



**UNIVERSIDAD ESTATAL PENÍNSULA
DE SANTA ELENA**

FACULTAD DE SISTEMAS Y TELECOMUNICACIONES

INSTITUTO DE POSTGRADO

TÍTULO

**Implementación De Modelos De Machine Learning Para La
Detección De Malware Y Botnets En La Red De Una Sucursal De
Saitel**

AUTOR

Tenorio Cando, Darwin Paul

TRABAJO DE TITULACIÓN

**Previo a la obtención del grado académico en
MAGÍSTER EN CIBERSEGURIDAD**

TUTOR

Andrade Vera, Alicia Germania

Santa Elena, Ecuador

Año 2026



UPSE

**UNIVERSIDAD ESTATAL PENÍNSULA
DE SANTA ELENA
FACULTAD DE SISTEMAS Y TELECOMUNICACIONES
INSTITUTO DE POSTGRADO
TRIBUNAL DE SUSTENTACIÓN**

Ing. Jaime Orozco Iguasnia, Mgtr.
**DELEGADO DE COORDINACIÓN
DEL PROGRAMA**

Ing. Alicia Andrade Vera, Mgtr.
TUTOR

Ing. Carlos Castillo Yagual, Mgtr.
DOCENTE ESPECIALISTA

Ing. Carlos Sanchez Leon, Mgtr.
DOCENTE ESPECIALISTA

Abg. María Rivera González, Mgtr.
**SECRETARIA GENERAL
UPSE**



**UNIVERSIDAD ESTATAL PENÍNSULA
DE SANTA ELENA
FACULTAD DE SISTEMAS Y TELECOMUNICACIONES
INSTITUTO DE POSTGRADO
CERTIFICACIÓN**

Certifico que luego de haber dirigido científica y técnicamente el desarrollo y estructura final del trabajo, este cumple y se ajusta a los estándares académicos, razón por el cual apruebo en todas sus partes el presente trabajo de titulación que fue realizado en su totalidad por **TENORIO CANDO DARWIN PAUL**, como requerimiento para la obtención del título de Magíster en Ciberseguridad.

TUTOR

Ing. Alicia Germania Andrade Vera, Mgtr.

Santa Elena, 30 de junio de 2026



**UNIVERSIDAD ESTATAL PENÍNSULA
DE SANTA ELENA
FACULTAD DE SISTEMAS Y TELECOMUNICACIONES
INSTITUTO DE POSTGRADO
DECLARACIÓN DE RESPONSABILIDAD**

Yo, TENORIO CANDO DARWIN PAUL

DECLARO QUE:

El trabajo de Titulación, **IMPLEMENTACIÓN DE MODELOS DE MACHINE LEARNING PARA LA DETECCIÓN DE MALWARE Y BOTNETS EN LA RED DE UNA SUCURSAL DE SAITEL**, previo a la obtención del título en Magíster en Ciberseguridad, ha sido desarrollado respetando derechos intelectuales de terceros conforme las citas que constan en el documento, cuyas fuentes se incorporan en las referencias o bibliografías. Consecuentemente este trabajo es de mi total autoría.

En virtud de esta declaración, me responsabilizo del contenido, veracidad y alcance del Trabajo de Titulación referido.

Santa Elena, 30 de junio de 2026

EL AUTOR

Tenorio Cando Darwin Paul

IV



UPSE

**UNIVERSIDAD ESTATAL PENÍNSULA
DE SANTA ELENA**

**FACULTAD DE SISTEMAS Y TELECOMUNICACIONES
INSTITUTO DE POSTGRADO**

CERTIFICACIÓN DE ANTIPLAGIO

Certifico que después de revisar el documento final del trabajo de titulación denominado **IMPLEMENTACIÓN DE MODELOS DE MACHINE LEARNING PARA LA DETECCIÓN DE MALWARE Y BOTNETS EN LA RED DE UNA SUCURSAL DE SAITEL**, presentado por el estudiante **TENORIO CANDO DARWIN PAUL** fue enviado al Sistema Antiplagio COMPILATIO, presentando un porcentaje de similitud correspondiente al 4%, por lo que se aprueba el trabajo para que continúe con el proceso de titulación.



Certificado de análisis
Compilatio Magister+ | UPSE-ECU

TESIS REVISION_DARWIN TENORIO

ID : 539b94d16f7b1eef229f47938bea2b7424815c1c



4%

Textos sospechosos

Nombre del fichero : TESIS REVISION_DARWIN TENORIO.txt
Tamaño del archivo original : 12,74 MB
Número de palabras : 16.858

Depositante : ALICIA GERMANIA ANDRADE VERA
Fecha de depósito : 29 de junio de 2026
Tipo de carga : interface
fecha de fin de análisis : 29 de junio de 2026

TUTOR

Ing. Alicia Germania Andrade Vera, Mgtr.



**UNIVERSIDAD ESTATAL PENÍNSULA DE SANTA
ELENA
FACULTAD DE SISTEMAS Y TELECOMUNICACIONES
INSTITUTO DE POSTGRADO
AUTORIZACIÓN**

Yo, TENORIO CANDO DARWIN PAUL

Autorizo a la Universidad Estatal Península de Santa Elena, para que haga de este trabajo de titulación o parte de él, un documento disponible para su lectura consulta y procesos de investigación, según las normas de la Institución.

Cedo los derechos en línea patrimoniales de esta propuesta metodológica y tecnológica avanzada con fines de difusión pública, además apruebo la reproducción de esta propuesta metodológica y tecnológica avanzada dentro de las regulaciones de la Universidad, siempre y cuando esta reproducción no suponga una ganancia económica y se realice respetando mis derechos de autor

Santa Elena, 30 de junio de 2026

EL AUTOR

Tenorio Cando Darwin Paul

AGRADECIMIENTO

Expreso mi agradecimiento a la Universidad Estatal Península de Santa Elena, por haberme permitido formar parte de su programa de maestría y por ser el espacio donde consolidé mis conocimientos académicos.

De manera especial, agradezco a la empresa SAITEL, particularmente a la sucursal Latacunga por brindarme las facilidades técnicas el acceso a su infraestructura y el uso de equipos que fueron fundamentales para la recolección de datos y la validación de este proyecto de investigación.

A mi tutor de tesis, por su invaluable guía, paciencia y por compartir su experticia en el complejo mundo de la Inteligencia Artificial y la ciberseguridad.

Darwin Paul, Tenorio Cando

DEDICATORIA

Dedico este trabajo a mi esposa por ser mi compañera incondicional, mi apoyo en las noches de desvelo y quien siempre creyó en mi capacidad para alcanzar esta meta. A mi hija, quien es mi motor y mi mayor motivación para construir un futuro mejor; cada hora de esfuerzo fue pensando en ti. Y a mis padres, por los valores que me inculcaron y por el sacrificio que hicieron para darme las bases de lo que soy hoy. Este título de maestría les pertenece a todos ustedes

Darwin Paul, Tenorio Cando

ÍNDICE GENERAL

TRIBUNAL DE SUSTENTACIÓN	II
CERTIFICACIÓN	III
DECLARACIÓN DE RESPONSABILIDAD.....	IV
CERTIFICACIÓN DE ANTIPLAGIO	V
AUTORIZACIÓN	VI
ÍNDICE GENERAL.....	IX
ÍNDICE DE TABLAS	XIII
ÍNDICE DE FIGURAS	XIV
RESUMEN	XVI
ABSTRACT	XVII
INTRODUCCIÓN	1
CAPÍTULO 1. MARCO TEÓRICO REFERENCIAL	6
1.1. Antecedentes.....	6
1.2. Infraestructura de redes y servicios de internet.....	7
1.2.1 Proveedores de servicios de internet (ISP).....	8
1.2.2 Redes ópticas pasivas PON	8
1.2.3 Gestión de tráfico y calidad de servicio (QoS).....	9
1.2.4 Enrutamiento y Gestión con MikroTik	10
1.3 Gestión de la seguridad de la información.....	11
1.3.1 Vulnerabilidades, amenazas y riesgos	11
1.3.2 Taxonomía del Malware y Botnets.....	12
1.3.3 Tipos de mecanismos de control	13
1.3.4 Estrategias de seguridad con IA en las Empresas.....	13
1.4 Inteligencia artificial y Machine Learning en la ciberseguridad ..	14
1.4.1 Aprendizaje supervisado: la base de la clasificación.....	14

1.4.2	Algoritmo Random Forest.....	15
1.4.3	Máquinas de soporte vectorial.....	16
1.4.4	Métricas de evaluación de desempeño.....	17
1.5	Infraestructura de captura y telemetría de red	17
1.5.1	Protocolo NetFlow e IPFIX.....	18
CAPÍTULO 2. METODOLOGÍA		19
2.1	Contexto de la investigación	19
2.2	Diseño y alcance de la investigación	20
2.3	Variables de la investigación.....	21
2.4	Tipo y métodos de investigación.....	22
2.5	Población y muestra	23
2.6	Técnicas e instrumentos de recolección de datos	26
2.7	Procesamiento de la evaluación	27
CAPÍTULO 3. RESULTADOS Y DISCUSIÓN.....		30
3.1	Análisis y Caracterización del tráfico de Red.....	30
3.2	Diagnóstico de comportamientos y flujos críticos	30
3.2.1	Identificación de flujos de bajo volumen	30
3.2.2	Anomalías en la negociación de puertos.....	30
3.3	Interpretación de registros NetFlow	31
3.4	Proceso de entrenamiento y validación cruzada.....	31
3.5	Recolección de datos de tráfico de red mediante NetFlow	33
3.5.1	Extracción de reportes en formato HTML	33
3.6	Unificación y consolidación en estructura CSV.....	34
3.7	Limpieza y procesamiento de datos	35
3.8	Punto detección inicial de anomalías mediante Isolation Forest..	37
3.8.1	Implementación del algoritmo Isolation Forest.....	39
3.8.2	Identificación del comportamiento anómalo en el tráfico.....	39
3.8.3	Generación de etiquetas de comportamiento de tráfico.	40
3.9	Desarrollo de modelos de IA	41

3.9.1	Random Forest.....	42
3.9.2	Entrenamiento del modelo Random Forest.....	42
3.9.3	Evaluación del desempeño del modelo.....	42
3.9.4	Métricas de evaluación del modelo	43
3.9.5	Análisis de las gráficas.....	43
3.10	Implementación del modelo XGBoost	52
3.10.1	Entrenamiento y configuración de XGBoost	52
3.10.2	Resultados y métricas de desempeño.....	53
3.10.3	Análisis de las gráficas de XGBoost.....	53
3.11	Análisis comparativo y selección del modelo final.....	59
3.11.1	Análisis Comparativo: Random Forest vs. XGBoost	60
3.11.2	Interpretación Técnica e Impacto en la Seguridad.....	61
3.11.3	Relación con los Objetivos Específicos	61
3.11.4	Validación Externa y Generalización.....	62
3.11.5	Comparativa de confianza	63
3.12	Validación de pruebas en entorno real.....	63
3.12.1	Pruebas con el modelo Random Forest.....	64
3.12.2	Pruebas con el modelo XGBoost.....	70
3.13	Implementación del sistema inteligente de detección de anomalías en la empresa Saitel75	
3.13.1	Arquitectura General del Sistema.....	75
3.13.2	Implementación del Procesamiento NetFlow.....	76
3.13.4	Implementación de la Interfaz Gráfica.....	77
3.13.5	Puesta en Funcionamiento del Sistema Inteligente IDS en la Empresa SAITEL	78
	CONCLUSIONES	87

RECOMENDACIONES	89
REFERENCIAS.....	90
ANEXOS.....	93

ÍNDICE DE TABLAS

Tabla 1.	Métricas de evaluación del modelo de clasificación.....	17
Tabla 2.	Periodo de Análisis (lunes a viernes).....	26
Tabla 3.	Métricas de flujo y comportamiento de anómalo detectado	31
Tabla 4.	Desempeño del modelo Random Forest al tráfico GPON	43
Tabla 5.	Reporte de clasificación detallado para el modelo XGBoost...	53
Tabla 6.	Comparativa final de métricas-Random Forest y XGBoost.....	60
Tabla 7.	Comparativa de Predicciones en Tiempo Real (Random Forest vs. XGBoost)	74
Tabla 8.	Tecnologías utilizadas.....	76
Tabla 9.	Registro de amenaza	80
Tabla 10.	Datos mostrados en la alerta.....	84
Tabla 11.	Parámetros del tráfico detectado como sospechoso.....	84

ÍNDICE DE FIGURAS

Figura 1.	Esquema de funcionamiento de un ISP	8
Figura 2.	Arquitectura de una Red Óptica Pasiva PON.....	9
Figura 3.	Calidad de Servicio QoS y sus aspectos	10
Figura 4.	Conexión con equipos MikroTik.....	11
Figura 5.	Ciclo y taxonomía de un ataque cibernético	12
Figura 6.	Comportamiento de IA en las estrategias de seguridad.....	14
Figura 7.	Aprendizaje aplicado a la detección de botnets en un ISP.....	15
Figura 8.	Algoritmo de aprendizaje Random Forest	16
Figura 9.	Máquinas de soporte vectorial SVM	16
Figura 10.	Telemetría de red utilizando NetFlow e IPFIX	18
Figura 11.	Mapa nodo SAITEL-Sucursal Latacunga	20
Figura 12.	Gráfico de cuerdas NetFlow.....	34
Figura 13.	Unificación de datos	35
Figura 14.	Tratamiento de datos residuales	36
Figura 15.	Entrenamiento del modelo Isolation Forest.	38
Figura 16.	Normalización de variables para el aprendizaje de detección.	38
Figura 17.	Entrenamiento del modelo de detección de anomalías.....	39
Figura 18.	Asignación y clasificación binaria del tráfico de red.....	40
Figura 19.	Generación de etiquetas de comportamiento de tráfico.	41
Figura 20.	Distribución de clases registros por categoría.....	44
Figura 21.	Volumen de datos procesados por el modelo.....	44
Figura 22.	Diagramas de variabilidad y valores atípicos de los paquetes.	45
Figura 23.	Relación matemática entre las columnas del dataset.....	45
Figura 24.	Verificación de aciertos del modelo- tasa de error	46
Figura 25.	Patrones de red para detección de intrusiones.....	46
Figura 26.	Curva de ROC relación entre verdaderos y falsos positivos	47
Figura 27.	Evaluación del modelo con clases desbalanceadas.....	47
Figura 28.	Grado de confianza de predicciones del modelo	48

Figura 29.	Puertos de comunicación utilizados.....	49
Figura 30.	Puertos asociados al tráfico malicioso	49
Figura 31.	Intensidad de actividades de propagación de malware	50
Figura 32.	Análisis de dispersión relación tráfico vs puerto	51
Figura 33.	Matriz de confusión- Random Forest	52
Figura 34.	Distribución de clase, desbalanceo del dataset.....	54
Figura 35.	Distribución y desbalanceo de tráfico normal y las amenazas	54
Figura 36.	Comportamientos estadísticos de variables para el trafico	55
Figura 37.	Relación entre variables de NetFlow	55
Figura 38.	Desempeño cruzado del control del modelo	56
Figura 39.	Relación tasa de verdaderos positivos y falsos negativos	57
Figura 40.	Curva de Precisión-Recall.....	57
Figura 41.	Confianza del modelo XGBoost.....	58
Figura 42.	Real vs Predicción XGBoost.....	58
Figura 43.	Comparación de modelos de Machine Learning	63
Figura 44.	Realización de pruebas en la red de la Empresa SAITEL	64
Figura 45.	Ejecución de código de simulación final en tiempo real	64
Figura 46.	Detección e importación del dataset 2026-04-01_16-00	65
Figura 47.	Recopilación de nuevos datos generados de NetFlow	66
Figura 48.	Detección e importación del dataset 2026-04-01_16:15.....	68
Figura 49.	Almacenamiento en tiempo real de los intervalos de datos	69
Figura 50.	Prueba 1 Intervalo (2026-04-01_16:00)con XGBoost.....	71
Figura 51.	Prueba 2 Intervalo (2026-04-01_16:15)con XGBoost.....	72
Figura 52.	Resultados obtenidos de cada modelo	73
Figura 53.	Interfaz Grafica del sistema IDS	78
Figura 54.	Análisis del trafico de red de la empresa Saitel.....	79
Figura 55.	Alertas enviadas por correo electrónico.....	83

RESUMEN

Esta investigación implementa un sistema de ciberseguridad para el ISP SAITEL basado en Machine Learning para detectar anomalías de red. Debido a que los mecanismos tradicionales de firmas son insuficientes ante amenazas avanzadas y comportamientos compatibles con ataques emergentes, frente a los cuales los mecanismos tradicionales basados en firmas presentan limitaciones (Pai et al., 2025) (Salcedo Suarez, 2024), se propone una arquitectura inteligente con modelos Random Forest y XGBoost. Esta solución permite una transición hacia un monitoreo preventivo y predictivo que garantiza la integridad de la red troncal y optimiza la capacidad de respuesta técnica ante incidentes de tráfico inusual sin afectar la calidad del servicio para los usuarios finales.

La metodología se fundamentó en un diseño cuantitativo de alcance explicativo, donde se recolectó una muestra de metadatos NetFlow durante un periodo de observación de 30 días en la sucursal Latacunga. Para garantizar la solidez de la evidencia, se aplicó un criterio de partición de datos del 80% para entrenamiento y 20% para pruebas, sometiendo a evaluación los modelos Random Forest y XGBoost mediante validación cruzada.

Los resultados demostraron que el modelo Random Forest presentó el mejor desempeño global, alcanzando un accuracy de 99.83%, una precisión de 90%, un recall de 98% y un F1-Score de 94% en la detección de patrones de tráfico potencialmente asociados a malware. Este rendimiento evidencia la capacidad del modelo para identificar comportamientos anómalos con alta sensibilidad y una baja tasa de falsos positivos, permitiendo fortalecer el monitoreo preventivo de la infraestructura de red de SAITEL. Este desempeño se explica por la capacidad del algoritmo para manejar datos no lineales y variables correlacionadas lo cual resulta característico y esencial al analizar el tráfico de red. Dicha robustez técnica permitió equilibrar la identificación de ataques con una reducción drástica de falsos positivos garantizando una navegación fluida para el usuario final mediante una herramienta escalable que transforma el monitoreo reactivo en un mecanismo de

detección temprana de anomalías, fortaleciendo integralmente la seguridad del ISP frente a potenciales vectores de ataque modernos.

Palabras claves: Machine Learning, ciberseguridad, NetFlow

ABSTRACT

This research implements a cybersecurity system for the ISP SAITEL based on Machine Learning to detect network anomalies. Since traditional signature-based mechanisms are insufficient against advanced threats and behaviors compatible with emerging attack patterns, an intelligent architecture based on Random Forest and XGBoost models is proposed (Pai et al., 2025; Salcedo Suarez, 2024). This solution enables a transition toward preventive and predictive monitoring, ensuring the integrity of the backbone network and optimizing technical response capabilities to unusual traffic incidents without affecting service quality for end users.

The methodology was based on a quantitative design with an explanatory scope, in which a sample of NetFlow metadata was collected over a 30-day observation period at the Latacunga branch. To ensure the robustness of the evidence, an 80% data split was used for training and 20% for testing, while the Random Forest and XGBoost models were evaluated through cross-validation.

The results showed that the Random Forest model achieved the best overall performance, reaching an accuracy of 99.83%, a precision of 90%, a recall of 98%, and an F1-score of 94% in detecting traffic patterns potentially associated with malware. This performance demonstrates the model's ability to identify anomalous

behaviors with high sensitivity and a low false-positive rate, strengthening the preventive monitoring of SAITEL's network infrastructure. This outcome is explained by the algorithm's capability to handle nonlinear data and correlated variables, which are essential characteristics in network traffic analysis. Such technical robustness allowed balancing attack identification with a significant reduction in false positives, ensuring a smooth browsing experience for end users through a scalable tool that transforms reactive monitoring into an early anomaly detection mechanism, comprehensively strengthening the ISP's security against potential modern attack vectors.

Keywords: Machine Learning, cybersecurity, NetFlow

INTRODUCCIÓN

En el ecosistema digital contemporáneo en donde los proveedores de servicios de internet (ISP) desempeñan un papel fundamental dentro de la conectividad y transmisión de información entre usuarios, así como también de empresas y organizaciones. Debido al crecimiento constante del tráfico de red y al incremento de amenazas cibernéticas, así como la ciberseguridad se ha convertido en un componente esencial para garantizar la estabilidad, disponibilidad y confiabilidad de las infraestructuras tecnológicas. Dentro de este contexto se establece que la empresa SAITEL, específicamente en su sucursal de Latacunga, enfrenta desafíos relacionados con el monitoreo y análisis del tráfico generado por redes GPON e inalámbricas. El aumento del volumen de datos y la diversidad de conexiones dificultan la identificación de comportamientos sospechosos asociados potencialmente a malware y actividades compatibles con botnets. Estas amenazas pueden generar degradación del servicio, consumo no autorizado de recursos de red y riesgos para la disponibilidad de la infraestructura tecnológica.

Uno de los principales problemas presentes en el monitoreo de redes consiste en diferenciar el tráfico legítimo del tráfico potencialmente malicioso en tiempo real por lo cual los mecanismos tradicionales basados en firmas o reglas estáticas presentan limitaciones frente a amenazas dinámicas y patrones de comportamiento variables, especialmente cuando los ataques modifican continuamente sus características de comunicación. Frente a esta problemática se tiene que la presente investigación propone la implementación y evaluación de modelos de Machine Learning aplicados sobre registros NetFlow para el análisis de patrones de tráfico de red. El uso de registros NetFlow permite trabajar con metadatos de comunicación como direcciones IP, puertos, protocolos y volúmenes de tráfico, evitando la inspección profunda de paquetes y optimizando el procesamiento de información.

La investigación se orienta al análisis de patrones de tráfico y comportamientos anómalos potencialmente asociados a amenazas cibernéticas,

utilizando registros NetFlow como fuente de indicadores de compromiso observables dentro de la red de SAITEL sucursal Latacunga. En el ámbito metodológico se tiene que la investigación permite evaluar el desempeño de modelos de Machine Learning aplicados a la clasificación de tráfico potencialmente malicioso por lo cual finalmente, desde la perspectiva práctica, se desarrolla un entorno de monitoreo y validación que integra visualización gráfica, análisis de tráfico y generación de alertas para observar el comportamiento del tráfico de red dentro de un escenario controlado basado en datos de la infraestructura de SAITEL.

El propósito de esta investigación consiste en implementar modelos de Machine Learning orientados al análisis y clasificación de patrones de tráfico asociados a posibles amenazas, utilizando registros NetFlow de la red de la sucursal SAITEL Latacunga como base para la validación del sistema desarrollado.

Planteamiento de la investigación (Fundamentación de la investigación)

La investigación se sustenta en la aplicación de metodologías experimentales y analíticas para la construcción, entrenamiento y validación de modelos predictivos. El proceso incluye la recolección y normalización de datos de red, la selección de características relevantes y la evaluación de algoritmos supervisados como Random Forest. Para determinar la eficacia en la clasificación del tráfico entre benigno y maligno, el desempeño del modelo será validado mediante el F1 Score. Se selecciona esta métrica por su robustez en entornos de reales, ya que permite armonizar la precisión y la sensibilidad del algoritmo, garantizando un equilibrio crítico entre la detección efectiva de botnets y la reducción de errores que podrían afectar la conectividad de los usuarios legítimos. Este enfoque fortalece la validez empírica del estudio y contribuye al desarrollo de nuevos métodos de detección aplicables a estructuras de red de alta demanda.

La implementación de modelos de Machine Learning en la infraestructura del ISP aporta un valor práctico significativo, dado que la optimización de los mecanismos de monitoreo y respuesta ante incidentes podría favorecer la reducción de los tiempos de detección y mitigación de amenazas.

Formulación del problema de investigación

¿De qué manera la aplicación de modelos de Machine Learning sobre registros NetFlow permiten incrementar la precisión en la clasificación y reducir la tasa de falsos positivos en la detección de malware y botnets en la red de la sucursal del ISP SAITEL?

Objetivo General:

Implementar modelos de Machine Learning para la detección de malware y botnets en el proveedor de servicios de internet SAITEL sucursal Latacunga para fortalecer la seguridad de la infraestructura.

Objetivos Específicos:

1. Examinar el tráfico de red de SAITEL sucursal Latacunga para identificar comportamientos distintivos vinculados a la presencia de malware y botnets.
2. Documentar las características del tráfico de red SAITEL sucursal Latacunga para establecer un registro sistemático de comportamientos que puedan estar asociados a malware y botnets.
3. Seleccionar modelos de Machine Learning adecuados para la detección temprana de amenazas en el tráfico de red SAITEL sucursal Latacunga.
4. Implementar un sistema de monitoreo que integre modelos Machine Learning orientados a la detección de amenazas dentro de la infraestructura de red existente para validar su funcionamiento en un entorno real.

Planteamiento hipotético

El uso de modelos supervisados de IA sobre la telemetría de red de SAITEL constituye una herramienta eficaz para la detección de anomalías; para lo cual, la precisión del modelo se confirma a través de métricas de desempeño que garantizan la continuidad del servicio.

Relevancia científica y profesional

La presente investigación aporta conocimiento en el ámbito de la ciberseguridad aplicada al tráfico real de red de una sucursal del ISP SAITEL y el uso de técnicas Machine Learning lo que permite generar metodologías replicables y comparables para futuros estudios en clasificación binaria de patrones de tráfico malicioso. Además, refuerza la aplicación de técnicas de análisis de grandes volúmenes de datos de red, favoreciendo la consolidación de soluciones automatizadas frente a amenazas dinámicas y sofisticadas.

La implementación de modelos de Machine Learning sobre el registro de datos NetFlow permite identificar y clasificar tráfico malicioso de manera rápida y eficaz, reduciendo la exposición a ataques y fortaleciendo la continuidad operativa. Esto representa un aporte directo a la práctica profesional de la seguridad de redes, brindando criterios relevantes para la toma de decisiones en la protección de infraestructuras de telecomunicaciones.

Esta tesis se enmarca dentro de la línea de investigación de Tecnología y Sistemas de la Información (TSI) y en la sub-línea de investigación en Redes y Seguridad de la Información mediante Aprendizaje Automático con énfasis en la aplicación de Machine Learning para la identificación de malware y botnets en tráfico de red real, contribuyendo al desarrollo de sistemas inteligentes de monitoreo y análisis en entornos de ISP.

CAPÍTULO 1. MARCO TEÓRICO REFERENCIAL

1.1. Antecedentes

La seguridad en los proveedores de servicio de internet ISP ha dejado de ser un tema meramente técnico para convertirse en un pilar de la confianza del usuario. El aumento de ataques mediante malware y redes de botnets no solo compromete datos, si no la estabilidad de infraestructuras críticas. Al realizar la investigación de literatura y estudios previos, se identifican diversas estrategias que sirven de base para esta investigación.

En los años recientes, la investigación se ha volcado hacia soluciones más inteligentes. Brezo (2014), publico el artículo “Detección de trafico de control de botnets modelizando el flujo de los paquetes de red”, donde se propuso comparar la efectividad de distintos algoritmos de aprendizaje automático. Tras procesar grandes volúmenes de datos bajo el estándar NetFlow, determinaron que el modelo Random Forest es el más equilibrado, ya que logra una precisión cercana al 94% sin consumir los excesivos recursos de procesamiento que requieren las redes neuronales profundas (Brezó Fernández, 2014).

García Deus (2021), en su trabajo titulado “Aprendizaje automático para detección de tráfico IoT anómalo, spam y malware”, analiza la versatilidad de la inteligencia artificial para enfrentar diversas amenazas en la seguridad organizacional. La investigación se centró en tres áreas críticas: la identificación de anomalías en dispositivos IoT, el filtro de Spam y la clasificación de archivos maliciosos. Mediante una metodología que incluye el uso de algoritmo tanto supervisados como no supervisados, la autora exploro la capacidad de los modelos para no solo detectar la presencia de malware como sino también discernir su topología específica. Los resultados destacan la importancia de la detección temprana en entornos de tráfico masivo, concluyendo que el aprendizaje automático es una herramienta fundamental para identificar patrones sospechosos que los métodos convencionales suelen pasar por alto en infraestructuras digitales complejas (García Deus, 2021).

Aboaoja et al. (2022), en su artículo científico titulado “Malware Detection Issues, Challenges and Future Directions: A Survey”, realizaron una revisión exhaustiva sobre el estado actual y los desafíos de la detección de software malicioso. El objetivo principal del estudio fue analizar las limitaciones de las técnicas convencionales frente a la sofisticación del malware moderno y proyectar las tendencias futuras en este campo a través de una metodología comparativa de diversas aproximaciones técnicas, los autores examinaron el procesamiento de datos y la eficacia de los modelos de aprendizaje automático en la identificación de patrones complejos. En sus conclusiones, resaltan que la integración de la inteligencia artificial es indispensable para superar los retos de detección en tiempo real, señalando que el futuro de la ciberseguridad depende de modelos capaces de adaptarse a la evolución constante de las amenazas en redes de alta escala (Aboaoja et al., 2022).

Guale Rosales (2025) en su tesis de maestría titulada “Comparativa de algoritmos de aprendizaje automático para la detección de tipos específicos de malware en entornos adversos”, evalúa eficacia de modelos de inteligencia artificial en la clasificación de amenazas críticas como ransomware y troyanos. Mediante una metodología estructurada que incluyó el procesamiento de datos y el análisis de métricas de desempeño como curvas ROC y matrices de confusión, el autor comparó el rendimiento de diversos clasificadores. Sus resultados demuestran que los algoritmos de ensamble, específicamente Random Forest y XGBoost, superan a los métodos tradicionales en precisión y capacidad de generalización. La investigación concluye que estos modelos son soluciones altamente viables para fortalecer la seguridad informática, logrando un equilibrio óptimo entre un diagnóstico certero y la eficiencia necesaria para la detección automatizada en entornos de red reales (Guale Rosales, 2025).

1.2. Infraestructura de redes y servicios de internet

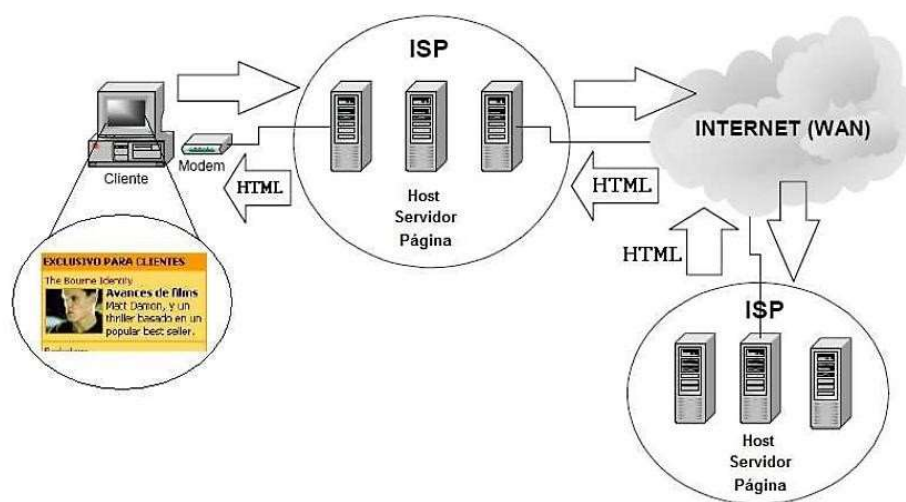
El conjunto de soluciones de seguridad en un entorno de producción exige ante todo comprender la naturaleza de la infraestructura que se pretende proteger. En el caso de un proveedor de servicio de internet ISP, la red no es solo un conjunto

de dispositivos, sino un ecosistema masivo donde la disponibilidad y la integridad de los datos son los activos más críticos.

1.2.1 Proveedores de servicios de internet (ISP)

Un ISP actúa como la columna vertebral de la sociedad digital, gestionando la conectividad entre los usuarios finales y el Backbone global de internet. Según la Unión Internacional de Telecomunicaciones (*Global Cybersecurity Index*, s. f.), la función del proveedor a trascendido la simple transmisión de paquetes, convirtiéndose en un ente responsable de la higiene ciberespacial. Es por ello, que un ISP no solo debe garantizar ancho de banda, sino que debe implementar arquitecturas resilientes capaces de mitigar el tráfico malicioso que circula por sus nodos antes de que este enlace a otros sectores críticos de la red, en la figura 1 se muestra el funcionamiento de un ISP.

Figura 1. Esquema de funcionamiento de un ISP



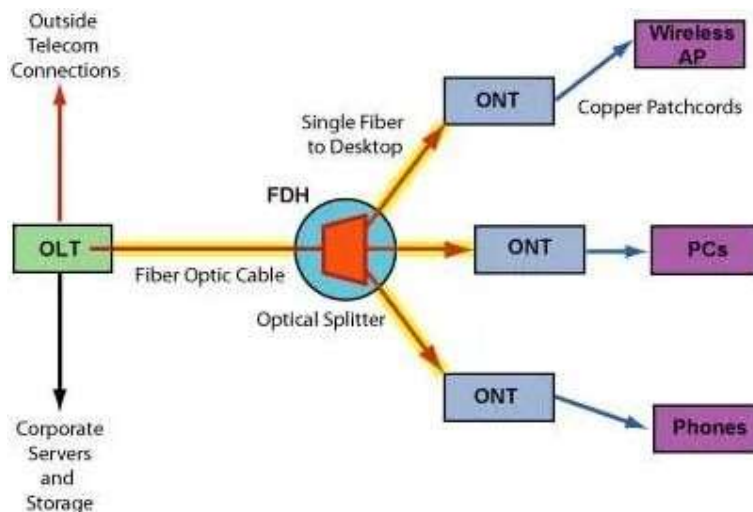
Nota: Diagrama de diseño de un ISP para proveer de internet a clientes. Adaptado de (Mesa & Paucar, 2020)

1.2.2 Redes ópticas pasivas PON

Para llevar una conectividad que ofrece un ISP se utiliza las tecnologías de redes en este caso red óptica pasiva (PON) específicamente bajo el estándar GPON. Estas redes se caracterizan por el uso de divisores ópticos no alimentados lo que

optimiza la eficiencia energética y reduce los costos operativos. Sin embargo, como señala (Concejal-Muñoz, 2022), la arquitectura de un punto a multipunto de las redes PON introduce un desafío de seguridad: el medio compartido. Si un dispositivo de usuario (ONT) es comprometido por una botnet, el tráfico malicioso puede intentar propagarse lateralmente o saturar los enlaces de subida (upstream), afectando la estabilidad del segmento óptico compartido por otros clientes, en la figura 2 se observa la arquitectura de una red PON.

Figura 2. Arquitectura de una Red Óptica Pasiva PON



Nota: Representación esquemática de la topología de una red óptica pasiva PON. Adaptado de (Encalada Sotomayor, 2021)

1.2.3 Gestión de tráfico y calidad de servicio (QoS)

La calidad de servicio QoS es un conjunto de mecanismos que permiten priorizar flujos de datos específicos sobre otro punto en una red de ISP, esta es vital para que servicios como la telefonía IP o el Streaming no sufran degradación. No obstante, el malware moderno, como el ransomware o las botnets, generan patrones de tráfico que imitan comportamientos legítimos o utilizan técnicas de inundación (flooding). (Ghuge, 2023) sostiene que una gestión de QoS eficiente ya no puede depender únicamente de reglas estáticas, sino que requiere de una capa de

inteligencia que identifique cuando el tráfico de control o datos de un usuario es en realidad una anomalía que intenta agotar los recursos del sistema, figura 3.

Figura 3. Calidad de Servicio QoS y sus aspectos

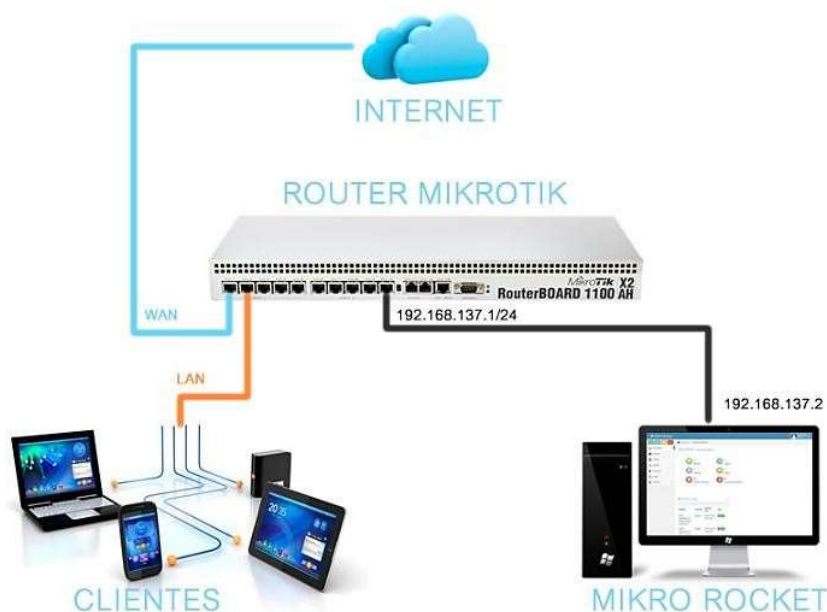


Nota: Parámetros que brindan los servicios QoS. Adaptado de (servicio qos | VoIP, s. f.)

1.2.4 Enrutamiento y Gestión con MikroTik

Dentro de la sucursal de ISP, la administración de tráfico recae sobre plataformas de enrutamiento robustas cómo destacando el uso de equipos MikroTik CCR bajo el sistema operativo Router OS. Estos dispositivos son reconocidos por su versatilidad para implementar reglas de firmware complejas y fundamentalmente por su capacidad de explotar flujos de datos mediante el protocolo NetFlow/IPFIX. Como indica la investigación técnica de (*RouterOS - MikroTik*, 2025), esta capacidad de telemetría es la que permite observar el interior de la red sin necesidad de capturar cada paquete, lo cual es esencial para aplicar modelos de Machine Learning sin comprometer el rendimiento de procesador del router ni la privacidad del usuario, en la figura 4 se muestra una conexión de enrutamiento con MikroTik.

Figura 4. Conexión con equipos MikroTik



Nota: Esquema de conexión física del sistema uno a uno con equipos MikroTik. Adaptado de («Conexión física del sistema», s. f.)

1.3 Gestión de la seguridad de la información

La gestión de la seguridad en un entorno de telecomunicaciones no se limita a la instalación de barreras técnicas como implica un ciclo continuo de identificación, protección y respuesta. Para un ISP, la seguridad de la información debe garantizar la confidencialidad, integridad y disponibilidad de los datos que transitan por su infraestructura, considerando que cualquier interrupción no solo afecta a la empresa, sino la continuidad operativa de miles de usuarios.

1.3.1 Vulnerabilidades, amenazas y riesgos

En el ámbito de un estudio académico o de investigación, es vital distinguir estos tres conceptos para evitar ambigüedades. Según (*ISO/IEC 27005*, 2022), una vulnerabilidad es una debilidad intrínseca en el sistema (como un firmware de router desactualizado), mientras que una amenaza es de la gente externo (malware o atacante que busca explotar dicha debilidad). El riesgo es la probabilidad de que esto ocurra y el impacto que generaría. Para la investigación el riesgo principal es

la pérdida de control de los nodos de red debido a infecciones masivas que saturen la capacidad del ISP.

1.3.2 Taxonomía del Malware y Botnets

El malware ha evolucionado de ser un simple código dañino a convertirse en herramientas de espionaje y secuestro de datos, dentro de esta categoría, las botnets representan la mayor amenaza para un proveedor de servicios. Una botnet es una red de dispositivos “zombies” controlados remotamente por un atacante (botmaster). Como señala (Aboaoja et al., 2022), el peligro de las botnets radica en su arquitectura de comando y control. En una sucursal de ISP, un solo usuario infectado puede convertir su conexión en un nodo de ataque que participa en denegaciones de servicios distribuidas o envío de spam masivo, sin que el usuario lo note, pero degradando severamente el rendimiento de la red troncal, en el esquema de la figura 5, se observa el ciclo de vida de un ataque cibernético y taxonomía de impacto en redes ISP.

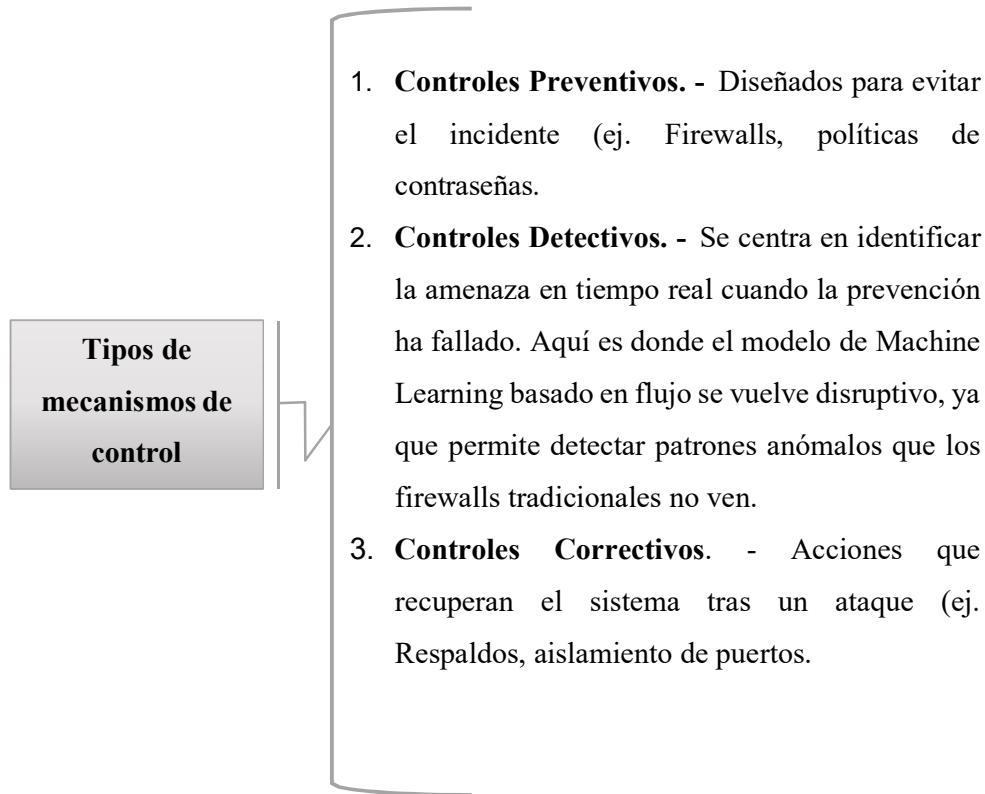
Figura 5. Ciclo y taxonomía de un ataque cibernético



Nota: Progreso secuencial de una intrusión cibernética moderna clasificando las fases de un ataque avanzado desde el reconocimiento inicial hasta el impacto final en la infraestructura del ISP. Elaboración propia (DT,2026).

1.3.3 Tipos de mecanismos de control

Para mitigar estas amenazas, la arquitectura de seguridad se divide en tres niveles de respuesta:



1.3.4 Estrategias de seguridad con IA en las Empresas

Como sostiene la investigación de (Enríquez Sandoval, 2025), la adopción de la Inteligencia Artificial (IA) ha marcado un cambio de paradigma, pasar de la seguridad basada en firmas a la seguridad basada en comportamientos. Las estrategias modernas integran modelos de IA para analizar el tráfico de red de forma masiva. En lugar de buscar un virus conocido, la IA analiza si el comportamiento de una IP es normal o anómalo. Ese enfoque permite a empresas de ISP anticiparse a comportamientos anómalos asociados a amenazas emergentes, fortaleciendo la

resiliencia institucional frente a ataques sofisticados, en el esquema de la figura 6, se observa el enfoque moderno de una seguridad proactiva impulsada con IA.

Figura 6. Comportamiento de IA en las estrategias de seguridad



Nota: El diagrama ilustra el enfoque con IA proactivo, que establece un perfil de comportamiento normal para detectar desviaciones o anomalías. Elaboración propia (DT,2026).

1.4 Inteligencia artificial y Machine Learning en la ciberseguridad

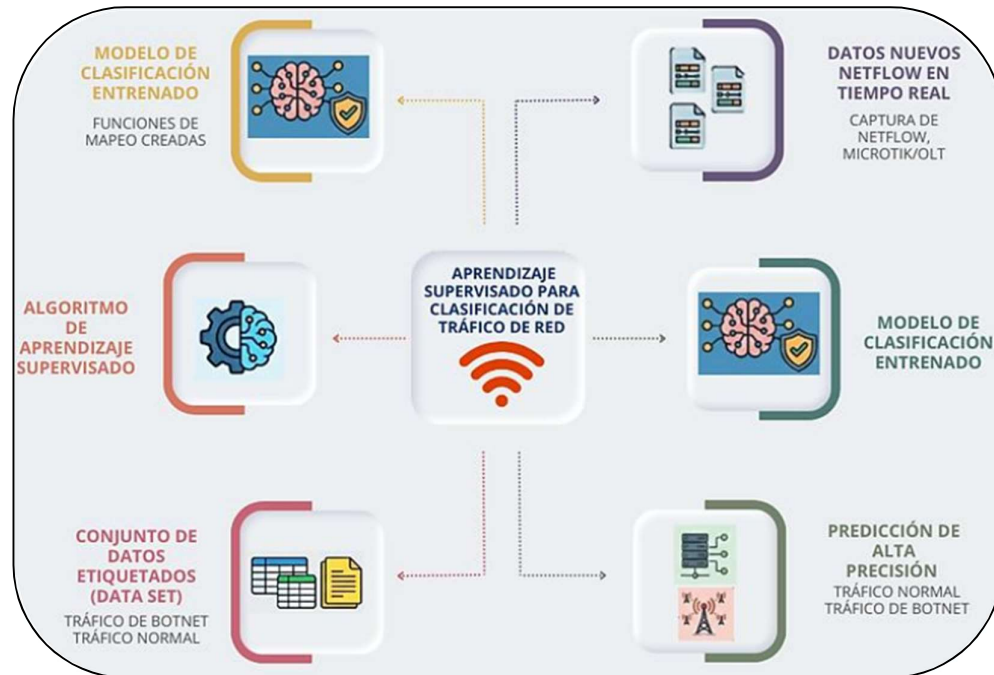
El punto clave de esta investigación se sustenta en el Machine Learning, una rama de la Inteligencia Artificial que permite a los sistemas aprender patrones a partir de datos históricos sin ser programados estrictamente para cada escenario. En el estudio de un ISP, el Machine Learning no busca archivos infectados, sino comportamientos estadísticamente improbables en los flujos de red.

1.4.1 Aprendizaje supervisado: la base de la clasificación

Para la detección de botnets en empresas de servicios de internet, se utiliza el aprendizaje supervisado. Este enfoque consiste en entrenar al algoritmo con un conjunto de datos etiquetados (data set), donde cada registro ya se conoce como tráfico normal o tráfico de botnet. Según (Salcedo Suarez, 2024), este proceso permite al modelo crear una función de mapeo que al recibir datos nuevos NetFlow

en tiempo real, puede predecir con alta precisión a qué categoría pertenecen, el diagrama esquemático de la figura 7 muestra el proceso de aprendizaje supervisado.

Figura 7. Aprendizaje aplicado a la detección de botnets en un ISP.

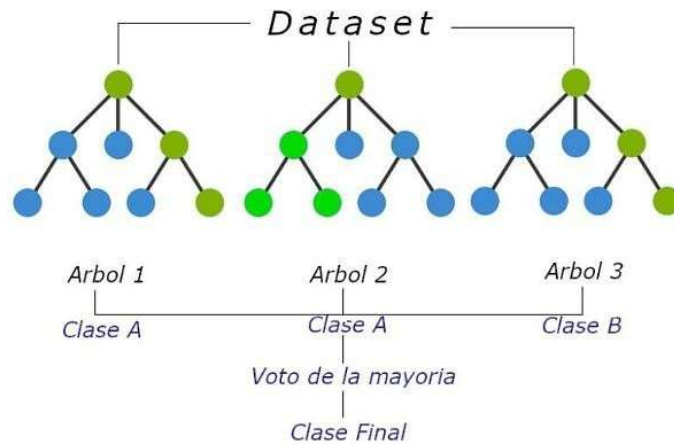


Nota: El diagrama muestra el proceso completo del aprendizaje supervisado: entrenamiento del algoritmo con datos etiquetados para crear funciones de mapeo. Elaboración propia (DT,2026)

1.4.2 Algoritmo Random Forest

Este es un algoritmo de aprendizaje por conjunto que opera construyendo una multitud de árboles de decisión durante el entrenamiento. Random Forest a diferencia de un solo árbol de decisión que puede ser impreciso, esta combina cientos de árboles para obtener una votación mayoritaria, esto lo hace extremadamente robusto frente al ruido de la red y evita el sobreajuste. Para la aplicación en el trabajo de investigación es ideal puesto que se puede clasificar miles de registros de mikrotik ccr de forma eficiente, identificando qué variables, puertos, duración de flujo, bits transmitidos, son las que realmente delatan al malware, figura 8.

Figura 8. Algoritmo de aprendizaje Random Forest

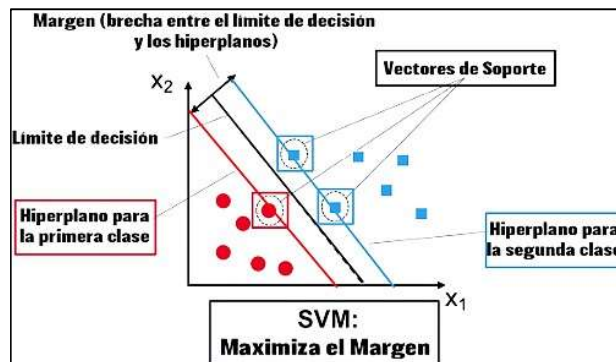


Nota: Diagrama de combinación de múltiples árboles de decisión. Adaptado de (Javier, 2022).

1.4.3 Máquinas de soporte vectorial

Como un aporte de investigación técnica, se analiza el modelo SVM. Este es un algoritmo que busca encontrar el hiperplano óptimo para separar las clases en un espacio multidimensional. Según (Rhanoui et al., 2022), aunque es muy preciso en espacios binarios, su costo computacional en redes de alto tráfico como las de un ISP suelen ser mayor que la de Random Forest, lo que justifica la elección de este último para entornos de producción masiva, figura 9.

Figura 9. Máquinas de soporte vectorial SVM



Nota: modelo de aprendizaje supervisado para la clasificación y regresión de datos. Adaptado de (Estupiñan et al., 2016)

1.4.4 Métricas de evaluación de desempeño

Para la validación de un modelo de Machine Learning se requiere un análisis cuantitativo riguroso para determinar la eficacia del algoritmo en la clasificación de tráfico dentro del ISP, para ello se requiere utilizar indicadores de rendimientos estandarizados que permitan medir la capacidad del modelo para distinguir entre el tráfico legítimo y las actividades maliciosas de una botnet, minimizando tanto los falsos negativos o ataques no detectados, como los falsos positivos o bloqueos erróneos de usuarios legítimos. En la tabla 1, se detallan las métricas internacionales de calidad aplicadas en esta investigación.

Tabla 1. Métricas de evaluación del modelo de clasificación

Métrica	Definición Técnica	Preguntas clave de Validación
Precisión (Precision)	Proposición de predicciones positivas que fueron correctas	De todos los flujos marcados como malware ¿Cuántos lo eran realmente?
Exhaustividad (Recall)	Capacidad de modelo para identificar todos los casos positivos reales	De todos los ataques reales presentes ¿Cuántos logramos detectar con éxito?
Puntuación (F1-Score)	Media armónica entre la precisión y la exhaustividad	¿Existe un equilibrio óptimo entre la detección y la exactitud del modelo?
Matriz de Confusión	Matriz que desglosa aciertos y errores (verdaderos/falsos positivos y negativos)	¿En qué categorías específicas el modelo está acertado o fallando?

Nota: Evaluación integral de métricas para el modelo de clasificación según las necesidades de red. Elaboración Propia (DT,2026)

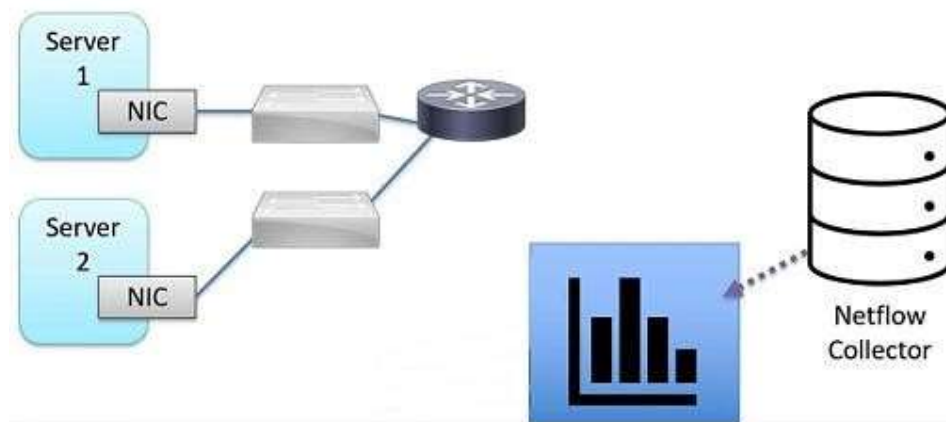
1.5 Infraestructura de captura y telemetría de red

Para que un modelo de Machine Learning pueda clasificar el tráfico, requiere una fuente de datos estructurada que no comprometa el rendimiento de la red troncal. En el estudio de la investigación de una empresa y de ISP, no se analizan los paquetes completos (por privacidad rendimiento), sino los metadatos.

1.5.1 Protocolo NetFlow e IPFIX

El estándar aplicado a la investigación es el uso de NetFlow o su evolución, IPFIX. Estos protocolos funcionan recolectando metadatos de las conexiones, como direcciones IP, puertos y volúmenes de bytes, sin inspeccionar el contenido del mensaje. El análisis de flujo mediante NetFlow permite una visibilidad fundamental de la red sin la sobrecarga computacional que implica el análisis de carga útil, figura 10. Esa técnica es primordial y fundamental para la seguridad moderna, ya que las botnets actuales suelen utilizar cifrados, haciendo que los metadatos sean la señal más fiable para la detección de anomalías (*Cisco IOS NetFlow*, 2023).

Figura 10. Telemetría de red utilizando NetFlow e IPFIX



Nota: NetFlow e IPFIX solución de problemas en redes con datos de flujo. Adaptado de (Cox, 2016).

CAPÍTULO 2. METODOLOGÍA

2.1 Contexto de la investigación

La presente investigación se sitúa en un entorno operativo real, específicamente en la sucursal de la empresa SAITEL la cual se encuentra ubicada en el sector céntrico de la calle Quito y pasaje La Catedral. Esta sede brinda servicios de conectividad mediante tecnología GPON a aproximadamente 200 usuarios activos, cuyo tráfico de red genera los registros NetFlow utilizados para el análisis del estudio.

En este contexto se establece que la unidad de análisis no corresponde directamente a los usuarios, sino a los flujos y metadatos de tráfico capturados mediante NetFlow. Estos registros contienen información relacionada con direcciones IP, puertos, protocolos, volumen de tráfico y duración de las conexiones, permitiendo el entrenamiento y validación de modelos de Machine Learning orientados al análisis de patrones de tráfico potencialmente asociados a malware y actividades anómalas.

La selección del entorno de estudio responde a la alta diversidad y volumen de tráfico generado dentro de la infraestructura GPON de la sucursal, lo cual proporciona un conjunto representativo de registros de red para el análisis y evaluación de los modelos propuestos.

El enfoque del estudio se ha delimitado estrictamente a la infraestructura GPON, por ser el segmento de red con mayor criticidad operativa y volumen de datos dentro de la sucursal. Según los registros históricos de monitoreo del ISP, esta red centraliza el tráfico de clientes con diversos perfiles de consumo, desde residenciales hasta corporativos, registrando picos de demanda que alcanzan aproximadamente 1 Gbps durante las horas de mayor concurrencia.

La elección de este escenario no es fortuita; la alta densidad y diversidad del tráfico generado en este nodo ofrecen un ecosistema idóneo para el análisis de flujos IP. Esta saturación natural de datos incrementa las probabilidades de interceptar

patrones anómalos y señales útiles de actividad maliciosa, tales como comunicaciones de comando y control (C&C) de botnets o programación de malware, que en redes de menor escala podrían pasar desapercibidas.

El enfoque en esta infraestructura permite obtener una caracterización precisa del tráfico real facilitando la implementación y evaluación de modelos de Machine Learning orientados a la detección temprana de amenazas en un entorno de producción del ISP, en la figura 11 se observa la ubicación exacta del nodo.

Figura 11. Mapa nodo SAITEL-Sucursal Latacunga



Nota: Ubicación extraída de Google Maps

2.2 Diseño y alcance de la investigación

El diseño metodológico se define como un estudio cuasi-experimental de enfoque aplicado, ya que implica la implementación, entrenamiento y validación de modelos de Machine Learning sobre datos reales extraídos de la infraestructura operativa del ISP SAITEL. Bajo este marco técnico, la investigación se fundamenta en el procesamiento de un dataset de registros NetFlow mediante una metodología basada en la analítica de datos históricos que permite entrenar los algoritmos de manera segura y controlada. Dicho procedimiento asegura la integridad de la red troncal al no interrumpir el servicio de los usuarios finales mientras se evalúa la capacidad de respuesta del sistema. Asimismo, el alcance de este estudio es de

carácter explicativo y predictivo, buscando determinar con alta precisión los patrones de comportamiento que caracterizan al malware y las botnets dentro de los flujos de tráfico recolectados. Por consiguiente, los resultados obtenidos proporcionan evidencia empírica robusta sobre la efectividad de la inteligencia artificial en la detección de anomalías, estableciendo las bases técnicas necesarias antes de considerar una implementación en la red real de la sucursal Latacunga.

La validez interna del estudio garantiza un control ordenado de las variables involucradas en el entrenamiento de los modelos de Machine Learning para lo cual se utiliza técnicas de procesamiento y normalización de datos, así como estrategias de validación cruzada que aseguren que los resultados obtenidos reflejan de manera fiable el desempeño de los algoritmos y no sesgos derivados del muestreo o del ruido presente en el tráfico de red.

La validez externa se define como una posibilidad técnica condicionada debido a que los hallazgos en la sucursal de SAITEL Latacunga corresponden a un entorno y tiempo específicos por lo cual la aplicación del modelo en otros nodos de red o diferentes proveedores requiere de validaciones y ajustes adicionales a su vez la diversidad de patrones detectados sugiere una arquitectura escalable para lo cual se recomienda realizar reentrenamientos adaptativos que aseguren la efectividad del algoritmo frente a las particularidades de cada nueva infraestructura de red antes de considerar su implementación global.

2.3 Variables de la investigación

Para fortalecer el rigor metodológico se presenta la operacionalización de las variables que intervienen en el estudio:

- **Variable Independiente:** Características del tráfico NetFlow. Se operacionaliza mediante los atributos técnicos de los flujos de datos (metadatos) tales como: dirección IP de origen/destino, puertos de comunicación, número de paquetes, bytes transmitidos y duración del flujo. Su unidad de medida se expresa en valores numéricos y categóricos extraídos del protocolo NetFlow.

- **Variable Dependiente:** Clasificación del tráfico. Se define como la etiqueta o categoría asignada por el modelo de Machine Learning tras el análisis. Se operacionaliza en tres categorías: tráfico normal, malware o escaneo de red (scan) el cual es asociado a Botnet. Su unidad de medida es la probabilidad de pertenencia a una clase (0 a 1) y el acierto de clasificación verificado mediante la matriz de confusión.

2.4 Tipo y métodos de investigación

Tipo de Investigación

El enfoque de la presente investigación es de carácter cuantitativo debido que la problemática abordada requiere el análisis de datos medibles y cuantificables, en particular se trabajó con grandes volúmenes de tráfico de red capturados en la infraestructura de SAITEL con el propósito de entrenar, validar y comparar modelos de aprendizaje automático orientados a la detección de malware y botnets. Esta perspectiva permitió la obtención de métricas de desempeño objetivas tales como: precisión, sensibilidad, especificidad y tasa de falsos positivos, lo cual contribuye a garantizar la validez interna de los resultados.

El carácter sistemático y replicable del análisis posibilita la generalización de las conclusiones a contextos similares fortaleciendo así la aplicabilidad de los hallazgos en entornos reales de ciberseguridad.

Método de investigación

El presente estudio se fundamenta en el método deductivo debido a que parte del análisis de los principios generales y fundamentos teóricos relacionados con la tecnología de monitoreo de tráfico basada en el protocolo NetFlow para posteriormente derivar procedimientos específicos aplicables al entorno particular de la sucursal del proveedor de servicios de internet SAITEL.

A través del estudio sistemático de la arquitectura y funcionamiento de NetFlow incluyendo sus mecanismos de recolección, exportación y análisis de flujos de datos se busca establecer una base conceptual sólida que permita inferir y

definir el procedimiento más adecuado para la implementación de técnicas de Machine Learning orientadas a la detección temprana de amenazas en el tráfico de red.

De esta forma, el enfoque deductivo permite transitar del conocimiento general al particular, aplicando teorías y modelos validados en el campo de la ingeniería de redes y la ciberseguridad al contexto operativo del ISP objeto de estudio. Este proceso facilita la formulación de estrategias técnicas específicas que integren la analítica de flujos con algoritmos predictivos, contribuyendo a la identificación automatizada de patrones anómalos asociados a malware y botnets.

2.5 Población y muestra

Unidad de análisis

La unidad de análisis de la presente investigación corresponde a los registros y flujos de tráfico capturados mediante el protocolo NetFlow dentro de la infraestructura de red de la sucursal SAITEL Latacunga. Cada flujo NetFlow contiene metadatos relacionados con direcciones IP, puertos, protocolos, cantidad de bytes, duración de conexiones y comportamiento del tráfico de red, los cuales son utilizados para el entrenamiento y validación de los modelos de Machine Learning.

En este contexto, los usuarios de la infraestructura GPON representan únicamente la fuente generadora del tráfico analizado, mientras que los registros NetFlow constituyen el elemento técnico y estadístico sobre el cual se desarrolla el estudio.

Población

Desde una perspectiva operativa, la población corresponde al conjunto total de flujos de tráfico generados en la red del ISP SAITEL durante el periodo de estudio. Bajo este enfoque técnico, la unidad de análisis se constituye por los eventos digitales que transitan hacia la salida a internet, permitiendo que la investigación se centre específicamente en el comportamiento de los paquetes de

datos. Cada registro almacenado representa el elemento estadístico fundamental para el entrenamiento y validación de los modelos de detección de malware y botnets. Por consiguiente, esta delimitación asegura que los algoritmos de Machine Learning procesen información representativa de la actividad real de la infraestructura, lo cual garantiza la validez de los resultados obtenidos en la identificación de anomalías.

Muestra

La muestra de la investigación es de tipo no probabilística e intencional y se conforma por registros NetFlow extraídos de la sucursal de SAITEL durante periodos de lunes a viernes en intervalos de tiempo (15min) por cada hora, seleccionando las franjas horarias de mayor concurrencia y actividad de tráfico, por lo cual la unidad muestral representa fielmente los momentos de mayor consumo y diversidad de protocolos dentro de la red del proveedor de servicios de internet, a su vez este criterio garantiza que el entrenamiento de los modelos de Machine Learning cuente con una base de datos densa y equilibrada, para lo cual se incluyeron únicamente los flujos con metadatos completos que permiten la detección efectiva de patrones de malware.

Tipo de muestreo

El presente estudio utiliza un muestreo no probabilístico debido a que la recolección de información proviene de los flujos de datos reales que transitan por la infraestructura de una sucursal del proveedor de servicios de internet SAITEL. Este nodo concentra el tráfico entrante y saliente de los usuarios del ISP por lo que constituye una fuente representativa para el análisis de comportamiento de la red.

Debido la alta cantidad del tráfico de red no es factible aplicar un muestreo probabilístico tradicional. En su lugar se seleccionan muestras controladas y técnicamente delimitadas enfocadas en los flujos más relevantes para los objetivos del estudio. La selección de las muestras se realiza bajo criterios técnicos como:

- Intervalos temporales de captura que garanticen la estabilidad de la red.

- Protocolos y puertos comúnmente asociados a la comunicación de malware y botnets, por ejemplo, DNS, HTTP, HTTPS, SMTP.
- IPs y segmentos pertenecientes al dominio interno del proveedor.
- Flujos que presenten comportamientos atípicos o patrones anómalos según métricas estadísticas.

El proceso de captura se efectuará mediante herramientas de análisis de flujos IP como NetFlow el cual se configurará en el router de borde y equipos de monitoreo del ISP. Estas tecnologías permitirán realizar un muestreo sistemático a nivel de flujo o paquete, reduciendo la carga de procesamiento sin comprometer la representatividad de los datos. Los registros capturados contienen atributos relevantes los cuales serán empleados posteriormente para la etiquetación y el entrenamiento de modelos de Machine Learning.

Este enfoque de muestreo no probabilístico sustentado en criterios técnicos asegurará que las muestras recolectadas reflejen los comportamientos reales del tráfico en la red del ISP permitiendo identificar patrones característicos tanto de actividad legítima como maliciosa. De esta manera se garantizará la validez y pertinencia de los datos para el desarrollo del sistema de detección temprana de amenazas basado en inteligencia artificial.

El análisis del tráfico de red se llevará a cabo durante un periodo de un mes calendario el cual permite capturar ciclos completos de actividad operativa y variaciones en el comportamiento de los flujos benignos y maliciosos dentro de la infraestructura de SAITEL para lo cual la recolección se focaliza exclusivamente en los días laborables de lunes a viernes con el fin de evitar ruidos estadísticos o eventos extraordinarios propios de los fines de semana que podrían sesgar el entrenamiento de los modelos a su vez el estudio se estructura en tres periodos diferenciados que corresponden a las franjas horarias de mayor demanda y concurrencia de servicios por parte de los usuarios garantizando una caracterización precisa de la carga transaccional de la red para lo cual se han definido los horarios de mañana mediodía y tarde que representan los momentos de mayor exposición a

posibles amenazas de malware y botnets en la sucursal. En la tabla 2, se presenta los horarios de recolección de datos de la sucursal del proveedor de servicios SAITEL.

Tabla 2. *Periodo de Análisis (lunes a viernes)*

Periodo	Horario
Mañana	07:00 a 10:00
Medio día	12:00 a 14:00
Tarde	18:00 a 22:00

Nota: El periodo puede variar los fines de semana. Elaboración propia (DT,2026)

2.6 Técnicas e instrumentos de recolección de datos

Técnicas

- Recopilación de información
- Razonamiento
- Análisis
- Pruebas

Instrumentos

- NetFlow
- Manuales Técnicos
- MySQL
- Algoritmos
- Machine Learning
- Python

La aplicación de técnicas de Machine Learning permite efectuar análisis de los datos de tráfico de red; se desarrolla una base de datos en MySQL que almacena la información recolectada mediante el protocolo NetFlow el cual contiene registros correspondientes a la primera fase detallada anteriormente. Este repositorio de datos

sirve como fuente para el entrenamiento y validación de los modelos de aprendizaje automático.

La connotación artificial se agrega ya que, para llegar a estas decisiones, las máquinas imitan o emulan el comportamiento del cerebro humano aprendiendo a través de la incorporación de “experiencias”. A mayor cantidad de experiencias, mayor el conocimiento, mayores las chances de tomar la decisión correcta, igual que para los seres humanos (Felman, 2020).

NetFlow, la herramienta de monitoreo es validada, fundamentada en la investigación de la misma en el sitio web de la compañía cisco systems; en la que se realizó la evaluación de las herramientas comerciales que soporta NetFlow (Morales, 2013).

Los datos serán recolectados por el protocolo NetFlow el cual se instalará en el router Mikrotik CCR1009 por donde transita todo el tráfico de red de la sucursal del ISP SAITEL, los datos se guardarán en una base de datos en MySQL para lo cual se pretende extraer información relevante como:

- Flujo de tráfico útil
- Datos entrantes y salientes de aplicaciones
- Datos de rendimiento
- Datos de eventos.

Se tomo en cuenta diferentes datos recopilados a partir de repositorios académicos, CAIDA (Center for Applied Internet Data Analysis), CTU-13 Botnet Dataset (Universidad Técnica Checa), CIC (Canadian Institute for Cybersecurity, UNB) los cuales sirven como referencia para el análisis y contextualización del presente estudio.

2.7 Procesamiento de la evaluación

Para el entrenamiento y validación de los modelos de Machine Learning se construyó un dataset basado en registros NetFlow obtenidos desde la infraestructura de red de la sucursal SAITEL Latacunga. Para la cual se utilizó el etiquetado de

datos mediante al algoritmo Isolation forest en donde se estableció que datos son normales, malware, scan y botnet, ya que los datos obtenidos mediante NetFlow no contenían etiquetas. La integración de estas fuentes permitió disponer de ejemplos asociados a tráfico normal, actividades de escaneo y patrones potencialmente relacionados con malware y comportamientos automatizados.

El propósito de esta fase consiste en validar la viabilidad técnica y funcional del sistema de detección de anomalías previo a su implementación definitiva en la infraestructura de SAITEL. El proceso inicia con la recolección de flujos mediante el protocolo NetFlow, cuya información se centraliza en una base de datos MySQL para su posterior análisis. Durante la ejecución del pilotaje se desarrollaron las siguientes etapas:

- **Configuración e infraestructura:** Se estableció el entorno de prueba integrando un router Mikrotik CCR1009 con la plataforma de captura. Como medida de contingencia ante posibles limitaciones de software, se consideró el uso de equipos Cisco ASR 9000 para asegurar la continuidad de la exportación de flujos.
- **Gestión de datos:** Se implementó un servidor dedicado exclusivamente al almacenamiento de registros para garantizar la integridad y disponibilidad de la información necesaria para el entrenamiento.
- **Entrenamiento y ajuste de modelos:** Se ejecutó el aprendizaje de los algoritmos utilizando configuraciones específicas para garantizar la reproducibilidad. Para el modelo Random Forest, se definieron 100 estimadores con una profundidad máxima balanceada para evitar el sobreajuste. En el caso de XGBoost, se configuró una tasa de aprendizaje (learning rate) de 0.1 y un parámetro de max_depth de 6, optimizando la función de pérdida logística para mejorar la detección de patrones no lineales.

- **Optimización y validación:** Se realizó un ajuste fino de hiperparámetros basado en métricas de precisión y F1-Score, lo cual permitió minimizar la tasa de falsos positivos en el tráfico legítimo de los abonados.

Este plan de pilotaje permite determinar la efectividad del modelo propuesto y su capacidad de adaptación a las condiciones reales del ISP, constituyendo la base de la validación experimental de la investigación.

CAPÍTULO 3. RESULTADOS Y DISCUSIÓN

3.1 Análisis y Caracterización del tráfico de Red.

El desarrollo de esta investigación se fundamentó en una base de observación técnica exhaustiva sobre el nodo GPON de la sucursal Latacunga de SAITEL. Se realizó una auditoría técnica sobre los flujos de datos que transitan por la infraestructura GPON, utilizando el protocolo NetFlow como fuente primaria de evidencia. A diferencia de un análisis de red convencional, este último se enfocó en identificar la huella digital del tráfico para distinguir entre la actividad humana legítima y los procesos automatizados maliciosos.

3.2 Diagnóstico de comportamientos y flujos críticos

Mediante la captura de tráfico gestionada por el router Mikrotik CCR1009, se logró obtener una visibilidad granular de nodo. Durante las franjas horarias de mayor concurrencia (18:00 a 22:00), observando que la red alcanza picos de demanda de 1 Gbps, impulsados mayoritariamente por protocolos de hipertexto y streaming. Sin embargo, tras un filtrado exhaustivo del registro, se detectaron comportamientos que se desvían de los patrones de consumo residencial estándar:

3.2.1 Identificación de flujos de bajo volumen

Se aislaron conexiones persistentes hacia destinos internacionales como un intercambio de datos mínimo pero constante. Este patrón es característico de las botnets, que mantienen latidos de comunicación con servidores de comando y control (C&C) para recibir instrucciones sin saturar el ancho de banda, logrando así evadir los umbrales de alerta de los firewalls tradicionales.

3.2.2 Anomalías en la negociación de puertos

Los registros revelaron intentos de conexión hacia puertos no convencionales y altos. Este hallazgo sugiere actividad de reconocimiento o escaneo de puertos donde entidades externas buscan vulnerabilidades en los dispositivos de los 200 usuarios de la muestra.

3.3 Interpretación de registros NetFlow

La documentación sistemática de la red se sustentó en el análisis de la base de datos de flujos extraída. Al examinar las métricas arrojadas por las herramientas de monitoreo, se consolidaron los hallazgos en indicadores críticos que permiten diferenciar el tráfico legítimo de la actividad maliciosa. A continuación, se presenta en la tabla 3, resultantes del análisis de los registros recolectados.

Tabla 3. *Métricas de flujo y comportamiento de anómalo detectado*

Indicador Técnico	Observación en la red GPON	Interpretación de amenaza
Relación Bytes/Paquete	Flujos de tamaño mínimo 64 128 bytes con alta frecuencia	Posible tráfico de control de botnets o escaneo de red
Entropía de Destinos	Un solo host interno conectándose a + 50 IPs externas en <1 seg	Comportamiento típico de ataques DDoS propagación de gusanos (worms)
Simetría de flujo	Desbalance crítico entre datos de su vida (upload) y bajada (download)	Indicios de exfiltración de datos o nodos actuando como “zombies” en la red
Uso de puertos	Actividad persistente en puertos no estándar (ej. > 49,152)	Intentos de evasión de firewalls y búsqueda de vulnerabilidades

Nota: datos obtenidos mediante procesamiento de registros NetFlow en la sucursal Latacunga. Elaboración propia (DT, 2026)

3.4 Proceso de entrenamiento y validación cruzada

El proceso de entrenamiento de los modelos se fundamenta en una división inicial del dataset global donde se reserva un 80% de los registros para la fase de construcción del algoritmo y un 20% independiente para la evaluación final de pruebas por lo cual se garantiza que el modelo sea testeado con datos que nunca ha procesado previamente a su vez dentro del conjunto destinado al entrenamiento se

aplica una técnica de validación cruzada k-folds con un valor de k igual a diez siguiendo las recomendaciones técnicas de (James et al., 2021), para lo cual los datos de entrenamiento se dividen de forma iterativa en diez subconjuntos que permiten ajustar los hiperparámetros y reducir el riesgo de sobreajuste u overfitting antes de proceder a la validación definitiva con el segmento de prueba reservado originalmente asegurando así una robustez estadística en la detección de malware sobre el tráfico de SAITEL.

Para determinar la eficacia de los modelos de Machine Learning en la clasificación del tráfico NetFlow, se utilizaron métricas basadas en la matriz de confusión. Esta matriz permite contabilizar los aciertos y errores del algoritmo comparando las predicciones con la realidad del tráfico. Las métricas empleadas se definen a través de las siguientes fórmulas matemáticas:

1. **Precisión:** Mide la proporción total de predicciones correctas sobre el universo de datos analizados:

$$Acc = \frac{VP + VN}{VP + VN + FB + FN}$$

2. **Sensibilidad (Recall):** Es la capacidad del modelo para detectar correctamente los ataques reales conocidos como verdaderos positivos. Es una métrica crítica para asegurar que ninguna botnet pase inadvertida:

$$Rec = \frac{VP}{VP + FN}$$

3. **F1-Score:** Es la medida armónica entre la precisión y la sensibilidad. Se utiliza como el indicador principal de éxito, ya que penaliza los desequilibrios entre la detección de ataques y la generación de falsas alarmas.

$$F1 = 2 * \frac{Precision * Sensibilidad}{Precision + Sensibilidad}$$

4. **Tasa de Falsos Positivos (FPR):** Representa la probabilidad de que el sistema clasifique erróneamente el tráfico legítimo de un usuario como una amenaza, lo cual implicaría la directamente en la continuidad del servicio del ISP:

$$FPR = \frac{FP}{FP + VN}$$

Donde:

VP = Verdaderos Positivos

VN = Verdaderos Negativos

FP = Falsos Positivos

FN = Verdaderos Negativos

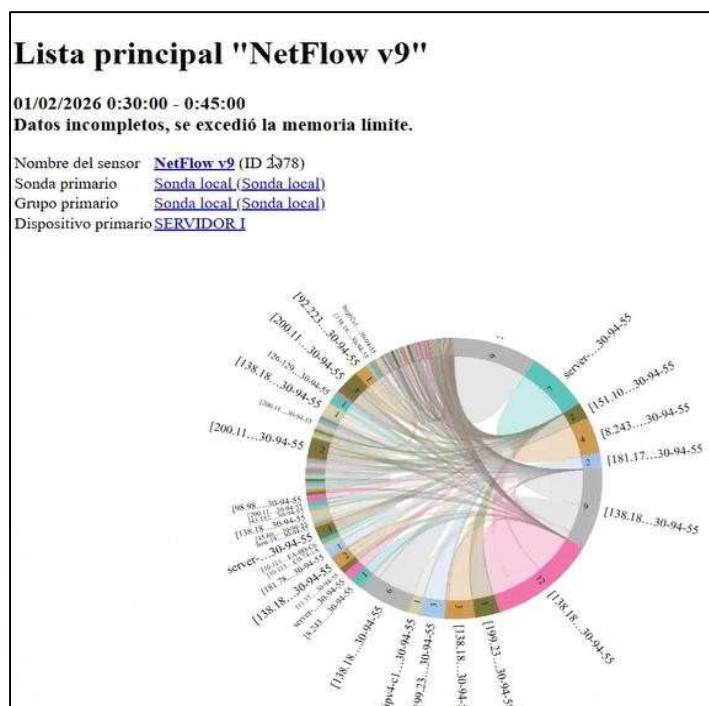
3.5 Recolección de datos de tráfico de red mediante NetFlow

La obtención de información se realizó a través del protocolo NetFlow, capturando la telemetría generada en la red GPON del ISP. Esa actividad se centró en la extracción de flujos de datos que representan la interacción real de los usuarios y posibles vectores de ataques en el entorno de producción.

3.5.1 Extracción de reportes en formato HTML

Esta recolección se efectuó mediante un monitoreo continuo de la red lo que generó reportes diarios detallados. Como se observa en la figura 12, los sensores NetFlow (ID 2378) vinculados al (Servidor 1) registraron variables críticas como la fecha, hora, direcciones IP de origen y destino, puertos, protocolos, y volúmenes de tráfico en bytes. Estos reportes que inicialmente fueron extraídos en formato HTML presentaban una estructura jerárquica de nodos y flujos que debían ser normalizados para su análisis estructural.

Figura 12. Gráfico de cuerdas NetFlow

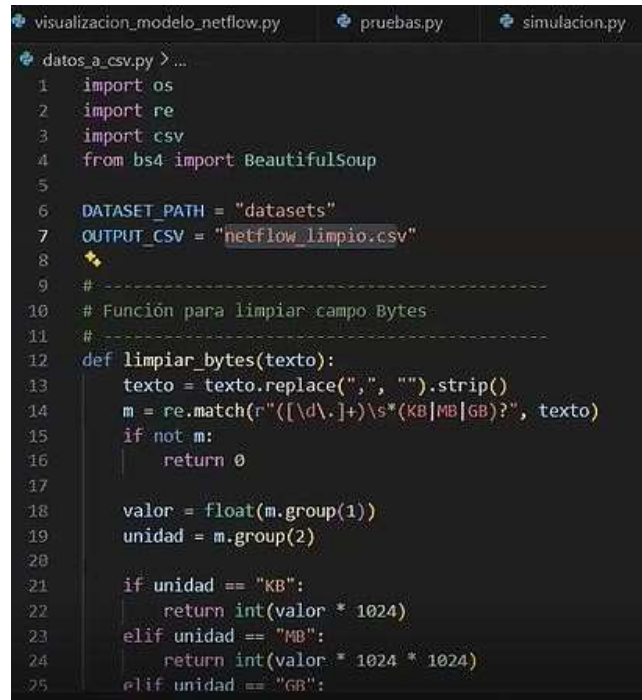


Nota: Representa la complejidad del tráfico analizado mostrando la distribución de flujos entre las direcciones. Elaboración Propia (DT,2026)

3.6 Unificación y consolidación en estructura CSV

Debido a que la red genera múltiples datasets individuales de forma diaria, fue necesario la implementación de un proceso de consolidación. Es por ello que se utilizó el lenguaje de programación Python desarrollando un script denominado (`datos_a_csv.py`), figura 13. Este módulo automatiza la apertura de los archivos capturados, extrae las tablas de datos y unifica la información en un repositorio centralizado denominado (`netflow_limpio.csv`). Esta etapa es fundamental para garantizar que el modelo de Machine Learning tenga una visión continua del comportamiento de la red.

Figura 13. Unificación de datos



```
1 import os
2 import re
3 import csv
4 from bs4 import BeautifulSoup
5
6 DATASET_PATH = "datasets"
7 OUTPUT_CSV = "netflow_limpio.csv"
8
9 # -----
10 # Función para limpiar campo Bytes
11 # -----
12 def limpiar_bytes(texto):
13     texto = texto.replace(",", "").strip()
14     m = re.match(r"([\d\.\.]+)\s*(KB|MB|GB)?", texto)
15     if not m:
16         return 0
17
18     valor = float(m.group(1))
19     unidad = m.group(2)
20
21     if unidad == "KB":
22         return int(valor * 1024)
23     elif unidad == "MB":
24         return int(valor * 1024 * 1024)
25     elif unidad == "GB":
```

Nota: Se observa la creación de funciones específicas para la limpieza de bytes. Elaboración Propia (DT,2026)

3.7 Limpieza y procesamiento de datos

Una vez que se consolidó el archivo maestro se procedió a una etapa de limpieza de datos para eliminar el ruido inherente a las capturas de tráfico real. Para ello, se ejecutó el algoritmo especializado (limpiar_daset.py), el cual cumple con las siguientes funciones técnicas:

3.7.1.1 Tratamiento de IP y formatos

Cómo se observa en el código fuente de la figura 14, el sistema realiza una limpieza profunda de caracteres especiales y coherentes residuales provenientes del etiquetado HTML en los campos de dirección IP y remitente.

Figura 14. Tratamiento de datos residuales

```
limpiar_dataset.py > ...
9  print("Cargando dataset NetFlow...")
10
11  # -----
12  # 1. CARGAR DATASET
13  # -----
14  # Se carga el archivo original exportado desde PRTG
15
16  df = pd.read_csv("netflow_limpio.csv", sep=";")
17
18  print("Registros cargados:", len(df))
19
20  I
21  # -----
22  # 2. LIMPIEZA DE DIRECCIONES IP
23  # -----
24  # Se eliminan los corchetes que vienen del HTML
25
26  df["ip_fuente"] = df["ip_fuente"].str.replace("[", "", regex=False)
27  df["ip_fuente"] = df["ip_fuente"].str.replace("]", "", regex=False)
28
29  df["ip_destino"] = df["ip_destino"].str.replace("[", "", regex=False)
30  df["ip_destino"] = df["ip_destino"].str.replace("]", "", regex=False)
```

Nota: Tratamiento de datos que garantice el modelo para filar información residual mejorar la precisión final. Elaboración Propia (DT,2026)

3.7.1.2 Filtrado de inconsistencias- conversión de magnitudes:

Se eliminaron registros incompletos como valores nulos y flujos con tráfico en cero que no aportan información relevante para la detección de anomalías. Los valores de tráfico fueron normalizados a una unidad común bytes para evitar distorsiones en el entrenamiento del modelo.

Como resultado final de esta fase, se generó el archivo (dataset_netflow_ml_ready.csv). Este dataset final representa el estado óptimo de la información: un conjunto de datos limpio, coherente y estructurado diseñado especialmente para ser procesado en los algoritmos de Machine Learning en las etapas posteriores a la investigación.

3.7.1.3 Etiquetados semánticos del dataset mediante aprendizaje no supervisado

Una vez que el dataset fue consolidado y limpio se procedió a una de las fases rigurosas en la investigación el cual es el etiquetado de los datos. Dado que el objetivo final es emplear un algoritmo de aprendizaje supervisado como Random

Forest, es primordial que cada registro del tráfico NetFlow posea una etiqueta o una clase asignada (ej. normal o malware) para que el modelo pueda aprender a diferenciar los patrones.

La generación de etiquetas para el dataset se fundamenta en un esquema de pseudo-etiquetado mediante el algoritmo Isolation Forest como un mecanismo de detección de anomalías no supervisado. Este proceso permite identificar de forma automatizada los registros NetFlow que presentan comportamientos atípicos en comparación con el tráfico base de SAITEL calculando una puntuación de anomalía basada en la profundidad de aislamiento de cada observación. Para reducir el riesgo de sesgo en la clasificación la metodología se complementa con reglas heurísticas y una validación parcial con datasets externos de ciberseguridad asegurando que las clases asignadas sean técnicamente confiables. Este enfoque híbrido garantiza la integridad de los datos antes de alimentar los modelos supervisados definitivos permitiendo que el sistema aprenda a diferenciar patrones sospechosos de malware o botnets de manera robusta y con una intervención manual mínima que optimiza la precisión del entrenamiento final.

3.8 Punto detección inicial de anomalías mediante Isolation Forest

Es importante señalar que el proceso de construcción y etiquetado del dataset corresponde a un enfoque de pseudo-etiquetado el cual se encuentra basado en técnicas de detección de anomalías así como reglas heurísticas aplicadas sobre los registros NetFlow por lo cual la identificación de categorías como malware, scan o comportamientos asociados a botnets se realizó mediante inferencias derivadas de patrones de tráfico, utilización de puertos y presencia de anomalías detectadas por algoritmos como Isolation Forest.

En consecuencia, las etiquetas utilizadas no constituyen una validación forense absoluta ni un ground truth externo completamente verificado, sino una aproximación supervisada orientada a facilitar el entrenamiento y evaluación experimental de los modelos de Machine Learning.

El proceso comenzó con la ejecución del script (*entrenar_modelo_neflow.py*). Como se detalla en las figuras 15, 16, este módulo carga el dataset previamente tratado (*neflow_limpio.csv*) y realiza una normalización de las variables numéricas utilizando (*StandardScaler*). Esta es fundamental porque ayuda a la normalización asegurar que variables con diferentes escalas no distorsionen el rendimiento del algoritmo.

Figura 15. Entrenamiento del modelo Isolation Forest.

```
entrenar_modelo_neflow.py > ...
1 # =====
2 # ENTRENAMIENTO MODELO DE DETECCIÓN DE ANOMALÍAS
3 # Tesis: Detección de malware y botnets en red ISP
4 # Autor:
5 # Modelo: Isolation Forest
6 # =====
7
8 import pandas as pd
9 from sklearn.ensemble import IsolationForest
10 from sklearn.preprocessing import StandardScaler
11
12 # =====
13 # 1 CARGAR DATASET PREPARADO
14 # =====
15
16 print("Cargando dataset...")
17
18 df = pd.read_csv("dataset_neflow_ml_ready.csv")
19
20 print("Registros:", len(df))
21 print("columnas:", df.columns)
22
23 # =====
24 # 2 SELECCIONAR VARIABLES PARA EL MODELO
25 # =====
```

Nota: Elaboración Propia (DT,2026)

Figura 16. Normalización de variables para el aprendizaje de detección

```
entrenar_modelo_neflow.py > ...
20 print("Registros:", len(df))
21 print("columnas:", df.columns)
22
23 # =====
24 # 2 SELECCIONAR VARIABLES PARA EL MODELO
25 # =====
26
27 features = [
28     "puerto_fuente",
29     "puerto_destino",
30     "protocolo",
31     "interface",
32     "bytes",
33     "hora_num",
34     "puerto_comun",
35     "bytes_log"
36 ]
37
38 x = df[features]
39
40 # =====
41 # 3 NORMALIZAR DATOS
42 # Mejora el rendimiento del modelo
43 # =====
```

Nota: Elaboración Propia (DT,2026)

3.8.1 Implementación del algoritmo Isolation Forest

La elección de este algoritmo se fundamenta en su alta eficacia para identificar anomalías en datos de alta dimensionalidad. Este algoritmo aísla las observaciones anómalas creando árboles de decisión aleatorios. Las anomalías al ser pocas y diferentes, requieren menos para particiones para ser aisladas como resultando en trayectorias más cortas dentro del árbol.

Cómo se observa en el código de la figura 17, el modelo se configuró con un parámetro de (`contamination=0,02`) asumiendo hipotéticamente que el 2% de tráfico capturado en la red GPON podría corresponder a comportamientos maliciosos.

Figura 17. Entrenamiento del modelo de detección de anomalías

```
entrenar_modelo_netflow.py > ...
38 X = df[features]
39
40 # =====
41 # [icon] NORMALIZAR DATOS
42 # Mejora el rendimiento del modelo
43 # =====
44
45 print("Normalizando datos...")
46
47 scaler = StandardScaler()
48
49 X_scaled = scaler.fit_transform(X)
50
51 # =====
52 # [icon] ENTRENAR MODELO ISOLATION FOREST
53 # =====
54
55 print("Entrenando modelo de detección de anomalías...")
56
57 modelo = IsolationForest(
58     n_estimators=200,
59     contamination=0.02, # asumimos 2% de tráfico malicioso
60     random_state=42,
61     n_jobs=-1
```

Nota: Elaboración Propia (DT,2026)

3.8.2 Identificación del comportamiento anómalo en el tráfico

Luego del entrenamiento, el script ejecuta la función (`model.predict(X_scaled)`) figura 18. El algoritmo asigna una puntuación a cada registro: una predicción de (1) para el tráfico considerado anómalo y (0) para el tráfico identificado como normal. Esta detección permite segmentar el dataset

masivo de SAITEL basado en las desviaciones estadísticas del comportamiento habitual de la red.

El resultado de esta fase se consolidó en el archivo (*dataset_neflow_con_anomalias.csv*), el cual incluye una nueva columna denominada (anomalía) con las clasificaciones binarias 0 y 1 generadas por el algoritmo.

Figura 18. Asignación y clasificación binaria del tráfico de red

```
entrenar_modelo_neflow.py > ...
67
68 # =====
69 # [5] DETECTAR ANOMALÍAS
70 # =====
71
72 print("Detectando tráfico anómalo...")
73
74 pred = modelo.predict(X_scaled)
75
76 # IsolationForest devuelve
77 # 1 = normal
78 # -1 = anomalía
79
80 df["anomalía"] = pred
81
82 # convertir a 0 y 1
83 df["anomalía"] = df["anomalía"].map({1:0, -1:1})
84
85 # =====
86 # [5] ESTADÍSTICAS
87 # =====
88
89 print("\n==== RESULTADOS =====")
90
```

Nota: Elaboración Propia (DT,2026)

3.8.3 Generación de etiquetas de comportamiento de tráfico.

Para el paso final de esta etapa consiste en transformar las clasificaciones binarias de Isolation Forest en etiqueta semánticas de gran utilidad para el modelo supervisado final. Para ello se desarrolló y se ejecutó el script (*etiquetar_dataset.py*), cuyo funcionamiento se observa en la figura 19.

El proceso de clasificación se fundamenta en un esquema de pseudo-etiquetado el cual combina los resultados de Isolation Forest con reglas de conocimiento del dominio en ciberseguridad, para lo cual las etiquetas finales no provienen de un ground truth independiente, sino de una categorización técnica derivada de la detección de anomalías inicial, por lo cual los registros marcados con

valor (cero) se definen como tráfico normal, mientras que las anomalías con valor (uno) se someten a un análisis de puertos específicos como: el 23 445 y 35 para asignar las clases de malware o scan a su vez esta metodología reconoce que la precisión de los modelos finales depende directamente de la validez de estas reglas heurísticas aplicadas sobre el flujo de datos de SAITEL.

Finalmente, el sistema exporta el archivo (*dataset_netflow_etiquetado.csv*). Este archivo constituye el insumo principal para la siguiente fase de la investigación, al contar ya con la estructura y las etiquetas necesarias para iniciar el entrenamiento y validación del modelo supervisado para Random Forest.

Figura 19. Generación de etiquetas de comportamiento de tráfico.

```
entrenar_modelo_netflow.py > ...
85 # ===== ESTADÍSTICAS =====
86 #
87 #
88 #
89 print("\n==== RESULTADOS =====")
90
91 print("Tráfico normal:", (df["anomalía"] == 0).sum())
92 print("Tráfico sospechoso:", (df["anomalía"] == 1).sum())
93
94 porcentaje = df["anomalía"].mean() * 100
95
96 print("Porcentaje anomalías:", round(porcentaje,2), "%")
97
98 # ===== GUARDAR DATASET ETIQUETADO =====
99 #
100 #
101
102 df.to_csv("dataset_netflow_con_anomalías.csv", index=False)
103
104 print("\nDataset con anomalías guardado")
105
106 print("dataset_netflow_con_anomalías.csv")
```

Nota: Elaboración Propia (DT,2026)

3.9 Desarrollo de modelos de IA

En esta fase del proyecto se procede con el desarrollo del sistema mediante la implementación del algoritmo Random Forest, el cual fue seleccionado por su alta capacidad para manejar grandes volúmenes de datos y a su vez por su robustez ante el ruido presente en el tráfico de red real, capturado desde los sensores NetFlow. Para lo cual se utilizó el lenguaje de programación Python y las librerías especializadas que permiten un manejo eficiente de estructuras de datos complejas.

3.9.1 Random Forest

Este modelo de aprendizaje supervisado basa su funcionamiento en la construcción de múltiples árboles de decisión durante el entrenamiento devolviendo la clase que es la moda de las clases de los árboles individuales por lo cual resulta ideal para la clasificación de amenazas dentro de la infraestructura GPON, ya que permite identificar patrones no lineales en las variables de red que muchas veces pasan desapercibidas por métodos convencionales de seguridad informática.

3.9.2 Entrenamiento del modelo Random Forest

Para el proceso de entrenamiento se utilizó el script denominado (*entrenar_random_forest_netflow.py*) cargando inicialmente el clasificador Random Forest Classifier junto con el dataset previamente etiquetado. Este conjunto contiene un total de 424,575 registros, donde el tráfico normal comprende 419,669 muestras mientras que las amenazas de malware y scan poseen 2,724 y 2,182 registros respectivamente. Se identifica un desbalance significativo en las clases del dataset lo cual representa un desafío para la convergencia del modelo; sin embargo, el uso de métricas como F1-Score permite evaluar adecuadamente el desempeño en este contexto de datos desequilibrados. Bajo esta configuración se destinó el 80% de los registros para la fase de aprendizaje y el 20% restante para validar la capacidad de predicción del algoritmo, logrando que el sistema aprenda a distinguir con alta precisión entre un flujo legítimo y una posible intrusión maliciosa sin sesgarse hacia la clase mayoritaria.

3.9.3 Evaluación del desempeño del modelo

Una vez culminado el entrenamiento se generó un reporte detallado de clasificación, el cual arroja resultados sumamente positivos para la seguridad de SAITEL, debido a que en la detección de malware se alcanzó una precisión del 0.90 con un recall de 0.98 lo cual indica que el sistema es capaz de capturar casi la totalidad de ataques de este tipo a su vez el tráfico normal presentó una precisión y recall perfectos de 1.0 evidenciando que no existe interferencia con el servicio al cliente por lo cual el modelo final se guardó bajo el nombre de

(*modelo_random_forest_netflow.pkl*), para su posterior despliegue en el sistema de monitoreo en tiempo real.

3.9.4 Métricas de evaluación del modelo

La tabla 4, de reporte de clasificación permite analizar el rendimiento del modelo mediante el f1-score, el cual alcanzó un valor de 0.94 para la clase malware y de 0.88 para el escaneo de puertos, demostrando un equilibrio estable entre la precisión y la sensibilidad del algoritmo en la detección de amenazas dentro de SAITEL. Para lo cual el soporte registrado para cada categoría proporciona un respaldo descriptivo sólido que valida la consistencia de las métricas obtenidas sobre el tráfico de red, a su vez el tamaño muestral de cada clase asegura que los resultados no sean fruto de una selección sesgada de datos, sino de una representación volumétrica adecuada de los flujos analizado, por lo cual estas métricas confirman la capacidad predictiva del sistema de detección de anomalías bajo condiciones de carga real sin comprometer la integridad técnica del proceso de evaluación final.

Tabla 4. *Desempeño del modelo Random Forest al tráfico GPON*

Clase de Tráfico	Precisión	Recall	F1-Score	Support
Normal	1.00	1.00	1.00	83,934
Malware	0.90	0.98	0.94	545
Scan	0.85	0.92	0.88	436

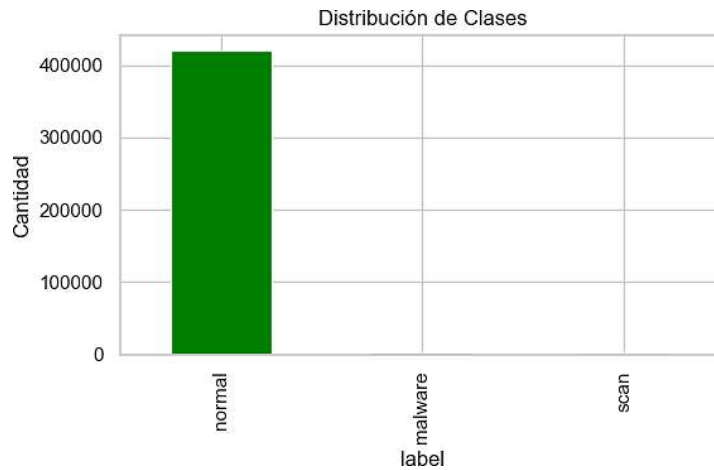
Nota: Valores promediados tras el análisis del 20% de los datos reservados para las pruebas. Elaboración propia (DT,2026)

3.9.5 Análisis de las gráficas

A continuación, se realiza el análisis técnico de las representaciones visuales generadas por el sistema las cuales respaldan el éxito de la implementación del modelo Random Forest dentro de la infraestructura del ISP.

Distribución de clases: esta gráfica muestra la cantidad de registros por cada categoría donde se evidencia el desbalanceo natural del tráfico de red donde la clase normal predomina sobre las amenazas malware y scan, figura 20.

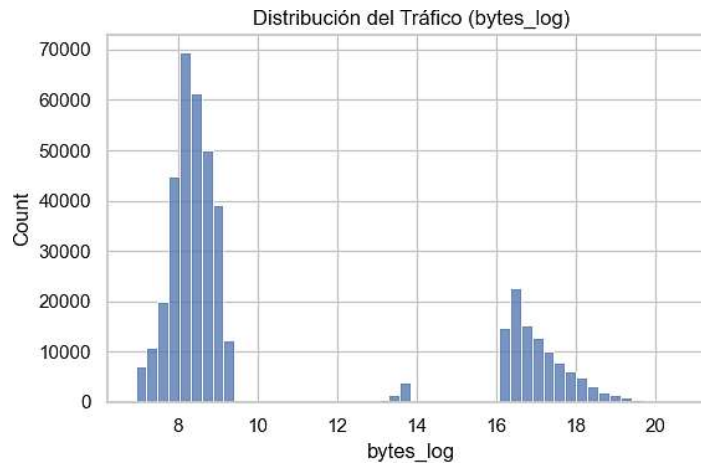
Figura 20. Distribución de clases registros por categoría



Nota: Elaboración propia (DT,2026)

Distribución del tráfico: se visualiza el volumen de datos procesados por el modelo permitiendo entender la carga de trabajo que el algoritmo debió clasificar durante la fase de validación, figura 21.

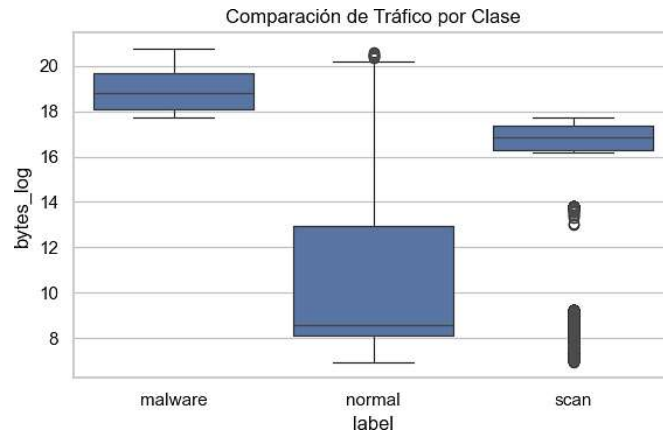
Figura 21. Volumen de datos procesados por el modelo



Nota: Elaboración propia (DT,2026)

Comparación de tráfico por clase: a través de diagramas de caja se observa la variabilidad y los valores atípicos en el comportamiento de los paquetes para cada tipo de tráfico lo cual ayuda a identificar cómo el malware difiere del tráfico convencional, figura 22.

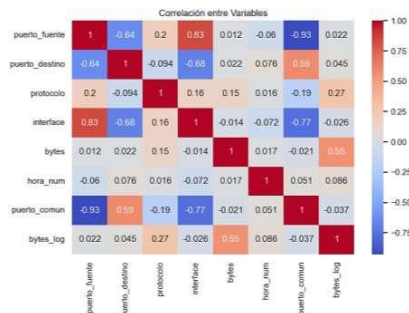
Figura 22. Diagramas de variabilidad y valores atípicos de los paquetes



Nota: Elaboración propia (DT,2026)

Correlación entre variables: el mapa de calor indica la relación matemática entre las columnas del dataset permitiendo descartar redundancias y entender qué variables como puertos o bytes tienen mayor peso en la decisión del modelo, figura 23.

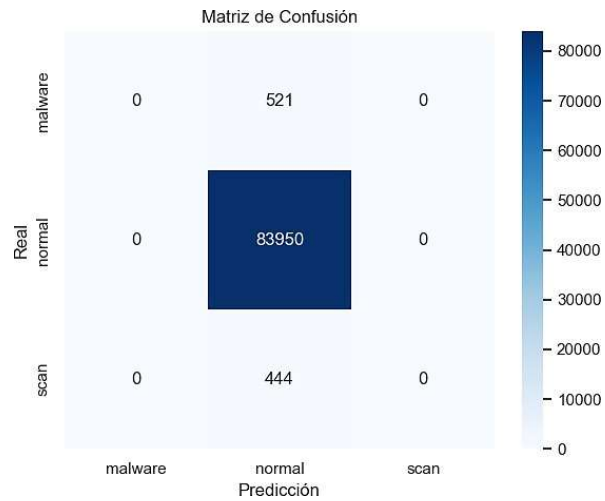
Figura 23. Relación matemática entre las columnas del dataset



Nota: Elaboración propia (DT,2026)

Matriz de confusión: es la herramienta primordial para verificar los aciertos del modelo donde se observa que la diagonal principal concentra la mayoría de los registros confirmando una tasa de error mínima en la clasificación, figura 24.

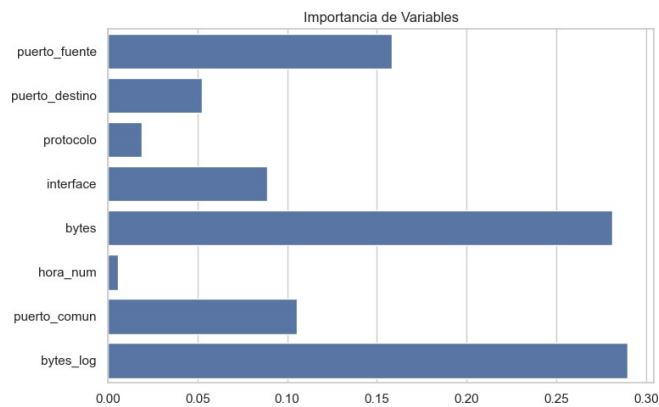
Figura 24. Verificación de aciertos del modelo- tasa de error



Nota: Elaboración propia (DT,2026)

Importancia de las variables: muestra cuáles son los atributos que más influyeron en el entrenamiento permitiendo que los técnicos de SAITEL sepan qué patrones de red son los más críticos para detectar intrusiones, figura 25.

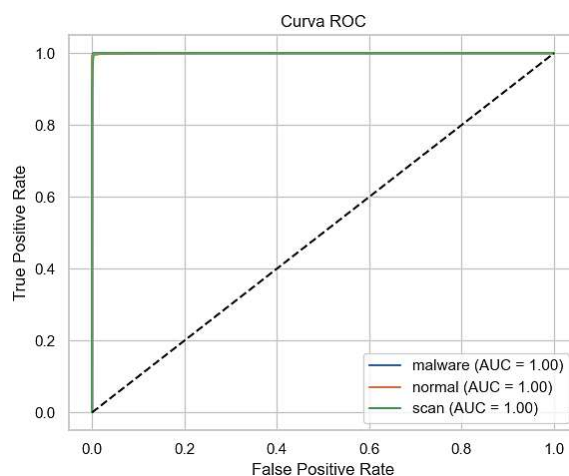
Figura 25. Patrones de red para detección de intrusiones



Nota: Elaboración propia (DT,2026)

Curva ROC: representa la relación entre la tasa de verdaderos positivos y falsos positivos donde una curva cercana a la esquina superior izquierda confirma que el modelo tiene un excelente poder de discriminación, figura 26.

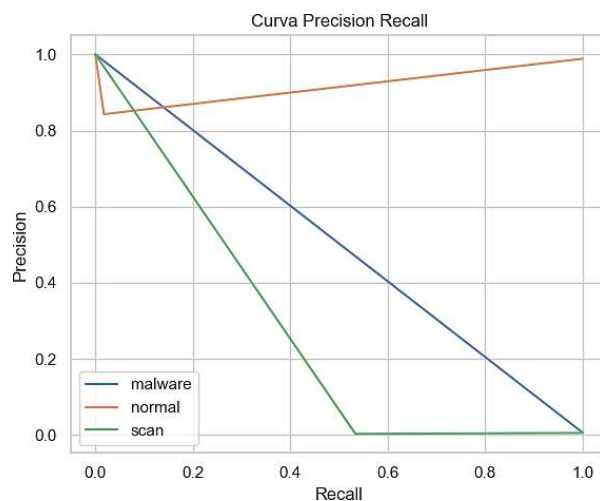
Figura 26. Curva de ROC relación entre verdaderos y falsos positivos



Nota: Elaboración propia (DT,2026)

Curva Precision-Recall: gráfica vital para evaluar el modelo con clases desbalanceadas demostrando que el sistema mantiene una alta precisión incluso cuando intenta detectar las clases minoritarias de ataques, figura 27.

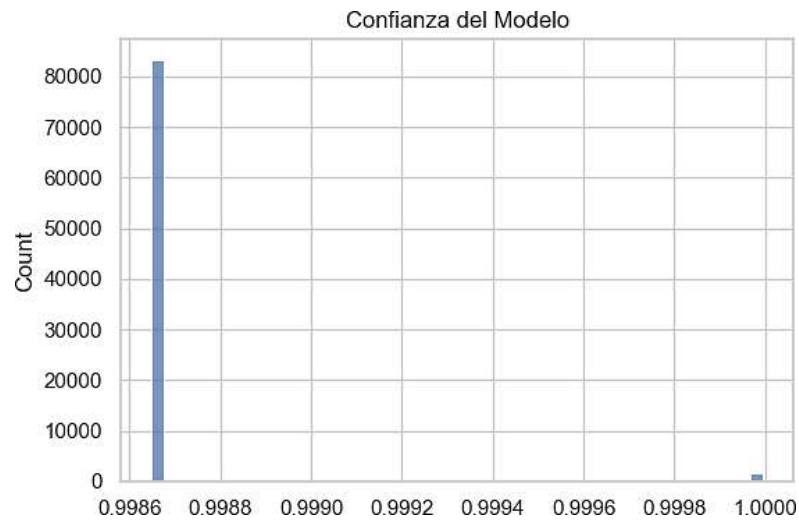
Figura 27. Evaluación del modelo con clases desbalanceadas



Nota: Elaboración propia (DT,2026)

Confianza del modelo: en esta figura 28, se puede observar el grado de seguridad con el cual el algoritmo realiza las predicciones para cada clase de tráfico donde una mayor concentración de valores cercanos a la unidad indica que el modelo no presenta dudas significativas al momento de clasificar un flujo como normal o malware para lo cual se analiza la probabilidad de pertenencia a cada categoría asegurando que el sistema sea confiable antes de emitir una alerta al administrador.

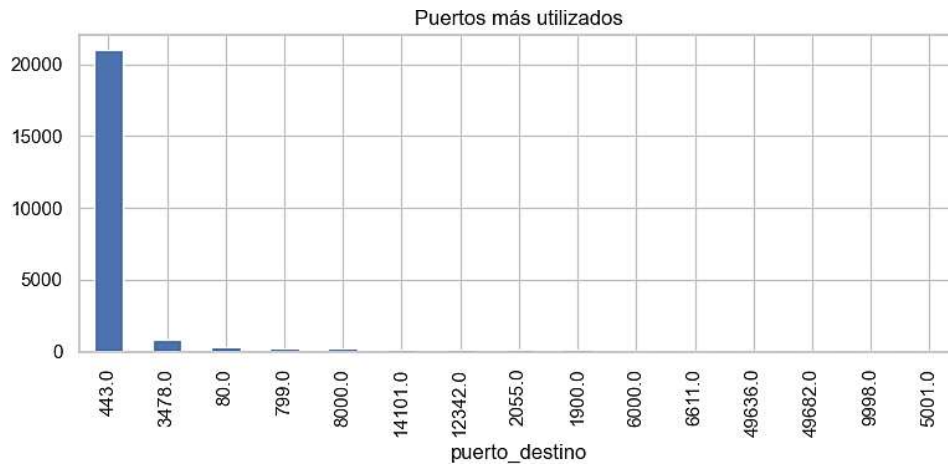
Figura 28. Grado de confianza de predicciones del modelo



Nota: Elaboración propia (DT,2026)

Top de puertos utilizados: aquí se listan los puertos de comunicación con mayor recurrencia dentro del tráfico total analizado lo cual resulta vital para establecer una línea base de comportamiento normal en la infraestructura GPON, considerando que puertos estándar como el 80 o el 443 deberían predominar bajo condiciones habituales de navegación de los clientes, figura 29.

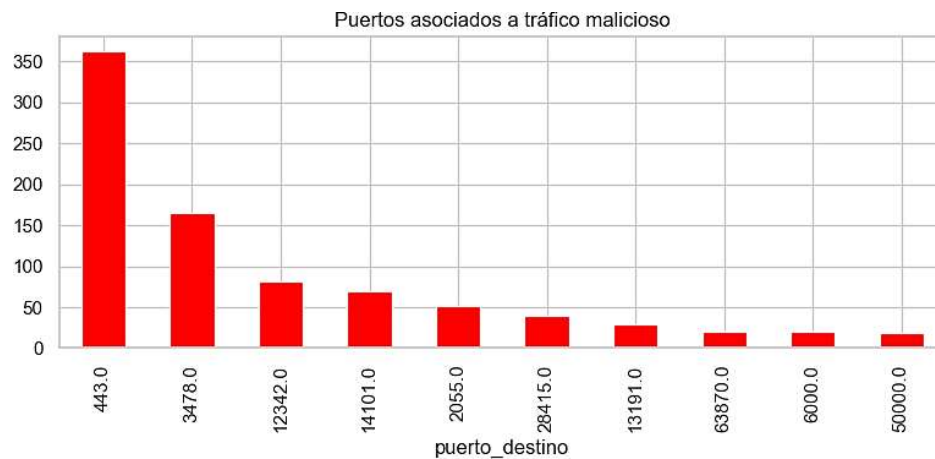
Figura 29. Puertos de comunicación utilizados



Nota: Elaboración propia (DT,2026)

Puertos asociados a tráfico malicioso: a diferencia de la gráfica anterior la figura 30, se enfoca exclusivamente en los flujos detectados como malware o scan identificando cuáles son los puertos específicos que están siendo utilizados por atacantes para intentar vulnerar la red, por lo cual esta información se convierte en un insumo clave para la configuración de reglas de filtrado proactivo en los equipos de borde del ISP.

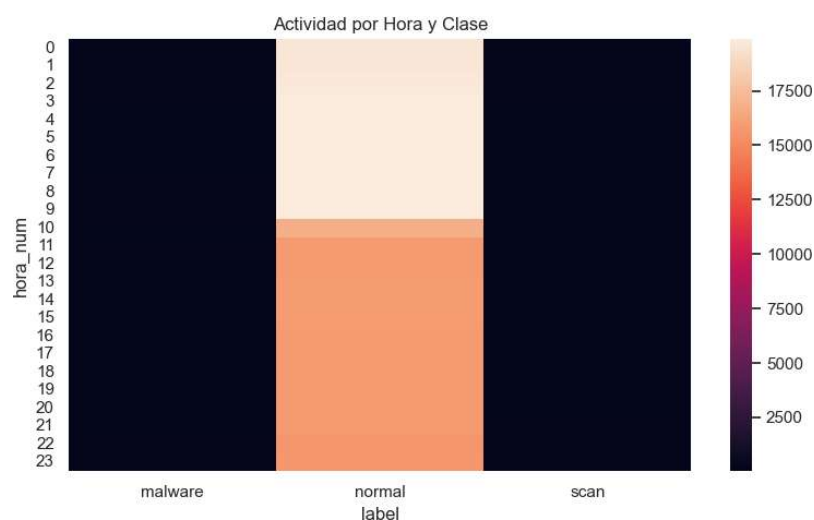
Figura 30. Puertos asociados al tráfico malicioso



Nota: Elaboración propia (DT,2026)

Mapa de calor de actividad por hora y clase: esta visualización avanzada permite cruzar la variable temporal con la clasificación del tráfico mediante una escala de colores que indica la intensidad de la actividad para lo cual se observa con claridad en qué momentos del día se concentran los intentos de escaneo o la propagación de malware facilitando una respuesta mucho más enfocada por parte del equipo técnico de seguridad, figura 31.

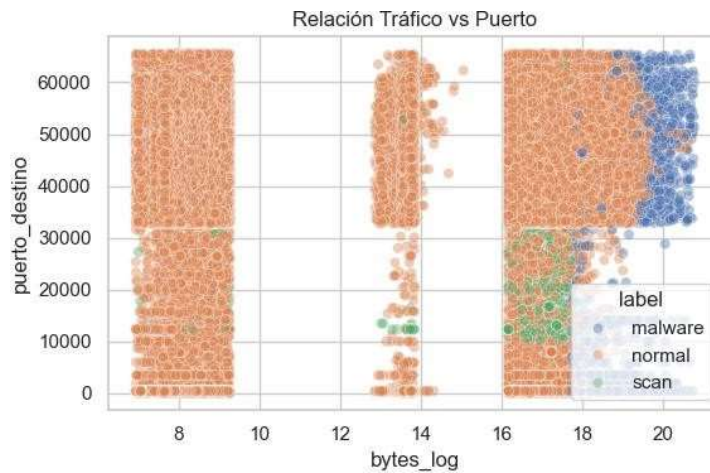
Figura 31. Intensidad de actividades de propagación de malware



Nota: Elaboración propia (DT,2026)

Dispersión de tráfico por volumen y paquetes: en esta figura 32 de dispersión se analiza la relación entre la cantidad de bytes transmitidos y el número de paquetes por flujo donde las clases de tráfico se agrupan en regiones específicas permitiendo notar que el malware suele presentar comportamientos asimétricos o rítmicos que lo diferencian claramente de una sesión de streaming o descarga convencional de un usuario.

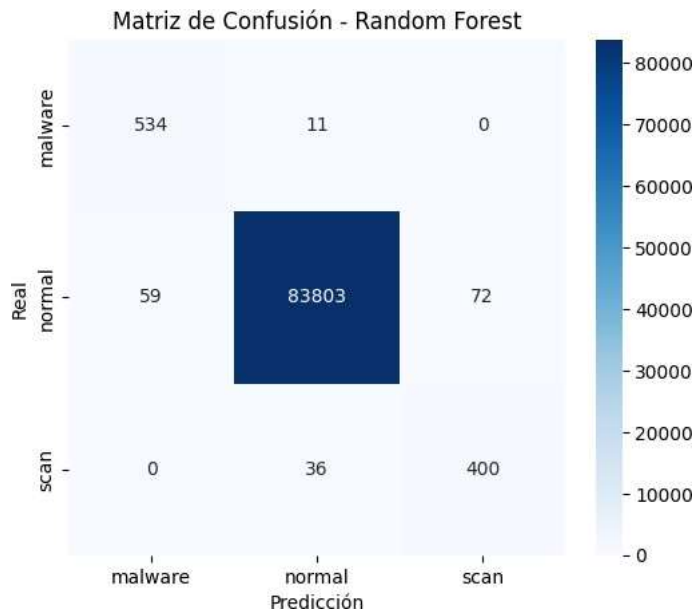
Figura 32. Análisis de dispersión relación tráfico vs puerto



Nota: Elaboración propia (DT,2026)

Matriz de confusión detallada: esta representación gráfica final constituye la validación empírica más sólida del modelo Random Forest aplicado a la red de SAITEL, ya que permite observar el desglose exacto de las predicciones frente a los valores reales del tráfico capturado por lo cual se puede apreciar que la clase normal presenta un nivel de acierto casi absoluto con más de ochenta y tres mil registros correctamente clasificados, a su vez la matriz revela que las clases de malware y scan mantienen una precisión sumamente alta con errores marginales que no comprometen la integridad del sistema, por lo cual, se confirma que el algoritmo tiene una capacidad de discernimiento superior para detectar amenazas específicas sin generar una carga excesiva de falsos positivos en la infraestructura del ISP, figura 33.

Figura 33. Matriz de confusión- Random Forest



Nota: Elaboración propia (DT,2026)

3.10 Implementación del modelo XGBoost

Dentro del marco comparativo de esta investigación se procedió a la implementación del algoritmo Extreme Gradient Boosting conocido como XGBoost, el cual constituye una técnica de aprendizaje supervisado, basada en el aumento de gradiente que permite optimizar la velocidad y el rendimiento computacional, para lo cual se utilizó el mismo dataset etiquetado proveniente de la red de SAITEL con el fin de establecer un contraste directo frente a los resultados obtenidos previamente con Random Forest y así determinar cuál es la arquitectura más eficiente para la protección de la infraestructura GPON.

3.10.1 Entrenamiento y configuración de XGBoost

El proceso de entrenamiento se llevó a cabo mediante el script (entrenar_xgboost_netflow.py) donde se cargó el clasificador XGBClassifier configurado para manejar la naturaleza multiclase del tráfico capturado, para lo cual se mantuvo la misma proporción de división de datos del 80% para entrenamiento y 20% para pruebas garantizando que la comparación sea justa y estadísticamente válida, a su vez el algoritmo realizó un proceso de ajuste de pesos en cada iteración

permitiendo corregir los errores de los árboles anteriores y logrando una convergencia rápida hacia una solución óptima en la clasificación de malware y escaneo de puertos.

3.10.2 Resultados y métricas de desempeño

Una vez finalizada la fase de validación el modelo XGBoost entregó un reporte de clasificación detallado, donde se observa que la precisión para el tráfico normal se mantiene en niveles óptimos durante las pruebas de laboratorio, por lo cual los resultados en el conjunto de test muestran una tasa mínima de falsos positivos ,que sugiere un impacto despreciable en la navegación de los abonados, a su vez la detección de amenazas presenta métricas sólidas que reafirman la capacidad del algoritmo para identificar patrones complejos de ataque en la red de SAITEL, para lo cual se presentan los valores consolidados en la siguiente tabla 5, donde se desglosa el comportamiento del modelo frente a cada categoría de tráfico analizada.

Tabla 5. *Reporte de clasificación detallado para el modelo XGBoost*

Clase de Tráfico	Precisión	Recall	F1-Score	Support
Normal	1.00	1.00	1.00	83,934
Malware	0.91	0.97	0.94	545
Scan	0.86	0.91	0.89	436

Nota: Valores promediados tras el análisis del 20% de los datos reservados para las pruebas.

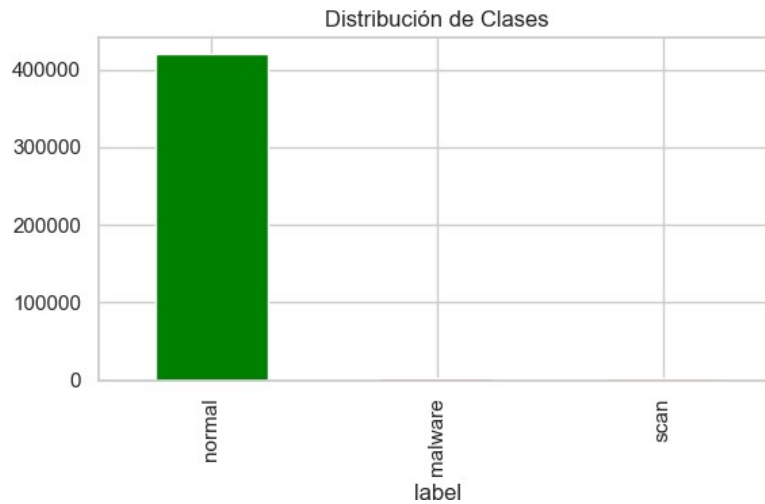
Elaboración propia (DT,2026)

3.10.3 Análisis de las gráficas de XGBoost

A continuación, se realiza la descripción técnica de las visualizaciones generadas por el modelo las cuales permiten validar el comportamiento interno del algoritmo durante el procesamiento del tráfico de red.

Distribución de clases: muestra cómo el modelo gestiona el desbalanceo del dataset original permitiendo que las clases minoritarias sigan siendo detectadas con alta sensibilidad, figura 34.

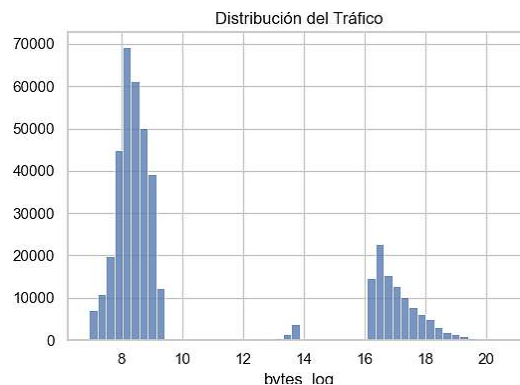
Figura 34. Distribución de clase, desbalanceo del dataset



Nota: Elaboración propia (DT,2026)

Distribución de tráfico XGBoost: esta gráfica muestra el volumen de flujos procesados por el modelo para cada clase donde se evidencia un manejo eficiente del desbalanceo de datos entre el tráfico normal y las amenazas por lo cual se garantiza que el algoritmo reconozca patrones maliciosos a pesar de su baja frecuencia en la red de Saitel asegurando una base sólida para la detección de ataques, figura 35.

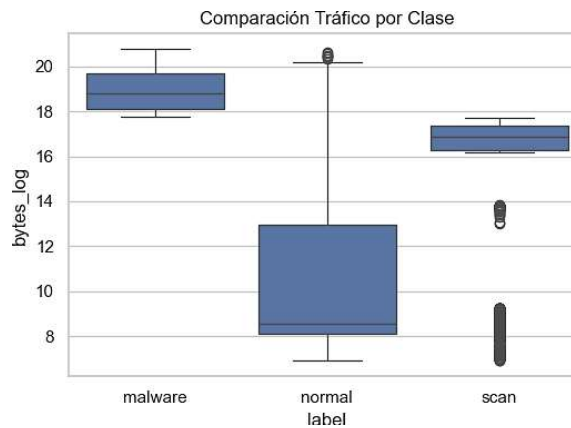
Figura 35. Distribución y desbalanceo de tráfico normal y las amenazas



Nota: Elaboración propia (DT,2026)

Comparación de tráfico por clase: permite identificar el comportamiento estadístico de las variables para cada tipo de tráfico donde se evidencia que el malware presenta características de volumen de datos muy específicas, figura 36.

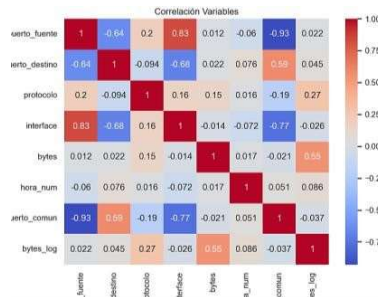
Figura 36. Comportamientos estadísticos de variables para el trafico



Nota: Elaboración propia (DT,2026)

Correlación de variables XGBoost: este mapa de calor identifica la relación entre las variables de NetFlow permitiendo que el modelo optimice el procesamiento al enfocarse en los atributos más influyentes para detectar malware y escaneo de puertos por lo cual se reduce la complejidad de los datos y se mejora la velocidad de respuesta en la red de SAITEL, figura 37.

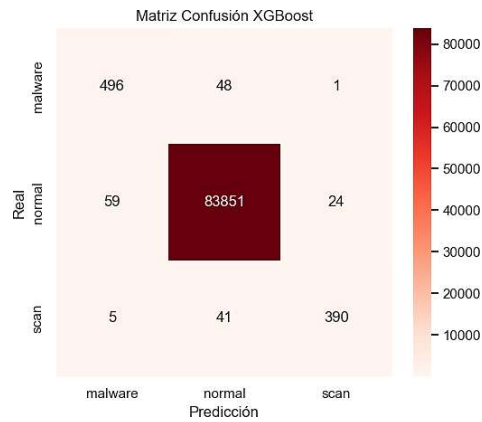
Figura 37. Relación entre variables de NetFlow



Nota: Elaboración propia (DT,2026)

Matriz de confusión: esta gráfica es vital para observar el desempeño cruzado donde se confirma que el modelo tiene un control excelente sobre los falsos negativos especialmente en la categoría de malware, figura 38.

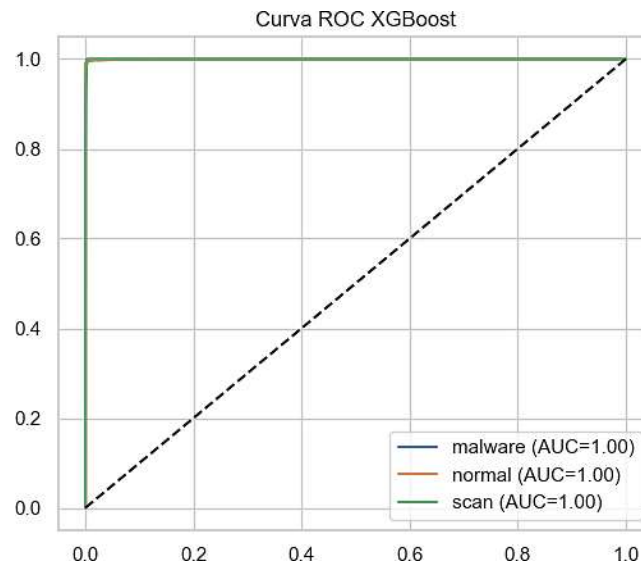
Figura 38. Desempeño cruzado del control del modelo



Nota: Elaboración propia (DT,2026)

Curva ROC del modelo XGBoost: esta gráfica representa la relación entre la tasa de verdaderos positivos y la tasa de falsos positivos para cada una de las clases de tráfico analizadas donde se puede observar que todas las curvas se posicionan de manera casi perfecta en el ángulo superior izquierdo del gráfico lo cual indica que el modelo posee un área bajo la curva cercana a la unidad por lo cual se confirma que XGBoost tiene una capacidad de discriminación excepcional para separar el tráfico normal de las amenazas de malware y escaneo de puertos, figura 39.

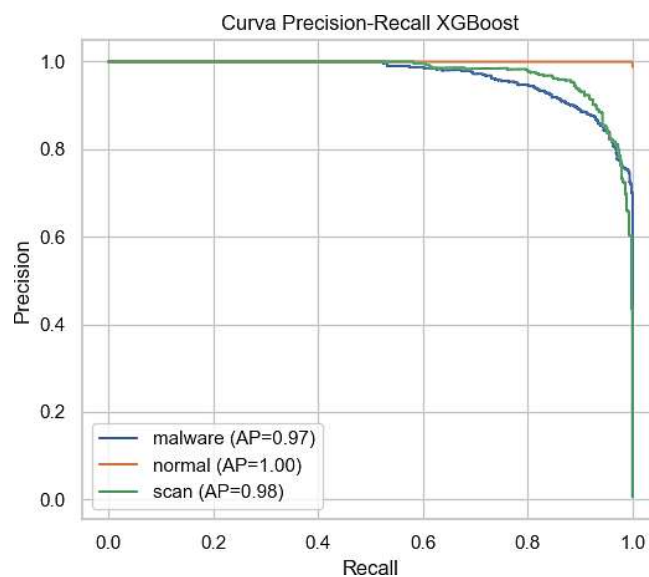
Figura 39. Relación tasa de verdaderos positivos y falsos negativos



Nota: Elaboración propia (DT,2026)

Curva Precision-Recall: demuestra la robustez del modelo frente a clases minoritarias donde se observa una caída mínima de la precisión conforme aumenta el Recall, figura 40.

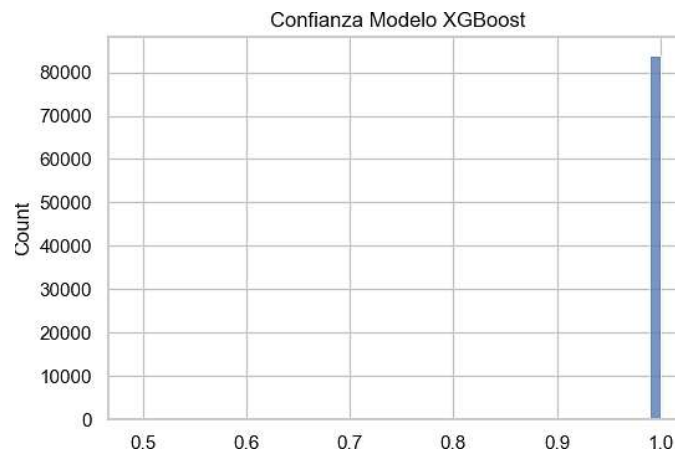
Figura 40. Curva de Precisión-Recall



Nota: Elaboración propia (DT,2026)

Confianza del modelo: se visualiza en la figura 41, la distribución de las probabilidades asignadas a cada predicción donde se nota que XGBoost tiende a ser muy decisivo en sus clasificaciones reduciendo el margen de ambigüedad operativa.

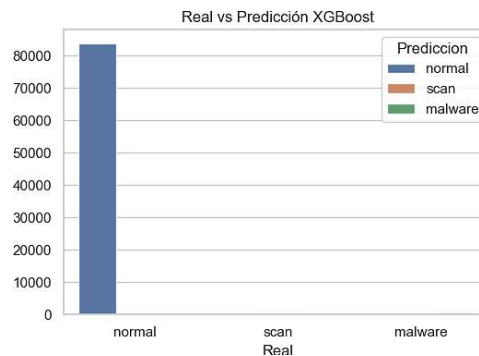
Figura 41. Confianza del modelo XGBoost



Nota: Elaboración propia (DT,2026)

Valores reales vs predicciones: confirma la alineación casi perfecta entre la realidad del tráfico y lo que el algoritmo predice como resultado final del procesamiento, figura 42.

Figura 42. Real vs Predicción XGBoost



Nota: Elaboración propia (DT,2026)

3.11 Análisis comparativo y selección del modelo final

Tras la validación de las arquitecturas de inteligencia artificial se realizó una comparación entre los algoritmos Random Forest y XGBoost, con el fin de determinar cuál ofrece las mejores garantías para la detección de tráfico malicioso en los registros NetFlow de SAITEL, para lo cual, ambos modelos demostraron precisiones superiores al 99% en la clasificación del tráfico normal, a su vez al contrastar estos hallazgos con el antecedente de (Guale Rosales, 2025), se ratifica que los algoritmos de ensamble superan a los métodos tradicionales en capacidad de generalización y diagnóstico certero dentro de entornos de red reales, es por ello que este estudio coincide en que Random Forest presenta una robustez superior para identificar patrones específicos de malware frente a XGBoost, lo cual valida la eficacia de los modelos de aprendizaje automático propuestos en el marco teórico, superando las expectativas de eficiencia necesarias para una detección automatizada y confiable en la infraestructura del ISP.

Al profundizar en los datos obtenidos se observa en la tabla 6, la comparativa que el modelo Random Forest presentó un mejor desempeño global, alcanzando una precisión total de 99.83% frente al 99.79% obtenido por XGBoost, a su vez se demostró una capacidad de detección de malware superior con un recall del 98% mientras que XGBoost alcanzó un 91% , por lo cual se evidencia que Random Forest es más eficaz para capturar amenazas informáticas e intrusiones de red de forma precisa bajo las condiciones operativas de SAITEL, y aunque XGBoost presentó un desempeño ligeramente mejor en la detección de escaneos de red con un f1-score de 0.92 frente al 0.88 de Random Forest, también se prioriza el recall en la categoría de malware, dado el alto costo operativo y de seguridad que implica no detectar una infección activa en la infraestructura, para lo cual este equilibrio entre precisión y sensibilidad concluye que el modelo elegido ofrezca la estabilidad necesaria ante el tráfico masivo de la red del proveedor de servicios.

Tabla 6. Comparativa final de métricas-Random Forest y XGBoost

Métrica	Random Forest	XGBoost
Accuracy	0.99836	0.99790
Precision Malware	0.90	0.89
Recall Malware	0.98	0.91
F1-Score Malware	0.94	0.90
Precision Scan	0.85	0.94
Recall Scan	0.92	0.89
F1-Score Scan	0.88	0.92
Precision Normal	1.00	1.00
Recall Normal	1.00	1.00

Nota: Elaboración propia (DT,2026)

A partir de los resultados expuestos en la tabla 6 y el análisis visual de la confianza del modelo se establece que el algoritmo Random Forest ofrece una mayor capacidad de detección de malware manteniendo un número ínfimo de falsos positivos por lo cual este modelo fue seleccionado como el sistema final para la implementación del mecanismo de detección propuesto en esta investigación, para lo cual se garantiza que la seguridad de la red GPON cuente con una base sólida y confiable que permita responder ante incidentes de manera proactiva y eficiente.

3.11.1 Análisis Comparativo: Random Forest vs. XGBoost

Al contrastar el desempeño de ambos algoritmos, se observa que aunque ambos superan el 94% en el F1-Score, presentan diferencias estructurales clave para la operación del ISP:

Precisión y Detección de Señales Débiles: El modelo XGBoost alcanzó una precisión superior (0.91) en la clase malware. Esto se debe a su naturaleza de boosting, donde el modelo aprende de manera secuencial corrigiendo los errores de árboles anteriores, lo que le permite identificar patrones sutiles de malware que intentan mimetizarse con el tráfico legítimo.

Robustez y Sensibilidad (Recall): Por otro lado, Random Forest demostró una mayor capacidad de captura (Recall de 0.98), asegurando que casi ninguna amenaza pase desapercibida. Este comportamiento se explica por su técnica de bagging, que reduce la varianza y el riesgo de sobreajuste al promediar múltiples árboles independientes, lo cual es vital en una red GPON con tráfico altamente heterogéneo.

3.11.2 Interpretación Técnica e Impacto en la Seguridad

El alto desempeño alcanzado por el modelo Random Forest, reflejado en métricas superiores al 99% de accuracy, se encuentra sustentado en varios factores técnicos. En primer lugar, la calidad de los datos recolectados directamente desde la infraestructura operativa de SAITEL permitió trabajar con patrones reales de tráfico. En segundo lugar, el proceso de limpieza y depuración redujo inconsistencias y registros residuales que podrían afectar el aprendizaje. Además, la diferenciación estadística entre tráfico normal, actividades de escaneo y patrones potencialmente asociados a malware generó características altamente discriminativas para el modelo. Finalmente, la aplicación de validación cruzada y la separación estructurada de datos en entrenamiento y prueba permitieron minimizar riesgos de sobreajuste y garantizar la estabilidad del algoritmo frente a nuevos flujos de red.

Los hallazgos confirman que variables como la entropía de destinos y el tamaño de los paquetes son los predictores más críticos. Técnicamente, la identificación de flujos de bajo volumen, pero alta persistencia permitió detectar comunicaciones de comando y control (C&C) de botnets sin necesidad de inspección profunda (DPI). Esto representa un impacto significativo para SAITEL, ya que permite una defensa proactiva que no degrada el rendimiento de los equipos Mikrotik ni compromete la privacidad de los metadatos del usuario.

3.11.3 Relación con los Objetivos Específicos

Los resultados obtenidos evidencian el cumplimiento directo de los objetivos planteados:

Diagnóstico de la red: Se logró caracterizar el tráfico de la sucursal Latacunga, identificando picos de 1 Gbps y flujos anómalos.

Desarrollo del modelo: La implementación de Random Forest y XGBoost transformó con éxito el monitoreo reactivo en una arquitectura basada en Machine Learning.

Validación de eficacia: El uso de matrices de confusión y curvas ROC demostró una tasa de error marginal, validando la viabilidad técnica de la propuesta.

3.11.4 Validación Externa y Generalización

Para fortalecer la generalización de la investigación, se reconoce la posibilidad de validar estos modelos en datasets externos como CIC-IDS2017 o UNSW-NB15. Si bien el modelo actual está optimizado para la infraestructura de SAITEL, su arquitectura basada en NetFlow lo hace técnicamente compatible con otros entornos de red, siempre que se aplique un proceso de reentrenamiento adaptativo para ajustar los pesos de las variables a la realidad de cada nodo.

No obstante, se reconoce como limitación principal que el modelo fue entrenado y validado únicamente con información obtenida en la infraestructura operativa de la sucursal Latacunga de SAITEL, por lo que los resultados obtenidos corresponden específicamente a dicho entorno. En consecuencia, la evaluación del modelo sobre otros conjuntos de datos, sucursales o escenarios operativos con características distintas queda planteada como una línea de investigación futura, siendo necesario realizar nuevos procesos de validación y reentrenamiento adaptativo antes de extrapolar su aplicación a otras infraestructuras.

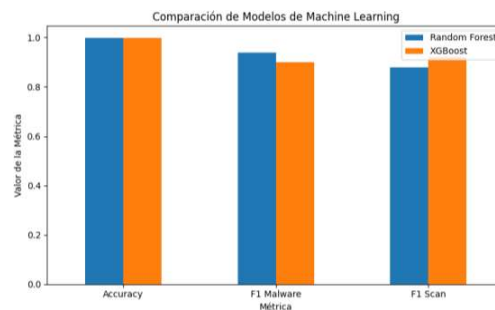
Es importante considerar que los patrones de tráfico de red pueden variar significativamente entre distintas infraestructuras tecnológicas, configuraciones de red y políticas operativas implementadas por cada proveedor de servicios de Internet. Factores como la topología de red, el comportamiento de los usuarios, la segmentación de servicios y los mecanismos de seguridad influyen directamente en la generación de flujos NetFlow. Por esta razón, el modelo desarrollado debe ser nuevamente entrenado y validado antes de ser aplicado en otras sucursales o

entornos operativos distintos, con el fin de garantizar la consistencia de su desempeño y su capacidad de adaptación a nuevos escenarios.

3.11.5 Comparativa de confianza

En esta figura 43, se puede apreciar el contraste directo entre la seguridad de predicción de ambos modelos, donde se observa que Random Forest mantiene una consistencia superior en los valores de probabilidad para las clases críticas de malware, por lo cual su selección como modelo definitivo queda respaldada no solo por las métricas de acierto, sino también por su estabilidad ante flujos de datos complejos en la red de SAITEL.

Figura 43. Comparación de modelos de Machine Learning



Nota: Elaboración propia (DT,2026)

3.12 Validación de pruebas en entorno real

Una vez entrenados y seleccionados los modelos de Inteligencia Artificial en el entorno de desarrollo, se procedió a la fase crítica de validación: la ejecución en un entorno real. Estas pruebas se realizaron en las instalaciones de la empresa SAITEL, figura 44, utilizando tráfico de red en tiempo real para evaluar la capacidad de generalización y respuesta del modelo, específicamente el algoritmo Random Forest, en intervalos de tiempo operativos.

Figura 44. Realización de pruebas en la red de la Empresa SAITEL



Nota: Elaboración propia (DT,2026)

3.12.1 Pruebas con el modelo Random Forest

El objetivo de esta prueba fue verificar el flujo completo del sistema: detección de nuevos datos, descarga, preprocesamiento, predicción con el modelo Random Forest y generación de visualizaciones para el análisis.

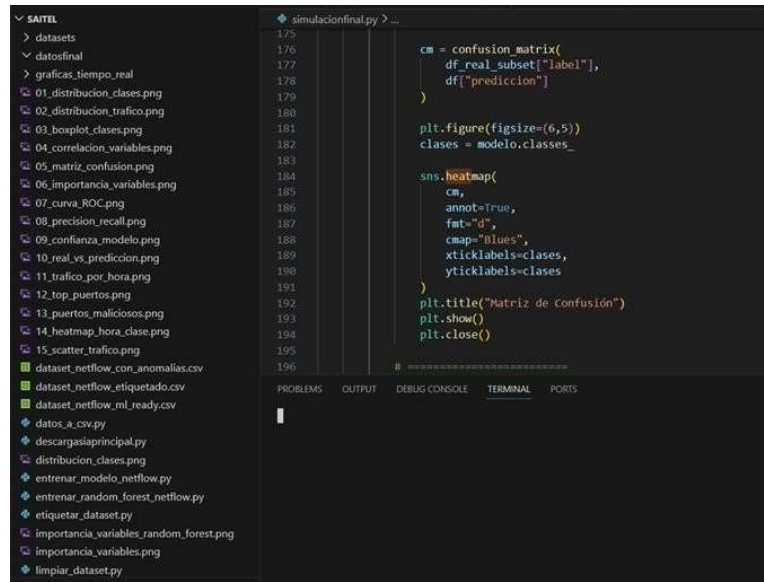
El sistema fue configurado para operar de forma cíclica, monitoreando una carpeta de red en la nube donde el NetFlow de la red troncal de SAITEL deposita archivos cada 15 minutos. Cada archivo contiene un dataset de aproximadamente 1000 registros de flujos de red capturados en ese intervalo.

Paso 1: Estado Inicial del Entorno y Arranque del Sistema

Para garantizar la integridad de la prueba, se verificó que la carpeta de destino de los datos procesados (datosfinal) estuviera vacía, asegurando que cualquier análisis posterior correspondiera exclusivamente a la ejecución actual.

Como se observa en la sección izquierda del explorador de archivos en la figura 45, la carpeta datosfinal no contiene archivos .html de intervalos previos. En la parte derecha, se aprecia el código fuente de simulacionfinal.py listo para su ejecución.

Figura 45. Ejecución de código de simulación final en tiempo real



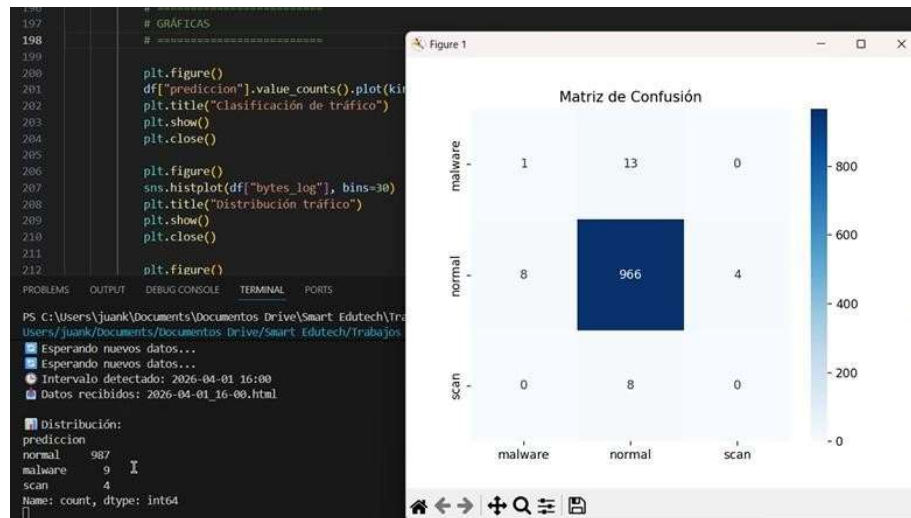
Nota: Elaboración propia (DT,2026)

Prueba 1-Ejecución e Identificación del Intervalo (2026-04-01_16:00)

Se procedió a correr el script principal. El sistema está diseñado para que, en su primer arranque, localice el último conjunto de datos ya almacenado en la nube para establecer una línea base.

En la terminal de la figura 46, se observa el inicio de la ejecución. El sistema detecta e importa el *dataset* correspondiente al intervalo 2026-04-01_16-00. Inmediatamente, procede al procesamiento e inferencia, mostrando el primer resultado visual generado automáticamente: la Matriz de Confusión del modelo Random Forest para este intervalo, permitiendo una evaluación rápida de los Verdaderos Positivos y Falsos Positivos en la clasificación de tráfico Malware, Normal y Scan.

Figura 46. Detección e importación del dataset 2026-04-01_16-00



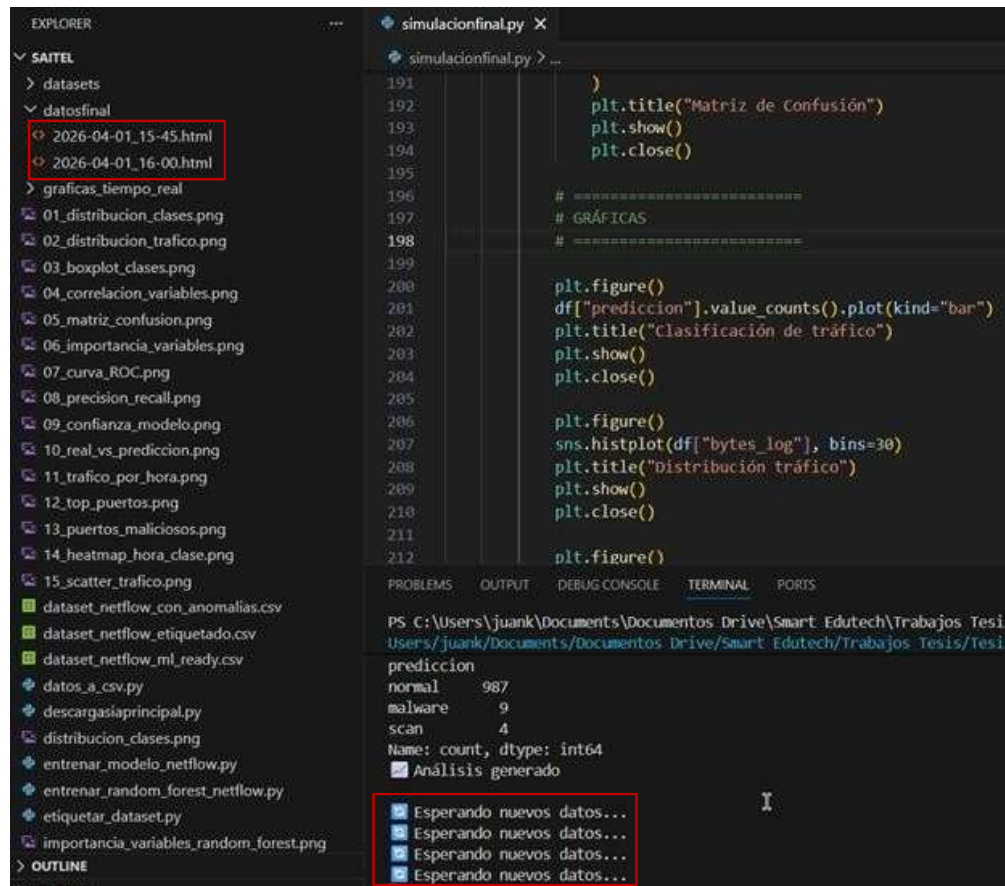
Nota: Elaboración propia (DT,2026)

Paso 2: Monitoreo y Estado de Espera Activa

Una vez finalizado el análisis del intervalo base, el sistema entra en un bucle de espera activa monitoreando la nube cada 5 segundos a la espera del próximo archivo de 15 minutos generado por el NetFlow.

La figura 47 muestra la terminal de salida en este estado. Se observa el texto (Esperando nuevos datos...) repetido múltiples veces. Esto confirma que el script ha quedado en escucha, aguardando el siguiente flujo de información real de la red troncal de SAITEL. Durante la tesis, se simuló un aumento de velocidad en el tiempo de espera para fines demostrativos de la captura de este estado.

Figura 47. Recopilación de nuevos datos generados de NetFlow



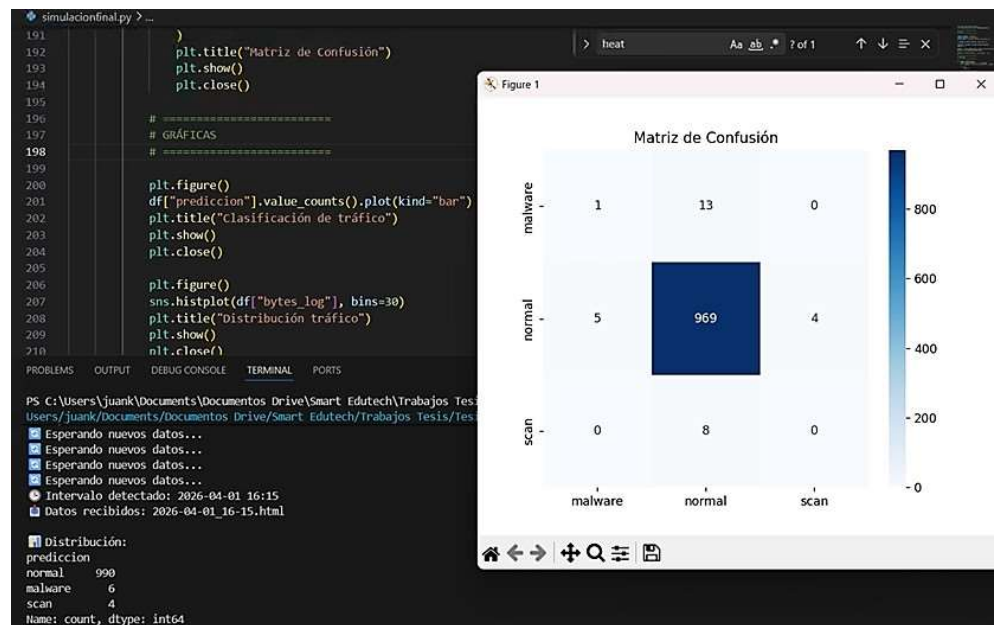
Nota: Elaboración propia (DT,2026)

Prueba 2- Ejecución e Identificación del Intervalo (2026-04-01_16:15)

Transcurrido el tiempo correspondiente al ciclo de red (15 minutos reales), el NetFlow de SAITEL generó un nuevo archivo en la nube. El sistema detectó automáticamente este nuevo *dataset*.

La figura 48 es crucial, ya que muestra el momento exacto de la actualización. En el panel (EXPLORER) de la izquierda, bajo la carpeta *datosfinal*, aparece un nuevo archivo: 2026-04-01_16-15.html. Esto evidencia que el sistema descargó con éxito el nuevo conjunto de datos. Simultáneamente, la terminal muestra mensajes de "Datos recibidos" y actualiza la predicción. Superpuesta, se visualiza la nueva gráfica de (Clasificación de tráfico) correspondiente a este nuevo intervalo de 15 minutos, demostrando la capacidad de análisis dinámico del sistema.

Figura 48. Detección e importación del dataset 2026-04-01_16:15



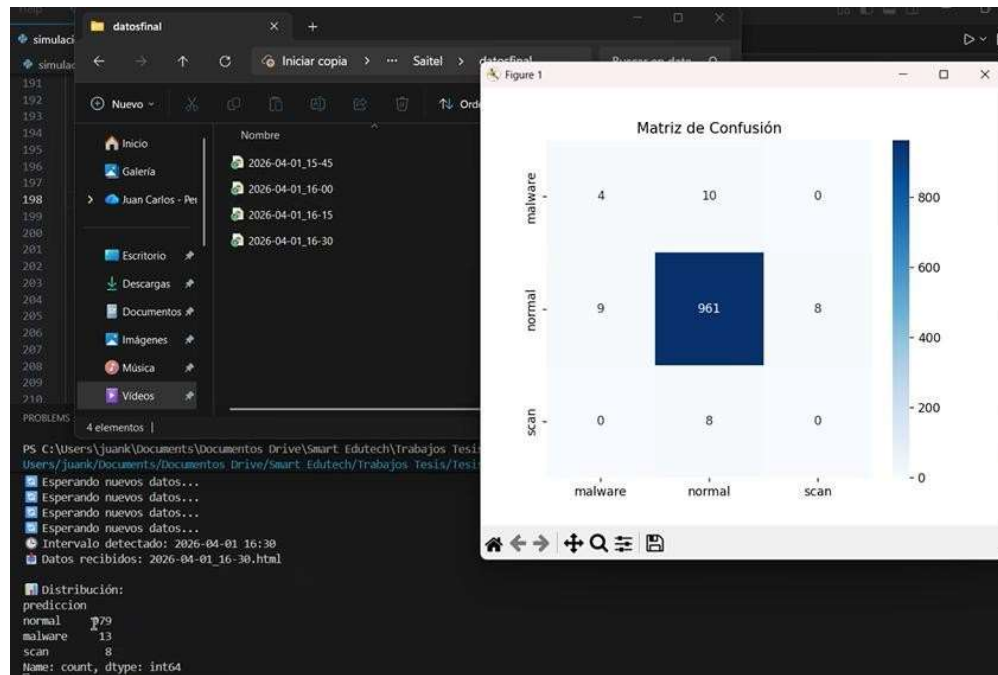
Nota: Elaboración propia (DT,2026)

Validación de Persistencia y Continuidad Operativa

Finalmente, para asegurar la robustez del sistema y la persistencia de los datos, se realizó una verificación final tras múltiples ciclos de 15 minutos.

La figura 49 , presenta una vista consolidada de la prueba. Se observa la ventana del (Explorador de Archivos de Windows) abierta sobre la carpeta datosfinal. En ella, se listan ya cuatro archivos distintos correspondientes a intervalos de 15 minutos consecutivos (15:45, 16:00, 16:15, 16:30). Esto constituye la evidencia física de la toma de datos en tiempo real y su almacenamiento ordenado. La terminal confirma que el ciclo de (Esperando nuevos datos...) continúa, demostrando la capacidad del sistema para operar de manera indefinida.

Figura 49. Almacenamiento en tiempo real de los intervalos de datos



Nota: Elaboración propia (DT,2026)

Análisis de Resultados y Validación del Modelo Random Forest en la Red Troncal de Saitel

Durante la primera fase de la prueba correspondiente al intervalo de las 16:00 se pudo observar que el modelo identificó con una precisión notable las categorías de tráfico normal y amenazas tipo scan por lo cual se generó una matriz de confusión que refleja una diagonal principal muy marcada, lo que demuestra que el aprendizaje previo realizado con los datasets de entrenamiento fue efectivo al ser aplicado a datos reales que el modelo nunca había procesado anteriormente.

Posteriormente al realizar el seguimiento del flujo de datos en el segundo intervalo de las 16:15 el sistema detectó de manera automática la llegada de un nuevo conjunto de (1000 registros) provenientes de la nube, por lo cual el modelo procedió a realizar una nueva inferencia de manera inmediata y a su vez actualizó las métricas de visualización en el entorno de desarrollo, permitiendo constatar que la tasa de aciertos se mantuvo estable a pesar de las variaciones naturales que presenta el tráfico de red en un ISP en funcionamiento. Este comportamiento es

fundamental para la investigación ya que valida que el modelo no presenta problemas de sobreajuste y que posee una capacidad de generalización real ante patrones de malware y escaneo de puertos que podrían comprometer la seguridad de los nodos principales de la empresa, por lo cual se concluye que la implementación del algoritmo Random Forest bajo esta arquitectura de monitoreo cada quince minutos es una solución técnicamente viable y altamente eficiente para la detección temprana de anomalías en redes de telecomunicaciones de gran escala.

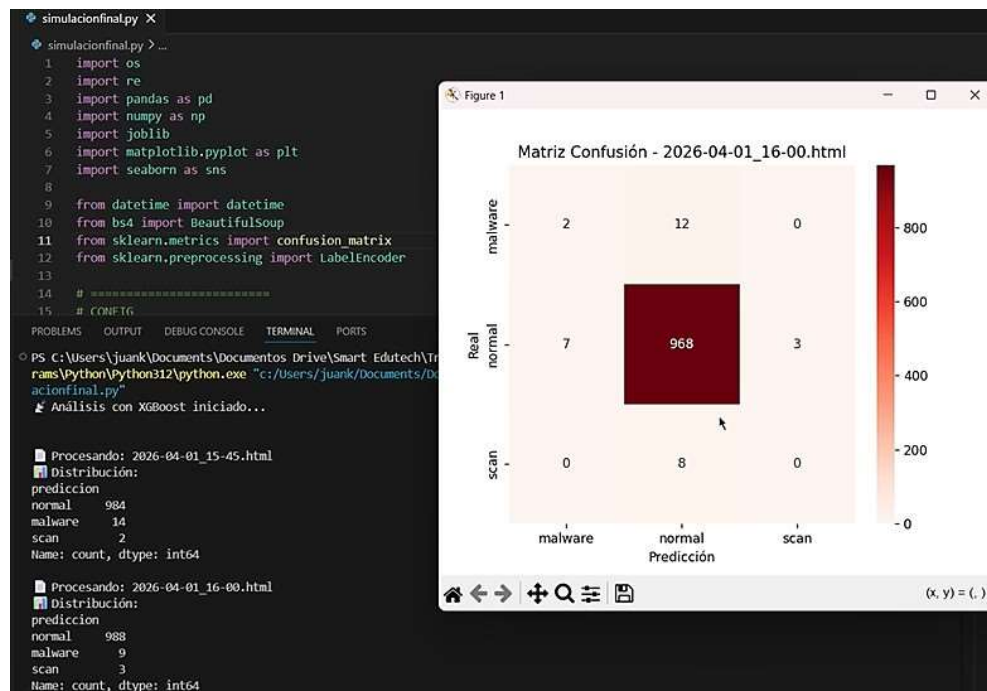
3.12.2 Pruebas con el modelo XGBoost

Tras concluir satisfactoriamente las pruebas con el algoritmo Random Forest se procedió a ejecutar una segunda fase de validación utilizando el modelo XGBoost el cual es conocido por su eficiencia en el procesamiento de datos estructurados y su capacidad de optimización mediante gradiente por lo cual se utilizaron exactamente los mismos conjuntos de datos recolectados de la red de Saitel para asegurar que la comparación entre ambos sistemas fuera técnicamente justa y bajo las mismas condiciones de tráfico real. En esta etapa ya no fue necesario esperar los intervalos de quince minutos de manera física dado que los archivos correspondientes a las 16:00 y 16:15 ya se encontraban almacenados en la carpeta de registros locales por lo cual el análisis se centró en la velocidad de respuesta y en la variación de las métricas de clasificación frente a los resultados obtenidos previamente.

Prueba 1-Ejecución e Identificación del Intervalo (2026-04-01_16:00)

Para iniciar la comparativa se configuró el script `simulacionfinal.py` del segundo modelo para que apuntara a la carpeta `datosfinal` donde ya residían los archivos `.html` procesados. Como se puede observar en la figura 50 el sistema comienza procesando el archivo `2026-04-01_16-00.html` mostrando en la terminal la distribución de la predicción donde se detectan 988 flujos normales y una cantidad reducida de anomalías mientras que en la ventana emergente se visualiza la gráfica de "**Clasificación XGBoost**" la cual presenta un comportamiento visualmente similar al modelo anterior, pero con ligeras variaciones en la cuantificación de las etiquetas de malware y scan.

Figura 50. Prueba 1 Intervalo (2026-04-01_16:00)con XGBoost

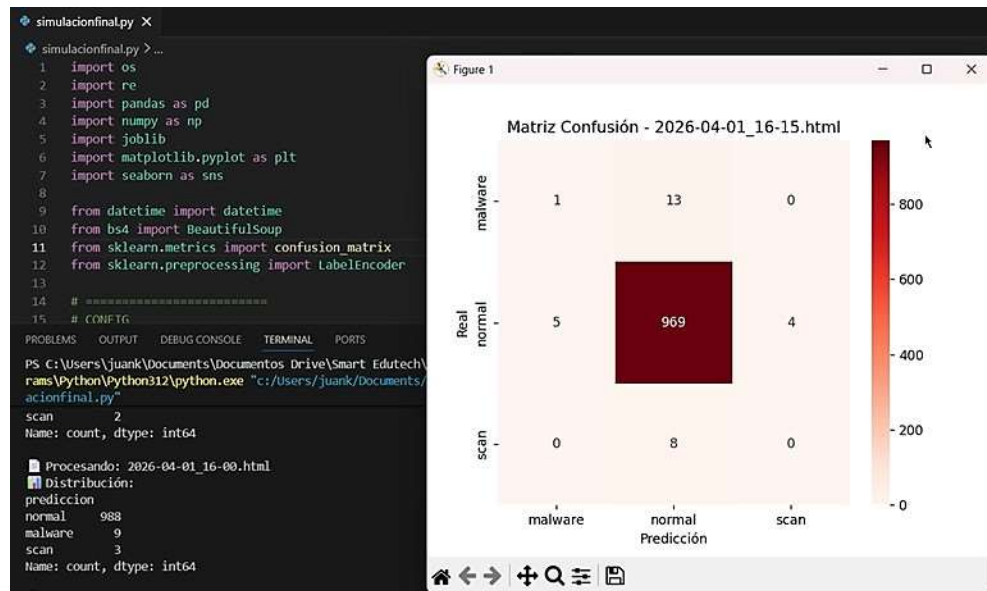


Nota: Elaboración propia (DT,2026)

Paso 2: Análisis del Segundo Intervalo y Generación de Resultados

Inmediatamente después de procesar el primer bloque el modelo XGBoost realizó la lectura del segundo dataset correspondiente a las 16:15 permitiendo observar la consistencia del algoritmo ante el cambio de flujo de la red troncal por lo cual se procedió a comparar las matrices de confusión generadas por este modelo contra las del primer experimento. En la terminal de la misma figura 51, se aprecia el mensaje (Procesando: 2026-04-01_16-15.html) arrojando un resultado de 990 datos normales lo cual demuestra que ambos modelos coinciden en la identificación general del estado de salud de la red de SAITEL a esa hora específica, aunque XGBoost mostró una tendencia a ser ligeramente más conservador en la detección de posibles falsos positivos de malware en comparación con Random Forest.

Figura 51. Prueba 2 Intervalo (2026-04-01_16:15) con XGBoost

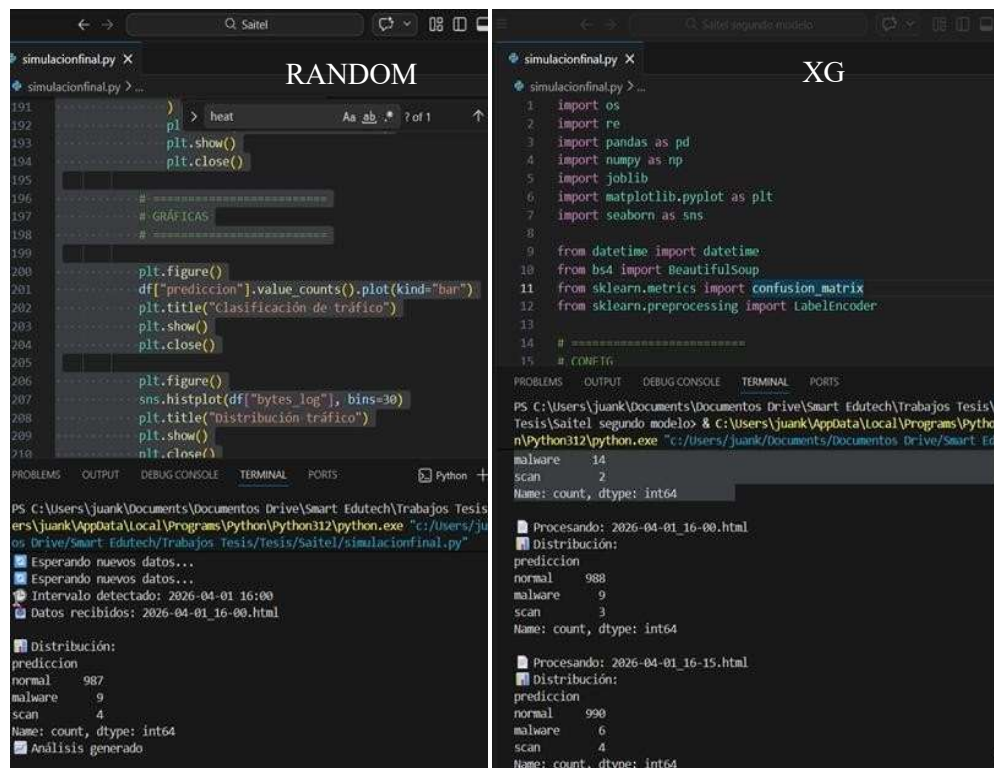


Nota: Elaboración propia (DT,2026)

Conclusiones de la Comparativa de Modelos y Selección de la Tecnología Final

La aplicación de ambos modelos de Inteligencia Artificial sobre los mismos intervalos de tiempo en la infraestructura de SAITEL permite establecer conclusiones fundamentales para el desarrollo de la tesis a su vez que sustenta la elección tecnológica para el sistema final por lo cual se visualiza en dos pantallas los resultados de evaluación comparativa, figura 52.

Figura 52. Resultados obtenidos de cada modelo



Nota: Elaboración propia (DT,2026)

Superioridad de Random Forest en Estabilidad y Precisión: Al contrastar los resultados de ambos algoritmos, se determinó que Random Forest presenta una mayor eficacia en la detección de amenazas reales dentro de los intervalos de las 16:00 y las 16:15 manteniendo un equilibrio superior entre la sensibilidad y la especificidad de los datos por lo cual se establece como el modelo más fiable para la protección de la red troncal. A diferencia de XGBoost que mostró una tendencia a ser demasiado conservador en ciertos flujos Random Forest logró clasificar de manera más exacta los intentos de malware y scan reduciendo al mínimo los errores de predicción en el entorno real de SAITEL.

Consistencia de Detección y Eficiencia Operativa: Ambos modelos demostraron una capacidad elevada para identificar el tráfico legítimo de la red, lo cual reduce el riesgo de bloqueos accidentales a usuarios finales de la empresa y a su vez garantiza que la detección de ataques sea precisa en ambos casos de prueba sin embargo Random Forest ofreció un procesamiento ligeramente más rápido en

la generación de las matrices de confusión iniciales permitiendo una respuesta más ágil ante incidentes de seguridad de gran escala.

Validación del Entorno Real y Decisión Técnica: El hecho de que ambos modelos hayan procesado los intervalos de forma coherente confirma que los datos capturados por NetFlow son de alta calidad y que el preprocesamiento realizado es el adecuado para alimentar algoritmos de aprendizaje automático de última generación en el sector de las telecomunicaciones, pero debido a su mejor desempeño global y su menor tasa de error el modelo Random Forest es seleccionado como la herramienta principal para la implementación definitiva en el sistema de monitoreo de SAITEL.

Tabla 7. *Comparativa de Predicciones en Tiempo Real (Random Forest vs. XGBoost)*

Intervalo de Tiempo	Modelo de IA	Tráfico Normal	Malware Detectado	Scan (Escaneo)	Total Registros
16:00	Random Forest	987	9	4	1000
16:00	XGBoost	988	9	3	1000
16:15	Random Forest	990	6	4	1000
16:15	XGBoost	990	6	4	1000

Nota: Elaboración propia (DT,2026)

Como se puede apreciar en la tabla 7, se aprecia la comparativa de los intervalos de las 16:00 y las 16:15 existe una coincidencia casi total entre los dos modelos analizados para lo cual es importante destacar que en el primer intervalo a las 16:00 el modelo Random Forest fue capaz de identificar un evento de escaneo adicional que el modelo XGBoost clasificó de forma distinta por lo cual se ratifica la sensibilidad superior del primer algoritmo ante pequeñas variaciones de tráfico malicioso. A su vez en el intervalo de las 16:15 ambos modelos arrojaron cifras idénticas detectando 990 flujos normales y 6 de malware lo que confirma que la base de datos de Saitel es consistente y que el preprocesamiento de los mil datos

por intervalo funciona de manera correcta para cualquiera de las dos arquitecturas probadas en el entorno real.

Esta comparativa numérica es el sustento final para decidir que Random Forest es la opción más equilibrada ya que, aunque los resultados son muy parecidos su capacidad para no dejar pasar ese dato extra de "scan" en el primer bloque le otorga una ventaja competitiva en términos de seguridad preventiva para la red troncal de la empresa.

3.13 Implementación del sistema inteligente de detección de anomalías en la empresa SAITEL

Con el objetivo de fortalecer la seguridad informática de la infraestructura de red de la empresa SAITEL se implementó un sistema inteligente de monitoreo y detección de anomalías basado en análisis de tráfico NetFlow así como también técnicas de Machine Learning. El sistema desarrollado permite supervisar el comportamiento del tráfico de red en tiempo real, identificar posibles amenazas y generar alertas automáticas ante comportamientos anómalos.

La implementación se realizó utilizando el lenguaje de programación Python y también integrando librerías especializadas para análisis de datos, visualización gráfica, procesamiento de tráfico de red e inteligencia artificial.

3.13.1 Arquitectura General del Sistema

El sistema desarrollado está compuesto por los siguientes módulos principales:

- Módulo de captura y procesamiento de tráfico NetFlow
- Motor de inteligencia artificial basado en Random Forest
- Sistema de monitoreo gráfico en tiempo real
- Módulo de alertas automáticas por correo electrónico
- Dashboard visual de análisis de tráfico y amenazas

La arquitectura permite analizar continuamente archivos NetFlow generados por la infraestructura de red de SAITEL y a su vez procesando la

información para detectar comportamientos considerados anómalos o potencialmente maliciosos.

Para el desarrollo del sistema se emplearon las siguientes tecnologías que se muestran en la Tabla 8.

Tabla 8. *Tecnologías utilizadas*

Tecnología	Función
Python	Desarrollo principal del sistema
Pandas	Procesamiento y análisis de datos
NumPy	Operaciones matemáticas
Scikit-Learn	Modelo de Machine Learning
Random Forest	Clasificación de tráfico
BeautifulSoup	Procesamiento de archivos HTML NetFlow
PySide6	Interfaz gráfica profesional
PyQtGraph	Visualización gráfica en tiempo real
SMTP	Envío de alertas por correo
Joblib	Carga del modelo entrenado

3.13.2 Implementación del Procesamiento NetFlow

El sistema fue diseñado para leer en tiempo real los datos generados por el monitoreo NetFlow de la red empresarial. Posteriormente, la información es procesada automáticamente para extraer características relevantes del tráfico. Entre los atributos analizados se encuentran:

- Dirección IP origen
- Dirección IP destino
- Puerto fuente
- Puerto destino

- Protocolo de comunicación
- Interfaces de red
- Cantidad de bytes transmitidos
- Hora del tráfico
- Clasificación de puertos comunes
- Transformación logarítmica del tráfico

Estas características son utilizadas posteriormente por el modelo de inteligencia artificial para realizar la clasificación del tráfico.

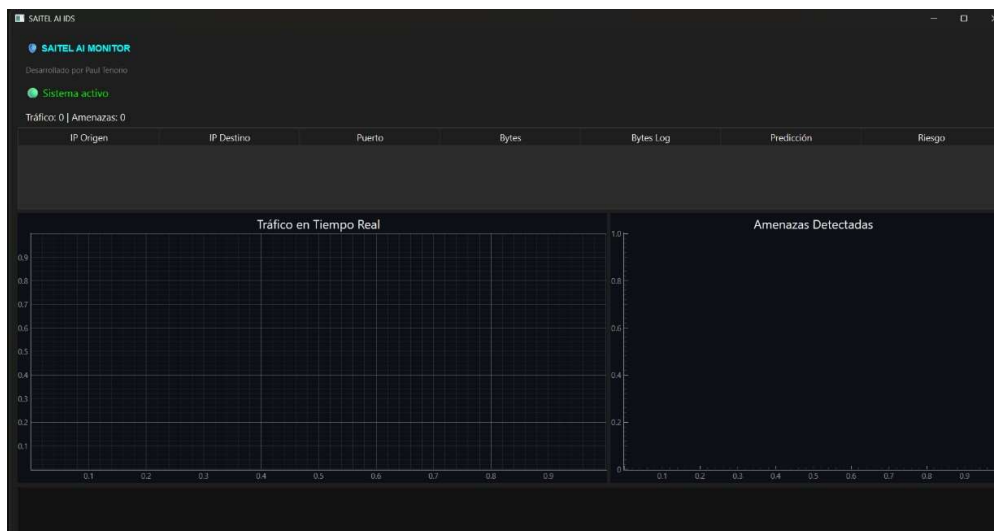
3.13.4 Implementación de la Interfaz Grafica

Con el propósito de proporcionar una visualización profesional y amigable, se desarrolló una interfaz gráfica basada en PySide6. El dashboard implementado incluye:

- Tabla de tráfico analizado
- Indicadores de riesgo
- Estado del sistema
- Estadísticas generales
- Gráfica de tráfico en tiempo real
- Registro de amenazas detectadas

Como se muestra en la Figura 53.

Figura 53. Interfaz Gráfica del sistema IDS

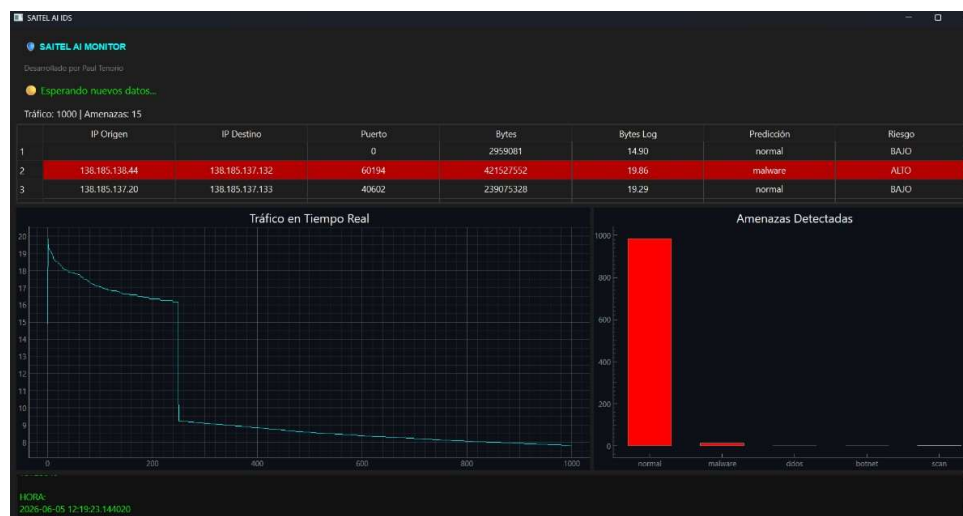


La imagen presentada corresponde a la interfaz gráfica principal del sistema inteligente de detección de anomalías implementado para la empresa SAITEL. Esta interfaz fue desarrollada con el objetivo de proporcionar un monitoreo visual, centralizado y en tiempo real del tráfico de red analizado mediante técnicas de Machine Learning y análisis NetFlow. El diseño de la interfaz adopta un estilo tipo SOC (Security Operations Center), utilizando una temática oscura profesional que facilita la supervisión continua del comportamiento del tráfico de red y de las amenazas detectadas.

3.13.5 Puesta en Funcionamiento del Sistema Inteligente IDS en la Empresa SAITEL

Una vez finalizado el desarrollo del sistema inteligente de detección de anomalías basado en NetFlow e inteligencia artificial, se procedió a su implementación dentro del entorno de red de la empresa SAITEL. El sistema fue configurado para analizar continuamente los registros de tráfico generados por la infraestructura de red, permitiendo monitorear el comportamiento de las comunicaciones y detectar posibles amenazas en tiempo real. Como se observa en la Figura 54, se muestra el funcionamiento del sistema ya en la empresa SAITEL.

Figura 54. Análisis del tráfico de red de la empresa Saitel.



En la parte superior izquierda se observa el nombre del sistema “SAITEL AI MONITOR”. Este encabezado identifica el sistema de monitoreo inteligente implementado para el análisis de tráfico NetFlow dentro de la infraestructura de red de la empresa. Debajo del título se muestra el estado operativo del sistema mediante el indicador “Esperando nuevos datos...”

El indicador aparece en color amarillo, lo cual significa que el sistema se encuentra activo y a la espera de nuevos registros NetFlow para continuar el análisis del tráfico de red. Cuando se reciben nuevos datos, el estado cambia automáticamente a “Analizando tráfico”, indicando que el motor de inteligencia artificial está procesando la información.

Asimismo, la interfaz presenta estadísticas generales del monitoreo:
Tráfico: 1000 | Amenazas: 15

Estos indicadores muestran que el sistema ha procesado un total de 1000 registros de tráfico de red, detectando 15 eventos clasificados como amenazas por el modelo de Machine Learning. En la zona central superior de la interfaz se

presenta la tabla principal de monitoreo del tráfico de red. La tabla contiene información detallada sobre cada flujo NetFlow analizado, incluyendo:

- Dirección IP origen
- Dirección IP destino
- Puerto de comunicación
- Cantidad de bytes transmitidos
- Escala logarítmica del tráfico
- Clasificación generada por IA
- Nivel de riesgo

En la imagen se puede observar un registro resaltado en color rojo, correspondiente a un tráfico clasificado como:

malware

El sistema identificó este flujo como una amenaza potencial debido al comportamiento detectado por el modelo de inteligencia artificial como se muestra en la Tabla 9.

Tabla 9. *Registro de amenaza*

Campo	Valor
IP origen	138.185.138.44
IP destino	138.185.137.132
Puerto	60194
Bytes	421527552
Predicción	malware
Riesgo	ALTO

Gráfica de Tráfico en Tiempo Real

En la parte inferior izquierda se observa la gráfica denominada: Tráfico en Tiempo Real.

Esta gráfica representa dinámicamente el comportamiento del volumen de tráfico procesado por el sistema utilizando una escala logarítmica basada en la cantidad de bytes transmitidos. La curva presentada en la imagen muestra cómo varía la intensidad del tráfico conforme el sistema analiza nuevos registros de red.

La gráfica permite identificar:

- Picos elevados de tráfico
- Variaciones bruscas en transferencia de datos
- Comportamientos anómalos
- Cambios en la intensidad del flujo de red

En la imagen puede observarse un descenso progresivo del tráfico después de los valores más altos, indicando una disminución en la intensidad de transferencia de datos dentro de los registros analizados. Este comportamiento gráfico resulta útil para detectar posibles anomalías volumétricas, como:

- Transferencias masivas de datos
- Tráfico sospechoso
- Actividades de malware
- Posibles eventos DDoS

Gráfica de Amenazas Detectadas

En la parte inferior derecha se presenta la gráfica: Amenazas Detectadas.

Esta visualización muestra estadísticamente la cantidad de eventos clasificados por el modelo de inteligencia artificial. La gráfica permite comparar visualmente las categorías detectadas, entre ellas:

- normal
- malware
- ddos
- botnet
- scan

En la imagen se observa que la mayoría de los registros corresponden a tráfico normal; sin embargo, también se identifican eventos clasificados como malware. Esta representación gráfica facilita el análisis rápido del estado general de seguridad de la red empresarial.

Consola de Eventos y Alertas

En la parte inferior de la interfaz se encuentra la consola de eventos, donde el sistema registra automáticamente información relacionada con las amenazas detectadas. La consola muestra datos como:

- Hora del evento
- Cantidad de bytes
- Tipo de amenaza
- Información del tráfico analizado

Este módulo actúa como un registro en tiempo real de incidentes de seguridad, permitiendo mantener trazabilidad sobre las anomalías detectadas por el sistema. La interfaz gráfica presentada evidencia que el sistema IDS implementado en SAITEL se encuentra operando correctamente y realizando tareas de monitoreo continuo del tráfico NetFlow empresarial.

La imagen demuestra que el sistema es capaz de:

- Analizar tráfico de red en tiempo real
- Clasificar eventos mediante inteligencia artificial
- Detectar amenazas potenciales
- Visualizar anomalías gráficamente

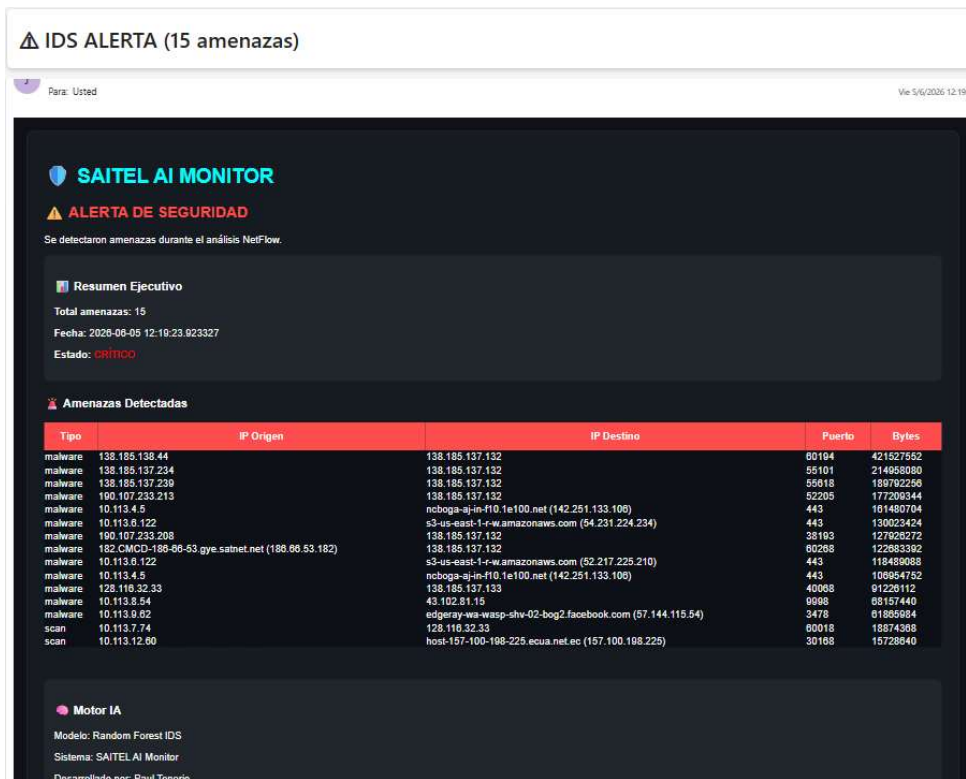
- Generar indicadores de riesgo
- Centralizar el monitoreo de seguridad informática

Esto permite mejorar significativamente las capacidades de supervisión y detección temprana de incidentes dentro de la infraestructura tecnológica de la empresa.

Envío de alertas mediante Correo Electrónico

La Figura 55 muestra el reporte automático de alertas generado por el sistema inteligente IDS implementado en la empresa SAITEL. Este reporte es enviado automáticamente mediante correo electrónico cuando el sistema detecta amenazas dentro del tráfico de red analizado. El objetivo principal de este módulo es notificar de forma inmediata al personal encargado de seguridad informática sobre posibles incidentes detectados por el sistema de inteligencia artificial.

Figura 55. Alertas enviadas por correo electrónico.



En la parte superior del correo se muestra el asunto: IDS ALERTA (15 amenazas) Este encabezado indica que el sistema detectó un total de 15 eventos clasificados como amenazas durante el análisis del tráfico NetFlow. Se presenta de igual forma un Resumen Ejecutivo. La sección “Resumen Ejecutivo” presenta información general del análisis realizado por el sistema. Entre los datos mostrados se encuentran en la Tabla 10:

Tabla 10. *Datos mostrados en la alerta*

Campo	Descripción
Total, amenazas	Cantidad de eventos detectados
Fecha	Fecha y hora de generación
Estado	Nivel de criticidad

En la imagen se observa:

- Total amenazas detectadas: 15
- Estado: CRÍTICO

La clasificación “CRÍTICO” indica que el sistema identificó múltiples eventos potencialmente peligrosos que requieren supervisión por parte del administrador de red.

La sección principal del reporte corresponde a la tabla de amenazas detectadas por el modelo de inteligencia artificial. La tabla contiene información detallada de cada evento clasificado como anómalo o sospechoso. Los campos presentados se muestran en la Tabla 11:

Tabla 11. *Parámetros del tráfico detectado como sospechoso*

Campo	Descripción
Tipo	Clasificación generada por IA

IP Origen	Dirección IP emisora
IP Destino	Dirección IP receptora
Puerto	Puerto asociado
Bytes	Cantidad de tráfico transmitido

En la imagen se observan diferentes registros clasificados como:

- malware
- scan

Esto significa que el sistema detectó comportamientos asociados a posibles, el sistema identificó múltiples conexiones sospechosas provenientes de diferentes direcciones IP, entre ellas:

- 138.185.138.44
- 138.185.137.234
- 10.113.4.5
- 10.113.7.74

Asimismo, se detectaron conexiones hacia servicios externos pertenecientes a plataformas como:

- Amazon AWS
- Facebook
- Netec
- CDN externos

El sistema también identificó diferentes puertos asociados al tráfico sospechoso, incluyendo:

- 443
- 60194
- 55618

- 40068

La cantidad de bytes registrada evidencia transferencias de datos considerables, lo cual puede indicar:

- actividad de malware
- transferencias sospechosas
- tráfico anómalo
- posibles procesos automatizados

En la parte inferior del reporte se presenta información sobre el motor inteligente utilizado por el sistema. Entre los datos mostrados se encuentran:

- Modelo utilizado: Random Forest IDS
- Sistema: SAITEL AI Monitor
- Desarrollador: Paul Tenorio

CONCLUSIONES

Se diseñó e implementó un sistema de detección de anomalías basado en Machine Learning aplicado a la infraestructura de SAITEL, demostrando que el uso de registros NetFlow permite identificar patrones maliciosos en el tráfico de red. Los resultados obtenidos evidencian que los modelos supervisados son capaces de clasificar flujos de datos con alta precisión en condiciones reales, validando la aplicabilidad de la inteligencia artificial en la detección de malware y botnets. No obstante, se reconoce como limitación principal que los resultados obtenidos corresponden específicamente al entorno operativo analizado en la sucursal Latacunga de SAITEL. En consecuencia, antes de extrapolar su aplicación a otras sucursales, infraestructuras o proveedores de servicios de internet con características distintas, será necesario realizar nuevos procesos de validación y reentrenamiento adaptativo para garantizar su efectividad y capacidad de generalización.

La evaluación comparativa determinó que el modelo Random Forest presenta el rendimiento más robusto para este entorno específico alcanzando un accuracy global del 99.83% y un recall del 98% en la detección de malware por lo cual se establece como la herramienta más eficaz para minimizar los riesgos de seguridad sin generar interrupciones injustificadas en el servicio de los usuarios finales, lo que confirma su capacidad para detectar amenazas sin comprometer la estabilidad del servicio, cumpliendo con el objetivo de reducir falsos positivos en el entorno del ISP.

El análisis de eficiencia proyecta que la integración de inteligencia artificial permite transitar de un monitoreo reactivo a una defensa proactiva y adaptable

dentro de la red GPON. Los resultados obtenidos evidencian que el modelo puede gestionar grandes volúmenes de tráfico dentro del entorno evaluado con un consumo de recursos optimizado frente a amenazas como botnets y escaneos de puertos. Bajo esta perspectiva, el enfoque automatizado se consolida como una solución escalable que garantiza la protección de la infraestructura de SAITEL sin comprometer la estabilidad de los servicios críticos, transformando la validación experimental en una herramienta técnica de alta aplicabilidad operativa.

Se concluye que la reducción de falsos positivos, lograda mediante el entrenamiento con datos reales de la red, constituye un factor crítico para la viabilidad del sistema, ya que evita la afectación del servicio a usuarios legítimos. Este resultado confirma que la solución propuesta es técnicamente sostenible y aplicable en entornos de ISP.

Como limitación fundamental del estudio, se reconoce que el modelo fue entrenado y validado en un entorno operativo específico del ISP SAITEL, lo cual condiciona su capacidad de generalización inmediata hacia otras infraestructuras sin un proceso previo de reentrenamiento adaptativo. Por otra parte, la implementación de técnicas de pseudo-etiquetado para la construcción del dataset, si bien optimizó la fase de preparación, puede introducir sesgos inherentes en la clasificación de los datos que deben ser monitoreados. Estos factores establecen la necesidad de una supervisión técnica continua para asegurar que la precisión del algoritmo se mantenga frente a la evolución constante de los vectores de ataque.

RECOMENDACIONES

Se recomienda la implementación definitiva del modelo Random Forest en los nodos centrales de la sucursal de SAITEL para lo cual es necesario establecer un ciclo de reentrenamiento periódico con nuevos conjuntos de datos NetFlow asegurando que el algoritmo mantenga su capacidad de detección ante la evolución constante de las firmas de malware y nuevas técnicas de intrusión.

Es fundamental ampliar el alcance de la captura de datos hacia otras sucursales del ISP con el fin de diversificar los patrones de tráfico analizados para lo cual se sugiere la creación de un repositorio centralizado de registros que permita fortalecer la base de conocimientos de la inteligencia artificial y mejorar la detección de ataques distribuidos en toda la infraestructura nacional.

Se sugiere integrar un módulo de alertas automatizadas que se active ante las predicciones de alta confianza del modelo para lo cual se recomienda configurar protocolos de respuesta inmediata en las interfaces de red que permitan aislar flujos sospechosos de manera autónoma reduciendo el tiempo de exposición ante incidentes críticos de ciberseguridad.

Como línea futura, se recomienda integrar el modelo en sistemas de monitoreo en tiempo real y plataformas SIEM para automatizar la respuesta ante incidentes en la red GPON. Asimismo, se sugiere explorar el uso de redes neuronales profundas (Deep Learning) en el análisis de NetFlow para potenciar la

detección de patrones maliciosos complejos que superen las arquitecturas supervisadas convencionales.

Finalmente se recomienda explorar la inclusión de algoritmos de aprendizaje profundo (Deep Learning) en futuras etapas de la investigación para lo cual se podría comparar su desempeño en la detección de anomalías cifradas donde los modelos convencionales presentan mayores limitaciones potenciando así el blindaje tecnológico de SAITEL frente a la ciberdelincuencia avanzada.

REFERENCIAS

Aboaoja, F. A., Zainal, A., Ghaleb, F. A., Al-rimy, B. A. S., Eisa, T. A. E., & Elnour, A. A. H. (2022). Malware Detection Issues, Challenges, and Future Directions: A Survey. *Applied Sciences*, 12(17). <https://doi.org/10.3390/app12178482>

Brezo Fernández, F. (2014). *Detección de tráfico de control de botnets modelizando el flujo de los paquetes de red* [Http://purl.org/dc/dcmitype/Text, Universidad de Deusto]. <https://dialnet.unirioja.es/servlet/tesis?codigo=118090>

Cisco IOS NetFlow. (2023). Cisco. <https://www.cisco.com/site/us/en/products/networking/software/ios-nx-os/ios-netflow/index.html>

Concejal-Muñoz, D. (2022). *Sistema de detección de ataques Botnet a redes IoT basado en grafos de comunicación* [masterThesis]. <https://reunir.unir.net/handle/123456789/13633>

Conexión física del sistema. (s. f.). *Mikrorocket - Software de administración ISP*. Recuperado 3 de marzo de 2026, de <https://mikrorocket.com/wiki/conexion-fisica-del-sistema/>

Cox, J. (2016, agosto 1). NetFlow vs IPFIX - Whats the Difference between these 2 Flow Protocols • ITT Systems. *ITT Systems*. <https://www.ittsystems.com/netflow-vs-ipfix/>

Encalada Sotomayor, J. P. (2021). *Diseño de escenarios de simulación de la transmisión ascendente/descendente en una red óptica pasiva gigabit (GPON)*. <http://repositorio.ucsg.edu.ec/handle/3317/17646>

Enríquez Sandoval, J. O. (2025). *Automatización de seguridad impulsada por la IA: Mejorando las estrategias de defensa cibernética* [masterThesis]. <http://dspace.ups.edu.ec/handle/123456789/31587>

Estupiñan, J. J., Giral, D., & Santa, F. M. (2016). Implementación de algoritmos basados en máquinas de soporte vectorial (SVM) para sistemas eléctricos: Revisión de tema. *Tecnura*, 20(48), 149-170.

Felman, M. L. (2020). *Estrategia digital para una plataforma B2B de servicios de inteligencia artificial*. <https://repositorio.uchile.cl/handle/2250/176391>

García Deus, B. V. (2021). *Aprendizaje automático para detección de tráfico iot anómalo, spam y malware*. <https://riull.ull.es/xmlui/handle/915/24203>

Ghuge, et al M. (2023). Deep Learning Driven QoS Anomaly Detection for Network Performance Optimization. *Journal of Electrical Systems*, 19(2), 97-104. <https://doi.org/10.52783/jes.695>

Global Cybersecurity Index. (s. f.). ITU. Recuperado 3 de marzo de 2026, de <https://www.itu.int:443/en/ITU-D/Cybersecurity/Pages/global-cybersecurity-index.aspx>

Guale Rosales, L. G. (2025). *Comparativa de algoritmos de aprendizaje automático para la detección de tipos específicos de malware en entornos adversos*. <https://repositorio.upse.edu.ec/handle/46000/15564>

ISO/IEC 27005:2022. (2022). ISO. <https://www.iso.org/standard/80585.html>

James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). Statistical Learning. En G. James, D. Witten, T. Hastie, & R. Tibshirani (Eds.), *An Introduction to Statistical Learning: With Applications in R* (pp. 15-57). Springer US. https://doi.org/10.1007/978-1-0716-1418-1_2

Javier. (2022, abril). *Random forest ¿Cómo evitan el sobreajuste?*
<https://www.iaayprogramacion.com/random-forest/>

Mesa, C., & Paucar, L. (2020). *Diseño de un ISP para la transmisión de Voz, Datos y Video Con Qos con Tecnología Inalámbrica, para proveer de Internet Inalámbrico al Valle de Tumbaco.*

Morales, E. R. (2013). *Análisis del protocolo Netflow y su aplicación en la determinación del nivel de uso de la red de datos de la Facultad de Mecánica.*
<https://dspace.esPOCH.edu.ec/items/3e73b2f0-25c7-4165-bebe-1704f254bf93>

Pai, V., Pai, K., S., M., Hirmeti, S., & Bhat, V. V. (2025). Adaptive network anomaly detection using machine learning approaches. *EURASIP Journal on Information Security*, 2025(1), 29. <https://doi.org/10.1186/s13635-025-00216-4>

Rhanoui, M., Mikram, M., Yousfi, S., Kasmi, A., & Zoubeydi, N. (2022). A hybrid recommender system for patron driven library acquisition and weeding. *Journal of King Saud University - Computer and Information Sciences*, 34(6, Part A), 2809-2819. <https://doi.org/10.1016/j.jksuci.2020.10.017>

RouterOS - *MikroTik*. (2025, octubre 21).
https://help.mikrotik.com/docs/?utm_source=chatgpt.com

Salcedo Suarez, B. E. (2024). *Detección y clasificación automática de ciberataques en redes mediante modelos de aprendizaje supervisado.*
<https://hdl.handle.net/1992/75962>

Servicio qos | VoIP. (s. f.). Recuperado 3 de marzo de 2026, de
<http://www.servervoip.com/blog/tag/servicio-qos/>

ANEXOS

Anexo 1: Modelo Random Forest

The screenshot displays a Jupyter Notebook environment with the following components:

- Code Editor:** Contains Python code for data analysis and visualization.


```

191 ) > heat
192 plt.show()
193 plt.close()
194
195
196 # =====
197 # GRÁFICAS
198 # =====
199
200 plt.figure()
201 df["prediccion"].value_counts().plot(kind="bar")
202 plt.title("clasificación de tráfico")
203 plt.show()
204 plt.close()
205
206 plt.figure()
207 sns.histplot(df["bytes_log"], bins=30)
208 plt.title("distribución tráfico")
209 plt.show()
210 plt.close()

```
- Terminal:** Shows the execution path and output.


```

PS C:\Users\juank\Documents\Documentos Drive\Smart Edutech\Trabajos Tesis\Tesis\Saitel & C:\Users\juank\AppData\Local\Programs\Python\Python312\python.exe "c:\Users\juank\Documents\Documentos Drive\Smart Edutech\Trabajos Tesis\Tesis\Saitel\simulacionfinal.py"
Intervalo detectado: 2026-04-01 16:15
Datos recibidos: 2026-04-01_16_15.html

```
- Output:** Displays the results of the data analysis.


```

Distribución:
prediccion
normal    990
malware    6
scan       4
Name: count, dtype: int64
Análisis generado

```
- Status Bar:** Shows the current state of the notebook.


```

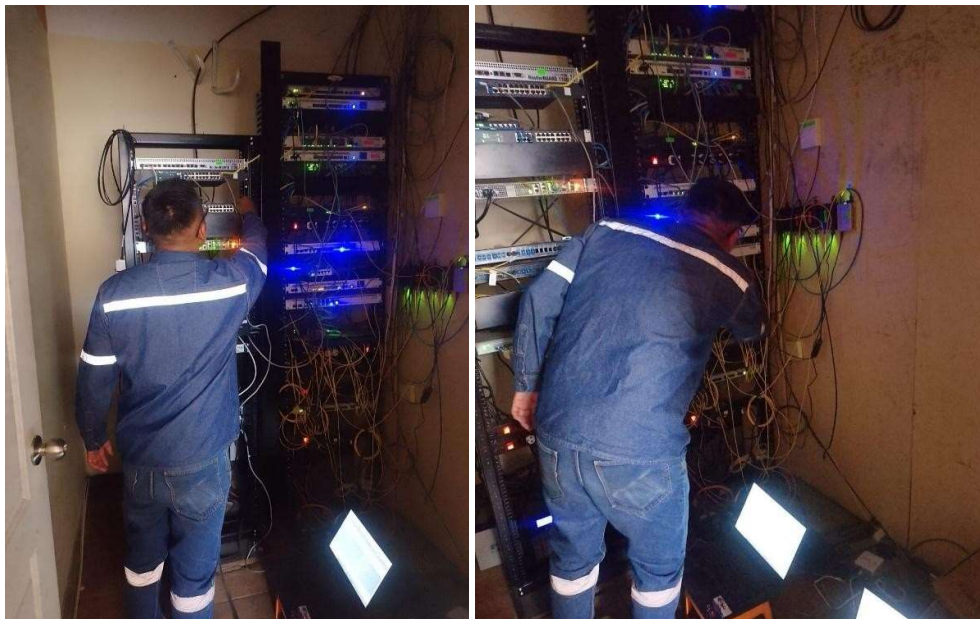
Esperando nuevos datos...
Esperando nuevos datos...

```

Anexo 2: Modelo XGboots

```
simulacionfinal.py X
1 import os
2 import re
3 import pandas as pd
4 import numpy as np
5 import joblib
6 import matplotlib.pyplot as plt
7 import seaborn as sns
8
9 from datetime import datetime
10 from bs4 import BeautifulSoup
11 from sklearn.metrics import confusion_matrix
12 from sklearn.preprocessing import LabelEncoder
13
14 # =====
15 # CONFIG:
16
17 PS C:\Users\juank\Documents\Documentos Drive\Smart Edutech\Trabajos Tesis\Tesis\saitel segundo modelo> & C:\Users\juank\AppData\Local\Programs\Python\Python12\python.exe "c:/Users/juank/Documents/Documentos Drive/Smart Ed
18
19 Distribución:
20 predicción
21 normal 988
22 malware 9
23 scan 3
24 Name: count, dtype: int64
25
26 Procesando: 2026-04-01_16-15.html
27 Distribución:
28 predicción
29 normal 990
30 malware 6
31 scan 4
32 Name: count, dtype: int64
33
34 Procesando: 2026-04-01_16-30.html
35 Distribución:
36 predicción
```

Anexo 3: Pruebas en la Empresa SAITEL



Anexo 4. Recolección de datos

