



ESCUELA UNIVERSITARIA DE POSGRADO

LOS MODELOS PREDICTIVOS DE IA Y SU RELACIÓN CON LA PROTECCIÓN
CONTRA AMENAZAS PERSISTENTES AVANZADAS (APTS) EN EL SISTEMA
BANCARIO DE LIMA, 2024

Línea de investigación:
Sistemas de información y optimización

Tesis para optar el Grado Académico de Doctor en Ingeniería de Sistemas

Autor

Espinoza Villavicencio, Héctor Martín

Asesor

Rodríguez Rodríguez, Ciro

ORCID: 0000-0003-2112-1349

Jurado

Cachay Boza, Orestes

Petrlik Azabache, Ivan Carlo

Ogosi Auqui, Jose Antonio

Lima - Perú

2026

LOS MODELOS PREDICTIVOS DE IA Y SU RELACIÓN CON LA PROTECCIÓN CONTRA AMENAZAS PERSISTENTES AVANZADAS (APTS) EN EL SISTEMA BANCARIO DE LIMA, 2024

INFORME DE ORIGINALIDAD

15%	12%	5%	6%
INDICE DE SIMILITUD	FUENTES DE INTERNET	PUBLICACIONES	TRABAJOS DEL ESTUDIANTE

FUENTES PRIMARIAS

1	Submitted to Universidad Nacional Federico Villarreal	2%
	Trabajo del estudiante	
2	www.coursehero.com	1%
	Fuente de Internet	
3	repositorio.ucv.edu.pe	1%
	Fuente de Internet	
4	repositorio.unfv.edu.pe	1%
	Fuente de Internet	
5	alicia.concytec.gob.pe	1%
	Fuente de Internet	
6	Submitted to Universidad Cesar Vallejo	<1%
	Trabajo del estudiante	
7	hdl.handle.net	<1%
	Fuente de Internet	
8	www.iksadamerica.org	<1%
	Fuente de Internet	
9	Submitted to Universidad San Marcos	<1%
	Trabajo del estudiante	
10	transportesynegocios.wordpress.com	<1%
	Fuente de Internet	
11	hemeroteca.unad.edu.co	<1%
	Fuente de Internet	



Universidad Nacional
Federico Villarreal

VRIN | VICERRECTORADO
DE INVESTIGACIÓN

ESCUELA UNIVERSITARIA DE POSGRADO

**LOS MODELOS PREDICTIVOS DE IA Y SU RELACIÓN CON LA PROTECCIÓN
CONTRA AMENAZAS PERSISTENTES AVANZADAS (APTS) EN EL SISTEMA
BANCARIO DE LIMA, 2024**

Línea de investigación:

Sistemas de información y optimización

Tesis para optar el Grado Académico de
Doctor en Ingeniería de Sistemas

Autor

Espinoza Villavicencio, Héctor Martín

Asesor

Rodríguez Rodríguez, Ciro

ORCID: 0000-0003-2112-1349

Jurado

Cachay Boza, Orestes

Petrlik Azabache, Ivan Carlo

Ogosi Auqui, Jose Antonio

Lima – Perú

2026

DEDICATORIA

A Dios y a mi madre por la vida, este objetivo es
logrado por su apoyo incondicional en el día a día.

AGRADECIMIENTO

Mi especial agradecimiento para mi asesor

Dr. Ciro Rodriguez Rodriguez

Por las sugerencias recibidas para el mejoramiento de este trabajo

Muchas gracias para todos

ÍNDICE

RESUMEN	iv
ABSTRACT	iv
I. INTRODUCCIÓN	1
1.1.Planteamiento del problema	2
1.2.Descripción del problema	4
1.3.Formulación del problema	5
1.3.1. Problema general	5
1.3.2. Problemas específicos	5
1.4.Antecedentes	6
1.4.1. Antecedentes nacionales	6
1.4.2. Antecedentes internacionales	11
1.5.Justificación de la investigación	17
1.5.1. Justificación práctica	17
1.5.2. Justificación teórica	17
1.5.3. Justificación social	18
1.6.Limitaciones de la investigación	18
1.6.1. Limitaciones prácticas	18
1.6.2. Limitaciones teóricas	19
1.6.3. Limitaciones metodológicas	19
1.7.Objetivos	19

1.7.1. Objetivo general	19
1.7.2. Objetivos específicos	19
1.8.Hipótesis	20
1.8.1. Hipótesis general	20
1.8.2. Hipótesis específicas.	20
II. MARCO TEÓRICO	21
2.1. Marco conceptual	21
2.1.1. Análisis del panorama actual de amenazas persistentes avanzadas (APTs) en el sistema bancario.	21
2.1.2. Evaluación del impacto de los modelos predictivos de IA en la resiliencia del sistema bancario frente a APTs.	24
2.1.3. Análisis del papel de los modelos predictivos de IA en la actualización y mejora de las estrategias de protección contra APTs en el sistema bancario.	27
2.1.4. Algoritmos	30
2.1.5. Inteligencia artificial	30
2.1.6. Datos	30
2.1.7. Sistema	31
2.1.8. Métricas	31
2.1.9. Patrones	31
2.1.10. Métodos	31
2.1.11. Protección	32
2.1.12. Capacidad	32

2.1.13. Modelos Predictivos de IA	32
2.1.14. Protección contra amenazas persistentes avanzadas (APT)	38
III. MÉTODO	44
3.1. Tipo de investigación	44
3.2. Población y muestra	45
3.3.Operacionalización de variables	47
3.4.Instrumentos	48
3.5.Procedimientos	49
3.6.Análisis de datos	49
3.7. Consideraciones éticas	50
IV. RESULTADOS	51
V. DISCUSIÓN DE RESULTADOS	94
VI. CONCLUSIONES	97
VII. RECOMENDACIONES	98
VIII. REFERENCIAS	99
IX. ANEXOS	109
Anexo A. Matriz de Consistencia	110
Anexo B. Instrumento de recolección de datos	111
Anexo C: Ficha de validación de instrumento por juicio de expertos	113

ÍNDICE DE TABLAS

Tabla 1 Operacionalización de las variables	48
Tabla 2 Alfa de Cronbach	49
Tabla 3 Análisis comparativo	57
Tabla 4 Frecuencia respecto a la eficacia de los algoritmos para abordar problemas.	62
Tabla 5 Frecuencia respecto a la confiabilidad de los modelos predictivos.	63
Tabla 6 Frecuencia respecto a la eficacia de los modelos predictivos ante cambios en los datos.	64
Tabla 7 Frecuencia de generalización de los algoritmos para prevenir amenazas en la seguridad bancaria.	65
Tabla 8 Frecuencia de actualización de métodos de entrenamiento para mejorar la detección de amenazas.	66
Tabla 9 Frecuencia de representatividad del conjunto de datos en la detección de amenazas reales.	67
Tabla 10 Frecuencia respecto a la revisión de métricas de rendimiento de los modelos predictivos.	68
Tabla 11 Frecuencia respecto al tiempo invertido en el entrenamiento de los modelos.	69
Tabla 12 Frecuencia de adecuación del tiempo de inferencia para la detección y respuesta a amenazas.	70
Tabla 13 Frecuencia respecto a la escalabilidad de los modelos predictivos ante mayor carga de trabajo.	71
Tabla 14 Frecuencia de facilidad de integración de modelos predictivos para la protección contra amenazas.	72

Tabla 15	Frecuencia respecto a la adaptación de los modelos predictivos a escenarios de ataque.	73
Tabla 16	Frecuencia respecto a la detección efectiva de ataques en comparación con otras amenazas.	74
Tabla 17	Frecuencia respecto a la generación de falsos positivos o negativos en los sistemas de detección.	75
Tabla 18	Frecuencia respecto a la rapidez en la detección de amenazas persistentes avanzadas.	76
Tabla 19	Frecuencia respecto a la respuesta oportuna del equipo ante ataques detectados.	77
Tabla 20	Frecuencia respecto a la recuperación del sistema tras un ataque.	78
Tabla 21	Frecuencia respecto a la resistencia del sistema frente a técnicas de evasión.	79
Tabla 22	Frecuencia respecto a la efectividad de los parches de seguridad frente a amenazas avanzadas.	80
Tabla 23	Frecuencia respecto a la actualización de modelos de seguridad frente a nuevas amenazas.	81
Tabla 24	Frecuencia respecto a la eficacia de las contramedidas implementadas frente a amenazas recientes.	82
Tabla 25	Frecuencia respecto a la adaptación rápida del sistema a nuevas amenazas.	83
Tabla 26	Frecuencia respecto a la detección y ajuste del sistema a la evolución de patrones de ataque.	84
Tabla 27	Correlación entre los modelos predictivos de IA y la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	85
Tabla 28	Correlación entre los modelos predictivos de IA y la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	86

Tabla 29	Correlación entre los modelos predictivos de IA y la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	87
----------	---	----

Tabla 30	Correlación entre los modelos predictivos de IA y la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	88
----------	--	----

ÍNDICE DE FIGURAS

Figura 1 Tendencia de Incidentes APTs en el Sector Bancario (2018-2023)	52
Figura 2 Impacto Financiero de APTs en el Sector Bancario por Año (2018-2023)	53
Figura 3 Frecuencia de tipos de Amenazas Persistentes Avanzadas (APT) en el Perú	55
Figura 4 Costo Promedio por Incidente (millones USD)	56
Figura 5 Gráfico de Resiliencia ante APTs por Componente	58
Figura 6 Distribución de Técnicas Utilizadas en APTs en el Sector Bancario (2023)	59
Figura 7 Frecuencia de uso de Softwares de Seguridad en la Nube en Bancos 2024	62
Figura 8 Productos de Ciberseguridad más comprados por Bancos	63
Figura 9 Frecuencia respecto a la eficacia de los algoritmos para abordar problemas.	67
Figura 10 Frecuencia respecto a la confiabilidad de los modelos predictivos.	68
Figura 11 Frecuencia respecto a la eficacia de los modelos predictivos ante cambios en los datos.	69
Figura 12 Frecuencia de generalización de los algoritmos para prevenir amenazas en la seguridad bancaria.	70
Figura 13 Frecuencia de actualización de métodos de entrenamiento para mejorar la detección de amenazas.	71
Figura 14 Frecuencia de representatividad del conjunto de datos en la detección de amenazas reales.	72
Figura 15 Frecuencia respecto a la revisión de métricas de rendimiento de los modelos predictivos.	73
Figura 16 Frecuencia respecto al tiempo invertido en el entrenamiento de los modelos.	74
Figura 17 Frecuencia de adecuación del tiempo de inferencia para la detección y respuesta a amenazas.	75

Figura 18 Frecuencia respecto a la escalabilidad de los modelos predictivos ante mayor carga de trabajo.	76
Figura 19 Frecuencia de facilidad de integración de modelos predictivos para la protección contra amenazas.	77
Figura 20 Frecuencia respecto a la adaptación de los modelos predictivos a escenarios de ataque.	78
Figura 21 Frecuencia respecto a la detección efectiva de ataques en comparación con otras amenazas.	79
Figura 22 Frecuencia respecto a la generación de falsos positivos o negativos en los sistemas de detección.	80
Figura 23 Frecuencia respecto a la rapidez en la detección de amenazas persistentes avanzadas.	81
Figura 24 Frecuencia respecto a la respuesta oportuna del equipo ante ataques detectados.	82
Figura 25 Frecuencia respecto a la recuperación del sistema tras un ataque.	83
Figura 26 Frecuencia respecto a la resistencia del sistema frente a técnicas de evasión.	84
Figura 27 Frecuencia respecto a la efectividad de los parches de seguridad frente a amenazas avanzadas.	85
Figura 28 Frecuencia respecto a la actualización de modelos de seguridad frente a nuevas amenazas.	86
Figura 29 Frecuencia respecto a la eficacia de las contramedidas implementadas frente a amenazas recientes.	87
Figura 30 Frecuencia respecto a la adaptación rápida del sistema a nuevas amenazas.	88
Figura 31 Frecuencia respecto a la detección y ajuste del sistema a la evolución de patrones de ataque.	89

RESUMEN

El presente estudio tuvo como objetivo determinar la relación entre los modelos predictivos de inteligencia artificial (IA) y la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima en 2024. La metodología adoptó un enfoque cuantitativo, de tipo correlacional y diseño no experimental de corte transversal. La muestra estuvo conformada por 79 empleados en ciberseguridad pertenecientes a diversas entidades bancarias. Los datos se recopilaron mediante un cuestionario estructurado basado en la escala Likert, que evaluó las dimensiones de protección frente a APTs. Para el análisis de los datos, se utilizó el coeficiente de correlación de Spearman (ρ), complementado con tablas descriptivas y gráficos generados en Microsoft Excel. Los resultados evidenciaron correlaciones positivas significativas entre los modelos predictivos de IA y las dimensiones estudiadas, destacando un coeficiente de 0.732 para la protección general contra APTs, lo que indica una relación alta. Asimismo, se observó que estos modelos permiten una evolución constante de las herramientas de seguridad cibernética y mejoran la resiliencia frente a amenazas dinámicas. Se concluye que la implementación de modelos predictivos de IA en el sector bancario contribuye significativamente a la detección, mitigación y prevención de APTs, fortaleciendo así la seguridad cibernética en un entorno crítico.

Palabras clave: Inteligencia artificial, Modelos predictivos, Ciberseguridad, Amenazas persistentes avanzadas

ABSTRACT

This study aimed to determine the relationship between predictive models of artificial intelligence (AI) and protection against advanced persistent threats (APTs) in the banking system of Lima in 2024. The methodology employed a quantitative approach, correlational type, and non-experimental cross-sectional design. The sample consisted of 79 cybersecurity specialists from various banking institutions. Data were collected through a structured questionnaire based on the Likert scale, which assessed the dimensions of detection, resilience, updating, and deployment of protection against APTs. Data analysis was conducted using Spearman's correlation coefficient (ρ), complemented by descriptive tables and graphs generated in Microsoft Excel. The results showed significant positive correlations between AI predictive models and the studied dimensions, with a coefficient of 0.732 for overall protection against APTs, indicating a strong relationship. Additionally, it was found that these models enable continuous evolution of cybersecurity tools and improve resilience against dynamic threats. It is concluded that the implementation of AI predictive models in the banking sector significantly contributes to the detection, mitigation, and prevention of APTs, thereby strengthening cybersecurity in a critical environment.

Keywords: artificial intelligence, predictive models, cybersecurity, advanced persistent threats

I. INTRODUCCIÓN

Los modelos predictivos de IA desempeñan un papel crucial en la protección contra amenazas persistentes en la ciberseguridad. Utilizando aprendizaje automático y aprendizaje profundo, estos modelos pueden analizar grandes cantidades de datos para identificar patrones y comportamientos maliciosos. A diferencia de los métodos tradicionales basados en firmas, los modelos predictivos de IA pueden detectar malware de día cero y adaptarse rápidamente a nuevas amenazas. Al combinar la heurística de comportamientos con la inteligencia de amenazas, la IA proporciona una detección más precisa y proactiva de vulnerabilidades, mejorando significativamente la seguridad de los usuarios en línea. (Varadaraj, 2024).

La creciente interconexión de dispositivos y la expansión del uso de internet han generado un entorno digital complejo y vulnerable, propenso a diversas ciberamenazas. Entre estas, las amenazas persistentes avanzadas (APTs) se destacan por su sofisticación y capacidad para evadir las defensas tradicionales durante períodos prolongados. Las APTs representan un riesgo significativo para la seguridad de la información, afectando tanto a individuos como a organizaciones. En este contexto, los modelos predictivos de inteligencia artificial (IA) han emergido como una solución innovadora y eficaz para enfrentar estos desafíos.

Los modelos predictivos de IA utilizan algoritmos de aprendizaje automático y aprendizaje profundo para analizar grandes volúmenes de datos en tiempo real, permitiendo la identificación de patrones y comportamientos anómalos que pueden indicar la presencia de actividades maliciosas. A diferencia de los métodos tradicionales de detección, que dependen principalmente de firmas conocidas y análisis retrospectivos, los modelos predictivos de IA pueden detectar y anticipar amenazas de día cero y otras formas de malware evolucionado. Esta capacidad predictiva no solo mejora la detección de amenazas, sino que también fortalece la capacidad de respuesta y mitigación, proporcionando una defensa más proactiva y robusta.

El propósito fundamental de este estudio fue examinar minuciosamente la relación existente entre la adopción de modelos predictivos basados en inteligencia artificial (IA) y la optimización de la efectividad en la defensa contra amenazas persistentes avanzadas. Mediante un análisis exhaustivo, se buscó evidenciar cómo la incorporación de la IA en las estrategias de ciberseguridad podría generar un cambio sustancial en la capacidad de identificación temprana y en la agilidad de respuesta frente a las amenazas digitales, proporcionando un enfoque holístico y vanguardista en la protección de los entornos digitales.

1.1.Planteamiento del problema

La creciente integración de la Inteligencia Artificial (IA) en diversos ámbitos ha transformado de manera significativa la gestión de datos personales, lo que genera desafíos críticos en términos de ciberseguridad. En particular, el uso de modelos predictivos de IA para analizar y procesar grandes volúmenes de datos expone a las instituciones, especialmente en el sector bancario, a riesgos de seguridad como la invasión de la privacidad y la discriminación. Además, se complica la garantía de transparencia y rendición de cuentas en estos sistemas. Esta problemática es especialmente relevante en un contexto global, incluidos los países en desarrollo, donde las regulaciones sobre protección de datos todavía están en proceso de evolución. La falta de marcos regulatorios robustos puede llevar a un uso indebido de la tecnología, incrementando la vulnerabilidad ante amenazas persistentes avanzadas (APT) en el ámbito financiero (Madrid, 2024).

El impacto de la Inteligencia Artificial y Big Data en la protección de datos personales ha suscitado una creciente preocupación respecto a cómo estas tecnologías afectan los derechos fundamentales de los trabajadores, particularmente en el sector bancario. La implementación de IA permite no solo el control y la vigilancia dentro del entorno laboral, sino que también extiende su influencia, generando perfiles de comportamiento y tomando decisiones

automatizadas que pueden comprometer la privacidad y la dignidad de los empleados. La capacidad de estas tecnologías para prever comportamientos plantea serios retos para la ciberseguridad, haciendo que las organizaciones sean más vulnerables a las APT. Es fundamental establecer un equilibrio entre la adopción de estas innovaciones tecnológicas y el respeto por los derechos fundamentales, abogando por una regulación adecuada que prevenga el avance tecnológico a expensas de la democracia y la libertad individual en el entorno bancario y más allá (Molina, 2023).

La inteligencia artificial está cada vez más presente en la vida cotidiana, generando un impacto considerable en la ciberseguridad. El avance de la digitalización ha provocado la acumulación de enormes volúmenes de datos personales, que son utilizados en modelos predictivos de IA para influir en el comportamiento de los usuarios. Este fenómeno plantea serias preocupaciones sobre la autonomía individual y la protección de datos, especialmente en contextos donde la desinformación y las amenazas persistentes avanzadas se entrelazan, aprovechando la vulnerabilidad de estos sistemas. En el sector bancario, la regulación y protección de estos datos se vuelve crucial para mitigar los riesgos asociados, ya que, si no se gestionan adecuadamente, estos sistemas de IA pueden amplificar las amenazas en línea y erosionar la confianza en las plataformas digitales, lo que podría tener consecuencias graves en la estabilidad financiera y la seguridad de los usuarios (Melo, 2022).

La inteligencia artificial (IA) presenta un dilema crucial en el ámbito de la ciberseguridad, actuando tanto como herramienta defensiva como vector de ataque. En particular, la IA generativa ha facilitado un aumento alarmante en la sofisticación de ataques de ingeniería social, como el 'phishing' y los 'deepfakes', que engañan a los usuarios y comprometen datos críticos. Recientes informes indican un incremento del 60% en los ataques de 'phishing' impulsados por IA, así como un aumento récord en los incidentes de ciberseguridad en España. Este panorama pone de manifiesto la creciente vulnerabilidad en la

protección de datos y sistemas, destacando la necesidad urgente de implementar modelos predictivos de IA que anticipen y mitiguen amenazas persistentes avanzadas, mejorando así la defensa ante un entorno de amenazas en constante evolución (BBVA, 2024).

La inteligencia artificial ha demostrado ser una herramienta poderosa, no solo para mejorar los sistemas de seguridad, sino también para sofisticar las amenazas digitales avanzadas. En el contexto español, la creciente implementación de modelos predictivos de IA plantea serias preocupaciones sobre su efectividad en la protección contra amenazas persistentes avanzadas. Aunque la IA generativa y los modelos predictivos pueden fortalecer las defensas cibernéticas al identificar patrones y anomalías, también pueden ser explotados por ciberdelincuentes para desarrollar ataques más elaborados, como el 'phishing' dirigido y la creación de 'deepfakes'. La capacidad de la IA para analizar grandes volúmenes de datos con alta precisión es una ventaja en la detección de amenazas, pero al mismo tiempo, incrementa el riesgo de que estas tecnologías se utilicen en contra de las mismas defensas que deberían proteger. (Pardiñas, 2020).

1.2.Descripción del problema

La creciente adopción de la inteligencia artificial (IA) en el sistema bancario de Lima ha transformado significativamente la gestión de datos y la seguridad en las operaciones financieras. Sin embargo, esta evolución tecnológica también ha planteado desafíos críticos en términos de ciberseguridad, especialmente frente a las amenazas persistentes avanzadas (APT). Los bancos en Lima están cada vez más expuestos a ataques cibernéticos sofisticados que utilizan IA generativa para realizar fraudes, como el 'phishing' dirigido y la creación de 'deepfakes', que buscan manipular a los usuarios y comprometer la integridad de los sistemas financieros.

Además, aunque los modelos predictivos de IA han sido implementados para mejorar la detección y prevención de estos ataques, su efectividad está en entredicho debido a la complejidad y evolución constante de las amenazas. La capacidad de los delincuentes para explotar vulnerabilidades en estos modelos resalta la necesidad urgente de fortalecer las medidas de seguridad, adaptando los sistemas de IA a un entorno de amenazas cada vez más dinámico. Este problema es particularmente relevante en Lima, donde la infraestructura bancaria aún enfrenta desafíos en términos de ciberseguridad, y donde la protección de datos financieros sensibles es una prioridad clave para mantener la confianza del público y la estabilidad del sistema financiero.

1.3. Formulación del problema

1.3.1. Problema general

¿Cómo los modelos predictivos de IA se relacionan con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?

1.3.2. Problemas específicos

- ¿Cómo los modelos predictivos de IA se relacionan con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?
- ¿Cómo los modelos predictivos de IA se relacionan con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?
- ¿Cómo los modelos predictivos de IA se relacionan con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?

1.4. Antecedentes

1.4.1. Antecedentes nacionales

Córdova et al. (2024) el presente estudio tuvo como propósito examinar la influencia de la inteligencia artificial (IA) en los procesos de evaluación crediticia dentro del Sistema Financiero peruano, evaluando su capacidad para incrementar la exactitud y optimizar la eficiencia en la administración del riesgo crediticio. La metodología empleada es de enfoque descriptivo y exploratorio, con un análisis documental sobre la regulación de la Superintendencia de Banca, Seguros y AFP (SBS) en ciberseguridad y la adopción de IA en América Latina. La muestra incluye información del Sistema Financiero peruano, donde los cuatro principales bancos concentran el 73.4% de los activos. Los resultados indican que la IA puede optimizar el acceso al crédito y fortalecer la estabilidad financiera, aunque enfrenta desafíos como altos costos y la necesidad de especialización. Se concluyó que, la integración de IA en la evaluación crediticia representa una oportunidad para reducir la morosidad y mejorar la toma de decisiones financieras, siempre que se implementen regulaciones adecuadas para proteger la seguridad y confidencialidad de los datos.

Sifuentes (2024) en el marco de su maestría, el estudio se propuso identificar el impacto de la innovación tecnológica en la satisfacción de los clientes de una entidad bancaria pública en Nuevo Chimbote durante el año 2023. Se adoptó un enfoque metodológico básico, con un diseño no experimental de corte transversal y un alcance correlacional-causal. La investigación abarcó una muestra de 384 usuarios, aplicando la técnica de la encuesta mediante un cuestionario estructurado. Los hallazgos obtenidos a través del análisis de regresión lineal evidenciaron un valor F de 218,290, un coeficiente de determinación r^2 de 0.424 y un nivel de significancia p de 0.00, claramente inferior al umbral de 0.01, lo cual valida una correlación significativa. En conclusión, se determinó que la innovación tecnológica incide de manera relevante en la satisfacción de los usuarios del banco público en Nuevo Chimbote durante 2023.

Espinoza (2024) en el contexto de su maestría, el estudio se propuso evaluar la influencia de la inteligencia artificial (IA) en la eficiencia de la gestión de incidentes de soporte en la Gerencia de Sistemas del sector justicia. La investigación se desarrolló bajo un enfoque cuantitativo, con un diseño metodológico no experimental y de tipo transversal. Se aplicaron encuestas estructuradas a una muestra compuesta por 45 profesionales de la Gerencia de Sistemas. Los resultados evidenciaron una correlación significativa entre el uso de la IA y la mejora en la resolución de incidentes, con un coeficiente Rho de Spearman de 0.844. Se identificó que la automatización y la optimización de procesos son factores determinantes. En conclusión, la integración de la IA contribuye a la reducción de los tiempos de respuesta, mejora la satisfacción de los usuarios y disminuye los costos operativos, demostrando su efectividad en la gestión de incidentes.

Maurtua (2024) en el marco de su maestría, el estudio tuvo como propósito analizar la relación entre la inteligencia artificial y la modernización de la gestión administrativa en el Gobierno Regional de Piura durante el año 2024. Se utilizó una metodología de tipo básica, bajo un enfoque cuantitativo, con un diseño no experimental, correlacional y de corte transversal. La investigación abarcó una muestra de 396 colaboradores del sector público. Los resultados descriptivos revelaron que el 53,8 % de los participantes percibió un nivel bajo de implementación de inteligencia artificial, mientras que el 43,6 % evaluó la gestión administrativa como deficiente. El análisis estadístico, con un nivel de significancia de 0.000, mostró una correlación moderadamente positiva entre las variables, con un coeficiente de 0.628. Se concluyó que la inteligencia artificial incide de manera significativa en la modernización de la gestión administrativa en la región.

Sánchez (2024) en el contexto de su maestría, el estudio se propuso analizar el efecto de la inteligencia artificial (IA) en la transformación de los modelos de negocio de las pequeñas y medianas empresas (PYMES) en Lima Metropolitana durante 2024. La investigación adoptó

una metodología mixta, integrando enfoques cuantitativos y cualitativos, con la participación activa de diversas PYMES y especialistas del sector. La muestra abarcó organizaciones de distintos tamaños y sectores, lo que permitió obtener una perspectiva amplia sobre la adopción de la IA. Los hallazgos evidenciaron que la IA potencia la eficiencia operativa, personaliza la oferta de servicios, redefine las estrategias de mercado y las estructuras organizacionales, además de facilitar la automatización de procesos y optimizar la toma de decisiones mediante análisis predictivo. Asimismo, se observó un impacto significativo en la interacción con los clientes, en los modelos de generación de ingresos y en la gestión de la cadena de suministro. En conclusión, se determinó que la IA es un factor determinante para la innovación y el desarrollo sostenible de las PYMES, reforzando su competitividad y apoyando el logro de los Objetivos de Desarrollo Sostenible (ODS).

Chávez et al. (2023) el estudio se propuso evaluar de manera exhaustiva los beneficios y desafíos asociados con la aplicación de la inteligencia artificial (IA) en la seguridad de la información dentro del sector cibernético. Para ello, se realizó una revisión sistemática de literatura basada en artículos y documentos publicados entre 2019 y 2023 en las bases de datos Scopus, Scielo y Alicia. Los hallazgos revelaron que la IA aporta notables ventajas, tales como la detección en tiempo real de ciberataques, la prevención proactiva de amenazas y la disminución de la carga operativa de los equipos de seguridad. No obstante, se identificaron limitaciones significativas, entre las que destacan la complejidad en la implementación, la dependencia de grandes volúmenes de datos, la posible desconfianza debido a la toma de decisiones autónoma y las preocupaciones en torno a la privacidad de los datos de los usuarios. En conclusión, aunque la IA presenta un potencial considerable para optimizar la protección de la información sensible en las organizaciones, resulta fundamental mitigar estas barreras para garantizar una implementación eficaz y segura.

Manrique (2023) la investigación se centró en llevar a cabo un exhaustivo análisis de la literatura científica con el propósito de evaluar la integración del Big Data en la ciberseguridad mediante la utilización de inteligencia artificial (IA) durante el periodo comprendido entre 2014 y 2023. Para ello, se realizó una revisión detallada de artículos académicos provenientes de diversas bases de datos reconocidas, tales como Dialnet, Scielo, Google Académico, Redalyc, Proquest, ScienceDirect y Scopus. Se seleccionaron 29 artículos pertinentes, los cuales fueron sistemáticamente organizados en una tabla con el fin de facilitar un análisis estructurado. Los resultados revelaron que la aplicación del Big Data contribuye significativamente al entrenamiento y optimización de algoritmos de aprendizaje automático, lo que permite una detección temprana y una prevención proactiva de ciberataques mediante el procesamiento y análisis de grandes volúmenes de datos. En conclusión, el estudio determinó que la sinergia entre Big Data e IA, especialmente a través de técnicas de Machine Learning, potencia sustancialmente la capacidad de identificar y gestionar amenazas cibernéticas, destacando el rol fundamental de estas tecnologías en el fortalecimiento de la seguridad digital.

Mejía (2023) la investigación titulada "Inteligencia artificial y sus efectos en la protección de datos personales en la ley N° 29733, Lima 2023" tuvo como propósito examinar el impacto de la inteligencia artificial (IA) en la salvaguardia de los datos personales de acuerdo con la Ley N° 29733 en Lima. El estudio se basó en un enfoque cualitativo, empleando métodos hermenéuticos e inductivos, con un diseño de teoría fundamentada de nivel descriptivo. El ámbito de la investigación se enfocó en Lima, utilizando técnicas como entrevistas y análisis documental. Se aplicaron instrumentos específicos, como una guía de entrevistas y un esquema de análisis documental. Los hallazgos evidenciaron que la IA plantea desafíos significativos en la protección de datos personales, resaltando la importancia de una gestión eficiente para mitigar posibles impactos adversos. Finalmente, el estudio formuló recomendaciones orientadas a optimizar la protección de los datos personales en un entorno influenciado por la

inteligencia artificial, aportando una perspectiva valiosa para la mejora de la gestión y el entendimiento de estos efectos.

Villacorta et al. (2023) se llevó a cabo una exhaustiva revisión sistemática con el propósito de evaluar la influencia de la inteligencia artificial (IA) en la gestión de servicios de tecnología de la información (TI) dentro de las organizaciones. Para este análisis, se utilizaron motores de búsqueda académicos como SciELO, Scopus, Google Académico y World Wide Science, aplicando criterios de estandarización de los documentos en cuanto al idioma y el año de publicación. Los hallazgos revelaron que la implementación de la IA en la gestión de servicios de TI genera un impacto positivo al optimizar la eficiencia operativa, mejorar la calidad del servicio, disminuir los costos operativos y apoyar la toma de decisiones basadas en datos. No obstante, también se identificaron desafíos significativos, tales como el riesgo de pérdida de empleos debido a la automatización de procesos y la percepción negativa de los usuarios en áreas de atención al cliente, donde la IA puede ser percibida como menos empática y humana. En conclusión, la IA puede ofrecer beneficios significativos en la gestión de servicios de TI, pero también presenta desafíos como la sustitución de roles humanos y posibles vulneraciones a la privacidad y seguridad. Además, los beneficios de la digitalización no se distribuyen equitativamente, favoreciendo a las organizaciones con mayores habilidades técnicas y organizacionales.

Montalvo (2022) en su maestría estableció como objetivo la implementación de estándares de seguridad para la protección de datos de tarjetas de pago en las entidades financieras. La metodología utilizada fue cualitativa, con un enfoque interpretativo y diseño de investigación acción. Las técnicas empleadas fueron entrevistas semiestructuradas, observación y análisis documental, utilizando triangulación para validar los resultados. No se especificó una muestra en el texto. Las conclusiones del estudio señalaron que, para salvaguardar la información de las tarjetas de pago, es fundamental adoptar estándares de

seguridad basados en los principios de disponibilidad, confiabilidad e integridad de los datos. Estos estándares deben estar respaldados por mecanismos robustos que aborden la seguridad durante las etapas de almacenamiento, procesamiento y transmisión de la información, asegurando una alineación coherente con la infraestructura tecnológica existente. Asimismo, se destacó la relevancia de los sistemas biométricos y las soluciones de autoaprendizaje como componentes críticos para mitigar el riesgo de fraude cibernético en cajeros automáticos (ATM), canales electrónicos y puntos de venta, proporcionando una defensa proactiva y adaptable frente a las amenazas emergentes.

1.4.2. Antecedentes internacionales

Farias y Centeno (2024) investigaron modelos de inteligencia artificial (IA) para la prevención de ataques cibernéticos en organizaciones. El objetivo del estudio fue analizar diferentes modelos de IA, El estudio evaluó la efectividad del aprendizaje automático y el análisis predictivo en la detección de patrones y anomalías en el tráfico de red y el comportamiento del usuario. Se utilizó una metodología basada en la revisión de literatura y el análisis de casos reales para examinar la aplicación práctica y los desafíos asociados con estos modelos. Los hallazgos evidenciaron que los modelos de inteligencia artificial (IA) son efectivos para la detección temprana y la mitigación de amenazas cibernéticas. No obstante, también se identificaron limitaciones, como la necesidad de contar con datos de alta calidad y los desafíos relacionados con la gestión de falsos positivos. En conclusión, el estudio subrayó la importancia crítica de la calidad de los datos en las estrategias de ciberseguridad basadas en IA y enfatizó la relevancia de la colaboración continua entre los equipos de seguridad para optimizar y fortalecer estos modelos predictivos.

Pedraza (2023) investigó el impacto actual y los desafíos futuros de la inteligencia artificial (IA) en la sociedad. El objetivo del trabajo fue analizar cómo la IA afecta diversos sectores laborales, su implementación en la vida humana y los riesgos, desafíos y oportunidades asociados. Utilizó un enfoque metodológico basado en una revisión exhaustiva de la literatura y estudios de casos para evaluar estos aspectos. Los resultados indicaron que la IA ha transformado significativamente la condición humana, influenciando desde la contratación laboral hasta la seguridad de datos y la interacción social. Además, destacó que, aunque la IA mejora la productividad y reduce riesgos en actividades, también presenta peligros como la discriminación predictiva y el aumento de brechas sociales. Las conclusiones subrayaron que, a pesar de sus avances tecnológicos, la IA no puede igualar la capacidad intelectual humana debido a su falta de conciencia y procesos de socialización. Se enfatizó la necesidad de un enfoque ético y regulado en la implementación de la IA para maximizar sus beneficios y mitigar sus riesgos.

Castro (2019) realizó una investigación con el objetivo de analizar el tratamiento de datos personales mediante herramientas de procesamiento automatizado de datos y explorar los desafíos que estas herramientas suponen para la privacidad y protección de datos personales. Utilizando un enfoque analítico, se examinó el impacto de la economía de datos en actividades como el comercio de datos personales, la mercadotecnia digital y el uso de plataformas digitales. Además, se llevó a cabo un análisis de los procesos de analítica de datos para entender su alcance y el riesgo de sobreexposición de la privacidad. El estudio incluyó una revisión de la legislación nacional e internacional, comparando instrumentos jurídicos relevantes para evaluar el avance en la protección de estos derechos y su relación con otros derechos humanos. Los resultados destacaron la necesidad de fortalecer las medidas de protección de datos y privacidad, y se emitieron recomendaciones para mejorar la interoperabilidad entre las actividades estatales y la industria privada. En conclusión, la

investigación subrayó la importancia de una protección robusta de los datos personales frente a los retos que plantea la analítica de big data.

Sáenz (2023) en su artículo planteó como objetivo analizar cómo el régimen jurídico europeo responde a los avances digitales como el Big Data, la inteligencia artificial y la tecnología blockchain, con énfasis en la protección de los usuarios, la competencia económica y la seguridad jurídica en las relaciones con consumidores y usuarios. La metodología empleada incluye una revisión normativa y un análisis de la legislación vigente en la Unión Europea respecto a la transparencia, el uso de datos y la regulación de los algoritmos en los procesos automatizados. La muestra abarca las normativas, políticas y directrices europeas más relevantes en este ámbito. Los resultados muestran que, aunque hay un progreso significativo, aún existen vacíos legales en cuanto a la regulación de decisiones automatizadas y contratos inteligentes. Se concluyó que hay la necesidad urgente de fortalecer el marco normativo para garantizar la seguridad jurídica y proteger adecuadamente a los usuarios en el entorno digital.

Dakalbab et al. (2024) en su artículo plantearon como objetivo analizar el uso de técnicas de Inteligencia Artificial (IA) en los mercados financieros mediante una Revisión Sistemática de Literatura (SLR). La metodología consistió en revisar 143 artículos publicados entre 2015 y 2023 que implementaron IA en el trading financiero. La muestra abarcó investigaciones sobre ocho mercados financieros y diferentes tipos de activos, considerando enfoques de análisis técnico y fundamental junto con diversas técnicas de IA. Los resultados indican que el análisis técnico es más utilizado que el fundamental, y que el 16% de los estudios automatiza completamente el proceso de trading. Además, se identificaron 40 técnicas de IA, con un predominio de modelos de deep learning. Se concluyó que, la aplicación de IA en los mercados financieros es un campo prometedor, con múltiples enfoques predictivos ya desarrollados, por lo que se brindan recomendaciones y orientación para futuras investigaciones.

Jada y Mayayise (2024) realizaron una revisión sistemática de la literatura (SLR) para evaluar el impacto de las tecnologías basadas en inteligencia artificial (IA) en la ciberseguridad organizacional y determinar su efectividad en comparación con los enfoques tradicionales de ciberseguridad. Utilizaron el diagrama de flujo PRISMA para guiar el proceso de revisión, incluyendo artículos revisados por pares de 2018 a 2023 de bases de datos como EBSCO Host, Google Scholar, Science Direct, ProQuest y SCOPUS, y sintetizaron 73 artículos. Los resultados mostraron que la IA puede impactar la ciberseguridad en todo su ciclo de vida, proporcionando beneficios como automatización, inteligencia de amenazas y mejora de la defensa cibernética. Sin embargo, también presenta desafíos como ataques adversariales y la necesidad de datos de alta calidad, lo que puede llevar a la ineficiencia de la IA. Estos hallazgos afirman la influencia positiva de la IA en la ciberseguridad, mejorando la efectividad y resiliencia. Las conclusiones subrayan que, aunque la IA introduce vulnerabilidades adicionales y requiere consideraciones regulatorias más estrictas, su incorporación en la ciberseguridad organizacional tiene un efecto predominantemente beneficioso, estableciendo una base para futuras investigaciones y ayudando a las organizaciones a tomar decisiones informadas sobre la implementación de soluciones de IA.

Kaur et al. (2023) llevaron a cabo una revisión sistemática de la literatura (SLR) para examinar el estado actual de la investigación sobre aplicaciones de inteligencia artificial (IA) en ciberseguridad. El objetivo fue identificar las técnicas de IA aplicadas en el dominio de la ciberseguridad y las actividades que se han beneficiado de esta tecnología. Utilizando la base de datos de Scopus, se analizaron 236 estudios primarios seleccionados de un total de 2395 artículos relacionados, publicados entre 2010 y febrero de 2022. El estudio analizó la literatura en términos de una taxonomía de IA en ciberseguridad, frecuencia de publicaciones por año y región geográfica, tipo de contribución en ciberseguridad y tipo de técnica de IA utilizada. Los resultados revelaron un aumento en el número de publicaciones, destacando la necesidad de

mejorar la adquisición y representación de datos históricos para implementar soluciones de ciberseguridad basadas en IA. En conclusión, este estudio clasificó los estudios primarios para integrar el estado de la literatura en el área y propuso direcciones futuras de investigación para abordar los problemas emergentes y facilitar la adopción exitosa de la IA en ciberseguridad.

De la Hoz et al. (2023) en su artículo de estudio plantearon como objetivo analizar cómo la inteligencia artificial contribuye a la gestión de los procesos en la auditoría financiera. La metodología utilizada es un estudio descriptivo basado en la revisión de libros y artículos científicos sobre auditoría financiera y nuevas tecnologías. Los resultados indican que la inteligencia artificial ha mejorado el análisis de grandes volúmenes de datos, la evaluación de riesgos, la automatización de tareas repetitivas y la detección de fraudes, reduciendo el tiempo y el esfuerzo en estos procesos. Se concluye que, la inteligencia artificial ha optimizado la auditoría financiera al hacerla más eficiente y precisa, aunque no puede reemplazar completamente el juicio y el escepticismo profesional del auditor humano.

Kalogiannidis et al. (2024) examinaron la eficacia de las tecnologías de inteligencia artificial (IA) en la evaluación predictiva de riesgos y su contribución a la continuidad del negocio. El objetivo de esta investigación fue comprender cómo diversos componentes de la IA, como el procesamiento del lenguaje natural (NLP), la analítica de datos impulsada por IA, el mantenimiento predictivo con IA y la integración de IA en la planificación de respuesta a incidentes, mejoran la evaluación de riesgos y apoyan la continuidad del negocio en un entorno donde las empresas enfrentan múltiples riesgos, incluidos desastres naturales, ciberataques y fluctuaciones económicas. Se utilizó un diseño transversal y un método cuantitativo para recopilar datos de una muestra de 360 especialistas en tecnología. Los resultados mostraron que las tecnologías de IA tienen un impacto significativo en la continuidad del negocio y la evaluación predictiva de riesgos. En particular, se descubrió que el NLP mejoró la precisión y rapidez de los procedimientos de evaluación de riesgos. La integración de la IA en los planes

de respuesta a incidentes fue especialmente efectiva, disminuyendo considerablemente las interrupciones en la empresa y mejorando la recuperación de eventos imprevistos. Se recomendó que las empresas inviertan en habilidades de IA, especialmente en áreas como el NLP para la evaluación automatizada de riesgos, la analítica de datos para la detección oportuna de riesgos, el mantenimiento predictivo para la eficiencia operativa y la planificación de respuesta a incidentes mejorada con IA para la gestión de crisis.

Sarker (2022) presentó una visión comprensiva sobre la modelización basada en inteligencia artificial (IA), destacando su papel crucial en la cuarta revolución industrial (Industria 4.0). El objetivo del estudio fue analizar las técnicas y capacidades de la IA que pueden desarrollar sistemas inteligentes y automatizados en diversas áreas de aplicación. Utilizó un enfoque exhaustivo para explorar diez categorías populares de técnicas de IA, incluyendo aprendizaje automático, aprendizaje profundo, procesamiento del lenguaje natural, descubrimiento de conocimiento y modelización de sistemas expertos. Los resultados revelaron que las técnicas de IA han demostrado ser beneficiosas en numerosos campos como la inteligencia empresarial, las finanzas, la atención médica, el reconocimiento visual, las ciudades inteligentes, el Internet de las cosas (IoT) y la ciberseguridad. La investigación también destacó la necesidad de entrenar algoritmos de aprendizaje complejos con datos y conocimientos específicos de la aplicación objetivo para mejorar la toma de decisiones inteligente. En las conclusiones, se subrayaron los futuros aspectos de la IA hacia la automatización y los sistemas de computación inteligentes, y se señalaron varios problemas de investigación dentro del alcance del estudio, proporcionando una guía de referencia para investigaciones y desarrollos futuros en los dominios de aplicación relevantes.

1.5. Justificación de la investigación

1.5.1. Justificación práctica

Desde una perspectiva práctica, la implementación de modelos predictivos de IA en la detección y mitigación de APTs transformó las estrategias de ciberseguridad utilizadas por las organizaciones. Este estudio evaluó la efectividad de estos modelos y ofreció recomendaciones prácticas para su implementación y optimización. Los resultados de esta investigación fueron aplicados directamente por profesionales de la ciberseguridad, mejorando la capacidad de respuesta ante amenazas y fortaleciendo las defensas de las infraestructuras digitales. Al proporcionar una guía basada en evidencias, esta tesis equipó a las organizaciones con herramientas necesarias para enfrentar de manera proactiva las amenazas cibernéticas en un panorama digital cada vez más complejo.

1.5.2. Justificación teórica

La integración de IA en la ciberseguridad representó una evolución significativa en la manera de abordar las amenazas digitales. Este estudio se fundamentó en teorías de aprendizaje automático y aprendizaje profundo, explorando cómo estos enfoques mejoraron la detección y mitigación de APTs. A través de la revisión de la literatura existente y el análisis de casos prácticos, se aportó al cuerpo teórico de la ciberseguridad, proporcionando nuevas perspectivas sobre la efectividad y las limitaciones de los modelos predictivos de IA. Esta investigación llenó un vacío en el conocimiento existente, ofreciendo una comprensión más profunda de cómo la IA revolucionó la protección contra amenazas persistentes avanzadas.

1.5.3. Justificación social

En la era digital, la protección de la información fue crucial para la seguridad y el bienestar de la sociedad. Las amenazas persistentes avanzadas (APTs) no solo afectaron a grandes corporaciones y entidades gubernamentales, sino que también pusieron en riesgo la privacidad y los datos personales de los individuos. La creciente dependencia de servicios en línea y la interconectividad de los dispositivos aumentaron la exposición a estas amenazas, afectando potencialmente la economía y la estabilidad social. Al investigar cómo los modelos predictivos de inteligencia artificial (IA) mejoraron la defensa contra APTs, se contribuyó a la creación de un entorno digital más seguro, protegiendo tanto a individuos como a organizaciones de las consecuencias devastadoras de los ciberataques.

1.6.Limitaciones de la investigación

1.6.1. Limitaciones prácticas

Los modelos predictivos de IA para combatir APTs en el sector bancario enfrentaron desafíos relacionados con la necesidad de grandes volúmenes de datos históricos, los cuales, en muchos casos, resultaron insuficientes o estuvieron sujetos a restricciones de sensibilidad, limitando su capacidad para detectar nuevas amenazas. Además, la integración con sistemas preexistentes presentó dificultades debido a la falta de compatibilidad tecnológica y los elevados costos asociados. La naturaleza evolutiva de las APTs demandó actualizaciones frecuentes de los modelos, lo que complicó mantener su efectividad frente a tácticas emergentes.

1.6.2. Limitaciones teóricas

Las limitaciones teóricas incluyeron la dificultad de modelar tácticas complejas que explotaron comportamientos humanos, como el phishing, y la incapacidad de anticipar técnicas de ataque novedosas. Los modelos tendieron a sobre ajustarse a los datos históricos, lo que redujo su capacidad de generalización frente a ataques nuevos. Además, la interpretación de modelos basados en redes neuronales profundas resultó especialmente compleja, lo que generó desconfianza en el uso de sus resultados para la toma de decisiones.

1.6.3. Limitaciones metodológicas

La investigación estuvo limitada por el tamaño y la representatividad de la muestra seleccionada entre los empleados del sector bancario. Una muestra insuficientemente amplia o poco diversa afectó la generalización de los resultados a contextos similares. Asimismo, la necesidad de una actualización continua de los modelos para adaptarse a nuevas amenazas representó un desafío metodológico, ya que estos no siempre demostraron ser resilientes ante cambios rápidos en el entorno de amenazas o en la infraestructura bancaria.

1.7.Objetivos

1.7.1. Objetivo general

Determinar si los modelos predictivos de IA se relacionan con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

1.7.2. Objetivos específicos

- Determinar si los modelos predictivos de IA se relacionan con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

- Determinar si los modelos predictivos de IA se relacionan con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.
- Determinar si los modelos predictivos de IA se relacionan con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

1.8.Hipótesis

1.8.1. Hipótesis general

Los modelos predictivos de IA se relacionan con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

1.8.2. Hipótesis específicas.

- Los modelos predictivos de IA se relacionan con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.
- Los modelos predictivos de IA se relacionan con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.
- Los modelos predictivos de IA se relacionan con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

II. MARCO TEÓRICO

2.1. Marco conceptual

2.1.1. Análisis del panorama actual de amenazas persistentes avanzadas (APTs) en el sistema bancario.

2.1.1.1. Recopilación de datos sobre incidentes de APTs en el sistema bancario de Lima entre 2020 y 2024. Durante 2022, Perú registró más de 15 mil millones de intentos de ciberataques, evidenciando una creciente amenaza en el sector bancario. Expertos en ciberseguridad señalaron que los delitos motivados financieramente representaron el mayor volumen de incidentes, con un 73.9% del total. Además, se identificó un incremento en el uso de malware sofisticado y técnicas de phishing dirigidas a usuarios de servicios financieros. La falta de protocolos de seguridad robustos en algunas instituciones permitió que estos ataques lograran vulnerar sistemas internos. Ante esta situación, organismos reguladores recomendaron reforzar la ciberseguridad y promover campañas de concientización para clientes y empleados. (Saldarriaga, 2024)

En 2023, la Policía Nacional del Perú registró un promedio de 107 intentos de ciberataques por minuto, evidenciando una creciente amenaza en el sector bancario. Expertos en ciberseguridad señalaron que Perú se encontraba entre los 20 países más atacados a nivel mundial, siendo las entidades financieras los principales objetivos. Se reportó un aumento en los ataques dirigidos a plataformas de banca digital, afectando transacciones en línea y servicios de pago. Las técnicas utilizadas por los ciberdelincuentes incluyeron la explotación de vulnerabilidades en software bancario y el uso de redes de bots para ataques distribuidos. Como respuesta, los bancos comenzaron a implementar autenticación multifactor y monitoreo en tiempo real para detectar actividades sospechosas. (Delgado, 2023)

En octubre de 2024, Interbank sufrió un acceso no autorizado a sus sistemas, resultando en la filtración de datos personales y financieros de numerosos clientes. Este ataque provocó

interrupciones en servicios digitales, incluyendo la billetera electrónica Plin. La Superintendencia de Banca, Seguros y AFP (SBS) inició acciones para determinar posibles infracciones y supervisar las medidas de seguridad implementadas por el banco. Según informes preliminares, los atacantes explotaron una vulnerabilidad en los servidores de autenticación, permitiéndoles acceder a información confidencial. La filtración de datos generó preocupación entre los usuarios, quienes reportaron intentos de fraude en sus cuentas. Interbank reforzó sus sistemas de seguridad y ofreció asesoría a los clientes afectados para minimizar los riesgos de suplantación de identidad. (Matos, 2024)

2.1.1.2. Identificación de patrones y tendencias en los ataques APT registrados en el sistema bancario durante el período analizado. Durante el período analizado, los ataques APT dirigidos al sistema bancario en Perú han evolucionado en términos de frecuencia, sofisticación y objetivos estratégicos. Se han identificado patrones claros en estos incidentes, lo que permite comprender mejor la naturaleza de estas amenazas y las estrategias utilizadas por los atacantes.

- **Aumento sostenido de intentos de ataque:** Durante los últimos años, se ha observado un crecimiento acelerado en los intentos de ciberataques dirigidos al sector bancario en Perú. En 2022, se registraron más de 15 mil millones de intentos de ataque, con un predominio del uso de malware avanzado y técnicas de phishing para comprometer credenciales bancarias. Esta tendencia se intensificó en 2023, con un promedio de 107 intentos de ataque por minuto, posicionando a Perú como uno de los países más afectados a nivel global en términos de ciberseguridad bancaria. Estos datos reflejan una evolución de los actores APT, quienes han desarrollado campañas persistentes para infiltrarse en infraestructuras críticas del sector financiero.
- **Diversificación de métodos de ataque:** Los atacantes han dejado de depender únicamente de técnicas tradicionales como el phishing y han comenzado a explotar

vulnerabilidades en software bancario y plataformas digitales. En el caso del ciberataque a Interbank en 2024, los atacantes lograron acceder a información confidencial mediante la explotación de fallos en los servidores de autenticación. Este tipo de ataque evidencia una tendencia en la que los grupos APT buscan comprometer directamente la infraestructura bancaria en lugar de solo atacar a los usuarios. Además, la filtración de datos financieros ha sido un objetivo recurrente, derivando en intentos de fraude y suplantación de identidad.

- **Uso intensivo de técnicas de ingeniería social y malware especializado:** La combinación de ingeniería social y software malicioso ha sido una táctica recurrente en los ataques APT recientes. Se ha observado un aumento en el uso de troyanos bancarios y herramientas de acceso remoto que permiten a los atacantes mantener el control de los sistemas comprometidos durante largos períodos sin ser detectados. Las campañas de phishing y vishing han seguido siendo utilizadas para engañar a empleados y clientes, facilitando la obtención de credenciales de acceso. Estos métodos reflejan un cambio en la estrategia de los atacantes, quienes buscan maximizar el impacto de sus incursiones a través de la manipulación humana y el uso de software avanzado.
- **Impacto en la seguridad y respuesta de las instituciones financieras:** El creciente número de ataques ha obligado a las instituciones bancarias a reforzar sus estrategias de seguridad. La implementación de autenticación multifactor, sistemas de monitoreo en tiempo real y políticas de ciberseguridad más estrictas han sido algunas de las medidas adoptadas para mitigar los riesgos. Sin embargo, a pesar de estos esfuerzos, la sofisticación de los ataques APT sigue en aumento, lo que sugiere la necesidad de estrategias más dinámicas y adaptativas. La filtración de datos y el fraude financiero siguen siendo preocupaciones clave, lo que indica que los bancos deben seguir invirtiendo en medidas de prevención y detección avanzada.

Estos patrones reflejan cómo los ataques APT en el sistema bancario de Perú han pasado de intentos masivos a operaciones dirigidas y sofisticadas, representando una amenaza constante para la seguridad financiera del país.

2.1.2. Evaluación del impacto de los modelos predictivos de IA en la resiliencia del sistema bancario frente a APTs.

2.1.2.1. Análisis de casos en los que se han implementado modelos predictivos de IA para la detección y mitigación de APTs en el sistema bancario.

- **Implementación de IA en Interbank:** Durante 2023, una entidad financiera en Lima fortaleció su infraestructura de ciberseguridad mediante el uso de modelos predictivos de inteligencia artificial para la detección temprana de amenazas persistentes avanzadas (APTs). Estos modelos, basados en machine learning, fueron entrenados para analizar patrones irregulares en el tráfico de red y detectar anomalías en tiempo real. En octubre de 2024, esta tecnología permitió identificar un intento de acceso no autorizado a su sistema de banca en línea, mitigando el riesgo de filtración de datos sensibles y evitando interrupciones en sus servicios digitales. (Matos, 2024).
- **Uso de IA en la detección de ciberataques:** Desde 2022, otra institución bancaria ha implementado tecnologías avanzadas para mejorar la detección y respuesta ante ataques dirigidos. A través de un sistema basado en deep learning, ha logrado analizar transacciones en tiempo real y detectar comportamientos sospechosos asociados a tácticas empleadas en APTs. Gracias a esta estrategia, en 2023 se bloqueó un intento de infiltración en el sistema de transferencias interbancarias, reduciendo el impacto potencial de un ataque dirigido y fortaleciendo la resiliencia de la infraestructura financiera. (Delgado, 2023).

- **Estrategias de prevención mediante IA:** En 2024, se adoptó un enfoque basado en inteligencia artificial para la prevención de ataques cibernéticos dirigidos a clientes y sistemas internos de otra entidad bancaria. La integración de modelos de aprendizaje automático permitió correlacionar eventos de seguridad en tiempo real, lo que facilitó la detección de patrones asociados con intentos de APTs antes de que los atacantes lograran comprometer información crítica. Con esta medida, se reforzó la capacidad de respuesta ante amenazas y se logró una mejor protección de los activos digitales. (Andina, 2024).

2.1.2.2. Comparación de la respuesta y recuperación del sistema bancario ante APTs en escenarios con y sin el uso de modelos predictivos de IA. La resiliencia del sistema bancario ante amenazas persistentes avanzadas (APTs) es clave para su seguridad y estabilidad. La implementación de modelos predictivos de inteligencia artificial (IA) ha mejorado significativamente la detección, respuesta y recuperación ante estos ataques. Este análisis compara la eficacia del sistema bancario en escenarios con y sin IA predictiva.

A. Escenario sin el uso de modelos predictivos de IA. En ausencia de IA, el sistema bancario depende de métodos tradicionales como firewalls y antivirus. Sus principales limitaciones incluyen:

- **Detección tardía:** Identificación reactiva de amenazas, lo que permite que los ataques permanezcan ocultos por más tiempo.
- **Tiempo de respuesta prolongado:** Dependencia de intervención manual, retrasando la mitigación.
- **Impacto financiero y reputación:** Mayor riesgo de pérdidas y daño a la confianza de los clientes.

- Recuperación lenta: Dependencia de auditorías manuales y análisis forenses, aumentando los tiempos de restauración.

B. Escenario con el uso de modelos predictivos de IA con IA predictiva, la respuesta ante APTs mejora significativamente en los siguientes aspectos.

- Detección proactiva: Análisis en tiempo real para identificar patrones sospechosos antes de que se materialicen ataques.
- Respuesta automatizada: Bloqueo de transacciones sospechosas y restricción de accesos sin intervención humana.
- Reducción del impacto financiero: Mitigación rápida de ataques, minimizando pérdidas.
- Recuperación eficiente: Recomendaciones automatizadas para acelerar la restauración del sistema.

C. Análisis comparativo.

Tabla 1

Análisis comparativo

Aspecto	Sin IA Predictiva	Con IA Predictiva
Detección de APTs	Reactiva, tardía	Proactiva, en tiempo real
Tiempo de respuesta	Lento, manual	Rápido, automatizado
Impacto financiero	Elevado	Optimizado
Recuperación del sistema	Lenta	Ágil
Confiabilidad	Riesgo elevado	Mayor seguridad

2.1.3. Análisis del papel de los modelos predictivos de IA en la actualización y mejora de las estrategias de protección contra APTs en el sistema bancario.

2.1.3.1. Evaluación de cómo los modelos predictivos de IA han contribuido a la adaptación y mejora de las medidas de ciberseguridad frente a APTs. La creciente complejidad y frecuencia de las Amenazas Persistentes Avanzadas (APTs) han impulsado a las instituciones financieras a reforzar sus estrategias de ciberseguridad. En este contexto, los modelos predictivos basados en Inteligencia Artificial (IA) han emergido como herramientas esenciales para la detección y mitigación proactiva de estas amenazas. (Hidalgo, 2020).

Las capacidades de los modelos predictivos de IA en ciberseguridad:

- Análisis en tiempo real: La IA permite examinar vastas cantidades de datos en tiempo real, identificando patrones y comportamientos anómalos que podrían indicar actividades maliciosas.
- Automatización de respuestas: Más allá de la detección, la IA puede automatizar la respuesta a incidentes de seguridad, acelerando la mitigación y reduciendo el impacto de los ataques.
- Reducción de falsos positivos: Al comprender el contexto y aprender de datos previos, los modelos de IA disminuyen la incidencia de alertas falsas, permitiendo que los equipos de seguridad se concentren en amenazas genuinas y optimicen el uso de recursos.
- Actualización dinámica de estrategias de seguridad: Las entidades bancarias que han integrado modelos predictivos de IA han demostrado una mayor capacidad para adaptar sus protocolos de seguridad en función de nuevas amenazas detectadas, manteniendo sus sistemas protegidos de forma continua.

Las instituciones bancarias, debido a la sensibilidad de la información que manejan, han integrado modelos predictivos de IA para fortalecer sus defensas contra APTs. Estos modelos analizan continuamente el tráfico de red y las transacciones, detectando anomalías que podrían señalar intentos de intrusión o actividades fraudulentas. (Hidalgo, 2020).

La incorporación de modelos predictivos de IA en las estrategias de ciberseguridad ha mejorado significativamente la capacidad de las instituciones bancarias para detectar y responder a APTs. Sin embargo, es esencial mantener una actualización constante de estos modelos y fomentar la colaboración entre entidades para compartir información sobre amenazas emergentes y mejores prácticas (Lucena, 2024).

El 40.5% de los participantes indica que los modelos de seguridad se actualizan "casi siempre" y el 39.2%, "siempre", mientras que un 20.3% señala que ocurre "a veces", evidenciando la necesidad de mayor consistencia en estos procesos, lo que evidencia la necesidad de establecer procesos más consistentes y dinámicos para enfrentar eficazmente las amenazas emergentes.

2.1.3.2. Identificación de los principales desafíos en la implementación de modelos predictivos de IA en la protección contra APTs en el sistema bancario de Lima, 2024. El uso de modelos predictivos de inteligencia artificial (IA) en la ciberseguridad bancaria ha permitido detectar y mitigar amenazas persistentes avanzadas (APTs) de manera más efectiva. Sin embargo, su implementación presenta diversos desafíos técnicos, operativos y regulatorios, especialmente en el contexto del sistema bancario de Lima en 2024.

- Disponibilidad y calidad de datos: La eficacia de los modelos de IA depende de la cantidad y calidad de los datos utilizados para su entrenamiento. En el sector bancario, garantizar bases de datos limpias, actualizadas y representativas es un reto,

especialmente cuando existen restricciones legales sobre la recopilación y el almacenamiento de información sensible.

- Evolución constante de las APTs: Los atacantes utilizan técnicas cada vez más sofisticadas para evadir los sistemas de detección basados en IA. La implementación de modelos predictivos requiere una actualización constante para adaptarse a las nuevas tácticas de ataque, lo que implica altos costos en investigación y desarrollo.
- Regulaciones y cumplimiento normativo: La adopción de IA en ciberseguridad debe alinearse con normativas nacionales e internacionales, como la Ley de Protección de Datos Personales en Perú y estándares como ISO/IEC 27001. Cumplir con estas regulaciones puede ralentizar la implementación de nuevas tecnologías en los bancos.
- Falta de personal capacitado: La implementación de IA en ciberseguridad requiere especialistas en aprendizaje automático, análisis de datos y ciberseguridad. En Lima, la escasez de profesionales capacitados en estas áreas ha dificultado la integración eficiente de estas tecnologías en el sector bancario.
- Costo de implementación y mantenimiento: Si bien los modelos predictivos pueden optimizar la ciberseguridad, su desarrollo e implementación requieren una inversión significativa en infraestructura tecnológica, software y capacitación del personal, lo que representa un desafío para muchas instituciones financieras. (Lucena, 2024).

La implementación de modelos predictivos de IA en la ciberseguridad bancaria de Lima enfrenta desafíos en términos de calidad de datos, evolución de amenazas, regulación, capacitación y costos. Para superar estos obstáculos, es fundamental que las instituciones adopten estrategias de actualización continua, inversión en talento especializado y cumplimiento normativo riguroso.

2.1.4. Algoritmos

Un algoritmo es un conjunto de reglas definidas que permiten resolver un problema mediante operaciones sistemáticas y finitas. Se compone de tres partes fundamentales: la entrada o input, que son los datos iniciales; el procesamiento, donde se aplican las instrucciones para realizar cálculos lógicos; y la salida, que es el resultado obtenido. (UNIR FP, 2022).

2.1.5. Inteligencia artificial

La inteligencia artificial (IA) se define como la facultad de los sistemas computacionales para emular habilidades humanas complejas, tales como el razonamiento lógico, el aprendizaje adaptativo y la capacidad creativa. Esta tecnología permite que los sistemas interactúen de manera autónoma con su entorno, analizando y procesando datos mediante sensores para tomar decisiones informadas y resolver problemas sin intervención humana directa. Aunque la IA no es un concepto nuevo —sus fundamentos datan de hace más de cinco décadas—, los avances recientes en poder computacional y la creciente disponibilidad de grandes volúmenes de datos han catalizado su evolución, permitiendo su integración en múltiples ámbitos de la vida cotidiana y profesional. (European Parliament, 2020).

2.1.6. Datos

Los datos son representaciones de variables cualitativas o cuantitativas, que se asignan mediante números, letras o símbolos. Tienen una base empírica, proviniendo de la realidad, y deben organizarse y contextualizarse para permitir análisis significativos. Por sí solos, los datos no comunican información, sino que requieren de un marco teórico y la relación con otras variables para adquirir relevancia. (Westreicher, 2020).

2.1.7. Sistema

Un sistema es un conjunto organizado de elementos interrelacionados que desempeñan la misma función o contribuyen a un fin común. Puede referirse a reglas, teorías, mecanismos o métodos estructurados de manera coherente. Además, en biología, un sistema puede designar un conjunto de órganos que cumplen funciones específicas dentro de un organismo (Real Academia Española [RAE], 2024).

2.1.8. Métricas

La métrica se refiere al arte que estudia la medida y estructura de los versos en la poesía, así como sus diversas clases y combinaciones. También abarca aspectos relacionados con el metro como unidad de longitud y su aplicación en la asignación del acento en diferentes lenguas (RAE, s.f.).

2.1.9. Patrones

Los patrones en la arquitectura de datos son modelos estructurales de alto nivel que representan cómo se organizan y agrupan los componentes para resolver problemas recurrentes en diferentes contextos. Estos patrones permiten aplicar principios de diseño generales y garantizan características como la integridad de datos, la disponibilidad del servicio y la escalabilidad, facilitando así el soporte de los requerimientos del negocio (Medium, 2024).

2.1.10. Métodos

Un método es una forma organizada de alcanzar un objetivo, que implica una secuencia estructurada de pasos aplicables en diversas áreas, como ciencias naturales, sociales o matemáticas. Se clasifica en diferentes tipos, como analítico, inductivo, deductivo y sintético, y es fundamental en el desarrollo de procedimientos sistemáticos para la investigación y resolución de problemas (Westreicher, 2024).

2.1.11. Protección

La protección de datos se refiere a los procesos, herramientas y estrategias diseñadas para salvaguardar información sensible de personas y organizaciones contra el acceso no autorizado y amenazas externas. Esta metodología busca prevenir la pérdida o corrupción de datos a través de copias de seguridad y protocolos de restauración, garantizando así la integridad de la información (SYDLE, 2024).

2.1.12. Capacidad

Las capacidades representan recursos fundamentales que facultan una actuación competente, al combinar de manera integral conocimientos, destrezas y actitudes que los estudiantes utilizan para abordar diferentes situaciones. Estas capacidades operan como componentes específicos y concretos dentro del contexto más amplio y holístico de las competencias, proporcionando las bases necesarias para el desarrollo de habilidades complejas y permitiendo la aplicación efectiva del aprendizaje en escenarios variados y desafiantes, que implican un conjunto más complejo de acciones y decisiones en contextos específicos (Ministerio de Educación [MINEDU], 2020).

2.1.13. Modelos Predictivos de IA

Los modelos predictivos de inteligencia artificial (IA) son herramientas que utilizan algoritmos para analizar datos históricos y realizar predicciones sobre eventos futuros o comportamientos que se aplican en diversas áreas como la salud, las finanzas, el marketing y la ingeniería. Se emplean frecuentemente técnicas de aprendizaje automático, como la regresión, los árboles de decisión y las redes neuronales, para analizar datos, identificar patrones y optimizar procesos en diversas aplicaciones (Sandua, 2024).

En marketing, los modelos predictivos de IA son algoritmos que analizan datos para identificar patrones en el comportamiento del cliente, permitiendo personalizar ofertas, mejorar experiencias y optimizar ventas mediante predicciones precisas. (Liberos et al., 2024).

Rincón y Vila (2021) destacan que el uso de modelos algorítmicos (matemáticos y datos comunes) de plataformas digitales para hacer predicciones pueden ser útiles en diferentes áreas. Liberos et al. (2024) mencionan que los beneficios del uso de la IA para predecir el comportamiento del cliente incluyen la mejora de la experiencia y retención del cliente mediante la personalización de ofertas, la optimización de la eficiencia en la planificación y distribución, la reducción de costos al evitar el exceso o falta de stock, y el aumento de las ventas adaptando los productos a las necesidades específicas del cliente.

Los modelos predictivos de IA tienen el potencial de transformar la previsión económica, permitiendo a economistas y políticos acceder a modelos más precisos. Utilizando técnicas de aprendizaje automático, la IA analiza datos económicos y factores externos para proporcionar predicciones más exactas sobre el crecimiento económico, inflación y tasas de interés. Estos pronósticos mejorados ayudan a gobiernos y empresas a tomar decisiones más informadas y a implementar estrategias más efectivas (Costa, 2023).

Liberos et al. (2024) indican que los modelos predictivos para predecir el comportamiento de los clientes con IA incluyen Árboles de decisión, que clasifican datos históricos para predecir resultados; Redes neuronales, que imitan la estructura cerebral para analizar grandes volúmenes de datos y ajustar predicciones; Regresión logística, utilizada para predecir probabilidades de acciones específicas de los clientes; y Máquinas de vectores de soporte (SVM), que identifican patrones en los datos basados en características específicas como edad e intereses. La elección del modelo depende de los datos disponibles y del objetivo de la predicción.

Calvillo et al. (2022) mencionan que existen diferentes métodos en los modelos predictivos que se utilizan en el campo médico, estos métodos se dividen en tres para predecir resultados y realizar diagnósticos. Esos métodos son: Gradient Boosting, Random Forest, y Support Vector Machine.

2.1.13.1. Algoritmos de Aprendizaje Automático. Soria et al. (2023) indican que los algoritmos son pasos matemáticos que guían a una máquina para aprender y alcanzar un objetivo. Un modelo de aprendizaje automático combina una estructura matemática y un algoritmo de aprendizaje, que se entrena con datos para ajustar parámetros y generar respuestas precisas. Estos modelos pueden utilizarse para predecir resultados futuros basados en datos previamente analizados, adaptándose a nuevos datos mediante la corrección de sus parámetros. Este estudio considera esta variable a la luz del autor Ropa-Carrión y Alama-Flores (2022), ya que administración y gestión suelen considerarse conceptos distintos. La gestión de los recursos y la organización interna son los objetivos principales de la administración; en cambio, la gestión tiene en cuenta no sólo estos factores, sino también los factores ambientales, el comportamiento organizativo y la formalización de los procedimientos para alcanzar los objetivos.

Los enfoques de aprendizaje automático (AM) resultan beneficiosos cuando se desea identificar automáticamente características relevantes entre numerosos candidatos (Tanaka et al., 2024). El aprendizaje automático (ML) es una rama de la ciencia de datos que capacita a las máquinas para aprender de forma autónoma, aplicando métodos estadísticos y matemáticos para identificar patrones y optimizar decisiones. Algoritmos como los utilizados por Google adaptan resultados basados en el comportamiento de los usuarios. Con avances tecnológicos, los algoritmos han mejorado en tareas como la comprensión de texto, reconocimiento de voz, identificación de imágenes y análisis de grandes volúmenes de datos. Estos algoritmos superan al humano en tareas repetitivas y optimización de decisiones, aunque todavía dependen de

parámetros definidos para alcanzar decisiones óptimas en situaciones complejas, según Soria et al. (2023).

- A. ***Tipos de algoritmos.*** Existen varios tipos de algoritmos de aprendizaje, como supervisado, no supervisado y reforzado, cuya elección depende del problema y los datos disponibles (Soria et al., 2023).
- B. ***Precisión de los modelos.*** Se refiere a la capacidad de un modelo para hacer predicciones correctas sobre nuevos datos. Se mide comúnmente como la proporción de predicciones correctas respecto al total de predicciones realizadas. La precisión puede variar dependiendo del tipo de modelo utilizado, los datos de entrenamiento, la calidad de las características seleccionadas y cómo se evalúa el modelo (por ejemplo, usando métodos como la validación cruzada) (Mirjalili y Raschka, 2020).
- C. ***Robustez frente a variaciones en los datos.*** Se refiere al destacar la capacidad de los algoritmos SVM y RF y evitar el sobreajuste. Esta robustez se refiere a la habilidad de estos algoritmos para mantener un rendimiento efectivo incluso cuando las condiciones del conjunto de entrenamiento varían, como en escenarios con datos equilibrados o desequilibrados (Garzón et al., 2023).
- D. ***Capacidad de generalización.*** Se refiere a la habilidad de los modelos de aprendizaje automático para aplicar lo aprendido a nuevos datos no vistos durante el entrenamiento, manteniendo su precisión en la predicción. Se destaca que los modelos de ensamblaje y redes neuronales no solo fueron precisos, sino que también tuvieron una alta capacidad de generalización, lo que significa que pueden hacer predicciones fiables en diferentes conjuntos de datos, más allá de los utilizados para entrenarlos (Párraga, 2024).

2.1.13.2. Entrenamiento y Validación. Se abordan como etapas del proceso de modelado en la evaluación de solicitudes. El conjunto de entrenamiento se utiliza para desarrollar los modelos, mientras que el conjunto de validación se emplea para comparar y

seleccionar el modelo más adecuado. Esta metodología permite garantizar que el modelo elegido tenga un buen desempeño en datos no vistos, mejorando su capacidad de generalización (Espinoza, 2020).

Generalmente, el entrenamiento se aborda mediante el uso de transferencia de aprendizaje, lo que permite utilizar redes neuronales preentrenadas con una base de datos más pequeña. Se realiza una comparativa entre distintas arquitecturas de redes neuronales, como Xception, ResNet y MobileNet, para seleccionar la más efectiva. La validación se lleva a cabo mediante pruebas experimentales en un prototipo con banda transportadora (Gante et al., 2022).

A. Métodos de entrenamiento. Los métodos de entrenamiento en algoritmos de aprendizaje automático incluyen aprendizaje supervisado, donde se entrena el modelo con datos etiquetados; aprendizaje no supervisado, que utiliza datos sin etiquetar para encontrar patrones; y aprendizaje por refuerzo, donde un agente toma decisiones en un entorno para maximizar recompensas (Soria et al., 2023).

B. Conjunto de datos utilizados para entrenamiento y prueba. Los conjuntos de datos utilizados para entrenamiento y prueba son definidos como esenciales para entrenar y evaluar los algoritmos de clasificación. El conjunto de entrenamiento permite al modelo "aprender" de las características de los documentos, mientras que el conjunto de prueba se usa para evaluar el rendimiento del modelo. Estos conjuntos permiten al algoritmo hacer predicciones precisas al clasificar nuevos documentos (Cedeño y Vargas, 2020).

C. Métricas de rendimiento. Las métricas de rendimiento se abordan como herramientas esenciales para evaluar la precisión y exactitud de las redes neuronales convolucionales (CNN). Estas métricas permiten detectar y corregir errores, falsas predicciones, subajustes y sobreajustes en el proceso de aprendizaje automático, contribuyendo a mejorar el rendimiento del software inteligente desarrollado (Zapeta et al., 2022).

D. Tiempo de entrenamiento. El tiempo de entrenamiento en algoritmos de aprendizaje automático puede variar significativamente según varios factores como la complejidad del modelo, la cantidad de datos, la potencia del hardware y la eficiencia del algoritmo utilizado. En general, modelos más complejos o grandes conjuntos de datos requieren más tiempo de entrenamiento (Soria et al., 2023).

2.1.13.3. Implementación y Despliegue. La implementación y despliegue se abordan como el proceso de desarrollar y poner en funcionamiento un modelo de predicción utilizando inteligencia artificial. La implementación incluye la creación del modelo mediante; por ejemplo, Python, mientras que el despliegue implica la utilización de herramientas como Flask para la API y Streamlit y Heroku para su distribución en la nube, asegurando que la solución sea accesible y funcional en un entorno real (Montesdeoca, 2024).

Una vez que el modelo ha sido entrenado y evaluado de manera satisfactoria, se puede implementar en producción para realizar predicciones o tomar decisiones en tiempo real sobre nuevos datos (Serra, 2023).

A. Tiempo de inferencia. El tiempo de inferencia en un entorno de producción para algoritmos de aprendizaje automático varía dependiendo del modelo, la complejidad de los datos y la capacidad del hardware utilizado. Generalmente, se busca que este tiempo sea lo más bajo posible para permitir decisiones rápidas en tiempo real, y puede oscilar desde milisegundos hasta varios segundos, dependiendo de los factores mencionados (Huyen, 2023).

B. Escalabilidad del modelo. La capacidad de poder gestionar un gran volumen de datos y que éstos puedan ser actualizados fácilmente (Miravé, 2023).

C. Facilidad de integración con sistemas existentes. Define esta facilidad como la capacidad de incorporar modelos de manera eficiente y sin necesidad de grandes modificaciones en el sistema, lo que optimiza la implementación en diversos contextos (Pérez, 2024).

D. Adaptabilidad a diferentes escenarios de ataque. La adaptabilidad de algoritmos de aprendizaje automático a diferentes escenarios de ataque se refiere a la capacidad de estos modelos para manejar y responder a distintas estrategias de manipulación o adversariales. Esto incluye la detección y mitigación de ataques como el envenenamiento de datos, ataques de adversarios y filtrado de características (Fundación Telefónica, 2021).

2.1.14. Protección contra amenazas persistentes avanzadas (APT)

La defensa frente a las amenazas persistentes avanzadas (APT) requiere un enfoque múltiple de la ciberseguridad para identificar, frustrar y contrarrestar este tipo de ataques. Las medidas de ciberseguridad incluyen la vigilancia continua de la red, el análisis del comportamiento de los usuarios, la división de los sistemas y la educación y formación del personal. Las organizaciones organizadas suelen lanzar APT, que son ataques dirigidos, con el espionaje o el sabotaje como posibles motivaciones (Mata, 2022).

Se trata de un conjunto de amenazas muy peligrosas que a menudo cuentan con apoyo financiero y son llevadas a cabo por organizaciones bien organizadas con sofisticadas capacidades de ataque. Los objetivos de estas amenazas suelen ser bastante específicos y tienen vínculos con el ciberespionaje de alto nivel (Moreno, 2022).

Las amenazas persistentes avanzadas (APT) pasan por las siguientes etapas en su ciclo de vida: El primer paso consiste en identificar el objetivo para poder adaptar el ataque; el segundo, en crear archivos maliciosos; el tercero, en hacer llegar el malware a la organización; el cuarto, en explotar las vulnerabilidades del malware; el quinto, en instalar el malware en los dispositivos; y el sexto, en llevar a cabo acciones sobre el objetivo, como robar información o utilizar la organización para ataques de mayor envergadura (Sierra, 2020).

Entre los objetivos de la APT figuran la protección de los modelos contra los ataques del aprendizaje automático adversario, la mejora de la interpretabilidad para aumentar la

fiabilidad de las predicciones y la elección de modelos eficaces que puedan funcionar en situaciones complejas y dinámicas (Hernández y Sánchez, 2023).

Integrar un sistema alineado con la tecnología blockchain y la inteligencia artificial en los dispositivos de borde es un enfoque común para protegerse contra las amenazas persistentes avanzadas (APT). Las amenazas persistentes avanzadas (APT) se describen como asaltos que socavan la seguridad del transporte de datos, comprometen los dispositivos de borde y hacen que las partes interesadas confíen menos entre sí de lo que lo harían en otras circunstancias (Rahman et al., 2023).

2.1.14.1. Detección y Prevención. Los sistemas de detección de intrusos (IDS) y los sistemas de prevención de intrusos (IPS) se centran en prevenir y detectar intrusiones. La función principal de un sistema de detección de intrusos (IDS) es la detección y notificación de amenazas, mientras que un sistema de prevención de intrusos (IPS) realiza ambas funciones automáticamente. Por ello, los IPS se denominan sistemas de detección y prevención de intrusiones (IDPS), ya que pueden tanto detectar como responder a las intrusiones (Ramírez, 2024).

Los sistemas de detección y prevención de intrusiones controlan el tráfico de la red. El propósito de las herramientas de análisis de registros es buscar en los registros del sistema cualquier comportamiento inusual. Las herramientas de gestión de la configuración se encargan de garantizar el cumplimiento de las normas de seguridad. Los objetivos y la arquitectura de la organización dictan la elección de las herramientas. Para salvaguardar los sistemas y las redes, es esencial integrarlas con reglas de seguridad sólidas (Arango, 2023).

A. Tasa de detección de ataques. Con el fin de comparar diversos enfoques de defensa, se aborda la tasa de detección de ataques. Esta tasa es una métrica fundamental para evaluar y comparar la eficacia de diversas medidas de seguridad, ya que indica lo bien que un mecanismo detecta las amenazas (Gómez, 2022).

- B. Falsos positivos y falsos negativos.** La precisión de un sistema de detección de intrusiones (IDS) a la hora de detectar intrusiones es una variable clave que puede afectar a la frecuencia y gravedad de los falsos positivos y negativos. Se cree que los sucesos que el sistema de detección de intrusiones (IDS) etiqueta erróneamente como amenazas se conocen como falsos positivos, mientras que las amenazas reales que el IDS no logra identificar se conocen como falsos negativos. La complejidad de las agresiones y la configuración de los criterios de identificación pueden afectar a la precisión del IDS y, por extensión, a las tasas de falsos positivos y negativos (Quiroz et al., 2020).
- C. Tiempo de detección.** El tiempo de detección de anomalías en la red es como un factor no estimado ni controlado debido a la falta de un sistema interno de monitoreo, indicando que la única opción para identificar anomalías es a través del proveedor de servicios de internet, lo cual puede demorar desde horas hasta semanas (Dalemborg, 2024).
- D. Tiempo de respuesta.** Limitar el alcance de un ciberataque depende de la rapidez con la que se pueda responder. Ejecutar esfuerzos sincronizados para detener el asalto, como cortar los contactos con el atacante, aislar los sistemas infectados y erradicar el malware, lo más rápidamente posible, es una definición implícita. Esto demuestra la importancia de que los equipos de respuesta a incidentes estén bien organizados y preparados para reducir los tiempos de respuesta (Arévalo, 2023).

2.1.14.2. Resiliencia del Sistema. Una de las estrategias más importantes en la ciberdefensa contra la denegación de servicio es la resistencia del sistema. La capacidad de un sistema para resistir y recuperarse de ciberataques, haciendo improbable el éxito de un adversario, es una definición implícita de resiliencia. A la luz del fracaso de la disuasión de represalias para lograr resultados significativos en esta área, es imperativo que los sistemas estén mejor equipados para la protección, defensa y recuperación (Cubeiro, 2021).

Según Camuñas (2024), la ciberseguridad es un ámbito en el que se debate la resiliencia de los sistemas; un método que puede utilizarse para mejorar la resiliencia son los honeypots. Una definición de resiliencia es la capacidad de un sistema para resistir, modificar y recuperarse de ciberataques y amenazas. Al obtener una visión más profunda de estos peligros, los honeypots fortalecen las defensas del sistema contra posibles debilidades.

A. Capacidad de recuperación post-ataque. Según Gómez (2022), la capacidad de recuperación post-ataque se enfoca en la prevención y mitigación del daño inicial causado por el ransomware Reventon. Al destacar cómo el ransomware infecta masivamente y utiliza tácticas de miedo para extorsionar a las víctimas, se sugiere que la capacidad de recuperación post-ataque dependería de la preparación previa y las respuestas rápidas y efectivas para evitar el pago del rescate y la propagación del ataque.

B. Capacidad de recuperación post-ataque. Según Gómez (2022), la capacidad de recuperación post-ataque se enfoca en la prevención y mitigación del daño inicial causado por el ransomware Reventon. Al destacar cómo el ransomware infecta masivamente y utiliza tácticas de miedo para extorsionar a las víctimas, se sugiere que la capacidad de recuperación post-ataque dependería de la preparación previa y las respuestas rápidas y efectivas para evitar el pago del rescate y la propagación del ataque.

C. Resistencia a técnicas de evasión empleadas por APT. La resistencia a las técnicas de evasión empleadas por APT se centra en la implementación de medidas de seguridad robustas, como la detección continua de amenazas, el uso de firewalls de próxima generación, sistemas de prevención de intrusiones y análisis de comportamiento (Jaén, 2021).

D. Efectividad de los parches de seguridad. Se infiere que la efectividad de los parches de seguridad se refiere a su capacidad para corregir debilidades y mejorar la seguridad del sistema frente a posibles ataques. Las pruebas de penetración ayudan a identificar

si los parches implementados son efectivos en la protección contra vulnerabilidades (Arango, 2023).

2.1.14.3. Actualización y Evolución. Las amenazas persistentes avanzadas (APT) han evolucionado significativamente en los últimos años, adaptándose a nuevas tecnologías y tácticas de defensa. Las actualizaciones en protección contra APT incluyen el uso de inteligencia artificial y machine learning para detectar patrones sospechosos, así como la implementación de soluciones de respuesta a incidentes para mitigar ataques en tiempo real (Mata, 2022).

Ramírez (2024) indica que la constante evolución de las tecnologías y el crecimiento exponencial de las amenazas cibernéticas han dado lugar a la necesidad de un enfoque proactivo y estratégico para defenderse contra los ataques maliciosos.

A. Frecuencia de actualizaciones de modelos. La frecuencia de actualizaciones de los modelos de Protección contra amenazas persistentes avanzadas (APT) varía según el proveedor y el tipo de software. Generalmente, las actualizaciones pueden ser semanales o mensuales, pero en respuesta a nuevas amenazas, algunas organizaciones optan por hacer actualizaciones en tiempo real o más frecuentes. Es crucial que las empresas monitoreen continuamente sus sistemas y se mantengan al día con las últimas actualizaciones de seguridad (Fundación Telefónica, 2021).

B. Eficiencia de las contramedidas. La eficacia de las contramedidas se destaca al señalar que la complejidad creciente de los ataques cibernéticos, incluidas las Amenazas Persistentes Avanzadas (APTs), requiere implementar estrategias de seguridad igualmente avanzadas. Se concluye que la eficiencia de estas contramedidas depende de su habilidad para anticipar y mitigar ataques sofisticados mediante la aplicación de técnicas de Inteligencia Artificial. (Larriva et al., 2023).

- C. *Adaptación a nuevas amenazas.*** La adaptación a nuevas amenazas de Protección contra Amenazas Persistentes Avanzadas (APT) implica una evolución constante de las estrategias de ciberseguridad. Esto incluye la implementación de tecnologías avanzadas de detección y respuesta, formación continua del personal, y el refuerzo de la protección de datos y sistemas críticos. La colaboración entre sectores y el intercambio de información sobre amenazas también son esenciales para fortalecer la defensa contra APTs. (Baqués, 2021).
- D. *Evolución de patrones de ataque identificados.*** Baqués (2021) indica que la evolución de patrones de ataque relacionados con la Protección contra Amenazas Persistentes Avanzadas (APT) incluye el uso de técnicas sofisticadas y enfoques multidimensionales que pueden abarcar desde phishing dirigido hasta la explotación de vulnerabilidades de software.

III. MÉTODO

3.1. Tipo de investigación

El estudio se enmarca en una investigación de tipo básico. La investigación básica, a diferencia de la aplicada, persigue ampliar el conocimiento científico sin un fin práctico inmediato (Carrasco, 2019). En este contexto, el estudio se centró en profundizar en la relación entre la gestión administrativa y la calidad de la atención, con el propósito de aportar una comprensión más sólida y fundamentada en este ámbito.

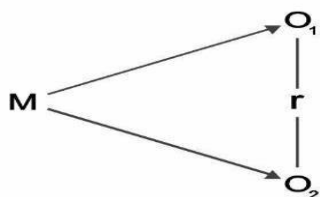
La investigación adoptó un enfoque cuantitativo, cuyo propósito fundamental fue simplificar el proceso de análisis a través de la recolección de datos numéricos. Para ello, se utilizaron instrumentos específicos que permitieron obtener información cuantificable, la cual fue analizada para responder de manera objetiva y precisa a las preguntas de investigación formuladas desde el inicio del estudio (Cabezas et al., 2018). Por esta razón, la presente indagación analizó la relación de las variables de manera cuantitativa.

El nivel o alcance fue correlacional. Para comprender el comportamiento de una variable observando sus relaciones con otras variables, los estudios cuantitativos correlacionales determinaron, en primer lugar, la fuerza de la relación entre las variables investigadas y, posteriormente, cuantificaron las correlaciones entre ellas (Cabezas et al., 2018). Por ello, esta indagación asoció las variables.

El diseño fue no experimental. La investigación, al no implicar la modificación de variables, permitió observar y analizar los acontecimientos en su entorno natural. Dado que el investigador no ejerció influencia sobre las variables independientes, no se crearon esas circunstancias, ya que estas existían previamente (Vera et al., 2018). Por lo tanto, no hubo manipulación o cambio en las variables de esta indagación.

El corte fue transversal. Este estudio tomó únicamente una instantánea en el tiempo para recopilar sus datos, proporcionando una descripción de las variables y examinando su aparición y conexión en un momento determinado (Cabezas et al., 2018). Por esta razón, esta investigación se llevó a cabo en un periodo específico.

Respecto al diseño asociativo, se consideró el análisis de Ñaupas et al. (2018):



Donde:

M: Muestra de la población.

O1: Observación de la variable 1

O2: Observación de la variable 2

r: Es coeficiente de correlación.

3.2. Población y muestra

3.1.1. Población

La población se definió como un conjunto de individuos agrupados con el propósito de realizar una investigación en profundidad. Este conjunto de datos estuvo estrechamente relacionado con las variables explorables del estudio (Cabezas et al., 2018). Para este estudio, se consideró como población a 100 empleados de instituciones bancarias ubicadas en Lima.

3.1.2. Muestra

Se entendió que los científicos realizaron generalizaciones sobre una población tras operar con una muestra, considerándola como un subconjunto representativo de la población total (Cabezas et al., 2018). En esta investigación, se empleó una muestra probabilística.

La muestra se seleccionó de forma aleatoria, lo que aseguró que todos los elementos de la población tuvieran la misma probabilidad de ser elegidos. Este tipo de muestreo permitió además calcular el error o grado de incertidumbre asociado a la muestra, para lo cual se aplicó la siguiente fórmula:

$$n = \frac{(Z)^2 (PQN)}{(e)^2 (N-1) + (Z)^2 PQ}$$

Donde:

Aplicando la fórmula correspondiente, se determinó como muestra de estudio la siguiente:

$$n = \frac{(1.96)^2 * 0.50 * 0.50 * 100}{(0.05)^2 (100-1) + (1.96)^2 * 0.50 * 0.50}$$

$$n = 79$$

Para este estudio se estableció una muestra de 79 empleados de instituciones bancarias en Lima.

3.3.Operacionalización de variables

Tabla 2

Operacionalización de las variables

VARIABLE	DIMENSIONES	INDICADORES
Modelos Predictivos de IA	Algoritmos de Aprendizaje Automático	Tipos de algoritmos Precisión de los modelos Robustez frente a variaciones en los datos Capacidad de generalización
	Entrenamiento y Validación	Métodos de entrenamiento Conjunto de datos utilizados para entrenamiento y prueba Métricas de rendimiento Tiempo de entrenamiento
	Implementación y Despliegue	Tiempo de inferencia en entorno de producción Escalabilidad del modelo Facilidad de integración con sistemas existentes Adaptabilidad a diferentes escenarios de ataque
Protección contra APT	Detección y Prevención	Tasa de detección de ataques Falsos positivos y falsos negativos Tiempo de detección Tiempo de respuesta
	Resiliencia del Sistema	Capacidad de recuperación post-ataque Resistencia a técnicas de evasión empleadas por APT Efectividad de los parches de seguridad
	Actualización y Evolución	Frecuencia de actualizaciones de modelos. Eficiencia de las contramedidas Adaptación a nuevas amenazas Evolución de patrones de ataque identificados

3.4. Instrumentos

La investigación se desarrolló utilizando la encuesta como técnica principal, una metodología comúnmente aplicada en estudios de campo. Este enfoque se caracterizó por el uso de preguntas estructuradas, formuladas a partir de un proceso sistemático como la operacionalización de variables, lo cual permitió recopilar opiniones de manera ordenada y objetiva. La construcción de interrogantes vinculadas a los fenómenos analizados posibilitó la obtención de respuestas concretas y pertinentes (Cabezas et al., 2018). Así, la encuesta se consolidó como la herramienta idónea para alcanzar los objetivos planteados en la investigación.

El instrumento utilizado fue el cuestionario, definido como una herramienta conceptual y material diseñada para recopilar información y datos relevantes para el investigador. Dependiendo de los métodos aplicados, los cuestionarios pueden adoptar diferentes formatos y niveles de complejidad (Ñaupas et al., 2018). En esta investigación, el cuestionario permitió recoger información estructurada directamente relacionada con las variables y dimensiones del estudio.

Para determinar su confiabilidad y validez se realizó un juicio de expertos que validó el instrumento empleado para medir la confiabilidad se realizó el alfa de Cronbach el cual tuvo un valor de 0.809

Tabla 3

Alfa de Cronbach

Alfa de Cronbach	N de elementos
,809	23

3.5.Procedimientos

En la primera etapa, se calcularon las frecuencias de los datos recopilados, lo que permitió obtener una visión general sobre la distribución y dispersión de las variables. Este análisis inicial ofreció una descripción detallada de los valores esperados en cada dimensión del estudio.

Posteriormente, se utilizó el software SPSS versión 25 para realizar el análisis de correlación de Spearman. Este procedimiento permitió examinar las relaciones entre las dimensiones y las variables principales, identificando la existencia y la magnitud de posibles vínculos significativos.

En una tercera etapa, se discutieron e interpretaron los datos obtenidos en función de los objetivos de la investigación, ofreciendo respuestas claras a las preguntas planteadas. Finalmente, se elaboraron recomendaciones pertinentes, basadas en los hallazgos, con el propósito de aportar soluciones prácticas y teóricas al problema abordado

3.6.Análisis de datos

Los datos recopilados mediante la encuesta se organizaron en Excel y luego se importaron a SPSS para su análisis. Se realizaron evaluaciones descriptivas para resumir la información y análisis inferenciales para identificar relaciones y generar proyecciones basadas en la muestra.

En el análisis descriptivo, se utilizaron gráficos y tablas de frecuencias para visualizar los datos y caracterizar las variables del estudio. El análisis inferencial aplicó el coeficiente Rho de Spearman para evaluar la relación entre las variables y verificar la hipótesis.

3.7.Consideraciones éticas

La investigación siguió el reglamento de la Universidad Nacional Federico Villarreal, asegurando el cumplimiento de principios éticos y un enfoque responsable en todas las etapas del proceso de recolección y análisis de datos.

Se respetaron los derechos de autenticidad mediante la adecuada citación de los autores referenciados en el trabajo, siguiendo las Normas APA, 7ª edición. Además, se aseguró la confidencialidad de los datos proporcionados por los participantes y se evitó cualquier forma de manipulación o distorsión de los resultados obtenidos. Este enfoque ético contribuyó a la elaboración de discusiones, conclusiones y recomendaciones fundamentadas y respetuosas del marco normativo y académico.

IV. RESULTADOS

4.1. Objetivo 1: Análisis de modelos predictivos de inteligencia artificial (IA)

Las Amenazas Persistentes Avanzadas (APTs) son ataques sofisticados y dirigidos que comprometen la seguridad de las entidades bancarias. Para mitigar estos riesgos, se emplean modelos predictivos de inteligencia artificial (IA) capaces de identificar anomalías y prevenir ataques antes de que se materialicen. Este análisis compara diferentes enfoques de IA en términos de efectividad, tiempo de respuesta y capacidad de adaptación a nuevas amenazas.

Tabla 4

Tabla comparativa de Modelos predictivos de inteligencia artificial

Modelo Predictivo	Principales Técnicas	Ventajas	Desventajas
Redes Neuronales Artificiales (ANNs)	Deep Learning, Redes Convolucionales (CNN), Redes Recurrentes (RNN)	Alta precisión en detección de anomalías complejas. Capacidad de autoaprendizaje.	Alto costo computacional. "Caja negra" difícil de interpretar.
Random Forest (Bosques Aleatorios)	Ensamble de árboles de decisión	Interpretabilidad, robustez ante datos desbalanceados.	Más lento en grandes volúmenes de datos. Puede sobreajustarse.
Support Vector Machines (SVM)	Máquinas de vectores de soporte con kernels	Precisión en clasificación binaria y detección de anomalías.	Ineficiente en grandes volúmenes de datos. Difícil ajuste de hiperparámetros.
Gradient Boosting (XGBoost, LightGBM, CatBoost)	Algoritmos de boosting para optimización	Alta precisión, rapidez en entrenamiento y predicción.	Puede ser difícil de ajustar en problemas muy complejos.
Análisis de Comportamiento del Usuario (UEBA)	Modelos de detección de anomalías con aprendizaje no supervisado	Detección en tiempo real de actividades inusuales.	Alta tasa de falsos positivos si no está bien calibrado.
Sistemas de Detección de Anomalías (ADS)	Modelos no supervisados como Autoencoders, Clustering (K-Means, DBSCAN)	Buen desempeño en detección de patrones desconocidos.	Puede requerir grandes volúmenes de datos históricos para ser preciso.
Redes Generativas Antagónicas (GANs)	Modelos generativos para detectar ataques adversariales	Capacidad de detectar ataques sofisticados generados por IA.	Difícil entrenamiento y alto costo computacional.

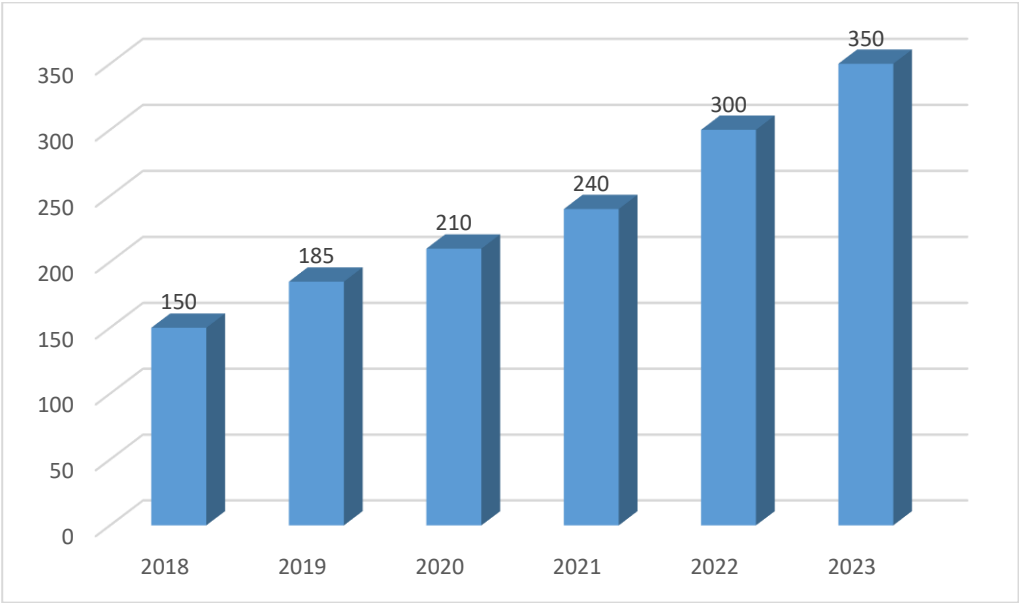
4.2. Objetivo 2: Análisis de amenazas persistentes avanzadas (ATPs) en el sistema bancario

La tabla muestra la tendencia creciente de incidentes de APTs en el sector bancario a lo largo de los últimos 6 años. Como se puede observar, el porcentaje de APTs sobre el total de ciberincidentes también ha ido aumentando, lo que indica que estas amenazas se están convirtiendo en un problema cada vez más significativo.

Tabla 5
Número de Incidentes de APTs Reportados en el Sector Bancario (2018-2023)

Año	Número de Incidentes Reportados	Porcentaje de APTs sobre el Total de Ciberincidentes (%)
2018	150	25%
2019	185	28%
2020	210	30%
2021	240	32%
2022	300	35%
2023	350	37%

Figura 1
Tendencia de Incidentes APTs en el Sector Bancario (2018-2023)

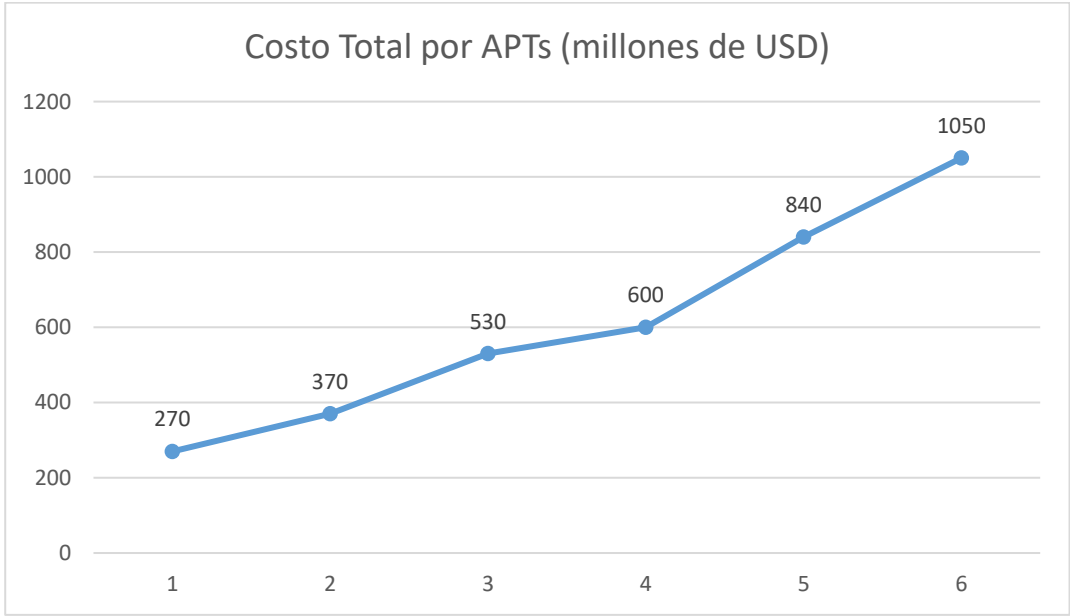


El gráfico ilustra un aumento progresivo en el impacto financiero de los ataques APTs en el sector bancario, con un notable salto en 2023 donde el costo total superó los mil millones de dólares.

Tabla 6
Impacto Financiero de APTs en el Sector Bancario por Año (2018-2023)

Año	Costo Total por APTs (millones de USD)
2018	270
2019	370
2020	530
2021	600
2022	840
2023	1050

Figura 2
Impacto Financiero de APTs en el Sector Bancario por Año (2018-2023)



El análisis estadístico de las Amenazas Persistentes Avanzadas (APT) detectadas en el sistema bancario peruano durante el año 2024 evidencia una creciente sofisticación de los ataques cibernéticos. Como se observa en los gráficos presentados, los tipos de APT más frecuentes fueron el Phishing Avanzado (APT) con 52 casos (representando el 22% del total), seguido de la Exfiltración de Datos Bancarios con 48 casos (20%) y el Uso de Malware Personalizado con 39 casos (17%).

Estas cifras demuestran que los atacantes cibernéticos priorizan técnicas orientadas al robo de información sensible y al acceso inicial a través de engaños personalizados, especialmente dirigidos a personal clave de las entidades financieras.

Un dato relevante es el incremento de ataques a través de Movimiento Lateral (35 casos) y Comandos y Control Remoto (C2) (32 casos), lo cual confirma que, una vez dentro de la infraestructura bancaria, los ciberatacantes buscan mantener un acceso prolongado, evadir detección y expandirse dentro de la red interna.

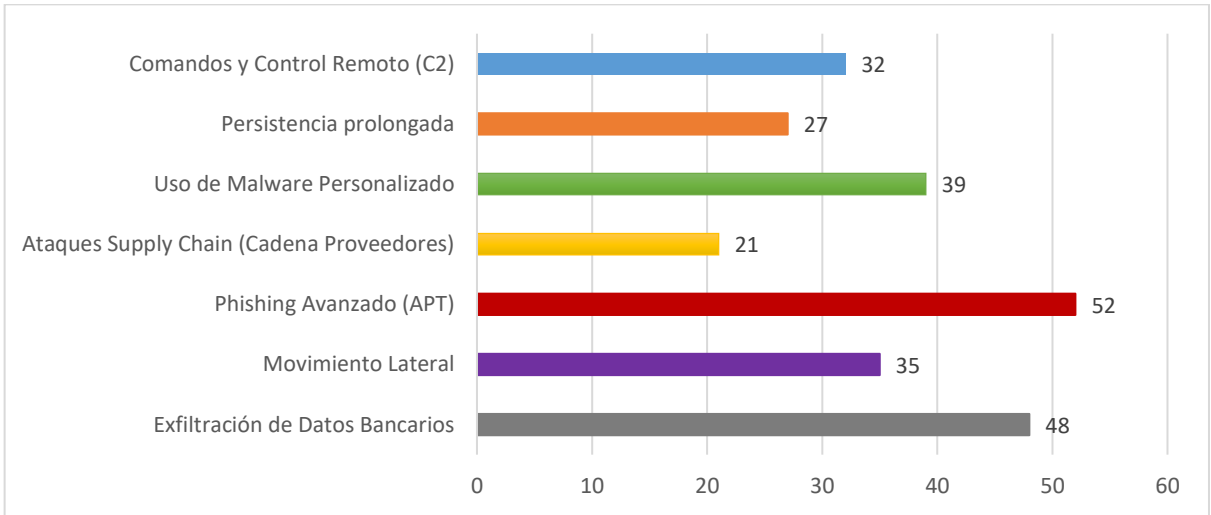
Por otro lado, los Ataques Supply Chain (21 casos) reflejan la creciente vulnerabilidad de los bancos frente a la dependencia de proveedores de software y servicios externos, lo cual se ha convertido en un vector importante de ingreso de amenazas APT.

La evolución y persistencia de este tipo de amenazas evidencia que las estrategias tradicionales de ciberseguridad ya no son suficientes, y que los modelos predictivos basados en Inteligencia Artificial (IA) se presentan como una herramienta clave para la detección y respuesta anticipada frente a estas amenazas.

Tabla 7
Tipos de Amenazas Persistentes Avanzadas (APT) en el Perú – Sector Bancario 2024

Tipo de APT Detectada	Frecuencia Estimada (casos)	Descripción
Exfiltración de Datos Bancarios	48	Robo sigiloso de información confidencial (clientes, transferencias, claves).
Movimiento Lateral	35	Expansión dentro de redes bancarias para capturar datos o infectar más sistemas.
Phishing Avanzado (APT)	52	Correos dirigidos a gerentes o área TI para acceso inicial.
Ataques Supply Chain (Cadena Proveedores)	21	Compromiso de proveedores de software bancario para infectar a bancos.
Uso de Malware Personalizado	39	Creación de malware específico para bancos peruanos.
Persistencia prolongada	27	APTs que se mantuvieron dentro de los sistemas por más de 3 meses sin ser detectados.
Comandos y Control Remoto (C2)	32	Uso de servidores externos para controlar dispositivos dentro del banco.

Figura 3
Frecuencia de tipos de Amenazas Persistentes Avanzadas (APT) en el Perú

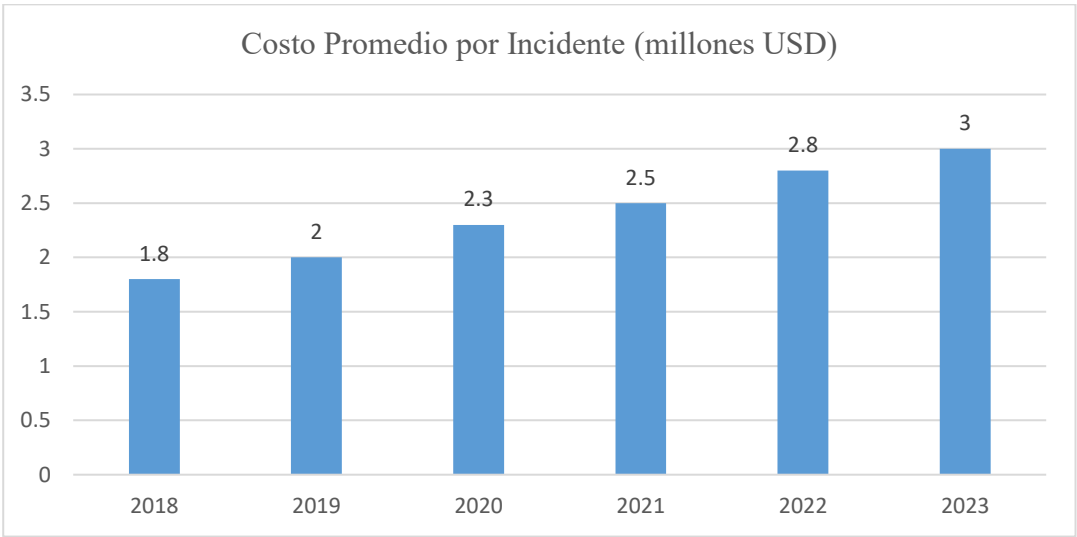


Esta tabla muestra cómo los costos promedio por incidente de APT en el sector bancario han aumentado a lo largo de los años. En 2023, las pérdidas estimadas alcanzaron más de mil millones de dólares, lo que subraya el impacto financiero grave de estos ataques.

Tabla 8
Costo Promedio por Incidente (millones USD)

Año	Costo Promedio por Incidente (millones USD)
2018	1,8
2019	2
2020	2,3
2021	2,5
2022	2,8
2023	3

Figura 4
Costo Promedio por Incidente (millones USD)



4.3. Objetivo 3: Resiliencia ante APTs por Componente

El gráfico de radar presentado ilustra el nivel de resiliencia del sistema frente a amenazas persistentes avanzadas (APTs) evaluando seis componentes clave de seguridad: detección, respuesta, recuperación, prevención, concientización del personal y gestión de vulnerabilidades. Los valores asignados a cada eje representan una escala hipotética del 0 al 10, donde 10 indica el máximo nivel de preparación o desempeño en ese aspecto.

Como se observa, los componentes de detección (8) y respuesta (7) obtienen las puntuaciones más altas, lo cual indica una capacidad sólida para identificar rápidamente intrusiones y actuar frente a ellas de manera eficaz. Esto sugiere la existencia de herramientas avanzadas de monitoreo, tales como sistemas SIEM o IDS, y protocolos bien establecidos para la gestión de incidentes. Por otro lado, la recuperación (6) y la prevención (7) muestran un nivel aceptable de resiliencia, evidenciando la capacidad del sistema para restablecer operaciones tras un ataque, así como esfuerzos preventivos como parches de seguridad y segmentación de red.

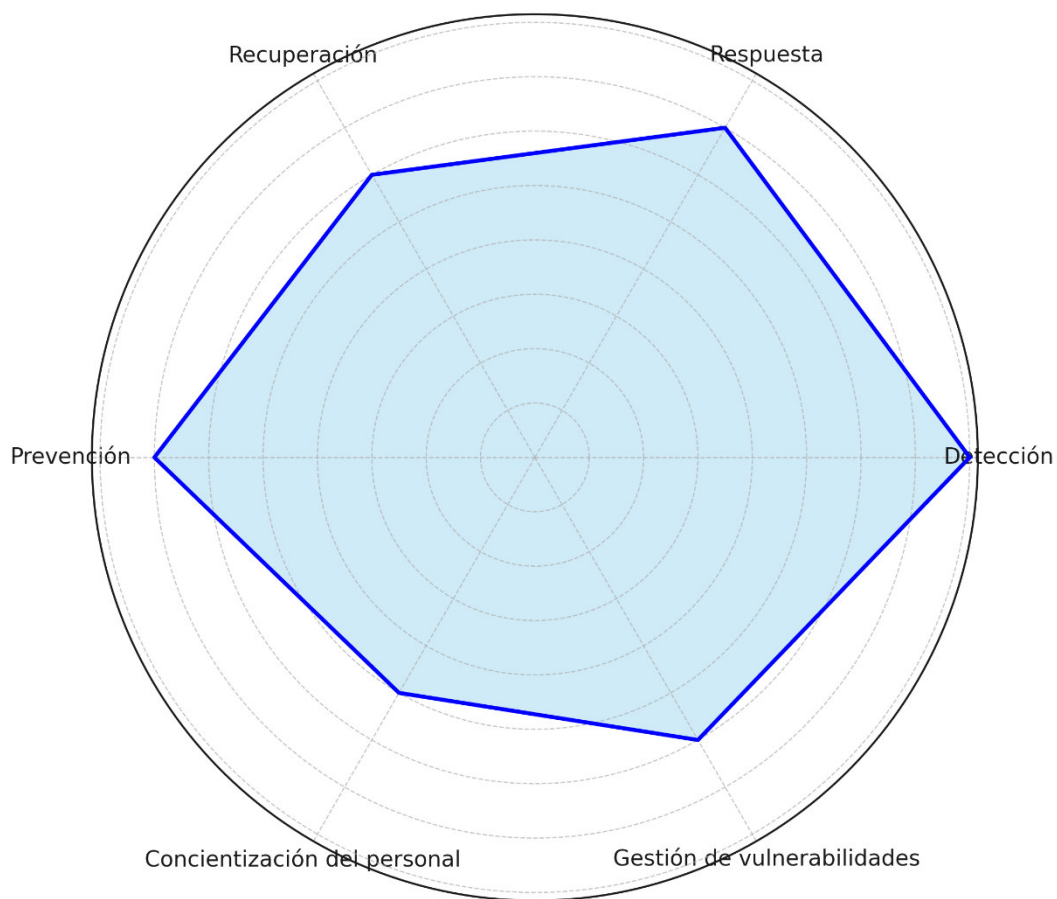
Sin embargo, se identifican oportunidades de mejora en los ejes de concientización del personal (5) y gestión de vulnerabilidades (6). Estas áreas representan un riesgo potencial, ya que los atacantes suelen explotar el error humano o las fallas no corregidas del sistema. Un fortalecimiento de la capacitación continua del personal, así como la implementación de programas de gestión de vulnerabilidades más robustos, permitirían incrementar la resiliencia general del sistema.

En resumen, el análisis evidencia una infraestructura con buena capacidad técnica, pero que requiere reforzamiento en la cultura organizacional de seguridad y en la gestión proactiva de amenazas. Esta evaluación visual es clave para tomar decisiones estratégicas y priorizar recursos en las áreas más vulnerables del sistema.

Figura 5

Gráfico de Resiliencia ante APTs por Componente

Nivel de Resiliencia ante APTs por Componente



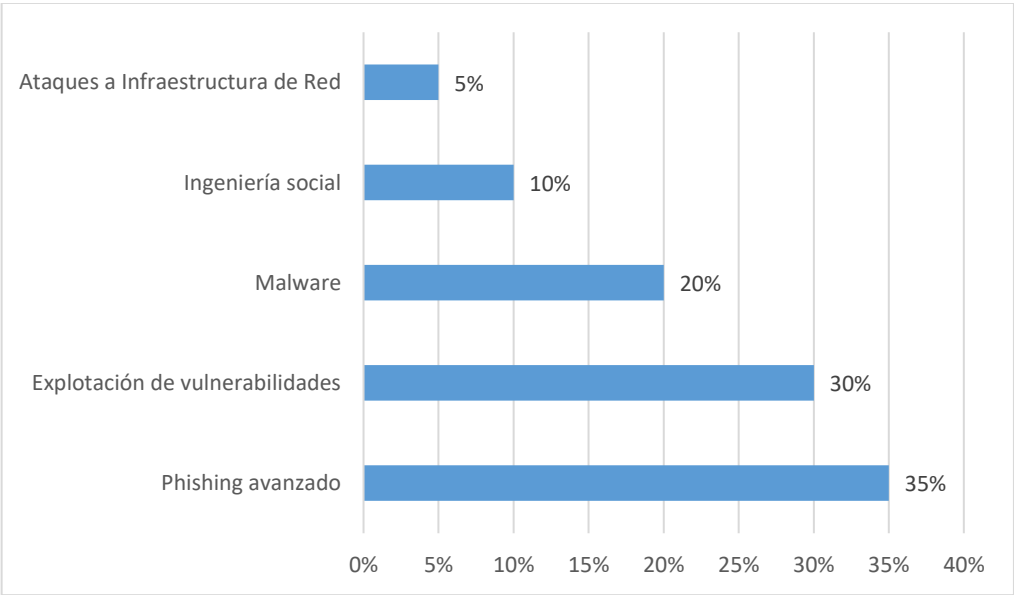
4.4. Objetivo 4: Actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs)

El gráfico de barras muestra que las técnicas de phishing (35%) y explotación de vulnerabilidades (30%) son las más utilizadas en los ataques APTs. Esto sugiere que los atacantes prefieren métodos de infiltración basados en ingeniería social y fallos en la seguridad de los sistemas.

Tabla 9
Distribución de Técnicas Utilizadas en APTs en el Sector Bancario (2023)

Técnica	Porcentaje de Uso (%)
Phishing avanzado	35%
Explotación de vulnerabilidades	30%
Malware	20%
Ingeniería social	10%
Ataques a Infraestructura de Red	5%

Figura 6
Distribución de Técnicas Utilizadas en APTs en el Sector Bancario (2023)



4.4.1. Principales enfoques utilizados para combatir APTs en el sector bancario

En el contexto actual de la ciberseguridad bancaria, los modelos predictivos de inteligencia artificial (IA) desempeñan un papel crucial en la detección, análisis y mitigación de las Amenazas Persistentes Avanzadas (APTs). Estas amenazas representan un desafío significativo para las instituciones financieras, debido a su capacidad para evadir mecanismos de defensa tradicionales y comprometer sistemas críticos.

El siguiente cuadro presenta una visión estructurada de los principales enfoques utilizados para combatir APTs en el sector bancario, categorizados en cinco áreas clave:

1. Técnicas de Modelos Predictivos: Algoritmos de IA utilizados para identificar patrones sospechosos y predecir ataques.
2. Ataques APTs en el Sector Bancario: Principales estrategias empleadas por ciberdelincuentes para comprometer infraestructuras financieras.
3. Estrategias de Protección: Tecnologías y metodologías implementadas para mitigar riesgos y fortalecer la seguridad.
4. Impacto en la Ciberseguridad: Beneficios obtenidos en términos de detección temprana, reducción del tiempo de respuesta y mejora en la protección.
5. Tendencias Futuras: Innovaciones emergentes que definirán el futuro de la ciberseguridad bancaria.

Este análisis proporciona una visión integral de cómo la IA está transformando la seguridad bancaria, permitiendo a las instituciones anticiparse a los ataques y fortalecer sus sistemas ante las crecientes amenazas digitales.

Tabla 10

Análisis comparativo entre las técnicas de IA y las estrategias de protección contra APTs en instituciones financieras.

Categoría	Técnicas de Modelos Predictivos	Ataques APTs en el Sector Bancario	Estrategias de Protección	Impacto en la Ciberseguridad	Tendencias Futuras
Algoritmos de IA	Redes neuronales, Random Forest, SVM	Uso de exploits de día cero, ingeniería social	Modelos híbridos de detección	Mejora en detección temprana	Implementación de IA explicable (XAI)
Análisis de Datos	Minería de datos, análisis de correlaciones	Persistencia en redes bancarias, movimientos laterales	SIEM con IA, monitoreo en tiempo real	Reducción del tiempo de respuesta	Uso de gemelos digitales en ciberseguridad
Mecanismos de Defensa	Sistemas de detección de anomalías (ADS)	Ataques dirigidos a endpoints y credenciales	Zero Trust Architecture (ZTA)	Protección contra accesos no autorizados	Implementación de IA en sistemas de Zero Trust
Automatización de Respuestas	SOAR (Security Orchestration, Automation and Response)	Phishing avanzado y ataques a API bancarias	Respuesta automática con Machine Learning	Reducción del impacto de incidentes	Uso de agentes autónomos de seguridad
Evaluación de Riesgos	Modelos probabilísticos para predecir ataques	Exfiltración de datos a través de canales encriptados	Análisis de comportamiento del usuario (UEBA)	Identificación proactiva de amenazas	Integración de ciberseguridad con Blockchain

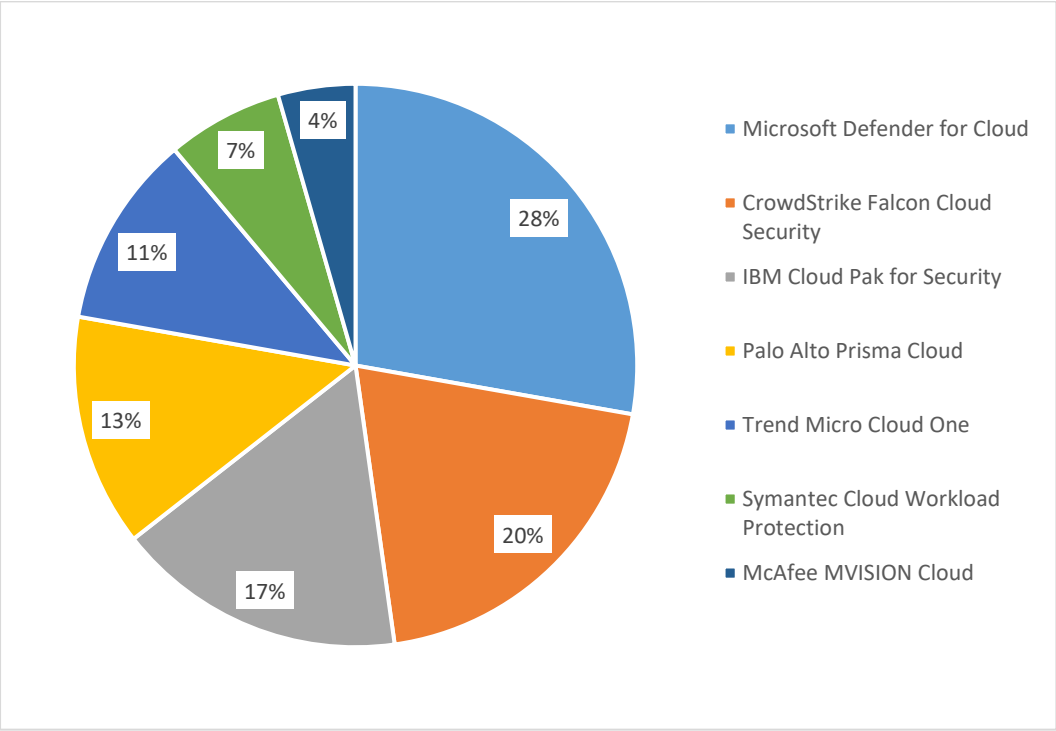
El uso de modelos predictivos de IA permite mejorar la detección, respuesta y prevención de ataques APTs en el sector bancario. Sin embargo, los ciberdelincuentes continúan evolucionando, lo que exige una adaptación constante de las estrategias de seguridad. Las tendencias futuras apuntan hacia la implementación de IA explicable, gemelos digitales, Zero Trust avanzado y Blockchain, marcando el camino para una ciberseguridad bancaria más robusta.

La Tabla 11 presenta los principales softwares de seguridad en la nube utilizados por las empresas a nivel mundial y en Latinoamérica, especialmente en el sector bancario. Estas soluciones se caracterizan por ofrecer protección integral frente a amenazas, vulnerabilidades, accesos no autorizados y ataques persistentes avanzados (APT) en entornos multinube, híbridos y SaaS. Estas herramientas son fundamentales para garantizar la seguridad de los datos y operaciones en las instituciones financieras.

Tabla 11
Frecuencia de uso de Softwares de Seguridad en la Nube en Bancos 2024

Software de Seguridad	Frecuencia de Uso
Microsoft Defender for Cloud	28%
CrowdStrike Falcon Cloud Security	20%
IBM Cloud Pak for Security	17%
Palo Alto Prisma Cloud	13%
Trend Micro Cloud One	11%
Symantec Cloud Workload Protection	7%
McAfee MVISION Cloud	4%

Figura 7
Frecuencia de uso de Softwares de Seguridad en la Nube en Bancos 2024

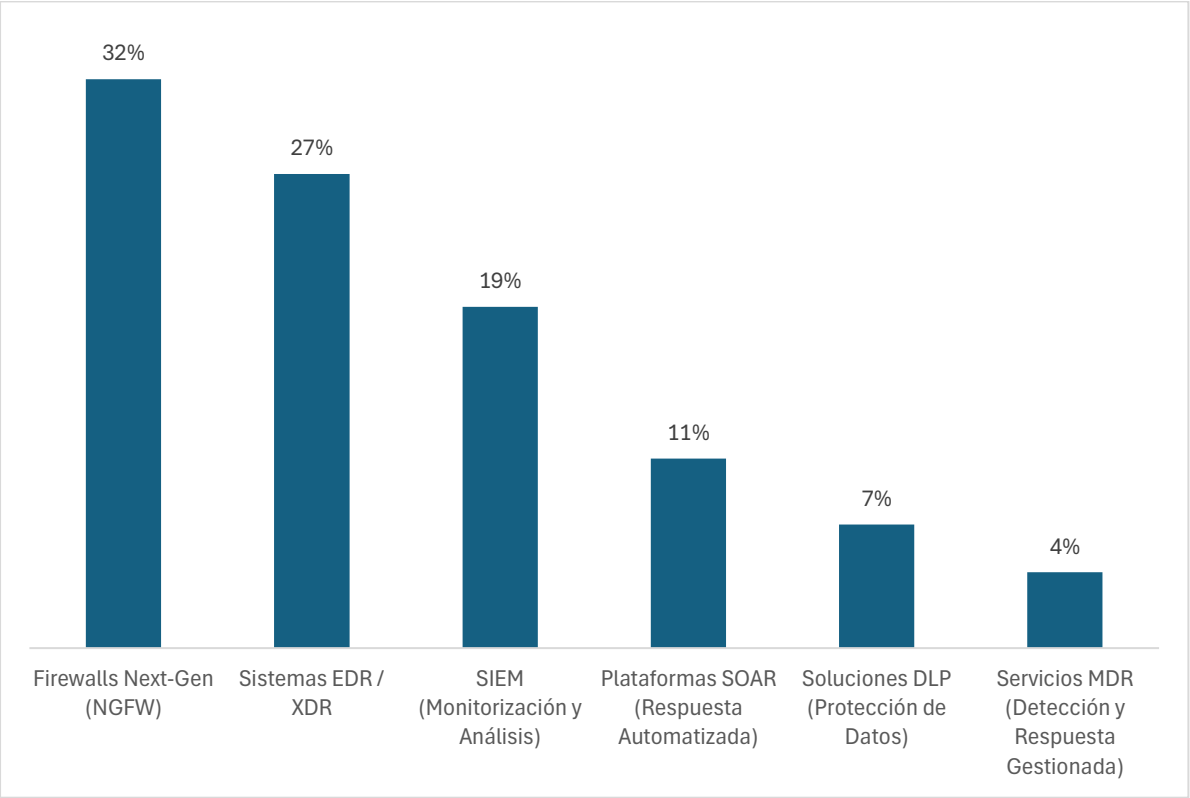


La Tabla 12 muestra los productos de ciberseguridad más adquiridos por los bancos en Perú y Latinoamérica durante el 2024. Estos productos buscan fortalecer las estrategias de protección frente a las amenazas cibernéticas, especialmente frente a los ataques de tipo APT, mediante soluciones tecnológicas de protección de redes, endpoints, identidades y análisis de amenazas.

Tabla 12
Productos de Ciberseguridad más comprados por Bancos

Producto de Ciberseguridad	Frecuencia de Compra
Firewalls Next-Gen (NGFW)	32%
Sistemas EDR / XDR	27%
SIEM (Monitorización y Análisis)	19%
Plataformas SOAR (Respuesta Automatizada)	11%
Soluciones DLP (Protección de Datos)	7%
Servicios MDR (Detección y Respuesta Gestionada)	4%

Figura 8
Productos de Ciberseguridad más comprados por Bancos



4.4.2. Cuadrante mágico de Gardner

En el contexto de la ciberseguridad bancaria, los modelos predictivos de inteligencia artificial (IA) desempeñan un papel crucial en la detección y mitigación de Amenazas Persistentes Avanzadas (APTs). Para analizar el posicionamiento de las principales soluciones en este campo, se ha construido un Cuadrante Mágico de Gartner, que clasifica a los proveedores en función de dos ejes fundamentales:

- 1. Capacidad de Ejecución (Eje Y): Evalúa la eficacia de la solución en la detección, respuesta y mitigación de amenazas.
- 2. Completitud de Visión (Eje X): Mide el grado de innovación y alineación estratégica con las necesidades futuras del sector bancario.

Las tablas presentadas a continuación contienen los valores utilizados para la construcción del gráfico, posicionando a cada proveedor en una de las cuatro categorías del cuadrante: Líderes, Competidores Desafiantes, Visionarios y Jugadores de Nicho. Este análisis proporciona una visión integral sobre el estado actual de las soluciones de ciberseguridad basadas en IA en el sistema bancario de Lima en 2024.

Tabla 13
Clasificación de proveedores en cuadrantes

Cuadrante	Proveedores	Características
Líderes	Microsoft Defender, CrowdStrike Falcon	Alto rendimiento, innovación en IA, gran adopción en la banca.
Competidores desafiantes	IBM QRadar, Palo Alto Cortex XDR	Fuerte ejecución, pero menos innovación en IA.
Visionarios	Trend Micro Apex One, Symantec Endpoint	Innovadores con IA, pero con menos adopción en banca.
Jugadores de nicho	McAfee MVISION	Focalizados en necesidades específicas, con menor alcance global.

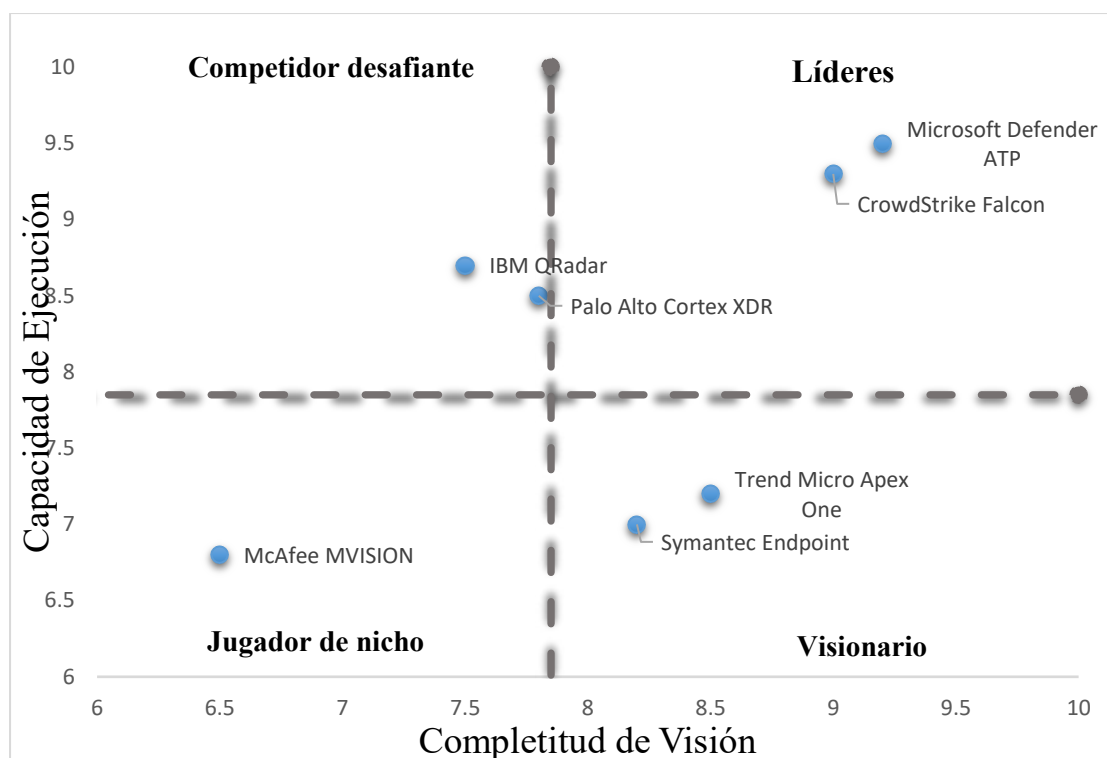
Tabla 14
Caracterización de los proveedores

Proveedor	Cuadrante	Puntos Fuertes	Debilidades
Microsoft Defender	Líder	Alta integración con sistemas bancarios, IA avanzada.	Depende del ecosistema de Microsoft.
CrowdStrike Falcon	Líder	Detección en tiempo real, gran precisión.	Costoso en comparación con otros.
IBM Security QRadar	Competidor Desafiante	Potente SIEM con análisis avanzado.	Interfaz compleja, curva de aprendizaje alta.
Palo Alto Cortex XDR	Competidor Desafiante	Respuesta automatizada, buen soporte.	Implementación costosa.
Trend Micro Apex One	Visionario	Fuerte en protección endpoint con IA.	No es líder en respuesta a incidentes.
Symantec Endpoint Security	Visionario	Protección multicapa contra APTs.	Menos integración con bancos.
McAfee MVISION	Jugador de Nicho	Arquitectura en la nube, gestión sencilla.	No cubre todo el ciclo de seguridad.

El Cuadrante Mágico evalúa soluciones de inteligencia artificial (IA) para la protección contra Amenazas Persistentes Avanzadas (APTs) en el sector bancario de Lima. Los proveedores se distribuyen en cuatro categorías, según su Capacidad de Ejecución (Eje Y) y su Completitud de Visión (Eje X).

Tabla 15
Valorización de los proveedores

Proveedor/Tecnología	Completitud de Visión (X)	Capacidad de Ejecución (Y)	Categoría
Microsoft Defender ATP	9,2	9,5	Líder
CrowdStrike Falcon	9	9,3	Líder
IBM QRadar	7,5	8,7	Competidor desafiante
Palo Alto Cortex XDR	7,8	8,5	Competidor desafiante
Trend Micro Apex One	8,5	7,2	Visionario
Symantec Endpoint	8,2	7	Visionario
McAfee MVISION	6,5	6,8	Jugador de dicho



Interpretación:

Líderes (Arriba a la derecha)

- Microsoft Defender y CrowdStrike Falcon dominan el mercado gracias a su IA avanzada y eficacia comprobada contra APTs, son las opciones más recomendadas para bancos con alta exigencia en ciberseguridad.

Competidores Desafiantes (Arriba a la izquierda)

- IBM QRadar y Palo Alto Cortex XDR tienen fuerte ejecución, pero aún no integran IA de manera tan avanzada como los líderes, son excelentes opciones para bancos con infraestructura ya consolidada.

Visionarios (Abajo a la derecha)

- Trend Micro Apex One y Symantec Endpoint innovan con IA en detección de amenazas, pero aún no tienen la adopción masiva de los líderes, son apuestas a futuro con potencial de crecimiento.

Jugadores de Nicho (Abajo a la izquierda)

- McAfee MVISION tiene capacidades específicas, pero no es una solución completa para grandes bancos, es útil para instituciones más pequeñas o con requerimientos específicos.

4.5. Análisis descriptivo

De acuerdo a los resultados, el 27,8% de los encuestados consideraron que los algoritmos utilizados en su institución abordan eficazmente los problemas "a veces". Un 44,3% señaló que esto ocurre "casi siempre", mientras que el 27,8% indicó que "siempre" se abordan eficazmente. Esto sugiere que la mayoría percibe una efectividad frecuente, aunque todavía hay margen de mejora para alcanzar una eficacia consistente.

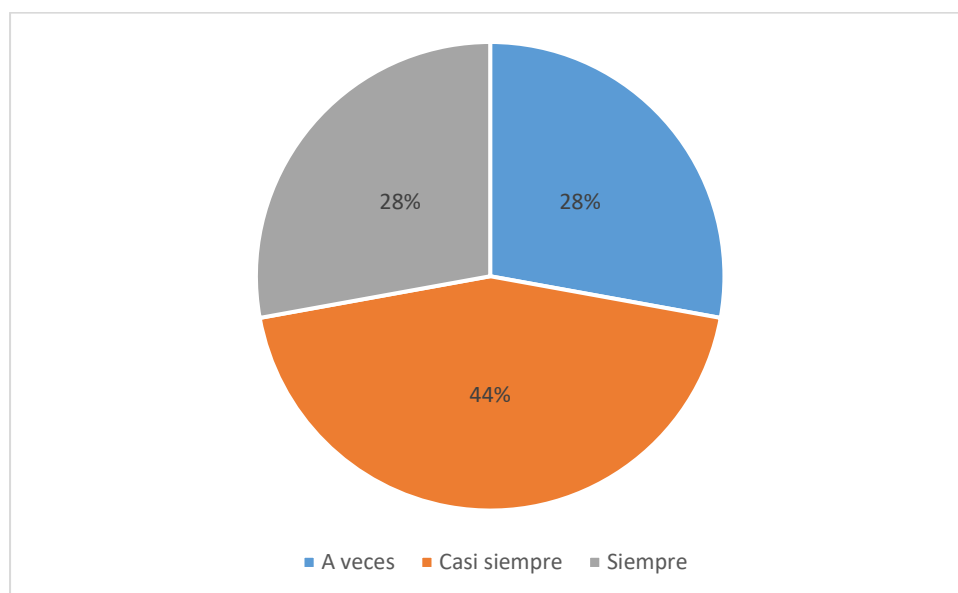
Tabla 16

Frecuencia respecto a la eficacia de los algoritmos para abordar problemas.

		Frecuencia	Porcentaje
Válido	A veces	22	27,8
	Casi siempre	35	44,3
	Siempre	22	27,8
	Total	79	100,0

Figura 9

Frecuencia respecto a la eficacia de los algoritmos para abordar problemas.



De acuerdo a los resultados, el 27,8% de los participantes opinaron que los resultados de los modelos predictivos son confiables "a veces". La mayoría, un 44,3%, percibió esta característica "casi siempre", mientras que el 27,8% afirmó que es así "siempre". Esto evidencia una percepción positiva generalizada respecto a la confiabilidad.

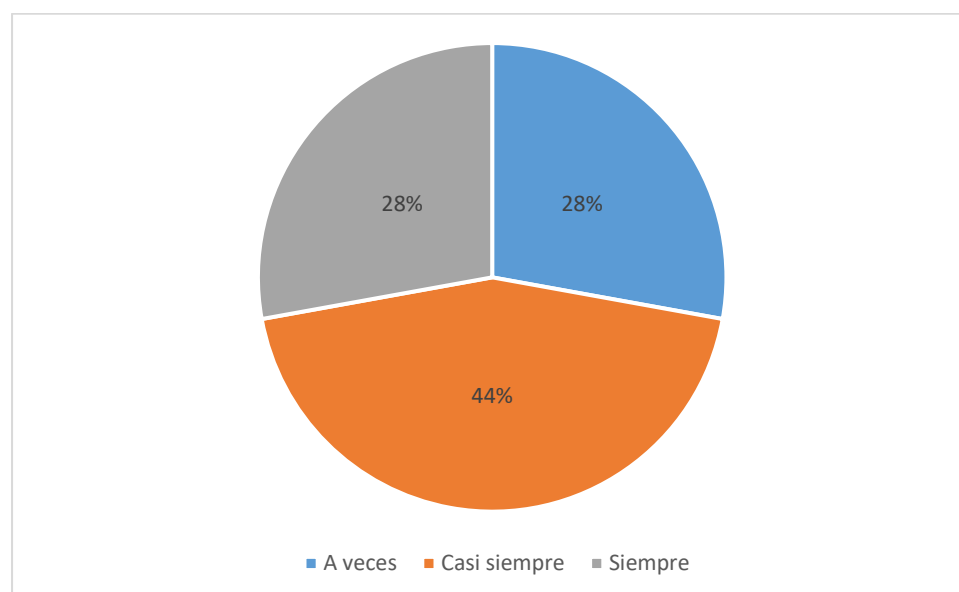
Tabla 17

Frecuencia respecto a la confiabilidad de los modelos predictivos.

		Frecuencia	Porcentaje	Porcentaje válido	Porcentaje acumulado
Válido	A veces	22	27,8	27,8	27,8
	Casi siempre	35	44,3	44,3	72,2
	Siempre	22	27,8	27,8	100,0
	Total	79	100,0	100,0	

Figura 10

Frecuencia respecto a la confiabilidad de los modelos predictivos.



De acuerdo a los resultados, el 26,6% de los encuestados consideraron que los modelos predictivos mantienen su eficacia ante cambios en los datos "a veces". Un 31,6% percibió esta eficacia "casi siempre", y el 41,8% la identificó "siempre". Este panorama resalta una percepción mayoritaria de adaptabilidad en los modelos, aunque con una proporción significativa que experimenta limitaciones.

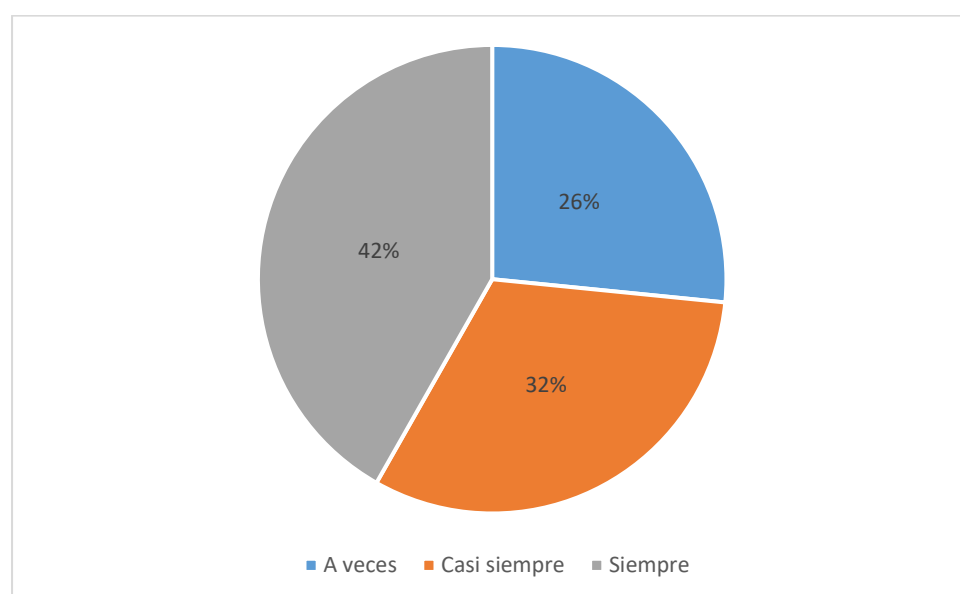
Tabla 18

Frecuencia respecto a la eficacia de los modelos predictivos ante cambios en los datos.

		Frecuencia	Porcentaje
Válido	A veces	21	26,6
	Casi siempre	25	31,6
	Siempre	33	41,8
	Total	79	100,0

Figura 11

Frecuencia respecto a la eficacia de los modelos predictivos ante cambios en los datos.



De acuerdo a los resultados, el 27,8% indicó que los algoritmos de aprendizaje automático pueden generalizar lo suficiente como para evitar nuevas amenazas "a veces". Un 40,5% expresó que esto ocurre "casi siempre", mientras que el 31,6% señaló que sucede "siempre". Esto sugiere que, si bien los algoritmos tienen una buena capacidad de generalización, aún existen casos en los que se requieren ajustes para garantizar una protección efectiva contra nuevas amenazas.

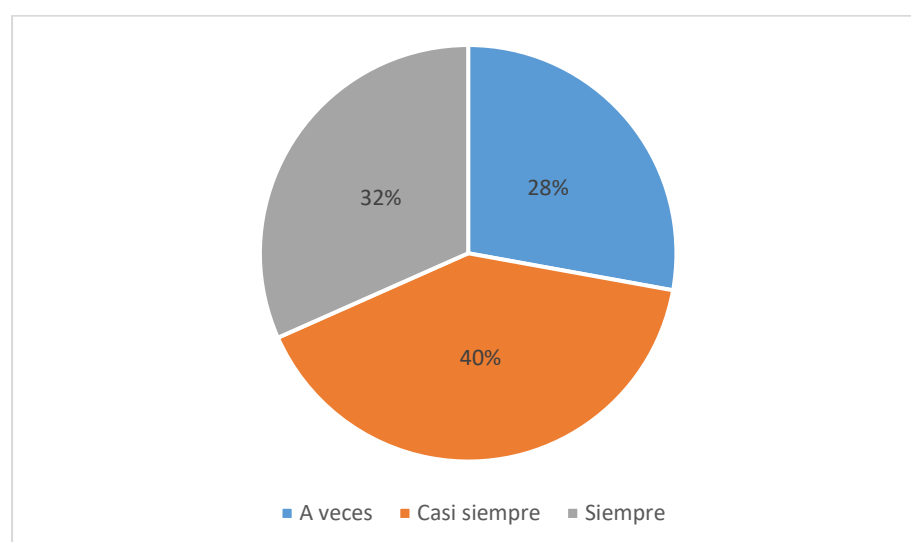
Tabla 19

Frecuencia de generalización de los algoritmos para prevenir amenazas en la seguridad bancaria.

		Frecuencia	Porcentaje
Válido	A veces	22	27,8
	Casi siempre	32	40,5
	Siempre	25	31,6
	Total	79	100,0

Figura 12

Frecuencia de generalización de los algoritmos para prevenir amenazas en la seguridad bancaria.



De acuerdo a los resultados, el 24,1% de los encuestados consideró que su institución actualiza regularmente los métodos de entrenamiento de los modelos predictivos para mejorar la detección de amenazas "a veces". Un 44,3% percibió que esto ocurre "casi siempre", mientras que el 31,6% afirmó que sucede "siempre". Esto refleja una tendencia positiva hacia la actualización de los métodos de entrenamiento, aunque aún existe margen para fortalecer esta práctica.

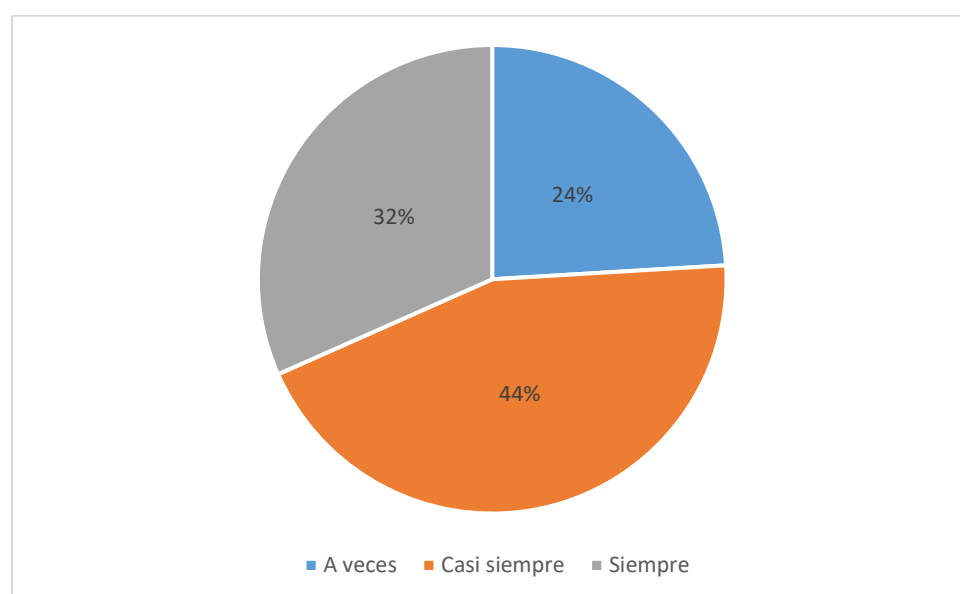
Tabla 20

Frecuencia de actualización de métodos de entrenamiento para mejorar la detección de amenazas.

		Frecuencia	Porcentaje
Válido	A veces	19	24,1
	Casi siempre	35	44,3
	Siempre	25	31,6
	Total	79	100,0

Figura 13

Frecuencia de actualización de métodos de entrenamiento para mejorar la detección de amenazas.



De acuerdo a los resultados, el 20,3% de los participantes percibió que el conjunto de datos utilizados para entrenar y probar los modelos refleja las amenazas reales para una protección efectiva "a veces". Un 40,5% consideró que esto ocurre "casi siempre", mientras que el 39,2% afirmó que sucede "siempre". Esto indica una percepción mayoritariamente positiva sobre la calidad de los datos, aunque aún existen casos en los que podrían mejorarse para una protección más efectiva.

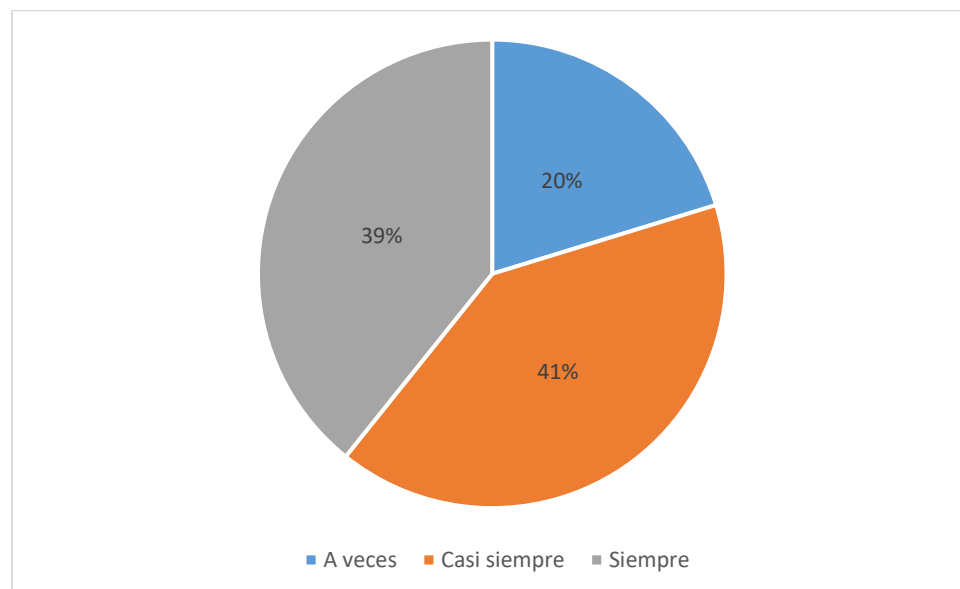
Tabla 21

Frecuencia de representatividad del conjunto de datos en la detección de amenazas reales.

		Frecuencia	Porcentaje
Válido	A veces	16	20,3
	Casi siempre	32	40,5
	Siempre	31	39,2
	Total	79	100,0

Figura 14

Frecuencia de representatividad del conjunto de datos en la detección de amenazas reales.



De acuerdo a los resultados, el 39,2% indicó que su equipo revisa las métricas de rendimiento para evaluar la efectividad de los modelos predictivos "a veces". Un 38,0% afirmó que esto ocurre "casi siempre", y un 22,8% señaló que se hace "siempre". Este resultado muestra una tendencia hacia la revisión frecuente, pero sugiere la necesidad de fortalecer su sistematicidad.

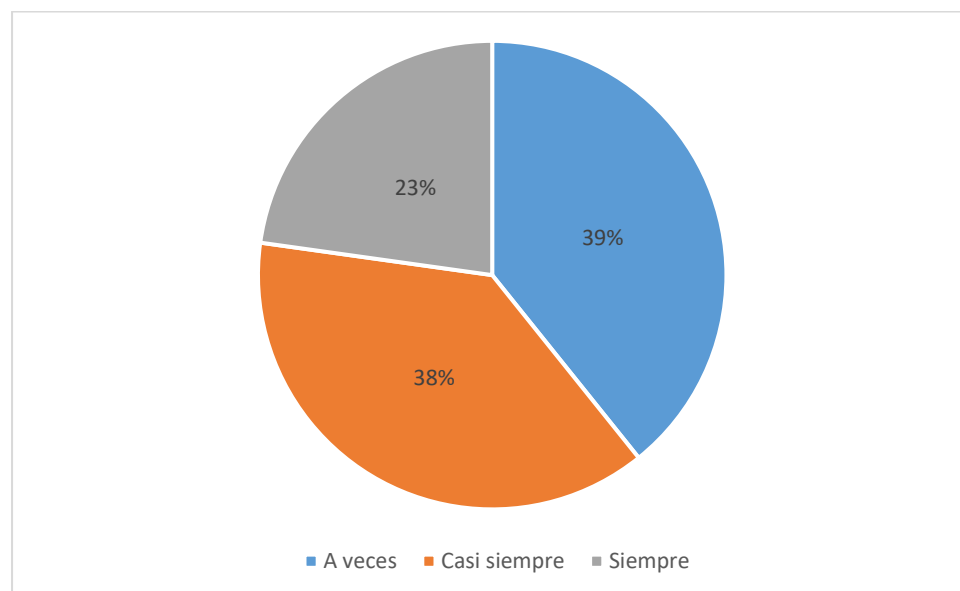
Tabla 22

Frecuencia respecto a la revisión de métricas de rendimiento de los modelos predictivos.

		Frecuencia	Porcentaje
Válido	A veces	31	39,2
	Casi siempre	30	38,0
	Siempre	18	22,8
	Total	79	100,0

Figura 15

Frecuencia respecto a la revisión de métricas de rendimiento de los modelos predictivos.



De acuerdo a los resultados, el 20,3% de los encuestados consideraron que el tiempo invertido en el entrenamiento de los modelos es razonable y productivo "a veces". Un 40,5% percibió esto "casi siempre", mientras que el 39,2% opinó que es así "siempre". Esto denota una percepción positiva, aunque con oportunidades para optimizar el uso del tiempo.

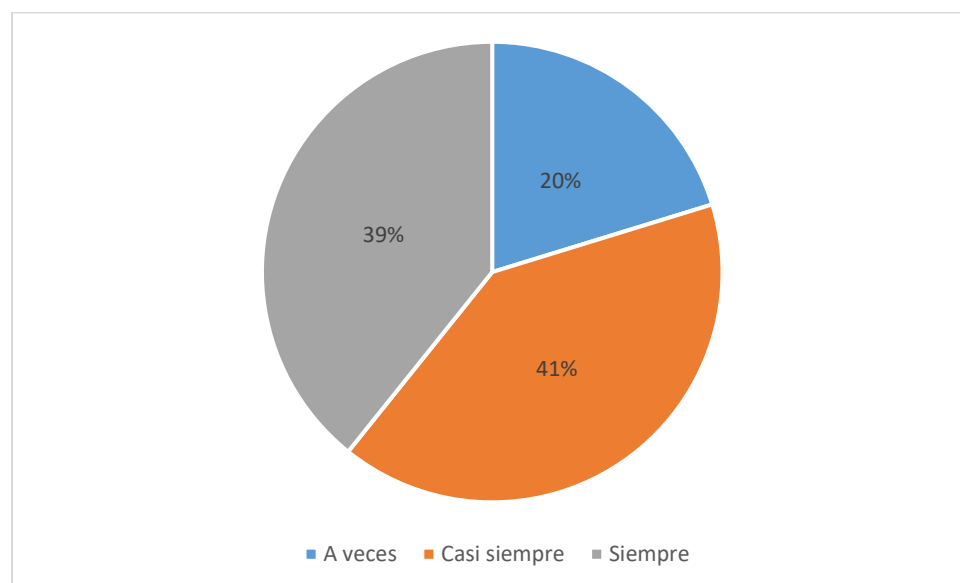
Tabla 23

Frecuencia respecto al tiempo invertido en el entrenamiento de los modelos.

		Frecuencia	Porcentaje
Válido	A veces	16	20,3
	Casi siempre	32	40,5
	Siempre	31	39,2
	Total	79	100,0

Figura 16

Frecuencia respecto al tiempo invertido en el entrenamiento de los modelos.



De acuerdo a los resultados, el 26,6% de los participantes señaló que el tiempo de inferencia de los modelos en producción es adecuado para detectar y responder a amenazas de manera efectiva "a veces". Un 39,2% indicó que es adecuado "casi siempre", mientras que el 34,2% afirmó que es así "siempre". Esto sugiere una percepción mayoritariamente favorable, aunque aún existen oportunidades de mejora en términos de rapidez y eficiencia para una protección más efectiva.

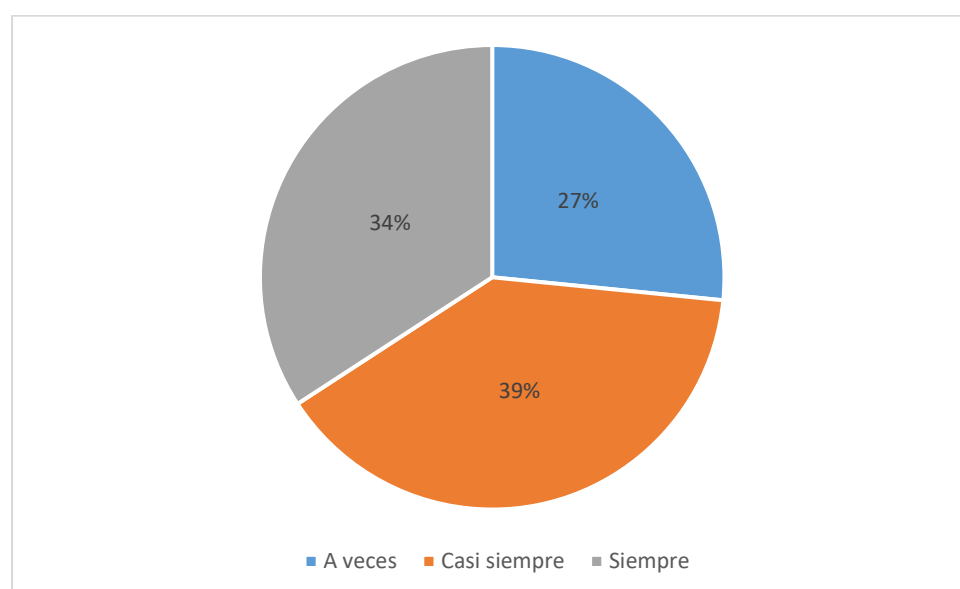
Tabla 24

Frecuencia de adecuación del tiempo de inferencia para la detección y respuesta a amenazas.

		Frecuencia	Porcentaje
Válido	A veces	21	26,6
	Casi siempre	31	39,2
	Siempre	27	34,2
	Total	79	100,0

Figura 17

Frecuencia de adecuación del tiempo de inferencia para la detección y respuesta a amenazas.



De acuerdo a los resultados, el 39,2% de los encuestados consideró que los modelos predictivos pueden escalar efectivamente ante un aumento en la carga de trabajo "a veces". Un 38,0% señaló que esto ocurre "casi siempre", y un 22,8% opinó que es así "siempre". Este panorama sugiere una necesidad de trabajar en la robustez de los modelos para manejar mayores demandas.

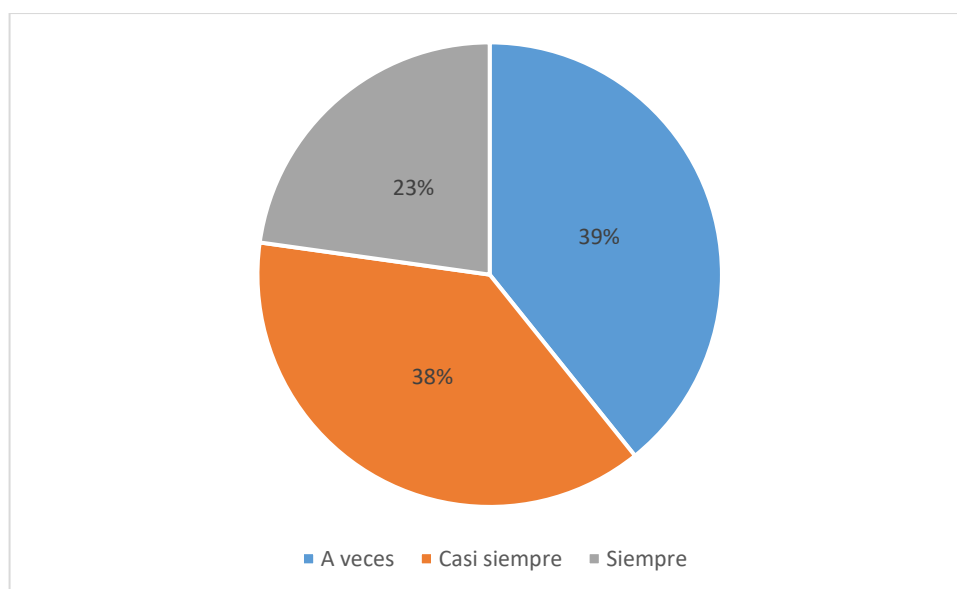
Tabla 25

Frecuencia respecto a la escalabilidad de los modelos predictivos ante mayor carga de trabajo.

		Frecuencia	Porcentaje
Válido	A veces	31	39,2
	Casi siempre	30	38,0
	Siempre	18	22,8
	Total	79	100,0

Figura 18

Frecuencia respecto a la escalabilidad de los modelos predictivos ante mayor carga de trabajo.



De acuerdo a los resultados, el 24,1% consideró que la integración de los modelos predictivos con los sistemas existentes para fortalecer la protección contra amenazas ha sido "a veces" sencilla. Un 39,2% indicó que esta integración se da "casi siempre", mientras que un 36,7% afirmó que sucede "siempre". Este resultado refleja avances significativos en la integración, aunque aún existen desafíos en ciertos contextos que podrían optimizarse.

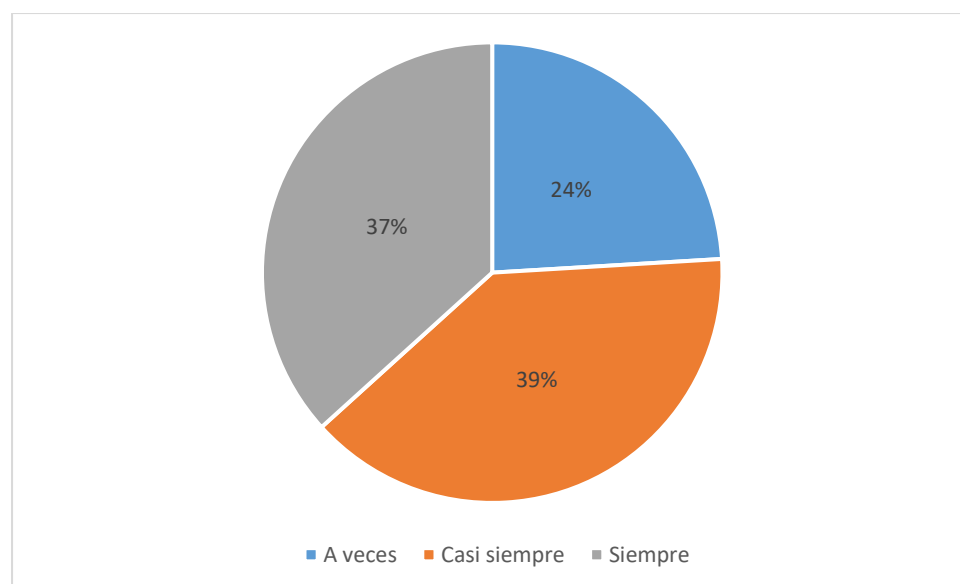
Tabla 26

Frecuencia de facilidad de integración de modelos predictivos para la protección contra amenazas.

		Frecuencia	Porcentaje
Válido	A veces	19	24,1
	Casi siempre	31	39,2
	Siempre	29	36,7
	Total	79	100,0

Figura 19

Frecuencia de facilidad de integración de modelos predictivos para la protección contra amenazas.



De acuerdo a los resultados, el 26,6% opinó que los modelos predictivos logran adaptarse a distintos escenarios de ataque "a veces". Un 36,7% señaló que esto ocurre "casi siempre", mientras que otro 36,7% expresó que sucede "siempre". Este equilibrio refleja una percepción de resiliencia en los modelos, pero con áreas de mejora frente a escenarios más complejos.

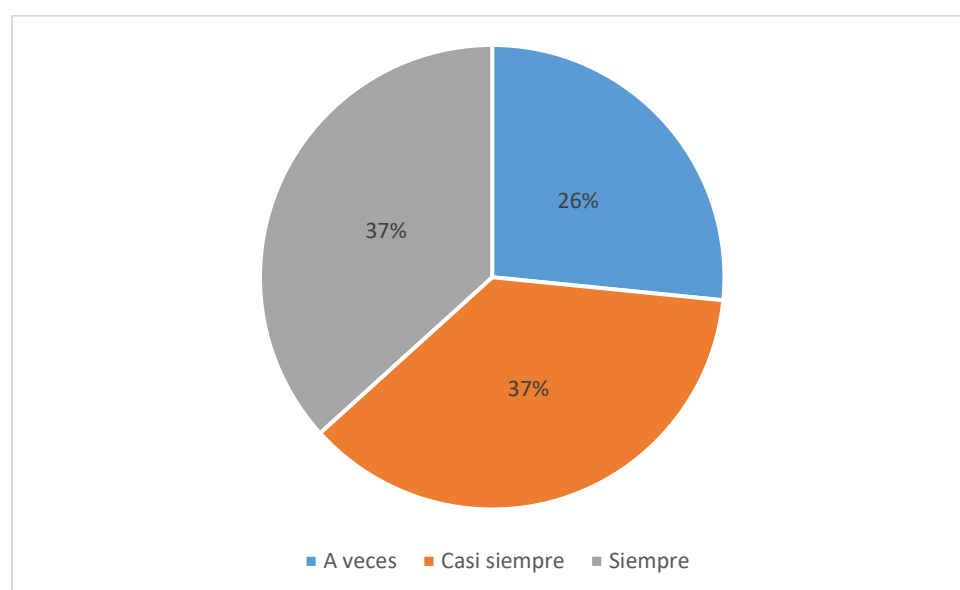
Tabla 27

Frecuencia respecto a la adaptación de los modelos predictivos a escenarios de ataque.

		Frecuencia	Porcentaje
Válido	A veces	21	26,6
	Casi siempre	29	36,7
	Siempre	29	36,7
	Total	79	100,0

Figura 20

Frecuencia respecto a la adaptación de los modelos predictivos a escenarios de ataque.



De acuerdo a los resultados, se evidencia que la efectividad del sistema para detectar ataques, en comparación con otras amenazas, es percibida como alta por la mayoría de los encuestados. El 41.8% considera que su sistema detecta ataques "casi siempre", mientras que un 22.8% indica que lo hace "siempre". Sin embargo, un 35.4% señala que esto ocurre "a veces", lo que sugiere que todavía existen oportunidades de mejora en la capacidad de detección.

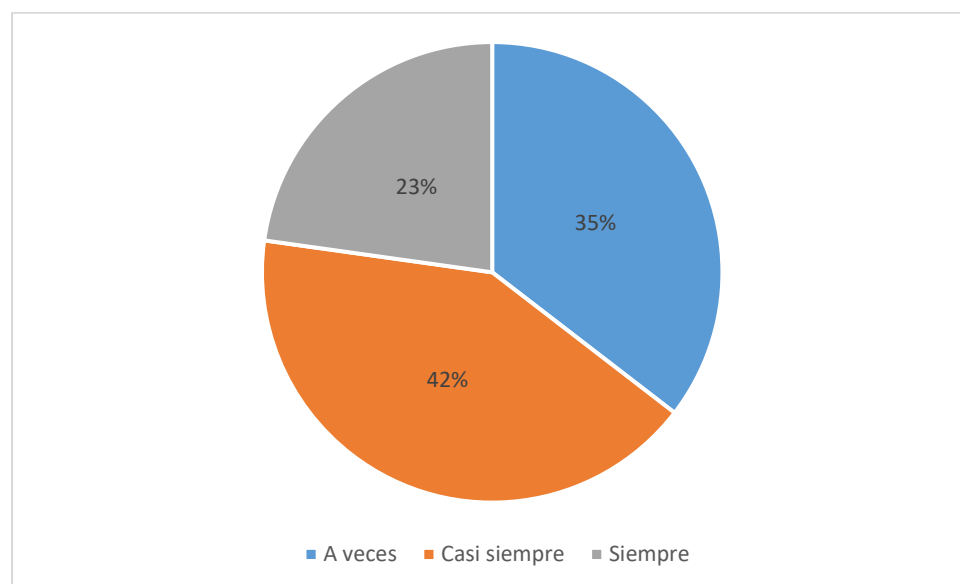
Tabla 28

Frecuencia respecto a la detección efectiva de ataques en comparación con otras amenazas.

		Frecuencia	Porcentaje
Válido	A veces	28	35,4
	Casi siempre	33	41,8
	Siempre	18	22,8
	Total	79	100,0

Figura 21

Frecuencia respecto a la detección efectiva de ataques en comparación con otras amenazas.



Los resultados reflejan que los falsos positivos o negativos generados por el sistema tienen un impacto relevante en su eficacia. El 39.2% de los participantes observa que estos eventos ocurren "casi siempre", mientras que un 29.1% menciona que se presentan "siempre". Aunque un 31.6% asegura que suceden "a veces", la frecuencia de estas incidencias podría estar influyendo en la percepción de la confiabilidad del sistema.

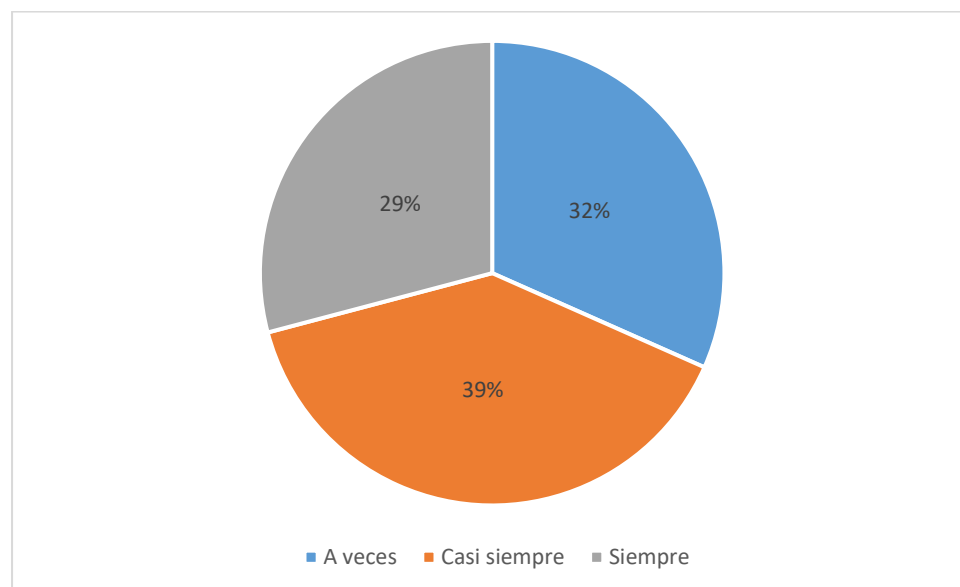
Tabla 29

Frecuencia respecto a la generación de falsos positivos o negativos en los sistemas de detección.

		Frecuencia	Porcentaje
Válido	A veces	25	31,6
	Casi siempre	31	39,2
	Siempre	23	29,1
	Total	79	100,0

Figura 22

Frecuencia respecto a la generación de falsos positivos o negativos en los sistemas de detección.



La rapidez del sistema para detectar amenazas persistentes avanzadas es considerada favorable por los encuestados. El 39.2% indica que el sistema actúa "casi siempre" con rapidez, y un 34.2% lo califica como "siempre" rápido. A pesar de esto, un 26.6% señala que dicha velocidad solo se cumple "a veces", destacando la necesidad de optimizar aún más el tiempo de respuesta.

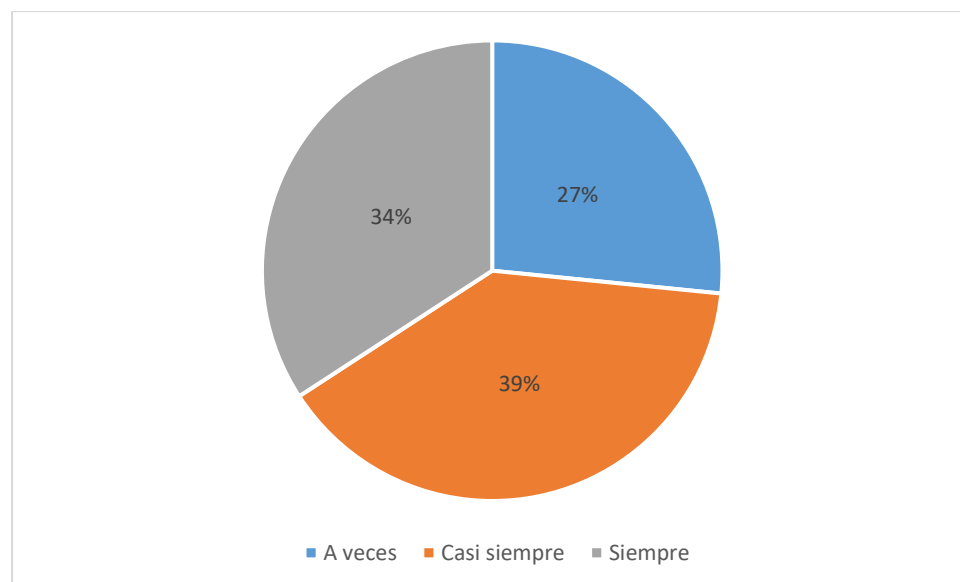
Tabla 30

Frecuencia respecto a la rapidez en la detección de amenazas persistentes avanzadas.

		Frecuencia	Porcentaje
Válido	A veces	21	26,6
	Casi siempre	31	39,2
	Siempre	27	34,2
	Total	79	100,0

Figura 23

Frecuencia respecto a la rapidez en la detección de amenazas persistentes avanzadas.



Respecto a la capacidad de respuesta del equipo ante los ataques detectados, los resultados muestran que el 36.7% de los encuestados cree que la reacción es oportuna "casi siempre", mientras que un 34.2% sostiene que esto ocurre "siempre". Por otro lado, el 29.1% menciona que esta capacidad de respuesta solo se da "a veces", lo que puede sugerir áreas de mejora en los procesos de intervención inmediata

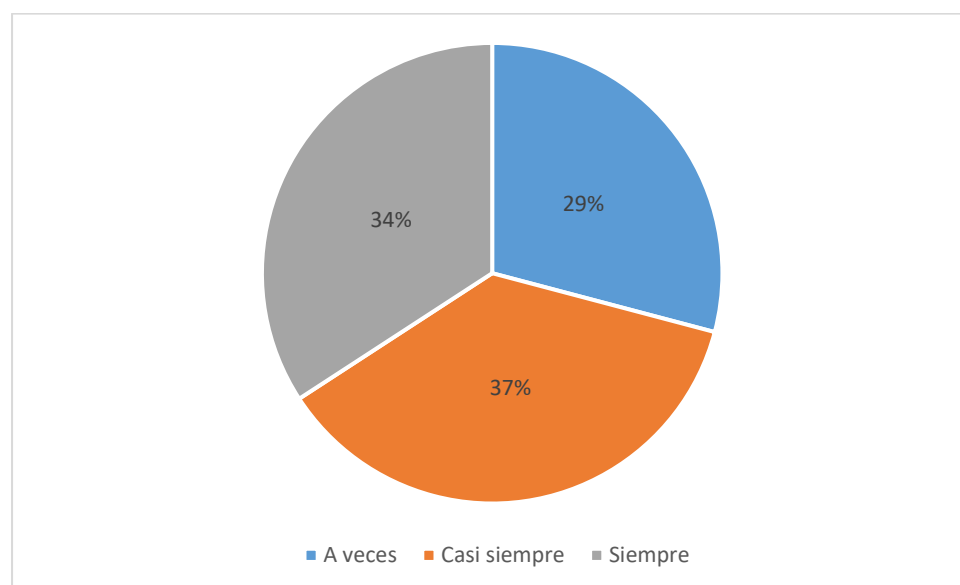
Tabla 31

Frecuencia respecto a la respuesta oportuna del equipo ante ataques detectados.

		Frecuencia	Porcentaje
Válido	A veces	23	29,1
	Casi siempre	29	36,7
	Siempre	27	34,2
	Total	79	100,0

Figura 24

Frecuencia respecto a la respuesta oportuna del equipo ante ataques detectados.



De acuerdo a los resultados, la capacidad del sistema para recuperarse después de un ataque es percibida como efectiva por una mayoría de los encuestados. El 34.2% considera que su sistema logra recuperarse "casi siempre", y un 30.4% lo califica como "siempre" eficiente en la recuperación. Sin embargo, un 35.4% menciona que esto ocurre "a veces", lo que sugiere que podrían fortalecerse las estrategias para asegurar una recuperación más consistente.

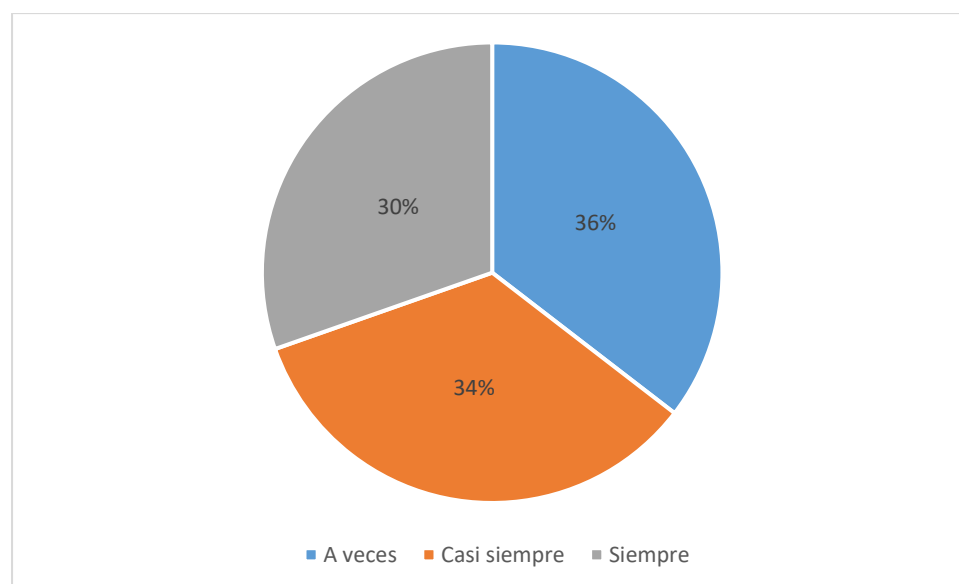
Tabla 32

Frecuencia respecto a la recuperación del sistema tras un ataque.

		Frecuencia	Porcentaje
Válido	A veces	28	35,4
	Casi siempre	27	34,2
	Siempre	24	30,4
	Total	79	100,0

Figura 25

Frecuencia respecto a la recuperación del sistema tras un ataque.



En cuanto a la resistencia del sistema frente a técnicas de evasión utilizadas por atacantes avanzados, el 39.2% de los encuestados opina que el sistema logra mantenerse resistente "casi siempre", mientras que un 27.8% afirma que esta resistencia se mantiene "siempre". No obstante, un 32.9% indica que dicha capacidad se manifiesta únicamente "a veces", lo que refleja una necesidad de adaptación continua frente a amenazas sofisticadas.

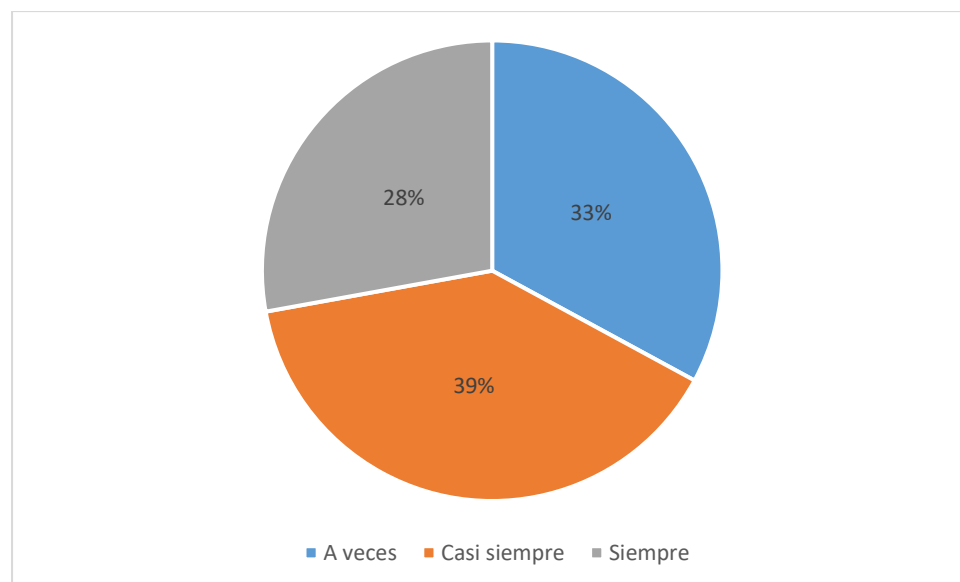
Tabla 33

Frecuencia respecto a la resistencia del sistema frente a técnicas de evasión.

		Frecuencia	Porcentaje
Válido	A veces	26	32,9
	Casi siempre	31	39,2
	Siempre	22	27,8
	Total	79	100,0

Figura 26

Frecuencia respecto a la resistencia del sistema frente a técnicas de evasión.



La efectividad de los parches de seguridad implementados para proteger el sistema contra amenazas avanzadas es valorada positivamente. El 35.4% considera que los parches son "siempre" efectivos, mientras que el 30.4% afirma que funcionan "casi siempre". Sin embargo, un 34.2% señala que su efectividad se observa solo "a veces", lo que podría estar relacionado con limitaciones en la aplicación o actualización de estas medidas.

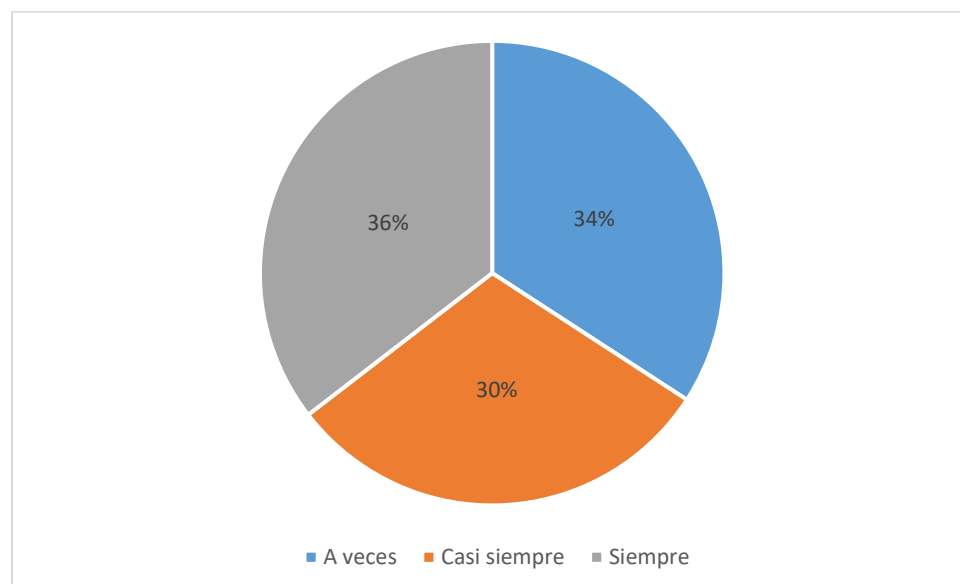
Tabla 34

Frecuencia respecto a la efectividad de los parches de seguridad frente a amenazas avanzadas.

		Frecuencia	Porcentaje
Válido	A veces	27	34,2
	Casi siempre	24	30,4
	Siempre	28	35,4
	Total	79	100,0

Figura 27

Frecuencia respecto a la efectividad de los parches de seguridad frente a amenazas avanzadas.



Sobre la frecuencia con la que se actualizan los modelos de seguridad para mantenerse al día con nuevas amenazas, el 40.5% de los participantes asegura que esto ocurre "casi siempre", y un 39.2% afirma que se realiza "siempre". No obstante, un 20.3% menciona que estas actualizaciones solo se llevan a cabo "a veces", lo que subraya la importancia de mantener procesos más regulares y dinámicos para abordar amenazas emergentes.

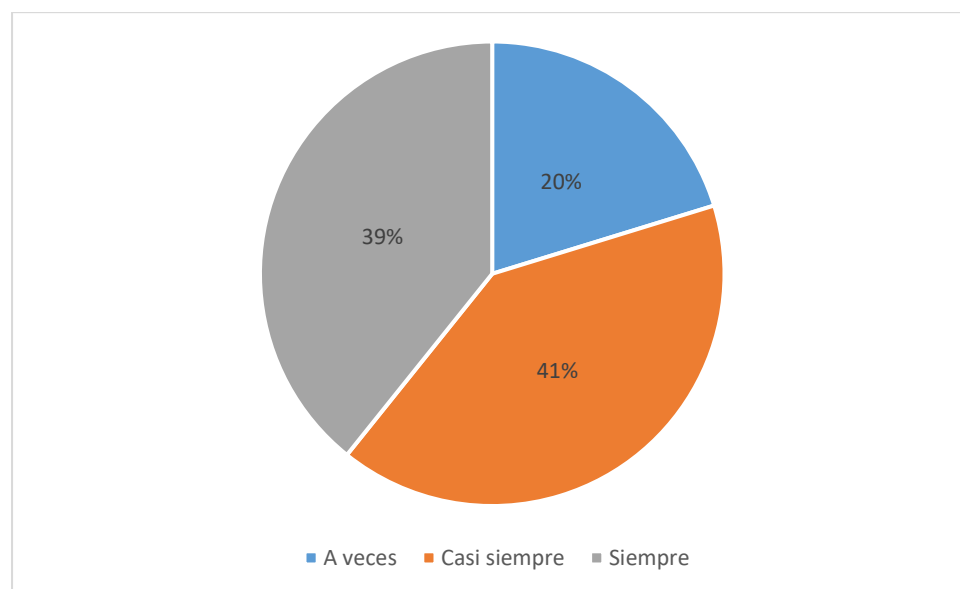
Tabla 35

Frecuencia respecto a la actualización de modelos de seguridad frente a nuevas amenazas.

		Frecuencia	Porcentaje
Válido	A veces	16	20,3
	Casi siempre	32	40,5
	Siempre	31	39,2
	Total	79	100,0

Figura 28

Frecuencia respecto a la actualización de modelos de seguridad frente a nuevas amenazas.



La eficacia de las contramedidas implementadas para proteger contra amenazas recientes es percibida de forma positiva por los encuestados. El 40.5% califica estas medidas como efectivas "casi siempre", mientras que un 32.9% las considera "siempre" eficaces. Sin embargo, un 26.6% opina que estas contramedidas solo funcionan "a veces", lo que evidencia áreas de mejora en la implementación de estrategias más robustas.

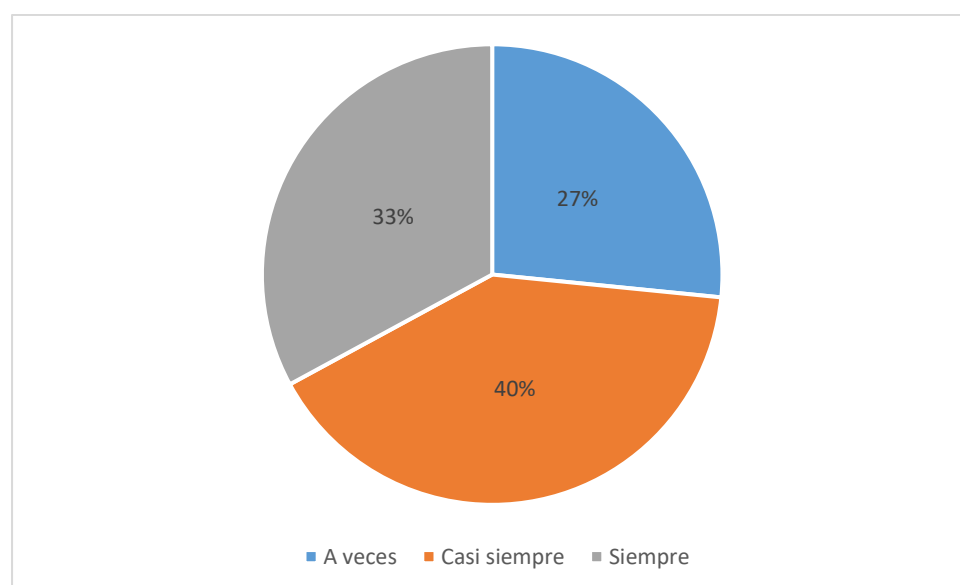
Tabla 36

Frecuencia respecto a la eficacia de las contramedidas implementadas frente a amenazas recientes.

		Frecuencia	Porcentaje
Válido	A veces	21	26,6
	Casi siempre	32	40,5
	Siempre	26	32,9
	Total	79	100,0

Figura 29

Frecuencia respecto a la eficacia de las contramedidas implementadas frente a amenazas recientes.



Respecto a la capacidad del sistema para adaptarse rápidamente a nuevas amenazas identificadas, el 41.8% de los encuestados asegura que esta adaptación ocurre "siempre", y el 40.5% señala que sucede "casi siempre". Solo un 17.7% indica que la adaptación se da "a veces", lo cual resalta un nivel positivo de flexibilidad en la mayoría de los casos, aunque existen oportunidades para fortalecer esta capacidad en escenarios menos frecuentes.

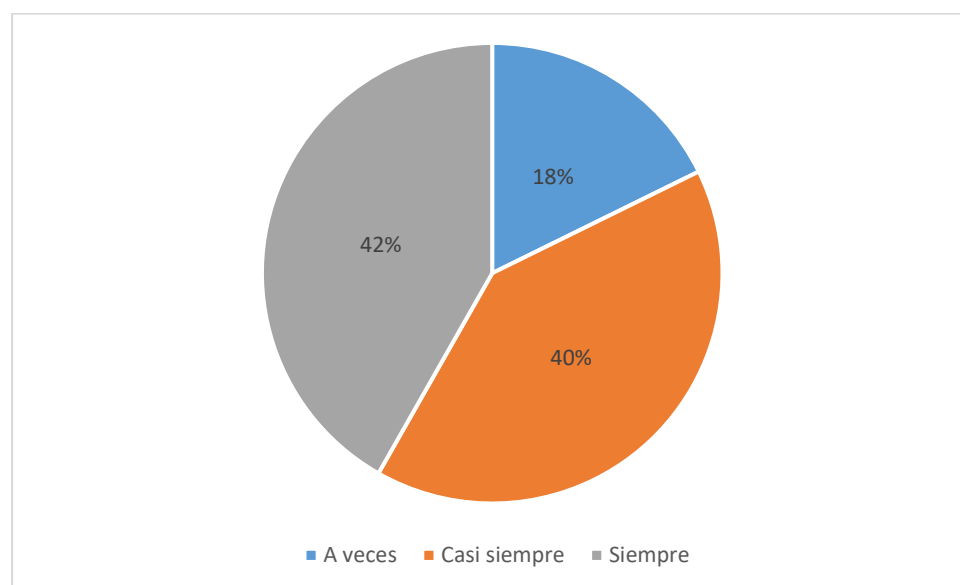
Tabla 37

Frecuencia respecto a la adaptación rápida del sistema a nuevas amenazas.

		Frecuencia	Porcentaje
Válido	A veces	14	17,7
	Casi siempre	32	40,5
	Siempre	33	41,8
	Total	79	100,0

Figura 30

Frecuencia respecto a la adaptación rápida del sistema a nuevas amenazas.



Finalmente, la efectividad del sistema para detectar y ajustarse a la evolución de los patrones de ataque es valorada por los participantes. El 40.5% opina que este ajuste se produce "casi siempre", y un 38.0% indica que se da "siempre". Sin embargo, un 21.5% menciona que esto ocurre solo "a veces", lo que sugiere la necesidad de un monitoreo continuo y mejoras en la capacidad de aprendizaje adaptativo del sistema.

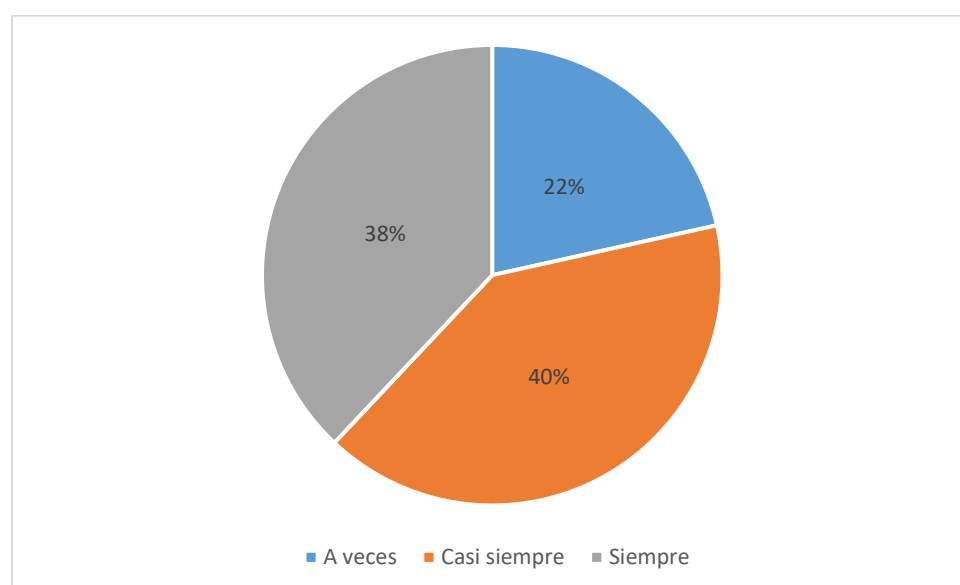
Tabla 38

Frecuencia respecto a la detección y ajuste del sistema a la evolución de patrones de ataque.

		Frecuencia	Porcentaje
Válido	A veces	17	21,5
	Casi siempre	32	40,5
	Siempre	30	38,0
	Total	79	100,0

Figura 31

Frecuencia respecto a la detección y ajuste del sistema a la evolución de patrones de ataque.



4.6. Análisis inferencial

4.6.1. Hipótesis general

H₀: Los modelos predictivos de IA no se relacionan con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

H₁: Los modelos predictivos de IA se relacionan con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

Tabla 39

Correlación entre los modelos predictivos de IA y la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024

			Protección contra APT
Rho de Spearman	Modelos	Coefficiente de correlación	,732**
	Predictivos de IA	Sig. (bilateral)	,000
		N	79

**Esta correlación es significativa al nivel $p < 0.01$ (99% de confianza).

Interpretación: Se acepta la hipótesis alterna Los modelos predictivos de IA se relacionan con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024 porque mediante el coeficiente de correlación de Spearman (ρ). se observa una significancia de 0.000, menor a 0.00, lo que sugiere una correlación significativa entre ambas variables. Además, se registra un coeficiente de correlación de 0.732**, indicando una alta correlación positiva entre ellas.

4.6.2. Hipótesis secundarias

a. Hipótesis específica 1

H₀: Los modelos predictivos de IA no se relacionarán con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

H₁: Los modelos predictivos de IA se relacionarán con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

Tabla 40

Correlación entre los modelos predictivos de IA y la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024

		Detección y Prevención	
Rho de Spearman	Modelos	Coefficiente de correlación	,429**
	Predictivos de	Sig. (bilateral)	,000
	IA	N	79

**Esta correlación es significativa al nivel $p < 0.01$ (99% de confianza).

Interpretación: Se acepta la hipótesis específica 1 alterna, que postula Los modelos predictivos de IA se relacionarán con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024, ya que el coeficiente de correlación de Spearman (ρ), se encontró una significancia de 0.000, lo que sugiere una correlación significativa entre ambas. Además, se observa un coeficiente de correlación de 0.429**, indicando una correlación positiva alta entre ellas.

b. Hipótesis específica 2

H₀: Los modelos predictivos de IA no se relacionarán con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

H₁: Los modelos predictivos de IA se relacionarán con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

Tabla 41

Correlación entre los modelos predictivos de IA y la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024

			Resiliencia del Sistema
Rho de Spearman	Modelos Predictivos de IA	Coefficiente de correlación	,479**
		Sig. (bilateral)	,000
		N	79

**Esta correlación es significativa al nivel $p < 0.01$ (99% de confianza).

Interpretación: Se presenta acepta hipótesis específica 2 alterna, que postula Los modelos predictivos de IA se relacionarán con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024, ya que el coeficiente de correlación de Spearman (ρ), se encontró una significancia de 0.000, lo que sugiere una correlación significativa entre ambas. Además, se observa un coeficiente de correlación de 0.479**, indicando una correlación positiva moderada entre ellas.

c. Hipótesis específica 3

H₀: Los modelos predictivos de IA no se relacionarán con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

H₁: Los modelos predictivos de IA se relacionarán con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.

Tabla 42

Correlación entre los modelos predictivos de IA y la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024

			Actualización y Evolución
Rho de Spearman	Modelos Predictivos de IA	Coefficiente de correlación	,750**
		Sig. (bilateral)	,000
		N	79

**Esta correlación es significativa al nivel $p < 0.01$ (99% de confianza).

Interpretación: Se acepta la hipótesis específica 3 alterna, que postula Los modelos predictivos de IA se relacionarán con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024 ya que el coeficiente de correlación de Spearman (ρ), se encontró una significancia de 0.000, lo que sugiere una correlación significativa entre ambas. Además, se observa un coeficiente de correlación de 0.750**, indicando una correlación positiva alta entre ellas.

V. DISCUSIÓN DE RESULTADOS

El estudio de Chávez et al. (2023) subrayó las ventajas de la inteligencia artificial (IA) en la seguridad de la información, destacando su capacidad para la detección en tiempo real y la reducción de la carga laboral en los equipos de seguridad. No obstante, también señaló desafíos en su implementación, especialmente en lo relacionado con la privacidad de los datos. En contraste, en el presente estudio se evidenció una alta correlación positiva ($\rho = 0.732$, significancia = 0.000) entre los modelos predictivos de IA y la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima. Esta diferencia pudo deberse a las particularidades del sector financiero, donde las herramientas de IA fueron entrenadas con datos específicos que permitieron identificar patrones de ataque con mayor precisión. Mientras que la investigación previa adoptó un enfoque general sobre la seguridad cibernética en diversas industrias, el presente análisis demostró que una implementación focalizada en el sector bancario superó varias de las barreras previamente señaladas. En particular, la IA no solo optimizó la detección de amenazas en tiempo real, sino que también mejoró la capacidad de respuesta y mitigación frente a ataques sofisticados. Además, los resultados obtenidos reforzaron la idea de que el uso de IA no solo fortaleció la seguridad bancaria, sino que también mejoró la capacidad de adaptación frente a amenazas emergentes. Esto sugiere que una integración estratégica de la IA en el sector financiero transformó la ciberseguridad, maximizando la protección ante ataques persistentes avanzados.

Manrique (2023), en su análisis sobre la integración de Big Data con IA para la detección de ciberataques, subrayó que estas tecnologías permitieron gestionar grandes volúmenes de datos, facilitando la identificación de patrones de comportamiento malicioso mediante algoritmos de aprendizaje automático. Este enfoque guardó relación con el hallazgo de que los modelos predictivos de IA en el sistema bancario de Lima tuvieron una correlación positiva alta ($\rho = 0.429$, significancia = 0.000) con la detección de APTs. Sin embargo, mientras

que el análisis previo se basó en una revisión teórica de la literatura, el presente estudio analizó empíricamente la relación entre los modelos predictivos de IA y la detección de APTs en el sector bancario. Los resultados evidenciaron una asociación significativa, lo que respalda la importancia de su implementación en la ciberseguridad financiera. Asimismo, los hallazgos sugieren que la inteligencia artificial puede desempeñar un papel clave en la optimización de la detección de amenazas avanzadas, fortaleciendo la seguridad operativa en las entidades bancarias.

Mejía (2023), al analizar la influencia de la IA en la protección de datos personales en el contexto de la Ley N.º 29733, identificó retos significativos relacionados con la gestión de riesgos y la necesidad de un marco regulatorio sólido. Si bien estos hallazgos resonaron con la correlación positiva moderada ($\rho = 0.479$, significancia = 0.000) encontrada en este estudio entre los modelos predictivos de IA y la resiliencia del sistema bancario ante APTs, los resultados actuales permitieron evidenciar cómo estos modelos se asociaron con un fortalecimiento de la capacidad del sector financiero para resistir y recuperarse de ataques avanzados. Mientras que el estudio previo se enfocó en el marco normativo y los desafíos éticos, el presente trabajo analizó la relación entre los modelos predictivos de IA y la resiliencia del sistema bancario frente a amenazas persistentes. Los hallazgos sugieren que estos modelos pueden contribuir a la anticipación y mitigación de riesgos cibernéticos, complementando las recomendaciones regulatorias propuestas en la investigación de Mejía.

Villacorta et al. (2023) señalaron que la inteligencia artificial (IA) en la gestión de servicios de tecnología de la información (TI) mejoró la eficiencia, redujo costos y optimizó la toma de decisiones. No obstante, también destacaron desafíos importantes, como la resistencia al cambio organizacional y el riesgo de reemplazo de tareas humanas, afectando potencialmente el empleo. De manera similar, en el presente estudio se identificó una correlación positiva alta ($\rho = 0.750$, significancia = 0.000) entre los modelos predictivos de IA

y la actualización continua de las estrategias de protección contra APTs. A diferencia del estudio mencionado, que abordó el impacto de la IA en los servicios generales de TI, este análisis se centró en su aplicación en la evolución de las estrategias de ciberseguridad en el sector bancario. Los hallazgos sugieren que los modelos predictivos no solo optimizaron procesos, sino que también favorecieron una adaptación constante de las herramientas de seguridad en función de amenazas emergentes, evidenciando su relevancia en entornos de alto riesgo cibernético.

VI. CONCLUSIONES

- 6.1. Se concluye que los modelos predictivos de IA se relacionaron significativamente con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024, evidenciándose una alta correlación positiva ($\rho = 0.732$) y una significancia estadística de 0.000.
- 6.2. Se concluye que los modelos predictivos de IA se relacionaron significativamente con la detección y prevención de APTs en el sistema bancario de Lima, 2024, con una correlación positiva alta ($\rho = 0.429$) y una significancia estadística de 0.000.
- 6.3. Se concluye que los modelos predictivos de IA se relacionaron significativamente con la resiliencia del sistema bancario frente a APTs, evidenciándose una correlación positiva moderada ($\rho = 0.479$) y una significancia estadística de 0.000.
- 6.4. Se concluye que los modelos predictivos de IA se relacionaron significativamente con la actualización y evolución de la protección contra APTs en el sistema bancario de Lima, 2024, demostrando una alta correlación positiva ($\rho = 0.750$) y una significancia estadística de 0.000.

VII. RECOMENDACIONES

- 7.1. A las entidades bancarias: Implementar programas de capacitación técnica en el uso y optimización de modelos predictivos de IA, con el objetivo de maximizar su eficacia en la protección contra amenazas persistentes avanzadas (APTs). Esto contribuirá a fortalecer la seguridad en el sistema bancario de Lima.
- 7.2. A los equipos de seguridad cibernética: Diseñar protocolos específicos basados en los modelos predictivos de IA para mejorar la detección temprana y la prevención eficiente de contramedidas frente a APTs. Además, promover simulacros regulares que evalúen la capacidad del sistema para responder a estas amenazas.
- 7.3. A las áreas de infraestructura tecnológica de las instituciones bancarias: Incorporar mejoras constantes en la infraestructura tecnológica que respalden la resiliencia del sistema frente a APTs, asegurando la implementación de herramientas que permitan monitorear y mitigar riesgos de manera continua.
- 7.4. A los responsables de innovación tecnológica en el sector bancario: Establecer planes de actualización periódica de los modelos predictivos de IA, asegurando que estos se mantengan alineados con las tendencias globales de evolución tecnológica y los nuevos tipos de APTs que surjan en el futuro. Esto garantizará la efectividad continua de las estrategias de protección.

VIII. REFERENCIAS

- Andina. (1 de marzo de 2023). *Más de 15 mil millones de intentos de ciberataques en Perú durante 2022*. <https://andina.pe/agencia/noticia-mas-15-mil-millones-intentos-ciberataques-peru-durante-2022-930795.aspx>
- Arango, O. (2023). *El ABC de la seguridad informática: Guía práctica para entender la seguridad digital*. Editorial Instituto Tecnológico Metropolitano. <https://repositorio.itm.edu.co/handle/20.500.12622/5901>
- Arévalo, E. (2023). *Cyber Kill Chain*. Universidad Piloto de Colombia. <https://repository.unipiloto.edu.co/bitstream/handle/20.500.12277/13079/Cyber%20Kill%20Chain%20Ataque%20y%20Defensa.pdf?sequence=1&isAllowed=y>
- Baqués, J. (2021). *De las guerras híbridas a la zona gris: La metamorfosis de los conflictos en el siglo XXI*. Editorial UNED. https://www.google.com.pe/books/edition/De_las_guerras_h%C3%ADbridas_a_la_zona_gris/odM2EAAAQBAJ?hl=es-419&gbpv=0
- Banco Bilbao Vizcaya Argentaria BBVA. (28 de junio de 2024). *La IA, en los dos lados de la ciberseguridad: Aliada y amenaza en el mundo digital*. <https://www.bbva.com/es/innovacion/la-ia-en-los-dos-lados-de-la-ciberseguridad-aliada-y-amenaza-en-el-mundo-digital/>
- Cabezas, E., Andrade, D., y Torres, J. (2018). *Introducción a la metodología de la investigación científica*.
- Calvillo, P., Cerdá, L., Fonfría, C., Carreres, A., Muñoz, C., Trilles, L., y Martí, L. (2022). Elaboración de modelos predictivos de la gravedad y la mortalidad en pacientes con COVID-19 que acuden al servicio de urgencias, incluida la radiografía torácica. *Radiología*, 64(3), 214-227. <https://doi.org/10.1016/j.rx.2021.09.011>

- Camuñas, A. (2024). *Honeypot: Resiliencia y conocimiento del adversario* [Tesis, Universitat Oberta de Catalunya]. Repositorio Institucional UOC. <http://hdl.handle.net/10609/150610>
- Carrasco, S. (2019). *Metodología de la investigación científica: Pautas metodológicas para diseñar y elaborar el proyecto de investigación*. Editorial San Marcos.
- Cedeño, D. y Vargas, M. (2020). Aprendizaje automático aplicado al análisis de sentimientos. *Revistas Académicas UTP*, 16(2). <https://doi.org/10.33412/idt.v16.2.2833>
- Córdova, L., Gonzales, L., Leiva, F., y Navarro, L. (2024). *Aplicación de la inteligencia artificial a la evaluación crediticia en el sistema financiero peruano* [Tesis de maestría, Universidad ESAN]. Repositorio ESAN. <https://hdl.handle.net/20.500.12640/4249>
- Costa, S. (2023). *Inteligencia artificial en la economía: Impulsando la transformación digital*. Editor Santos Costa. https://www.google.com.pe/books/edition/Inteligencia_artificial_en_la_econom%C3%ADa/vX3cEAAQBAJ?hl=es-419&gbpv=0
- Cubeiro, E. (2021). Unidades de ciberinteligencia y ciberguerra al servicio de Estados. *bie3: Boletín IEEE*, 23, 432-458. <https://dialnet.unirioja.es/servlet/articulo?codigo=8175397>
- Dakalbab, F., Abu, M., Nasir, Q. y Saroufil, T. (2024). Artificial intelligence techniques in financial trading: A systematic literature review. *Journal of King Saud University - Computer and Information Sciences*, 36(3). <https://doi.org/10.1016/j.jksuci.2024.102015>
- Dalembert, F. (2024). *Detección de SPAM en correos electrónicos mediante el uso de sistema inteligente para la facultad de sistemas y telecomunicaciones de la UPSE* [Tesis, Universidad Estatal Península de Santa Elena]. Repositorio Universidad Estatal Península de Santa Elena. <https://repositorio.upse.edu.ec/handle/46000/10929>

- De la Hoz, B., Manjarres, A., De la Hoz, A. y Luna, I. (2023). Inteligencia artificial como estrategia para gestionar los procesos de auditoría financiera. *Revista Estrategia Organizacional*, 13(1), 52-72.
- Delgado, A. (2 de noviembre de 2023). *Caso Interbank: Perú entre los países más afectados por cibercriminalidad a nivel mundial*. Infobae. <https://www.infobae.com/peru/2024/11/02/caso-interbank-peru-entre-los-paises-mas-afectados-por-cibercriminalidad-a-nivel-mundial/>
- Espinoza, D. (2024). *La inteligencia artificial y la gestión de incidentes de soporte en una gerencia de sistemas del sector justicia, 2023* [Tesis de maestría, Universidad César Vallejo]. Repositorio de la Universidad César Vallejo. <https://hdl.handle.net/20.500.12692/150154>
- Espinoza, J. (2020). Aplicación de algoritmos Random Forest y XGBoost en una base de solicitudes de tarjetas de crédito. *Ingeniería, investigación y tecnología*, 21(3). <https://doi.org/10.22201/fi.25940732e.2020.21.3.022>
- European Parliament. (8 de septiembre de 2020). *¿Qué es la inteligencia artificial y cómo se usa?* <https://www.europarl.europa.eu/topics/es/article/20200827STO85804/que-es-la-inteligencia-artificial-y-como-se-usa>
- Fundación Telefónica. (2021). *El año en que todo cambió*. Editorial Fundación Telefónica. https://www.google.com.pe/books/edition/Sociedad_Digital_en_Espa%C3%B1a_2020_2021/i34xEAAAQBAJ?hl=es-419&gbpv=0
- Gante, S., Lozano, Y., Maldonado, M., Flores, G. y Villegas, D. (2022). Integración de una cámara multiespectral y aprendizaje automático para clasificación de manzanas. *Ingeniería, investigación y tecnología*, 23(4). <https://doi.org/10.22201/fi.25940732e.2022.23.4.031>

- Garzón, J., Sánchez, N. y Londoño, D. (2023). Evaluación comparativa de los algoritmos de aprendizaje automático Support Vector Machine y Random Forest: Efectos del tamaño del conjunto de entrenamiento. *Ciencia e Ingeniería Neogranadina*, 33(2). <https://doi.org/10.18359/rcin.6996>
- Gómez, J. (2022). *Cripto-ransomware: Análisis y detección temprana basada en el uso de archivos trampa* [Tesis, Universidad de Granada]. Repositorio Institucional UG. <http://hdl.handle.net/10481/76803>
- Hernández, A. y Sánchez, J. (2023). Metodología de selección de modelos de analítica de datos para la detección de APT con enfoque en la confiabilidad y optimización. *Memorias*, 2(3). <https://erevistas.uacj.mx/ojs/index.php/memoriascyt/article/view/6141>
- Hidalgo, V. (2020). *El creciente impacto de Machine Learning en la Ciberseguridad*. PWC. <https://www.pwc.com/ia/es/publicaciones/perspectivas-pwc/El-creciente-impacto-de-Machine-Learning.html#:~:text=Nos%20encontramos%20en%20una%20era,informaci%C3%B3n%20y%20se%20comprometen%20equipos.>
- Huyen, C. (2023). *Diseño de sistemas de Machine Learning: Un proceso iterativo para aplicaciones listas para funcionar*. Editorial Marcombo. https://www.google.com.pe/books/edition/Dise%C3%B1o_de_sistemas_de_Machine_Learning/itDJEAAAQBAJ?hl=es-419&gbpv=0
- Jaén, F. (2021). *Ciberseguridad industrial e infraestructuras críticas*. Ra-Ma Editorial. [https://books.google.es/books?id=GoNTEAAAQBAJ&dq=Resistencia+a+t%C3%A9cnicas+de+evasi%C3%B3n+empleadas+por+Protecci%C3%B3n+contra+amenazas+persistentes+avanzadas+\(APT\)&lr=&hl=es&source=gbs_navlinks_s](https://books.google.es/books?id=GoNTEAAAQBAJ&dq=Resistencia+a+t%C3%A9cnicas+de+evasi%C3%B3n+empleadas+por+Protecci%C3%B3n+contra+amenazas+persistentes+avanzadas+(APT)&lr=&hl=es&source=gbs_navlinks_s)

- Larriva, X., Plaza, A., Jover, O., Sanchez, C. y Villagra, V. (2023). Simulador de APTs realistas avanzados basado en el marco de MITRE ATT&CK. *Jornadas Nacionales de Investigación en Ciberseguridad*, 601. https://oa.upm.es/81620/1/JNIC2023_APT.pdf
- Liberos, E., Ahumada, S., y Sánchez, M. (2024). *Inteligencia artificial para el marketing: Cómo la tecnología revolucionará tu estrategia*. ESIC Editorial. https://www.google.com.pe/books/edition/INTELIGENCIA_ARTIFICIAL_PARA_EL_MARKETIN/PjvuEAAAQBAJ?hl=es-419&gbpv=0
- Lucena, P. (2024). *Cómo la inteligencia artificial está revolucionando la ciberseguridad en 2024*. Centro Europeo de Estudios y Formación Superior. CESUMA. https://www.cesuma.mx/blog/como-la-inteligencia-artificial-esta-revolucionando-la-ciberseguridad-en-2024.html?utm_source=chatgpt.com
- Madrid, J. (2024). *El impacto de la Inteligencia Artificial en la protección de datos personales y el acceso a la información (DT-EG 8/2024)*. Escuela de Gobierno, Universidad Complutense de Madrid. <https://www.ucm.es/eg/file/el-impacto-de-la-inteligencia-artificial-en-protecci%C3%93n-de-datos-personales-y-acceso-a-la-informaci%C3%93n?ver>
- Mata, A. (2022). *E-Book - Kali Linux para Hackers: Técnicas y metodologías avanzadas de seguridad informática ofensiva*. RA-MA S.A. Editorial y Publicaciones. https://www.google.com.pe/books/edition/Kali_Linux_para_Hackers/7VW6EAAAQBAJ?hl=es-419&gbpv=0
- Matos, E. (30 de octubre de 2024). *SBS iniciará acciones para determinar infracciones contra Interbank tras filtración de datos y caída de sistema*. La República. <https://larepublica.pe/economia/2024/10/30/sbs-iniciara-acciones-para-determinar-infracciones-contra-interbank-tras-filtracion-de-datos-y-caida-de-sistema>

- Maurtua, R. (2024). *La Inteligencia Artificial y modernización de la gestión administrativa en la Región Piura, 2024* [Tesis de maestría, Universidad César Vallejo]. Repositorio de la Universidad César Vallejo. <https://hdl.handle.net/20.500.12692/147731>
- Medium. (1 de mayo de 2024). *Arquitectura de datos: Patrones*. <https://medium.com/data-world-portafolio/arquitectura-de-datos-patrones-7bd538dfec6e>
- Melo, V. (2022). Inteligencia artificial, desinformación y protección de datos personales. *In Itinere. Revista Digital de Estudios Humanísticos*, 12(1). <https://repositorio.uca.edu.ar/handle/123456789/16405>
- Ministerio de Educación MINEDU (6 de noviembre de 2020). *¿Qué son las capacidades?* <https://sites.minedu.gob.pe/curriculonacional/2020/11/06/que-son-las-capacidades/>
- Miravé, M. (2023). *Predicción del valor de mercado de una vivienda en la ciudad de Barcelona mediante la obtención de un conjunto de datos y el desarrollo de un algoritmo de aprendizaje automático* [Tesis, Universitat Oberta de Catalunya]. Repositorio Universitat Oberta de Catalunya. <https://openaccess.uoc.edu/handle/10609/148363>
- Mirjalili, V. y Raschka, S. (2020). *Python machine learning: Aprendizaje automático y aprendizaje profundo con Python, scikit-learn y TensorFlow*. Editorial Marcombo. https://www.google.com.pe/books/edition/Python_Machine_Learning/5EtOEAAAQBAJ?hl=es-419&gbpv=0
- Molina, O. (2023). Inteligencia artificial, Big Data y derecho a la protección de datos de las personas trabajadoras. *Revista Española de Derecho del Trabajo y de la Seguridad Social*, 89-117. <https://revistas.uma.es/index.php/REJLSS/article/view/16225/16779>
- Montalvo, C. (2022). *Estándar de seguridad para la protección de datos de tarjetas de pago en las entidades financieras, Lima 2021* [Tesis de maestría, Universidad César Vallejo].

Repositorio de la Universidad César Vallejo.

<https://hdl.handle.net/20.500.12692/84060>

Montesdeoca, S. (2024). *Desarrollo de una API de predicción de fraude en transacciones financieras aplicando inteligencia artificial* [Tesis, Universidad Técnica del Norte].

Repositorio Digital Universidad Técnica del Norte.

<http://repositorio.utn.edu.ec/handle/123456789/15665>

Moreno, M. (2022). *Gestión de incidentes de ciberseguridad*. RA-MA S.A. Editorial y Publicaciones.

https://www.google.com.pe/books/edition/Gesti%C3%B3n_de_incidentes_de_ciberseguridad/_824EAAAQBAJ?hl=es-419&gbpv=0

Ñaupas, H., Valdivia, M., Palacios, J. y Romero, H. (2018). *Metodología de la investigación cuantitativa-cualitativa y redacción de la tesis*. Ediciones U.

Pardiñas, S. (2020). *Inteligencia Artificial: Un estudio de su impacto en la sociedad* [Trabajo de fin de grado, Universidade da Coruña]. <https://ruc.udc.es/dspace/handle/2183/28479>

Párraga, D. (2024). Modelos Predictivos de Rendimiento Académico Universitario Mediante Aprendizaje Automático. *Revista Científica de Salud y Desarrollo Humano*, 5(2), 974-991. <https://doi.org/10.61368/r.s.d.h.v5i2.204>

Pérez, A. (2024). *Prototipado de visión por computadora para la detección de objetos en drones autónomos* [Tesis], Universitat Politècnica de Catalunya]. Repositorio Institucional UPC. <https://upcommons.upc.edu/handle/2117/411614>

Quiroz, S., Zapata, J. y Vargas, H. (2020). Predicción de ciberataques en sistemas industriales SCADA a través de la implementación del filtro Kalman. *TecnoLógicas*, 23(48). <https://doi.org/10.22430/22565337.1586>

Real Academia Española- RAE (23 de mayo de 2024). *Sistema*. En Diccionario del estudiante.

<https://www.rae.es/diccionario-estudiante/sistema>. <https://dle.rae.es/m%C3%A9trico>

Rahman, Z., Yi, X. y Khalil, I. (2023). Blockchain-Based AI-Enabled Industry 4.0 CPS Protection Against Advanced Persistent Threat. *IEEE Internet of Things Journal*, 10(8), 6769-6778. <https://doi.org/10.1109/JIOT.2022.3147186>

Ramírez, L. (2024). Tecnologías de defensa frente a inteligencia de amenazas y ciberataques. *InnDev*, 3(1), 127-141. <https://doi.org/10.69583/inndev.v3n1.2024.94>

Rincón, J. y Vila, M. (2021). Modelo Predictivo Multivariable En Tiempo Real Para Predecir El Desempeño De Los Estudiantes, En Programas Virtuales De Posgrado, Empleando Inteligencia Artificial. *American Journal of Distance Education*, 35(4), 307-328. <https://doi.org/10.1080/08923647.2021.1954839>

Ropa, B. y Alama, M. (2022). Gestión organizacional: Un análisis teórico para la acción. *Revista Científica de la UCSA*, 9(1), 81-103. <https://doi.org/10.18004/ucsa/2409-8752/2022.009.01.081>

Sáenz, B. (2023). Inteligencia artificial y tecnología blockchain: Transparencia e información como pilares de la protección del consumidor. *Almedina*, 505-536. <https://investigacion.unirioja.es/documentos/654c6c3eaaaaaa46358ca73f?lang=ca>

Saldarriaga, M. (18 de noviembre de 2024). *Ciberataques en el Perú: ¿Qué tan preparados estamos?* https://caretas.pe/nacional/ciberataques-en-el-peru-que-tan-preparados-estamos/?utm_source=chatgpt.com

Sánchez, F. (2024). *Inteligencia artificial y sus implicancias en la transformación de los modelos de negocios convencionales de las Pymes en Lima, 2024* [Tesis de maestría, Universidad César Vallejo]. Repositorio de la Universidad César Vallejo. <https://hdl.handle.net/20.500.12692/149234>

- Sandua, D. (2024). *La Evolución de la Inteligencia Artificial de Asistentes Virtuales a Robots Autónomos*. Editor David Sandua. https://www.google.com.pe/books/edition/LA_EVOLUCI%C3%93N_DE_LA_INTELIGENCIA_ARTIFIC/OZcNEQAAQBAJ?hl=es-419&gbpv=0
- Sarker, I. (2022). AI-based modeling: Techniques, applications and research issues towards automation, intelligent and smart systems. *SN Computer Science*, 3, Article 158. <https://doi.org/10.1007/s42979-022-01043-x>
- Serra, D. (2023). Intervenciones experienciales en procesos de aprendizaje automático. *Essay Gráfica*, 1-9. <https://doi.org/10.5565/rev/grafica.325>
- Sierra, P. (2020). Capítulo 45 arquitectura del ciclo de vida de acciones de monitorización en remoto de equipos informáticos. <https://acmpublicaciones.revistabarataria.es/wp-content/uploads/2023/05/45-Gomez-Sierra-Arquitectura-del-ciclo-de-vida-informatica-2019-2023-pp494-499.pdf>
- Sifuentes, F. (2024). *Influencia de la innovación tecnológica en la satisfacción de los usuarios del Banco de la Nación de Nuevo Chimbote -- 2023* [Tesis de maestría, Universidad César Vallejo]. Repositorio de la Universidad César Vallejo. <https://hdl.handle.net/20.500.12692/137057>
- Soria, E., Sánchez, M., Gamero, R., Castillo, B. y Cano, P. (2023). *Sistemas de Aprendizaje Automático*. Ra-Ma S.A. Editorial y Publicaciones. https://www.google.com.pe/books/edition/Sistemas_de_Aprendizaje_Autom%C3%A1tico/2OPGEAAAQBAJ?hl=es-419&gbpv=0
- SYDLE. (19 de febrero de 2024). *Protección de datos: ¿Qué es y cómo funciona en los contextos empresarial y personal?* <https://www.sydle.com/es/blog/proteccion-de>

datos-que-es-y-como-funciona-en-los-contextos-empresarial-y-personal-
635c9031dd972c188d314148

Tanaka, M., Akiyama, Y., Mori, K., Hosaka, I., Kato, K., Endo, K., Ogawa, T., Sato, T., Suzuki, T., Yano, T., Ohnishi, H., Hanawa, N. y Furuhashi, M. (2024). Predictive modeling for the development of diabetes mellitus using key factors in various machine learning approaches. *Diabetes Epidemiology and Management*, 13(1).
<https://doi.org/10.1016/j.deman.2023.100191>

Universidad Internacional de La Rioja UNIR FP. (20 de junio de 2022). *Algoritmo: Qué es, para qué sirve, ejemplos de algoritmos y cómo funciona*. Revista UNIR FP.
<https://unirfp.unir.net/revista/ingenieria-y-tecnologia/que-es-algoritmo/>

Westreicher, G. (19 de febrero de 2024). *Método: Qué es, tipos y ejemplo*. Economipedia.
<https://economipedia.com/definiciones/metodo.html>

Westreicher, G. (22 de diciembre de 2020). *Dato*. Economipedia.
<https://economipedia.com/definiciones/dato>

Zapeta, A., Galindo, G., Juan, H. y Martínez, M. (2022). Métricas de rendimiento para evaluar el aprendizaje automático en la clasificación de imágenes petroleras utilizando redes neuronales convolucionales. *Ciencia Latina Revista Científica Multidisciplinar*, 6(5), 4624-4637. https://doi.org/10.37811/cl_rcm.v6i5.3420

IX. ANEXOS

Anexo A. Matriz de Consistencia

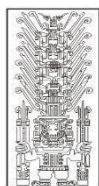
PROBLEMA	OBJETIVO	HIPÓTESIS	VARIABLES	DIMENSIONES	INDICADORES
Problema general ¿Los modelos predictivos de IA se relacionarán con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?	Objetivo general Determinar si los modelos predictivos de IA se relacionarán con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.	Hipótesis general Los modelos predictivos de IA se relacionan con la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024.	Modelos Predictivos de IA	Algoritmos de Aprendizaje Automático	- Tipos de algoritmos - Precisión de los modelos - Robustez frente a variaciones en los datos - Capacidad de generalización
Problemas específicos a. ¿Los modelos predictivos de IA se relacionarán con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?	Objetivos específicos a. Determinar si los modelos predictivos de IA se relacionarán con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	Hipótesis específicas a. Los modelos predictivos de IA se relacionarán con la detección y prevención de amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024		Entrenamiento y Validación	- Métodos de entrenamiento - Conjunto de datos utilizados para entrenamiento y prueba - Métricas de rendimiento - Tiempo de entrenamiento
b. ¿Los modelos predictivos de IA se relacionarán con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?	b. Determinar si los modelos predictivos de IA se relacionarán con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	b. Los modelos predictivos de IA se relacionarán con la resiliencia del sistema contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024		Implementación y Despliegue	- Tiempo de inferencia en entorno de producción - Escalabilidad del modelo - Facilidad de integración con sistemas existentes - Adaptabilidad a diferentes escenarios de ataque
c. ¿Los modelos predictivos de IA se relacionarán con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024?	c. Determinar si los modelos predictivos de IA se relacionarán con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	c. Los modelos predictivos de IA se relacionarán con la actualización y evolución de la protección contra amenazas persistentes avanzadas (APTs) en el sistema bancario de Lima, 2024	Protección contra APT	Detección y Prevención	- Tasa de detección de ataques - Falsos positivos y falsos negativos - Tiempo de detección - Tiempo de respuesta
				Resiliencia del Sistema	- Capacidad de recuperación post-ataque - Resistencia a técnicas de evasión empleadas por APT - Efectividad de los parches de seguridad
				• Actualización y Evolución	- Frecuencia de actualizaciones de modelos. - Eficiencia de las contramedidas - Adaptación a nuevas amenazas - Evolución de patrones de ataque identificados

Anexo B. Instrumento de recolección de datos

ITEMS		Nunca	Casi nunca	A veces	Casi siempre	Siempre
Variable Independiente Modelos Predictivos de IA						
Algoritmos de Aprendizaje Automático						
1	¿Con qué frecuencia los algoritmos de aprendizaje automático utilizados en su institución abordan eficazmente los problemas de seguridad o amenazas persistentes?					
2	¿Con qué regularidad los resultados de los sistemas del banco considera usted que son confiables?					
3	En su experiencia, ¿los modelos predictivos mantienen su eficacia cuando se enfrentan a amenazas que afectan los datos?					
4	¿Los algoritmos de aprendizaje automático pueden generalizar lo suficiente como para evitar que nuevas amenazas comprometan la seguridad bancaria?					
Entrenamiento y Validación						
5	¿Se actualiza regularmente los métodos de entrenamiento de los modelos predictivos de su institución para mejorar la detección de amenazas?					
6	¿El conjunto de datos utilizados para entrenar y probar los modelos en su organización refleja las amenazas reales para una protección efectiva?					
7	¿Con qué frecuencia su equipo revisa las métricas de rendimiento para evaluar la efectividad de los modelos predictivos?					
8	¿Considera que el tiempo invertido en el entrenamiento de los modelos es razonable y productivo?					
Implementación y Despliegue						
9	¿Qué tan a menudo considera que el tiempo de inferencia de los modelos en producción es adecuado para detectar y responder a amenazas de manera efectiva?					
10	¿En qué medida cree que los modelos predictivos pueden escalar efectivamente ante un aumento en la carga de trabajo?					

11	¿Los modelos predictivos se integran fácilmente con los sistemas existentes para fortalecer la protección contra amenazas?					
12	En su opinión, ¿con qué frecuencia los modelos predictivos logran adaptarse a distintos escenarios de ataque sin problemas?					
Variable dependiente Protección contra APT						
Detección y Prevención						
13	¿Qué tan a menudo considera que su sistema detecta ataques de manera efectiva en comparación con otras amenazas?					
14	¿Con qué frecuencia observa que su sistema de detección genera falsos positivos o negativos que afectan la eficacia de la protección?					
15	¿Cómo calificaría la rapidez con la que su sistema detecta ataques de amenazas persistentes avanzadas?					
16	¿Con qué frecuencia su equipo responde de manera oportuna a los ataques detectados por el sistema?					
Resiliencia del Sistema						
17	¿Qué tan efectiva considera que es la capacidad de su sistema para recuperarse después de un ataque?					
18	¿Con qué frecuencia su sistema se mantiene resistente frente a técnicas de evasión utilizadas por atacantes avanzados?					
19	¿Cómo calificaría la efectividad de los parches de seguridad aplicados en su sistema para proteger contra amenazas persistentes avanzadas?					
Actualización y Evolución						
20	¿Con qué regularidad se actualizan los modelos de seguridad para mantenerse al día con nuevas amenazas?					
21	¿Qué tan eficaces considera que son las contramedidas implementadas para proteger contra las amenazas más recientes?					
22	¿Con qué frecuencia su sistema se adapta rápidamente a nuevas amenazas identificadas en el entorno de seguridad?					
23	¿Con qué regularidad su sistema detecta y se ajusta a la evolución de los patrones de ataque?					

Anexo C: Ficha de validación de instrumento por juicio de expertos



UNIVERSIDAD NACIONAL FEDERICO VILLARREAL ESCUELA UNIVERSITARIA DE POSGRADO

I. DATOS GENERALES

1.1. Apellidos y Nombres: Tejada Estrada, Gina Coral

1.2. Grado académico: Doctor

1.3. Cargo e Institución donde labora: Docente de EUPG-UNFV

1.4. Nombre del instrumento motivo de evaluación: Encuesta

1.5. Título de la Investigación: LOS MODELOS PREDICTIVOS DE IA Y SU RELACIÓN CON LA PROTECCIÓN CONTRA AMENAZAS PERSISTENTES AVANZADAS (APTS) EN EL SISTEMA BANCARIO DE LIMA, 2024

1.6. Autor(a) del Instrumento: Espinoza Villavicencio, Héctor Martín

II. ASPECTOS DE VALIDACIÓN

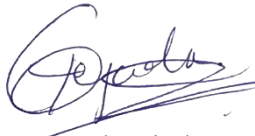
Indicadores	Criterios	Deficiente 0-20%	Baja 21-40%	Regular 41-60%	Buena 61%-80%	Muy buena 81%-100%
1. Claridad	Está formulado con lenguaje apropiado.					95%
2. Objetividad	Está expresado en conductas observables					95%
3. Actualidad	Adecuado al avance de la especialidad					95%
4. Organización	Existe una organización lógica					95%
5. Suficiencia	Comprende los aspectos en cantidad y calidad.					95%
6. Intencionalidad	Adecuado para valorar la investigación					95%
7. Consistencia	Basado en aspectos teóricos científicos.					95%
8. Coherencia	Entre lo descrito en dimensiones e indicadores					95%
9. Metodología	La formulación responde a la investigación					95%
10. Pertinencia	Es útil y adecuado para la investigación					95%

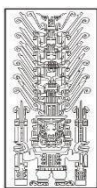
III. PROMEDIO DE VALORACIÓN: 95%

a) Deficiente ☐ b) Baja ☐ c) Regular ☐ d) Buena ☐ e) Muy Buena ☒

IV. OPINIÓN DE APLICABILIDAD: El Instrumento es aplicable en la investigación.

Lima, Febrero 2025


Dra. Gina Coral Tejada Estrada
(ORCID: 0000-0002-0023-5147)



UNIVERSIDAD NACIONAL FEDERICO VILLARREAL
ESCUELA UNIVERSITARIA DE POSGRADO

Ficha de Validación

(Juicio de Experto)

I. DATOS GENERALES

I.1. Apellidos y Nombres: Collazos Páucar, Edwin

I.2. Grado académico: Doctor

I.3. Cargo e Institución donde labora: Docente de EUPG-UNFV

I.4. Nombre del instrumento motivo de evaluación: Cuestionario

I.5. Título de la Investigación: LOS MODELOS PREDICTIVOS DE IA Y SU RELACIÓN CON LA PROTECCIÓN CONTRA AMENAZAS PERSISTENTES AVANZADAS (APTS) EN EL SISTEMA BANCARIO DE LIMA, 2024

I.6. Autor(a) del Instrumento: Espinoza Villavicencio, Héctor Martin

II. ASPECTOS DE VALIDACIÓN

Criterios	Indicadores	Deficiente 0-20%	Baja 21-50%	Regular 51-70%	Buena 71%-90%	Muy buena 91%-100%
1. Claridad	Está formulado con lenguaje apropiado.				90%	
2. Objetividad	Está expresado en conductas observables				90%	
3. Actualidad	Adecuado al avance de la especialidad				90%	
4. Organización	Existe una organización lógica				90%	
5. Suficiencia	Comprende los aspectos en cantidad y calidad.				90%	
6. Intencionalidad	Adecuado para valorar la investigación				90%	
7. Consistencia	Basado en aspectos teóricos científicos.				90%	
8. Coherencia	Entre lo descrito en dimensiones e indicadores				90%	
9. Metodología	La formulación responde a la investigación				90%	
10. Pertinencia	Es útil y adecuado para la investigación				90%	

III. PROMEDIO DE VALORACIÓN: 90%

a) Deficiente

☐

b) Baja

☐

c) Regular

☐

d) Buena

☒

e) Muy Buena

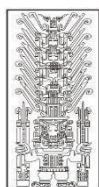
☐

IV. OPINIÓN DE APLICABILIDAD:

El Instrumento es aplicable en la investigación

Lima, Febrero 2025

Dr. Edwin Collazos Paucar
(ORCID: 0000-0002-3528-1354)



**UNIVERSIDAD NACIONAL FEDERICO VILLARREAL
ESCUELA UNIVERSITARIA DE POSGRADO**

**Ficha de Validación
(Juicio de Experto)**

I. DATOS GENERALES

1.1. Apellidos y Nombres: Soto Soto, Luis

1.2. Grado académico: Doctor en Ingeniería

1.3. Cargo e Institución donde labora: Docente UNFV y UNMSM

1.4. Nombre del instrumento motivo de evaluación: Cuestionario

1.5. Título de la Investigación: LOS MODELOS PREDICTIVOS DE IA Y SU RELACIÓN CON LA PROTECCIÓN CONTRA AMENAZAS PERSISTENTES AVANZADAS (APTS) EN EL SISTEMA BANCARIO DE LIMA, 2024

1.6. Autor(a) del Instrumento: Espinoza Villavicencio, Héctor Martín

II. ASPECTOS DE VALIDACIÓN

Criterios	Indicadores	Deficiente 0-20%	Baja 21-50%	Regular 51-70%	Buena 71%-90%	Muy buena 91%-100%
11. Claridad	Está formulado con lenguaje apropiado.				90%	
12. Objetividad	Está expresado en conductas observables				90%	
13. Actualidad	Adecuado al avance de la especialidad				90%	
14. Organización	Existe una organización lógica				90%	
15. Suficiencia	Comprende los aspectos en cantidad y calidad.				90%	
16. Intencionalidad	Adecuado para valorar la investigación				90%	
17. Consistencia	Basado en aspectos teóricos científicos.				90%	
18. Coherencia	Entre lo descrito en dimensiones e indicadores				90%	
19. Metodología	La formulación responde a la investigación				90%	
20. Pertinencia	Es útil y adecuado para la investigación				90%	

III. PROMEDIO DE VALORACIÓN: 90%

a) Deficiente

☐

b) Baja

☐

c) Regular

☐

d) Buena

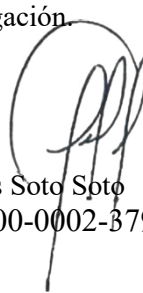
☒

e) Muy Buena

☐

V. OPINIÓN DE APLICABILIDAD: El instrumento es aplicable en la investigación.

Lima, Febrero 2025


 Dr. Luis Soto Soto
 (ORCID: 0000-0002-3799-645X)