prácticas en el contexto de las redes de telecomunicaciones.

1. Rendimiento de los modelos de aprendizaje automático

Los modelos de aprendizaje automático clásico (SVM, Random Forest, Naïve Bayes, Decision Tree) siguen siendo competitivos en escenarios donde la detección de *malware* depende de características estáticas extraídas de archivos ejecutables o del análisis de cabeceras de tráfico. Según Permana et al. (2020), Random Forest alcanzó precisiones del 95,6 % en la clasificación de *malware*, mientras que SVM logró un rendimiento similar con tiempos de inferencia inferiores.

Sin embargo, su capacidad de detección disminuye ante muestras de *malware* polimórfico o metamórfico, así como en escenarios con tráfico cifrado o técnicas avanzadas de evasión. En el contexto de redes de telecomunicaciones críticas, donde los flujos son heterogéneos y de alta velocidad, los modelos clásicos tienden a generar una alta tasa de falsos positivos ($FPR > 6\%$), lo cual puede saturar los sistemas de alerta en centros de monitoreo (Cruz et al., 2021).

Por ello, las organizaciones del sector han comenzado a incorporar modelos híbridos, que combinan la capacidad interpretativa de los métodos clásicos con la potencia de generalización del aprendizaje profundo.

2. Eficacia de los modelos de aprendizaje profundo

El análisis reveló que los modelos basados en redes neuronales profundas (DNN, CNN, RNN, LSTM) ofrecen un rendimiento sustancialmente superior en la detección de *malware* dinámico y en la clasificación de patrones complejos de tráfico de red. En particular, las CNN han sido ampliamente utilizadas para representar el código binario del *malware* como una imagen matricial y detectar patrones maliciosos ocultos.

Kolosnjaji et al. (2018) demostraron que una CNN entrenada sobre imágenes binarias de ejecutables alcanzó una precisión del 98,4 % en la detección de familias de *malware* sin necesidad de desensamblar el código. Por su parte, LSTM y GRU resultan eficaces en la detección de comportamientos maliciosos persistentes al capturar relaciones temporales entre eventos en secuencias de tráfico (Pektaş & Acarman, 2020).

Los modelos Autoencoder destacan en la detección de anomalías

mediante reconstrucción de características. Al ser entrenados exclusivamente con tráfico legítimo, logran identificar desviaciones con tasas de precisión superiores al 97 %, lo que reduce de manera significativa los falsos positivos en entornos de alta disponibilidad (Wang et al., 2022).

Por otro lado, las GNN y los *Transformers* se posicionan como modelos emergentes de referencia, capaces de analizar dependencias estructurales entre nodos de red, identificar flujos coordinados de ataque y correlacionar alertas entre dominios distribuidos (Hussain et al., 2023).

La Tabla 1 resume las métricas promedio obtenidas de las principales investigaciones revisadas, comparando los modelos según precisión (*accuracy*), sensibilidad (*recall*), y tasa de falsos positivos (FPR).

Tabla 1. Rendimiento comparativo de modelos de IA en detección de *malware*

Tipo de modelo	Ejemplo de algoritmo	Precisión (%)	<i>Recall</i> (%)	FPR (%)
ML clásico	Random Forest	95.6	94.2	6.1
ML clásico	SVM	94.9	92.8	5.8
DL – CNN	CNN 2D en binarios	98.4	97.2	2.3
DL – LSTM	Secuencias de tráfico	97.6	96.9	3.1
DL – <i>Autoencoder</i>	Tráfico legítimo	97.2	95.8	2.6
GNN / <i>Transformer</i>	Topología de red	98.9	98.1	1.9

3. Aplicabilidad en infraestructuras de telecomunicaciones

En redes de telecomunicaciones, los flujos de datos se caracterizan por su alta dimensionalidad, diversidad de protocolos y baja latencia, lo que dificulta la detección de comportamientos anómalos en tiempo real. Los modelos de IA deben adaptarse a este contexto mediante técnicas de reducción de dimensionalidad y optimización de inferencia en el borde (*edge inference*).

Diversas investigaciones (Shafiq et al., 2020; Al-Yaseen et al., 2020) han implementado modelos CNN-LSTM integrados en entornos *Software Defined Networking* (SDN y han logrado detectar

ataques de tipo *botnet*, *ransomware* y *denial-of-service* en menos de 200 ms de tiempo de respuesta. En modelos mostraron una reducción del 35 % en la latencia de detección y una mejora del 40 % en la preparación respecto a métodos convencionales en escenarios experimentales que simularon infraestructuras de telecomunicaciones 5G.

Asimismo, la combinación de modelos Autoencoder con sistemas de detección basados en *flow analysis* permitió aislar comportamientos sospechosos en entornos de virtualización NFV (*Network Function Virtualization*), garantizando continuidad operativa y mitigación proactiva de amenazas (ENISA, 2023).

En el caso de las redes IoT vinculadas a telecomunicaciones, se evidenció que los modelos GNN son especialmente eficaces al representar la topología como un grafo dinámico, identificando relaciones entre dispositivos comprometidos y eventos coordinados (Zhang et al., 2022).

4. Ventajas y desafíos de la aplicación de IA

Entre las principales ventajas de los modelos de IA en la detección de *malware* destacan:

- Alta precisión en la identificación de amenazas conocidas y desconocidas.
- Adaptabilidad ante cambios en los patrones de ataque.
- Capacidad de aprendizaje continuo mediante actualización de pesos y *datasets*.
- Reducción de falsos positivos, optimizando los recursos de los centros de operaciones de seguridad (SOC).

Sin embargo, los desafíos son igualmente significativos:

- Explicabilidad limitada: la mayoría de los modelos de DL son cajas negras difíciles de interpretar (Goodman & Flaxman, 2017).
- Dependencia de datos etiquetados: la escasez de *datasets* específicos de telecomunicaciones limita el entrenamiento de modelos especializados.
- Vulnerabilidad ante ataques adversariales, que pueden engañar a los modelos de IA modificando levemente los patrones de entrada (Huang et al., 2021).

- Altos requerimientos computacionales, lo que dificulta la implementación en sistemas en tiempo real con recursos limitados.

5. Integración en sistemas de seguridad de telecomunicaciones

El despliegue de modelos de IA para detección de *malware* se ha integrado de manera progresiva en arquitecturas de seguridad existentes, como los *Security Information and Event Management* (SIEM) y los *Intrusion Detection Systems* (IDS) basados en IA. En entornos de telecomunicaciones críticas, esta integración requiere una orquestación inteligente capaz de correlacionar eventos provenientes de múltiples capas (física, lógica y de aplicación).

Sistemas como IBM QRadar, Splunk y Cisco SecureX han comenzado a incorporar módulos de detección basados en aprendizaje profundo para mejorar la correlación de alertas en redes 5G (CISA, 2022). Estos enfoques híbridos permiten la detección en tiempo real y la respuesta automatizada, lo que reduce en más del 60 % el tiempo medio de detección (MTTD).

Además, se están desarrollando *frameworks* federados de IA que posibilitan el entrenamiento de modelos de forma colaborativa entre operadores de telecomunicaciones sin compartir datos sensibles, lo que garantiza la privacidad mediante técnicas de aprendizaje federado (Yang et al., 2019). Este enfoque resulta prometedor para la defensa colectiva ante ciberataques de gran escala.

Conclusiones

El estudio realizado demuestra que la IA constituye una herramienta esencial para fortalecer la seguridad de las infraestructuras críticas de telecomunicaciones, y ofrece soluciones de detección de *malware* más precisas, adaptativas y escalables que los métodos tradicionales.

Los resultados comparativos indican que los modelos basados en aprendizaje profundo, especialmente las CNN, LSTM y GNN, superan de forma significativa a los enfoques clásicos de *Machine Learning* en términos de precisión, tasa de falsos positivos y capacidad de detección de amenazas desconocidas. Los modelos híbridos que combinan análisis estático, dinámico y de flujo de red han mostrado resultados

sobresalientes en entornos de alta criticidad, que alcanzan tasas de detección superiores al 98 %.

Asimismo, la integración de la IA con sistemas de seguridad existentes, como los IDS y SIEM, permite la correlación automatizada de eventos y la respuesta en tiempo real, lo que mejora la capacidad de detección y mitigación ante ciberataques complejos.

Sin embargo, persisten desafíos relevantes. Entre ellos se destacan la escasa disponibilidad de datos representativos de infraestructuras críticas, la interpretabilidad limitada de los modelos de aprendizaje profundo y la vulnerabilidad ante ataques adversariales, que pueden manipular las predicciones mediante perturbaciones mínimas en los datos de entrada. Además, la exigencia computacional de algunos modelos limita su despliegue en sistemas con recursos restringidos o en nodos perimetrales.

En consecuencia, se recomienda avanzar hacia el desarrollo de modelos explicables (XAI) que permitan comprender las decisiones de los algoritmos, la creación de *datasets* específicos del sector de telecomunicaciones y la adopción de enfoques federados que favorezcan el entrenamiento colaborativo entre operadores sin comprometer la privacidad de los datos.

La aplicación de IA en la detección de *malware* no solo representa una mejora técnica, sino un paso estratégico hacia la ciberresiliencia integral de las telecomunicaciones, garantizando la continuidad operativa de los servicios que sustentan la economía digital y la seguridad nacional.

Referencias bibliográficas

- Alasmary, W., Alasmary, M., & Alhaidari, F. (2022). AI-based Malware Detection and Classification: A Survey. *Computers & Security*, 121, 102856. *Elsevier*. <https://doi.org/10.1016/j.cose.2022.102856>
- Al-Yaseen, W. L., Othman, Z. A., & Nazri, M. Z. (2020). Hybrid Intelligent Intrusion Detection System Using Machine Learning Techniques. *Journal of Information Security and Applications*, 55, 102607. <https://doi.org/10.1016/j.jisa.2020.102607>

CISA – *Cybersecurity and Infrastructure Security Agency*. (2022). *Securing Communications Infrastructure: Guidance for Critical Networks*. U.S. Department of Homeland Security. Disponible en: <https://www.cisa.gov>

Cruz, A., Márquez, J., & Ramírez, D. (2021). Machine Learning Models for Malware Detection in Telecommunication Networks. *IEEE Access*, 9, 178543–178557. <https://doi.org/10.1109/ACCESS.2021.3130201>

ENISA – European Union Agency for Cybersecurity. (2023). *Threat Landscape for the Telecom Sector 2023*. European Commission. Disponible en: <https://www.enisa.europa.eu>

García-Teodoro, P., Maciá-Fernández, G., Díaz-Verdejo, J., & Vázquez, E. (2023). Advanced Intrusion Detection in Telecommunications: AI-driven Approaches. *Computer Networks*, 235, 109000. <https://doi.org/10.1016/j.comnet.2023.109000>

Goodman, B., & Flaxman, S. (2017). European Union Regulations on Algorithmic Decision-Making and a “Right to Explanation”. *AI Magazine*, 38(3), 50–57. <https://doi.org/10.1609/aimag.v38i3.2741>

Huang, L., Joseph, A. D., Nelson, B., Rubinstein, B. I. P., & Tygar, J. D. (2021). Adversarial Machine Learning. *Communications of the ACM*, 64(7), 64–73. <https://doi.org/10.1145/3434228>

Hussain, F., Abbas, G., Fatima, R., & Wang, H. (2023). Graph Neural Networks for Network Threat Detection in 5G and IoT Environments. *IEEE Internet of Things Journal*, 10(8), 7254–7267. <https://doi.org/10.1109/JIOT.2023.3241073>

Kolosnjaji, B., Zarras, A., Webster, G., & Eckert, C. (2018). *Deep Learning for Classification of Malware System Call Sequences*. In *Proceedings of the 2018 IEEE International Conference on Big Data* (pp. 3787–3796). IEEE. <https://doi.org/10.1109/BigData.2018.8621968>

- Pektaş, A., & Acarman, T. (2020). Malware Classification Based on Deep Learning Architectures for Android. *Neural Computing and Applications*, 32(3), 817–833. <https://doi.org/10.1007/s00521-018-3930-1>
- Permana, A., Pradana, F., & Utami, R. (2020). Comparative Analysis of Machine Learning Algorithms for Malware Detection. *International Journal of Computer Applications*, 175(13), 14–22. <https://doi.org/10.5120/ijca2020920431>
- Sgandurra, D., & Lupu, E. C. (2020). Automated Malware Classification and Analysis Using Machine Learning. *Computers & Security*, 95, 101867. <https://doi.org/10.1016/j.cose.2020.101867>
- Shafiq, M., Tian, Z., Bashir, A. K., Du, X., & Guizani, M. (2020). Corrauc: A Deep Learning Model for Anomaly Detection in 5G Telecom Networks. *IEEE Transactions on Network Science and Engineering*, 7(4), 2348–2362. <https://doi.org/10.1109/TNSE.2020.3025241>
- Ucci, D., Aniello, L., & Baldoni, R. (2019). Survey of Machine Learning Techniques for Malware Analysis. *Computers & Security*, 81, 123–147. <https://doi.org/10.1016/j.cose.2018.11.001>
- Wang, J., Zhang, H., & Zhu, X. (2022). Deep Autoencoder-Based Network Anomaly Detection for Critical Infrastructure. *Expert Systems with Applications*, 201, 117138. <https://doi.org/10.1016/j.eswa.2022.117138>
- Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated Machine Learning: Concept and Applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), 12. <https://doi.org/10.1145/3298981>
- Zhang, Y., Pan, S., & Chang, X. (2022). Graph Representation Learning for Cybersecurity: A Review. *IEEE Transactions on Neural Networks and Learning Systems*, 33(5), 1805–1822. <https://doi.org/10.1109/TNNLS.2021.3114520>