



MÁSTER UNIVERSITARIO EN INGENIERÍA DE TELECOMUNICACIONES

TRABAJO FIN DE MÁSTER

Herramienta de Threat Intelligence para Sistemas de
Control Industrial (SCI)

Autor: Claudia Blanco García

Director: Juan Carlos Cortinas Abad

Madrid

Declaro, bajo mi responsabilidad, que el Proyecto presentado con el título
Herramienta de Threat Intelligence para Sistemas de Control Industrial (SCI)
en la ETS de Ingeniería - ICAI de la Universidad Pontificia Comillas en el
curso académico 2024/25 es de mi autoría, original e inédito y
no ha sido presentado con anterioridad a otros efectos.

El Proyecto no es plagio de otro, ni total ni parcialmente y la información que ha sido
tomada de otros documentos está debidamente referenciada.



Fdo.: Claudia Blanco García

Fecha: 14 / 06 / 2025

Autorizada la entrega del proyecto

EL DIRECTOR DEL PROYECTO



Fdo.: Juan Carlos Cortinas Abad

Fecha: 14 / 06 / 2025

AUTORIZACIÓN PARA LA DIGITALIZACIÓN, DEPÓSITO Y DIVULGACIÓN EN RED DE PROYECTOS FIN DE GRADO, FIN DE MÁSTER, TESIS O MEMORIAS DE BACHILLERATO

1º. Declaración de la autoría y acreditación de la misma.

El autor D. Claudia Blanco García DECLARA ser el titular de los derechos de propiedad intelectual de la obra: Herramienta de Threat Intelligence para Sistemas de Control Industrial, que ésta es una obra original, y que ostenta la condición de autor en el sentido que otorga la Ley de Propiedad Intelectual.

2º. Objeto y fines de la cesión.

Con el fin de dar la máxima difusión a la obra citada a través del Repositorio institucional de la Universidad, el autor **CEDE** a la Universidad Pontificia Comillas, de forma gratuita y no exclusiva, por el máximo plazo legal y con ámbito universal, los derechos de digitalización, de archivo, de reproducción, de distribución y de comunicación pública, incluido el derecho de puesta a disposición electrónica, tal y como se describen en la Ley de Propiedad Intelectual. El derecho de transformación se cede a los únicos efectos de lo dispuesto en la letra a) del apartado siguiente.

3º. Condiciones de la cesión y acceso

Sin perjuicio de la titularidad de la obra, que sigue correspondiendo a su autor, la cesión de derechos contemplada en esta licencia habilita para:

- a) Transformarla con el fin de adaptarla a cualquier tecnología que permita incorporarla a internet y hacerla accesible; incorporar metadatos para realizar el registro de la obra e incorporar “marcas de agua” o cualquier otro sistema de seguridad o de protección.
- b) Reproducirla en un soporte digital para su incorporación a una base de datos electrónica, incluyendo el derecho de reproducir y almacenar la obra en servidores, a los efectos de garantizar su seguridad, conservación y preservar el formato.
- c) Comunicarla, por defecto, a través de un archivo institucional abierto, accesible de modo libre y gratuito a través de internet.
- d) Cualquier otra forma de acceso (restringido, embargado, cerrado) deberá solicitarse expresamente y obedecer a causas justificadas.
- e) Asignar por defecto a estos trabajos una licencia Creative Commons.
- f) Asignar por defecto a estos trabajos un HANDLE (URL *persistente*).

4º. Derechos del autor.

El autor, en tanto que titular de una obra tiene derecho a:

- a) Que la Universidad identifique claramente su nombre como autor de la misma
- b) Comunicar y dar publicidad a la obra en la versión que ceda y en otras posteriores a través de cualquier medio.
- c) Solicitar la retirada de la obra del repositorio por causa justificada.
- d) Recibir notificación fehaciente de cualquier reclamación que puedan formular terceras personas en relación con la obra y, en particular, de reclamaciones relativas a los derechos de propiedad intelectual sobre ella.

5º. Deberes del autor.

El autor se compromete a:

- a) Garantizar que el compromiso que adquiere mediante el presente escrito no infringe ningún derecho de terceros, ya sean de propiedad industrial, intelectual o cualquier otro.
- b) Garantizar que el contenido de las obras no atenta contra los derechos al honor, a la intimidad y a la imagen de terceros.
- c) Asumir toda reclamación o responsabilidad, incluyendo las indemnizaciones por daños, que pudieran ejercitarse contra la Universidad por terceros que vieran infringidos sus derechos e intereses a causa de la cesión.
- d) Asumir la responsabilidad en el caso de que las instituciones fueran condenadas por infracción de derechos derivada de las obras objeto de la cesión.

6º. Fines y funcionamiento del Repositorio Institucional.

La obra se pondrá a disposición de los usuarios para que hagan de ella un uso justo y respetuoso con los derechos del autor, según lo permitido por la legislación aplicable, y con fines de estudio, investigación, o cualquier otro fin lícito. Con dicha finalidad, la Universidad asume los siguientes deberes y se reserva las siguientes facultades:

- La Universidad informará a los usuarios del archivo sobre los usos permitidos, y no garantiza ni asume responsabilidad alguna por otras formas en que los usuarios hagan un uso posterior de las obras no conforme con la legislación vigente. El uso posterior, más allá de la copia privada, requerirá que se cite la fuente y se reconozca la autoría, que no se obtenga beneficio comercial, y que no se realicen obras derivadas.
- La Universidad no revisará el contenido de las obras, que en todo caso permanecerá bajo la responsabilidad exclusiva del autor y no estará obligada a ejercitar acciones legales en nombre del autor en el supuesto de infracciones a derechos de propiedad intelectual derivados del depósito y archivo de las obras. El autor renuncia a cualquier reclamación frente a la Universidad por las formas no ajustadas a la legislación vigente en que los usuarios hagan uso de las obras.
- La Universidad adoptará las medidas necesarias para la preservación de la obra en un futuro.
- La Universidad se reserva la facultad de retirar la obra, previa notificación al autor, en supuestos suficientemente justificados, o en caso de reclamaciones de terceros.

Madrid, a 10 de junio de 2025

ACEPTA



Fdo. Claudia Blanco García

Motivos para solicitar el acceso restringido, cerrado o embargado del trabajo en el Repositorio Institucional:



MÁSTER UNIVERSITARIO EN INGENIERÍA DE TELECOMUNICACIONES

TRABAJO FIN DE MASTER

Herramienta de Threat Intelligence para Sistemas de
Control Industrial (SCI)

Autor: Claudia Blanco García

Director: Juan Carlos Cortinas Abad

Madrid

Agradecimientos

Quiero agradecer al departamento de Ciberseguridad de IESE (Iberdrola Energía Sostenible España) por hacerme sentir parte del equipo desde el primer día y por todo lo que me han enseñado; a Ana por incluirme siempre en cada paso del proyecto, a Darío y Julio por compartir conmigo sus conocimientos sobre ciberseguridad en entornos OT, a Fernando por fomentar mi participación, y a Juan Carlos por su guía y acompañamiento en la realización de este trabajo fin de máster.

THREAT INTELLIGENCE PARA SISTEMAS DE CONTROL INDUSTRIAL (SCI)

Autor: Blanco García, Claudia.

Director: Cortinas Abad, Juan Carlos.

Entidad Colaboradora: Iberdrola

RESUMEN DEL PROYECTO

Se ha desarrollado una herramienta de ciberinteligencia orientada al sector energético, integrando fuentes de datos en tiempo real, análisis contextual y modelos de machine learning para la predicción del riesgo o impacto potencial. El sistema evalúa ciberataques de forma automatizada e incluye una API para comunicación con el SOC, y una interfaz gráfica.

Palabras clave: SCI (Sistema de Control Industrial), OT (*Operational Technology*), CTI (*Cyber Threat Intelligence*), ML (*Machine Learning*), LLM (*Large Language Model*), SOC (*Security Operations Center*), API (*Application Programmable Interface*), OSINT (*Open Source Intelligence*), CVE (*Common Vulnerability Exposure*), web scraping.

1. Introducción

En las últimas décadas, los sistemas de control industrial (SCI) se han convertido en objetivos de ciberataques debido a su papel crítico en sectores estratégicos como el energético. A medida que estos ataques se vuelven más sofisticados y con mayor carga geopolítica, se hace imprescindible contar con herramientas capaces de contextualizar, priorizar y anticipar amenazas en tiempo real.

Este proyecto nace con la intención de desarrollar una solución de ciberinteligencia adaptada a las particularidades de los entornos OT, donde los activos industriales requieren una protección especializada frente a amenazas de ciberseguridad complejas.

2. Definición del proyecto

El objetivo principal de este trabajo es diseñar y construir una herramienta llamada O.C.A.S. (*OT Cyberrisk Assessment System*) que permita evaluar el riesgo de ciberataques en entornos industriales del sector energético. El sistema combina datos técnicos sobre amenazas en tiempo real (como IoCs, CVEs, actores implicados) con información contextual (como el estado geopolítico o las técnicas utilizadas por los atacantes).

Para ello, se ha planteado un flujo dual que extrae información tanto de noticias generales como de vulnerabilidades concretas, procesando los datos mediante scraping, normalización con modelos LLM, enriquecimiento técnico y contextual, y finalmente evaluación del riesgo con un modelo de machine learning entrenado sobre un histórico de más de 14.000 ciberataques. Esto se combina con la extracción de inteligencia de ciberseguridad OT a partir de fuentes OSINT.

3. Descripción del sistema

La herramienta se compone de varios módulos interconectados. El modelo de ML se entrena con un conjunto de datos procedentes de distintas fuentes (ciberataques

históricos, contexto geopolítico e información de adversario), con el impacto potencial o “riesgo pasado” como variable objetivo. Otro módulo se encarga de la ingesta en tiempo real de noticias de ciberataques generales y de otros asociados a vulnerabilidades o CVEs. Estas se procesan mediante web scraping y se normalizan mediante un LLM. La información se enriquece con contexto geopolítico, con información de técnicas y adversarios basada en MITRE, y con indicadores de compromiso (IoCs) usando diferentes fuentes como GreyNoise o Shodan. Toda esta información se pasa al modelo de regresión entrenado (con un R^2 de 0.9111) que predice el riesgo del ataque. Los resultados se publican en una API que se puede consultar fácilmente, y se almacenan en una base de datos de Elastic para su visualización mediante un dashboard de Kibana.

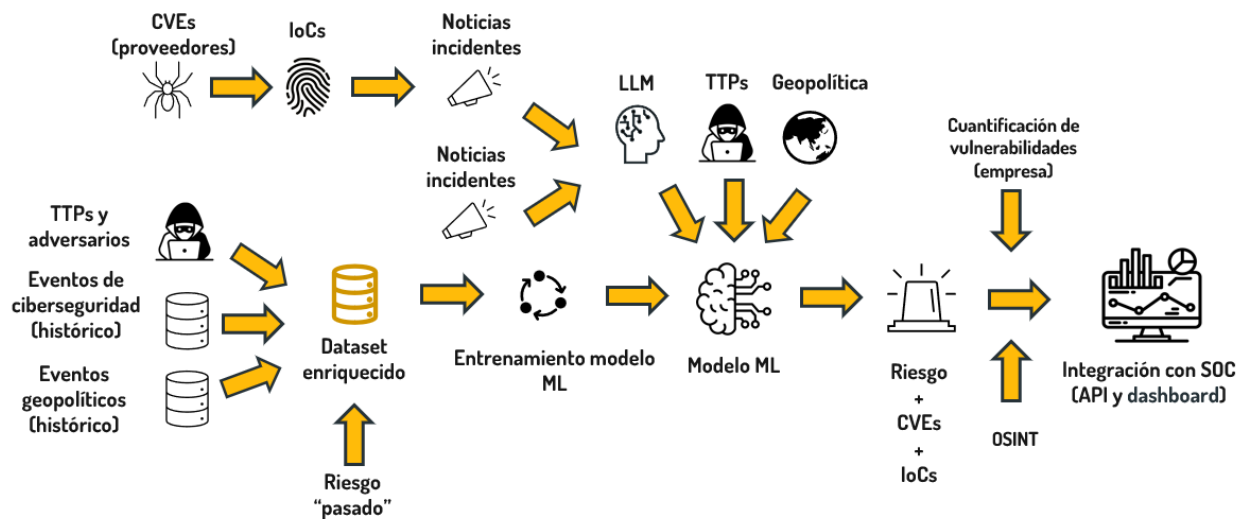


Ilustración 1. Esquema del sistema de O.C.A.S.

4. Resultados

- Se logró desarrollar un pipeline funcional en dos líneas paralelas de ingesta y análisis, uno centrado en noticias de ciberataques y otro en ciberinteligencia general.
- El modelo de ML mostró un rendimiento elevado y permitió detectar patrones complejos de riesgo.
- La herramienta demostró su utilidad operativa al poder integrarse con un SOC mediante una API desplegada en Render, aunque con ciertas limitaciones de velocidad atribuibles a la capacidad de procesamiento del entorno de despliegue en Render.
- Además, se ha integrado en la herramienta el sistema de cuantificación de vulnerabilidades propio de Iberdrola IESE.

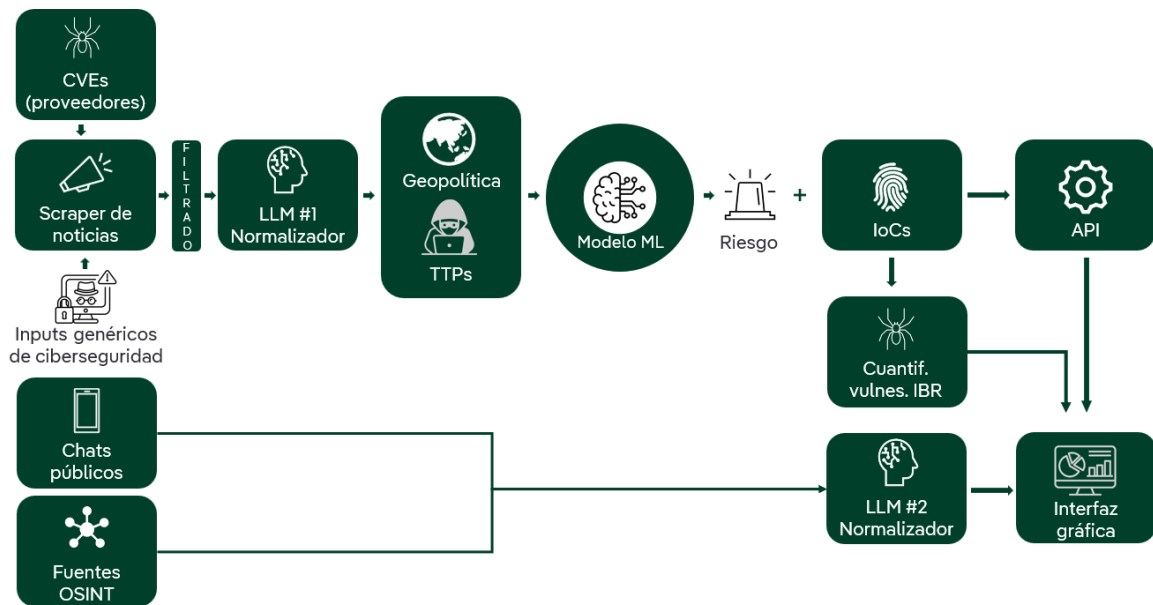


Ilustración 2. Pipeline final de la herramienta O.C.A.S.

5. Conclusiones

El sistema O.C.A.S. se ha mostrado como una herramienta útil, modular y adaptable para evaluar amenazas de ciberseguridad dirigidas a infraestructuras energéticas con componentes OT. Además de cumplir con los cuatro objetivos operativos inicialmente propuestos (entrenamiento del modelo de ML, desarrollo de la herramienta de ingesta y análisis, integración con el SOC y visualización en un dashboard), alcanza dos objetivos adicionales: la integración de la cuantificación de CVEs de la empresa y el diseño de una arquitectura modular. Esta última permite incorporar nuevas fuentes de datos y adaptar el sistema a futuro. La plataforma ofrece una ventaja operativa clara al combinar inteligencia técnica y contextual, y su estructura modular permite escalarla y adaptarla fácilmente a nuevos escenarios y sectores industriales.

6. Referencias

- [1] MITRE, «ATT&CK,» MITRE, 2025. [En línea]. Available: <https://attack.mitre.org/>
- [2] K. Knerler, I. Parker y C. Zimmerman, «11 Strategies of a World-Class Cybersecurity Operations Center,» MITRE, 2022.

THREAT INTELLIGENCE TOOL FOR INDUSTRIAL CONTROL SYSTEMS (ICS)

Author: Blanco García, Claudia.

Supervisor: Cortinas Abad, Juan Carlos.

Collaborating Entity: Iberdrola.

ABSTRACT

A cyber threat intelligence tool has been developed, specifically designed for the energy sector, integrating real-time data sources, contextual analysis, and machine learning models to predict risk or potential impact. The system automatically evaluates cyberattacks and includes an API for communication with the SOC, as well as a graphical user interface.

Keywords: ICS (Industrial Control Systems), OT (Operational Technology), CTI (Cyber Threat Intelligence), ML (Machine Learning), LLM (Large Language Model), SOC (Security Operations Center), API (Application Programmable Interface), OSINT (Open Source Intelligence), CVE (Common Vulnerability Exposure), web scraping.

1. Introduction

In recent decades, industrial control systems (ICS) have become frequent targets of cyberattacks due to their critical role in strategic sectors such as energy. As these attacks grow more sophisticated and geopolitically charged, the need for tools capable of contextualizing, prioritizing, and anticipating threats in real time becomes increasingly urgent.

This project aims to develop a cyber intelligence solution tailored to the unique characteristics of OT environments, where industrial assets require specialized protection against complex cybersecurity threats.

2. Project Definition

The main objective of this work is to design and build a tool called O.C.A.S. (OT Cyberrisk Assessment System) that can assess the risk of cyberattacks in industrial environments within the energy sector. The system combines technical data on real-time threats (such as IoCs, CVEs, and threat actors) with contextual information (such as geopolitical conditions or attacker techniques).

A dual pipeline has been designed to extract information from both general news and specific vulnerabilities, processing the data through scraping, normalization using LLMs, technical and contextual enrichment, and finally risk evaluation using a machine learning model trained on a historical dataset of over 14,000 cyberattacks. This is complemented by OSINT-based OT cyber threat intelligence extraction.

3. System Description

The tool is composed of several interconnected modules. The ML model is trained on a dataset from various sources (historical cyberattacks, geopolitical context, adversary information), with the potential impact or “past risk” as the target variable. Another module handles real-time ingestion of general cyberattack news and events linked to

vulnerabilities or CVEs. These are processed through web scraping and normalized via an LLM. The data is enriched with geopolitical context, MITRE-based technique and adversary information, and IoCs from sources like GreyNoise and Shodan. All this information is passed to a trained regression model (with an R^2 of 0.9111) that predicts the attack risk. The results are published through an easily accessible API and stored in an Elastic database for visualization via a Kibana dashboard.

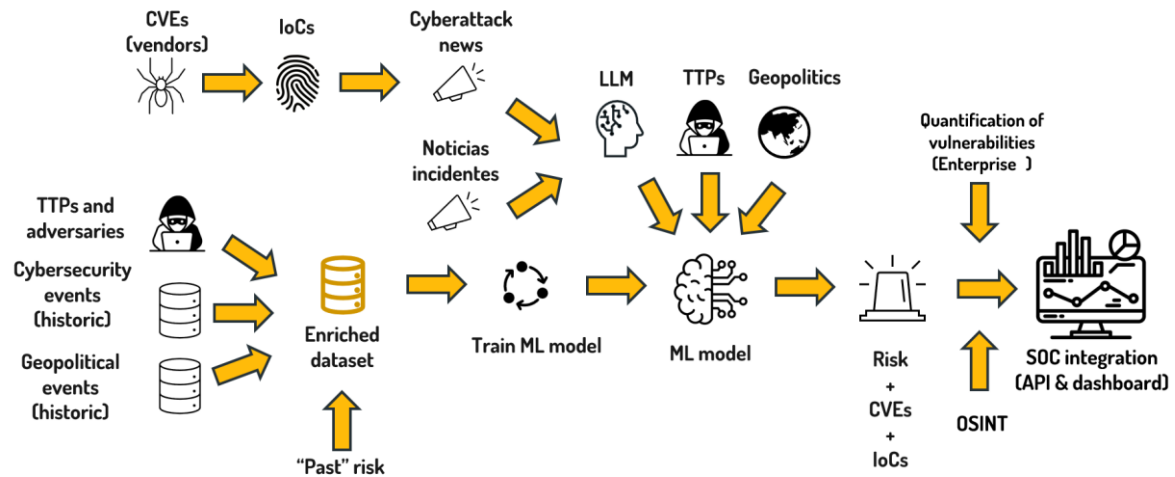


Ilustración 3. Diagram of the O.C.A.S. tool

4. Results

- A functional pipeline was successfully developed, with two parallel lines of ingestion and analysis—one focused on cyberattack news and the other on general cyber intelligence.
- The ML model showed strong performance and enabled the detection of complex risk patterns.
- The tool demonstrated operational value by integrating with a SOC through an API deployed on Render, although some performance limitations were observed due to the processing capacity of the Render environment.
- Additionally, Iberdrola IESE's proprietary vulnerability quantification system was integrated into the tool.

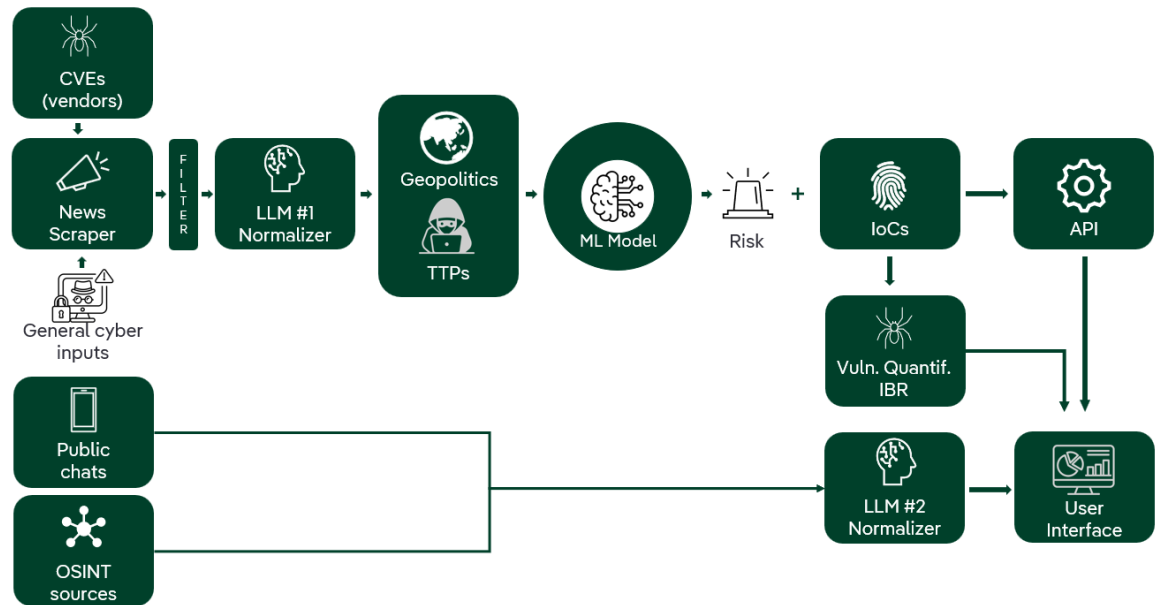


Ilustración 4. Final pipeline of the O.C.A.S. tool

5. Conclusions

The O.C.A.S. system has proven to be a useful, modular, and adaptable tool for assessing cybersecurity threats targeting energy infrastructures with OT components. In addition to meeting the four initially proposed operational objectives—training the ML model, developing the ingestion and analysis tool, integrating with the SOC, and visualizing data through a dashboard—it also achieves two additional goals: the integration of company-specific CVE quantification and the design of a modular architecture. The latter enables the incorporation of new data sources and enables future system adaptation. The platform provides a clear operational advantage by combining technical and contextual intelligence, and its modular structure allows for easy scaling and adaptation to new scenarios and industrial sectors.

6. References

- [1] MITRE, «ATT&CK,» MITRE, 2025. [En línea]. Available: <https://attack.mitre.org/>
- [2] K. Knerler, I. Parker y C. Zimmerman, «11 Strategies of a World-Class Cybersecurity Operations Center,» MITRE, 2022.

Índice de la memoria

Capítulo 1. Introducción	8
1.1 Motivación del proyecto.....	8
1.2 Solución propuesta	10
Capítulo 2. Descripción de las Tecnologías.....	11
2.1 Python.....	11
2.1.1 Jupyter Notebook.....	11
2.1.2 Visual Studio Code	11
2.1.3 Pandas	11
2.1.4 MITRE ATT&CK.....	12
2.1.5 Newspaper3k	12
2.1.6 Keras	13
2.1.7 FastAPI.....	13
2.1.8 Telethon.....	13
2.2 Herramienta de LLM: OpenRouter	13
2.3 Herramientas de Threat Intelligence	13
2.3.1 GreyNoise.....	13
2.3.2 SHODAN	14
2.4 Herramientas de visualización.....	15
2.4.1 Kibana y Elasticsearch.....	15
2.5 Otras herramientas.....	15
2.5.1 Render	15
Capítulo 3. Estado de la Cuestión	17
3.1 CTI (Cyber Threat Intelligence).....	17
3.1.1 Fuentes de CTI	17
3.1.2 Recogida de CTI.....	19
3.1.3 STIX.....	20
3.2 Soluciones tecnológicas y trabajos de investigación.....	22
3.2.1 Soluciones tecnológicas existentes	22
3.2.2 Trabajos de investigación relacionados.....	24

Capítulo 4. Definición del Trabajo	26
4.1 Justificación.....	26
4.1.1 Necesidad estratégica.....	26
4.1.2 Proyecto modular	26
4.2 Objetivos	27
4.3 Metodología.....	28
4.4 Planificación y Estimación Económica	29
4.4.1 Planificación.....	29
4.4.2 Estimación económica.....	32
Capítulo 5. El Sistema/Modelo Desarrollado	33
5.1 Modelo de Machine Learning para predicción del riesgo	34
5.1.1 Definición del riesgo	34
5.1.2 Pruebas de calibración.....	45
5.1.3 Especificación del modelo y entrenamiento	66
5.1.4 Pruebas con ciberataques recientes	71
5.2 Ingesta de noticias de ciberataques en tiempo real.....	71
5.2.1 Búsqueda automatizada de noticias	72
5.2.2 Scraping y filtrado.....	73
5.2.3 Normalización con LLM.....	73
5.3 Ingesta de información de contexto en tiempo real.....	75
5.3.1 Información de contexto geopolítico	75
5.3.2 Información de adversarios y TTPs.....	76
5.4 Aplicación del modelo entrenado.....	76
5.5 Enriquecimiento con indicadores de compromiso.....	77
5.5.1 IPs maliciosas: GreyNoise	78
5.5.2 Análisis en profundidad de IPs maliciosas: SHODAN.....	79
5.6 Ingesta de datos de inteligencia OT.....	79
5.6.1 Telegram.....	80
5.6.2 Feeds de OSINT.....	81
5.6.3 Normalización con LLM.....	81
5.7 Cuantificación de vulnerabilidades: modelo empresarial.....	82
5.8 Diseño de la comunicación con el SOC	84
5.9 Interfaz gráfica	85

Capítulo 6. Análisis de resultados.....	88
6.1 Resultados del modelo de Machine Learning	89
6.2 Resultados de la integración de fuentes en tiempo real.....	90
6.3 Resultados de la API	92
Capítulo 7. Conclusiones y Trabajos Futuros.....	93
7.1 Conclusiones	93
7.2 Trabajos futuros.....	94
Capítulo 8. Bibliografía.....	96
ANEXO A. Alineación con Objetivos de Desarrollo Sostenible (ODS)	102
ANEXO B. Códigos completos.....	105

Índice de ilustraciones

Ilustración 1. Esquema del sistema de O.C.A.S.	11
Ilustración 2. Pipeline final de la herramienta O.C.A.S.	12
Ilustración 3. Diagram of the O.C.A.S. tool	14
Ilustración 4. Final pipeline of the O.C.A.S. tool	15
Ilustración 5. Chat de Telegram sobre OpSpain [3]	9
Ilustración 6. MITRE ATT&CK ICS [1]	12
Ilustración 7. Interfaz de usuario de GreyNoise [4]	14
Ilustración 8. Ejemplo de búsqueda de dispositivos Siemens en Shodan	15
Ilustración 9. Visualización de datos en mapa de Kibana	16
Ilustración 10. Pirámide del dolor [5].....	19
Ilustración 11. Aspectos a considerar al recopilar CTI [2].....	19
Ilustración 12. Recolección de datos [2]	20
Ilustración 13. Indicadores STIX	21
Ilustración 14. Interfaz gráfica de Intel741 [6].....	22
Ilustración 15. Interfaz gráfica de Nozomi Networks OT & IoT Threat Intelligence [7] ...	23
Ilustración 16. Fortiguard [8]	24
Ilustración 17. Esquema de la herramienta.....	28
Ilustración 18. Gantt del proyecto	29
Ilustración 19. Esquema completo de la herramienta.....	31
Ilustración 20. Funcionamiento de la herramienta	33
Ilustración 21. Cuantificación del riesgo del peor caso #1	47
Ilustración 22. Cuantificación del riesgo del peor caso #2.....	49
Ilustración 23. Cuantificación del riesgo del mejor caso 1.....	50
Ilustración 24. Cuantificación del riesgo del mejor caso 2.....	52
Ilustración 25. Cuantificación del riesgo del caso intermedio #1.....	53
Ilustración 26. Cuantificación del riesgo del caso intermedio #2.....	55

Ilustración 27. Comparativa Industroyer e Industroyer2 [9]	57
Ilustración 28. Arquitectura del modelo Keras.....	70
Ilustración 29. Ejemplo de prueba válida del modelo con caso sintético.....	71
Ilustración 30. Proceso de ingesta de noticias de ciberataques en tiempo real	72
Ilustración 31. Ingesta de información de contexto en tiempo real.....	75
Ilustración 32. Aplicación del modelo entrenado	77
Ilustración 33. Predicciones del modelo en tiempo real.....	77
Ilustración 34. Enriquecimiento con IoCs	78
Ilustración 35. Funcionamiento de la clasificación de IPs de GreyNoise	78
Ilustración 36. Ingesta de datos de inteligencia OT.....	80
Ilustración 37. Palabras clave para filtrado de mensajes de Telegram.....	80
Ilustración 38. Palabras clave para OSINT	81
Ilustración 39. Prompt al LLM que normaliza la inteligencia OT	81
Ilustración 40. Ejemplo de entrada de JSON de inteligencia OT.....	82
Ilustración 41. Inicio de la API.....	84
Ilustración 42. Interfaz gráfica.....	86
Ilustración 43. Interfaz gráfica (funcionalidades).....	86
Ilustración 44. Formulario de cuantificación de CVEs en contexto.....	87
Ilustración 45. Panel de inteligencia OT	87
Ilustración 46. Flujo final de la herramienta.....	88
Ilustración 47. Distribución de la variable de riesgo	90
Ilustración 48. Objetivos de Desarrollo Sostenible	102

Índice de tablas

Tabla 1. Ejemplos de CTI [2]	18
Tabla 2. Coste de recursos humanos	32
Tabla 3. Coste de licencias y herramientas.....	32
Tabla 4. Coste de hardware	32
Tabla 5. Cuantificación final del riesgo	38
Tabla 6. Justificación de las métricas de riesgo geopolítico	39
Tabla 7. Detalle de métricas de riesgo geopolítico [9]	40
Tabla 8. Resumen de la cuantificación del riesgo	57
Tabla 9. Ponderaciones en iteración 2	60
Tabla 10. Ponderaciones en iteración 3	61
Tabla 11. Ponderaciones en iteración 4	62
Tabla 12. Resultados de las pruebas de calibración	65

Índice de fragmentos de código

Fragmento de código 1. Función de cuantificación del riesgo final.....	44
Fragmento de código 2. Función de cálculo del riesgo parcial asociado al contexto geopolítico	44
Fragmento de código 3. Definición del “worst case 1”	47
Fragmento de código 4. Definición del "worst case 2"	49
Fragmento de código 5. Definición del "best case 2"	51
Fragmento de código 6. Definición del “caso intermedio #1”	53
Fragmento de código 7. Definición del “caso intermedio #2”	54
Fragmento de código 8. Código de cuantificación (iteración 0)	59
Fragmento de código 9. Cambios en código de cuantificación (iteración 1)	60
Fragmento de código 10. Cambios en código de cuantificación (iteración 5)	63
Fragmento de código 11. Cambios en código de cuantificación (iteración 7)	64
Fragmento de código 12. Definición de embeddings	68
Fragmento de código 13. Definición de la red neuronal	68
Fragmento de código 14. Modelo definido	69
Fragmento de código 15. Interpretación de las métricas obtenidas.....	69
Fragmento de código 16. Obtención de noticias de ciberataques generales	73
Fragmento de código 17. Scraping de las noticias	73
Fragmento de código 18. TLD confiables	73
Fragmento de código 19. Consulta o prompt al primer LLM	74
Fragmento de código 20. Consulta a la API de ACLED.....	75
Fragmento de código 21. Obtención de técnicas asociadas a un actor.....	76
Fragmento de código 22. API creada	84
Fragmento de código 23. Publicación de datos en formato STIX.....	85

Capítulo 1. INTRODUCCIÓN

Los sistemas de control industrial (SCI) son fundamentales para la operación segura y eficiente de infraestructuras críticas, especialmente en el sector energético. Estos sistemas gestionan procesos como la generación, transmisión y distribución de electricidad, supervisando equipos como turbinas, generadores y redes de distribución. Su creciente digitalización e interconexión han aumentado su exposición a riesgos de ciberseguridad, a lo que se suma la connotación geopolítica que tienen la mayoría de los ataques al sector energético, y que cada vez es más relevante.

1.1 MOTIVACIÓN DEL PROYECTO

La motivación del proyecto se centra en dos líneas principales. En primer lugar, el **aumento en las últimas décadas de ciberataques a SCIs**, que demuestra las amenazas crecientes en este ámbito. En segundo lugar, la **necesidad de una herramienta de ciberinteligencia orientada a OT con una aproximación al sector energético**, ya que se trata del foco de muchos ciberataques con connotaciones políticas y, además, gran parte de los activos de las empresas de esta industria son activos OT.

Uno de los primeros ataques maliciosos a un sistema de control industrial ocurrió en el año 2010: fue el famoso ataque de Stuxnet, en el que se desarrolló un gusano informático muy sofisticado dirigido a controladores lógicos programables (PLCs) de Siemens que causó daños físicos en las centrifugadoras de la central nuclear de Natanz (Irán). Este fue uno de los primeros ataques a SCIs con contexto geopolítico, ya que fue un programa creado por Israel y Estados Unidos para retrasar el programa de energía nuclear de Irán. En los últimos años, han ocurrido otros ataques en el contexto del conflicto entre Rusia y Ucrania. Entre 2014 y 2015, varios grupos de hackers rusos crearon un malware llamado BlackEnergy dirigido a la red eléctrica ucraniana, provocando cortes de electricidad en el país que afectaron a cientos de miles de personas. Posteriormente, se ha desarrollado más malware

diseñado para atacar a infraestructuras críticas ucranianas, como Industroyer (2016, 2022), así como otras herramientas de ciberespionaje vinculadas al conflicto ruso-ucraniano. También han ocurrido otros ataques con daños físicos a las plantas industriales, como el ataque a una planta de fabricación de acero en Alemania en 2014, y otros que provocaron disrupciones en las operaciones de las plantas, como el ransomware EKANS/Snake entre 2019 y 2023. Esta variedad de ataques diseñados para afectar a SCIs en entornos OT demuestran la necesidad de las empresas que operan en estos ámbitos de comprender el contexto completo que afecta a su entorno en cuanto a ciberseguridad.

Más recientemente, “OpSpain” en marzo de 2025 hizo saltar las alarmas: España fue el objetivo de ciberataques rusos tras su apoyo a Zelensky en el contexto del conflicto entre Rusia y Ucrania.



Ilustración 5. Chat de Telegram sobre OpSpain [3]

Además de la política y los objetivos militares (que en muchos casos se centran en redes eléctricas y otras infraestructuras energéticas), el hacktivismo ha jugado un papel de igual importancia en la ciberseguridad, donde grupos como Anonymous han llevado a cabo ataques para reivindicar causas políticas o sociales. Estos ataques, motivados por razones que van desde la política hasta el deseo de reconocimiento (*bragging rights*), han demostrado la capacidad de los actores no estatales para influir en la ciberseguridad global.

El auge de la IA y los modelos de ML en la ciberseguridad, tanto la parte defensiva como la ofensiva, manifiesta la necesidad de incorporar estas herramientas en el día a día de la protección de los SCIs.

1.2 SOLUCIÓN PROPUESTA

La creciente sofisticación de los ataques dirigidos de entornos OT implica repercusiones que van más allá de interrupciones inmediatas de las operaciones, suponiendo un riesgo para la seguridad pública e incluso la estabilidad económica. El proyecto surge, por tanto, en un contexto de necesidad de concienciación y preparación ante posibles ataques con repercusiones tan graves. Se propone desarrollar una herramienta de inteligencia de ciberamenazas o ciberinteligencia orientada al sector energético y a sus sistemas de control industrial. El objetivo es cubrir las necesidades de inteligencia de este ámbito, combinando inteligencia de amenazas técnicas con inteligencia del contexto geopolítico.

El nombre de la herramienta es O.C.A.S. (*OT Cyberrisk Assessment System*), y hace honor a la leyenda de la Antigua Roma que cuenta que las ocas del templo de Juno salvaron a los últimos defensores de la ciudad de ser derrotados por los galos.

Este proyecto se plantea como una iniciativa de investigación para impulsar la aplicación de la inteligencia de amenazas (TI) en entornos OT. Se engloba en una fase de prueba y desarrollo, orientada a validar enfoques, evaluar herramientas y sentar las bases para futuras implementaciones.

Capítulo 2. DESCRIPCIÓN DE LAS TECNOLOGÍAS

En este capítulo se describen las tecnologías y herramientas específicas utilizadas en el desarrollo de la herramienta de inteligencia de amenazas de ciberseguridad orientada al sector energético para Iberdrola. Esta descripción tiene como objetivo facilitar la lectura y comprensión del proyecto.

2.1 PYTHON

Python es un lenguaje de programación de alto nivel, ampliamente utilizado por su simplicidad y legibilidad. Python ha sido la base para el desarrollo de este proyecto por su gran cantidad de librerías y módulos y por su capacidad para manejar grandes volúmenes de datos.

2.1.1 JUPYTER NOTEBOOK

Jupyter Notebook es una aplicación web que permite crear y compartir documentos que contienen código en vivo, ecuaciones, visualizaciones y texto narrativo. En este proyecto, se ha utilizado para el desarrollo y la prueba de scripts de análisis de datos, facilitando la visualización y la documentación del proceso.

2.1.2 VISUAL STUDIO CODE

Visual Studio Code es un editor de código fuente desarrollado por Microsoft. Soporta diferentes lenguajes de programación y tiene una amplia gama de extensiones. Aquí se ha utilizado como el entorno de desarrollo integrado (IDE) principal para escribir y depurar el código.

2.1.3 PANDAS

Pandas es una librería de Python que facilita la manipulación y el análisis de datos. Proporciona estructuras de datos y herramientas para trabajar con datos estructurados (tablas,

series temporales, etc.). Se ha utilizado desde el principio del proyecto para para la limpieza y el análisis de los datos recopilados, tanto en las pruebas iniciales (durante el análisis y la elección de fuentes de información para el desarrollo de la base de la herramienta) como en el núcleo final de la herramienta desarrollada.

2.1.4 MITRE ATT&CK

MITRE ATT&CK es una base de datos de tácticas y técnicas utilizadas por adversarios en ciberataques, basada en observaciones del mundo real. Se ha utilizado para mapear y analizar las amenazas detectadas y los autores de los ciberataques a sus países de origen, a sus técnicas empleadas y a su presencia o no en ataques a entornos industriales. Para esto último tiene especial relevancia el marco MITRE ATT&CK ICS, orientado a las tácticas y técnicas utilizadas por adversarios en ataques a sistemas de control industrial.

Initial Access	Execution	Persistence	Privilege Escalation	Evasion	Discovery	Lateral Movement	Collection	Command and Control	Inhibit Response Function	Impair Process Control	Impact
12 techniques	10 techniques	6 techniques	2 techniques	7 techniques	5 techniques	7 techniques	11 techniques	3 techniques	14 techniques	5 techniques	12 techniques
Drive-by Compromise	Autorun Image	Hardcoded Credentials	Exploitation for Privilege Escalation	Change Operating Mode	Network Connection Enumeration	Default Credentials	Adversary-in-the-Middle	Commonly Used Port	Activate Firmware Update Mode	Brute Force I/O	Damage to Property
Exploit Public-Facing Application	Change Operating Mode	Modify Program	Hooking	Exploitation for Evasion	Network Sniffing	Exploitation of Remote Services	Automated Collection	Connection Proxy	Alarm Suppression	Modify Parameter	Denial of Control
Exploitation of Remote Services	Command-Line Interface	Module Firmware		Indicator Removal on Host	Remote System Discovery	Hardcoded Credentials	Data from Information Repositories	Standard Application Layer Protocol	Block Command Message	Module Firmware	Denial of View
External Remote Services	Execution through API	Project File Infection		Masquerading	Remote System Information Discovery	Lateral Tool Transfer	Data from Local System		Block Reporting Message	Spoof Reporting Message	Loss of Availability
Internet Accessible Device	Graphical User Interface	System Firmware		Rootkit	Wireless Sniffing	Program Download	Detect Operating Mode		Block Serial COM	Unauthorized Command Message	Loss of Control
Remote Services	Hooking	Valid Accounts		Spoof Reporting Message		Remote Services	I/O Image		Change Credential		Loss of Productivity and Revenue
Replication Through Removable Media	Modify Controller Tasking			System Binary Proxy Execution		Valid Accounts	Monitor Process State		Data Destruction		Loss of Protection
Rogue Master	Native API						Point & Tag Identification		Denial of Service		Loss of Safety
Spearphishing Attachment	Scripting						Program Upload		Manipulate I/O Image		Manipulation of Control
Supply Chain Compromise	User Execution						Screen Capture		Modify Alarm Settings		Manipulation of View
Transient Cyber Asset							Wireless Sniffing		Rootkit		Theft of Operational Information
Wireless Compromise									Service Stop		
									System Firmware		

Ilustración 6. MITRE ATT&CK ICS [1]

2.1.5 NEWSPAPER3K

Newspaper3k es una librería de Python que permite extraer y analizar artículos de páginas web, como noticias. Es esencial para la herramienta desarrollada, que se basa en noticias de ciberataques para obtener inteligencia.

2.1.6 KERAS

Keras es una librería de Python que facilita la creación y el entrenamiento de modelos de aprendizaje automático o *machine learning*. Aquí se ha usado para desarrollar y entrenar un modelo de inteligencia artificial basado en una red neuronal capaz de identificar patrones y anomalías en los datos de amenazas.

2.1.7 FASTAPI

FastAPI es un framework web que facilita la construcción de APIs con Python. Aquí se ha utilizado para desarrollar la API que permite la interacción con la herramienta O.C.A.S., y también para la comunicación con el SOC.

2.1.8 TELETHON

Telethon es una librería de Python que permite desplegar un cliente asíncrono de Telegram. Se ha utilizado para monitorizar las conversaciones de grupos públicos de ciberseguridad industrial.

2.2 ***HERRAMIENTA DE LLM: OPENROUTER***

OpenRouter es una plataforma que permite a los usuarios acceder y utilizar modelos de inteligencia artificial desarrollados por diferentes empresas a través de una única interfaz unificada. Facilita la conexión con múltiples LLMs, como los de OpenAI, lo que permite integrarlos fácilmente en aplicaciones mediante una API común.

2.3 ***HERRAMIENTAS DE THREAT INTELLIGENCE***

2.3.1 GREYNOISE

GreyNoise es una plataforma que analiza el tráfico de Internet para identificar indicadores de compromiso (IoCs) como direcciones IP asociadas a determinados ataques, actores o

ubicaciones. Se ha utilizado para enriquecer datos en tiempo real de ciberataques y explotaciones de vulnerabilidades. En concreto, se ha utilizado su funcionalidad de búsqueda de IPs asociadas a la explotación de CVEs, utilizando GNQL (*GreyNoise Query Language*), el formato de consultas de su API.

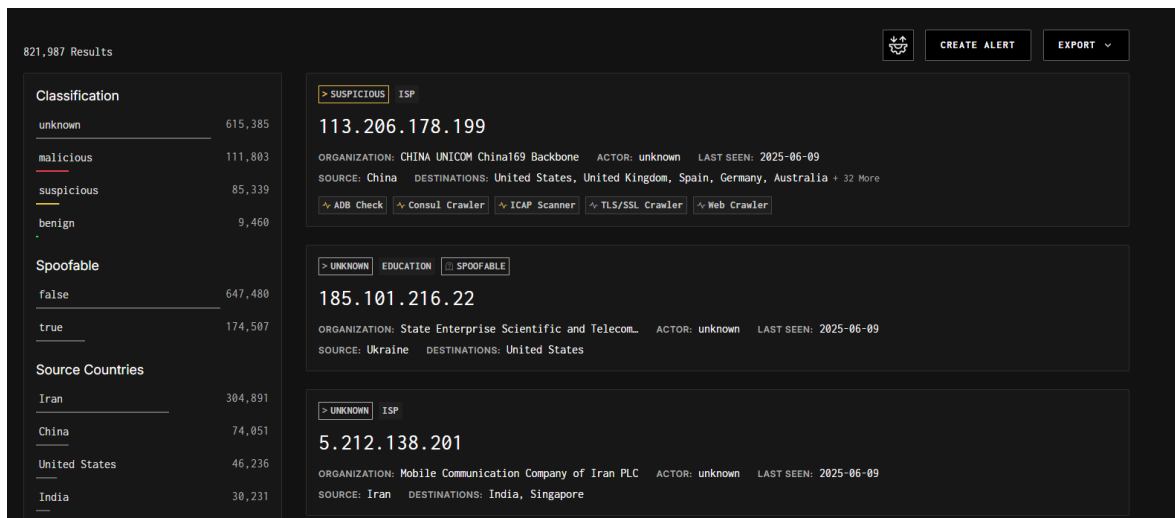


Ilustración 7. Interfaz de usuario de GreyNoise [4]

2.3.2 SHODAN

Shodan es un motor de búsqueda especializado que, a diferencia de los motores de búsqueda tradicionales que indexan páginas web, permite explorar dispositivos conectados a Internet. A través de escaneos constantes, recopila información sobre servidores, cámaras IP, routers, impresoras, sistemas de control industrial como PLCs y otros dispositivos del Internet de las Cosas (IoT). Sobre los dispositivos encontrados, proporciona detalles como direcciones IP, puertos abiertos o servicios en ejecución. Se ha utilizado la API de Shodan para hacer consultas automatizadas a sus datos.

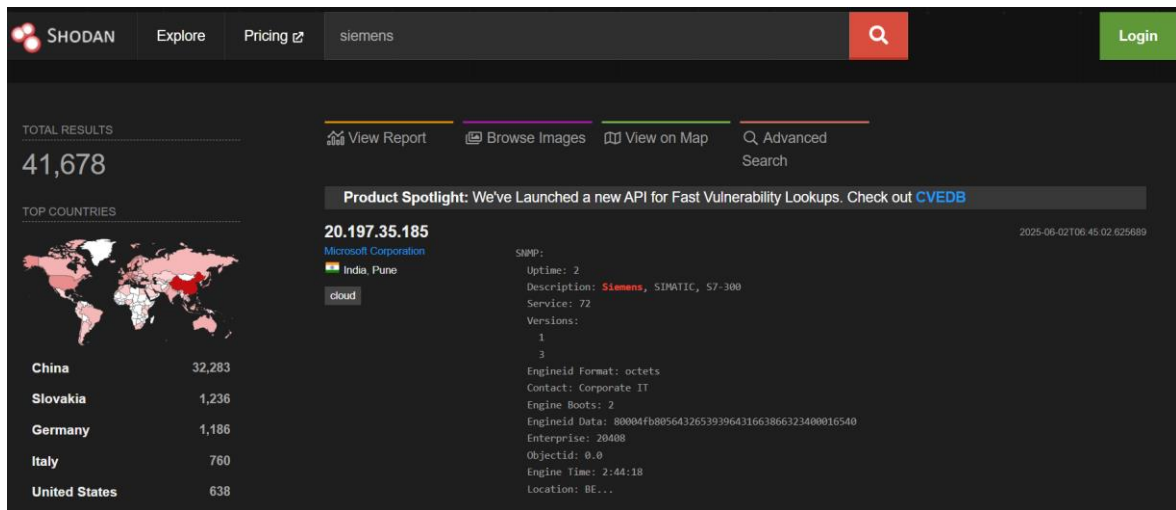


Ilustración 8. Ejemplo de búsqueda de dispositivos Siemens en Shodan

2.4 HERRAMIENTAS DE VISUALIZACIÓN

2.4.1 KIBANA Y ELASTICSEARCH

Kibana es una plataforma de código abierto que permite visualizar y explorar datos almacenados en Elasticsearch. Permite analizar grandes volúmenes de información en gráficos, diagramas y mapas, y es muy útil para datos en tiempo real. En este proyecto se ha utilizado para la visualización de datos de ciberataques en un mapa.

2.5 OTRAS HERRAMIENTAS

2.5.1 RENDER

Render es una plataforma gratuita que permite desplegar APIs (*Application Programmable Interfaces*) a partir de repositorios de código en otras plataformas como GitHub. Ofrece un despliegue rápido y seguro.

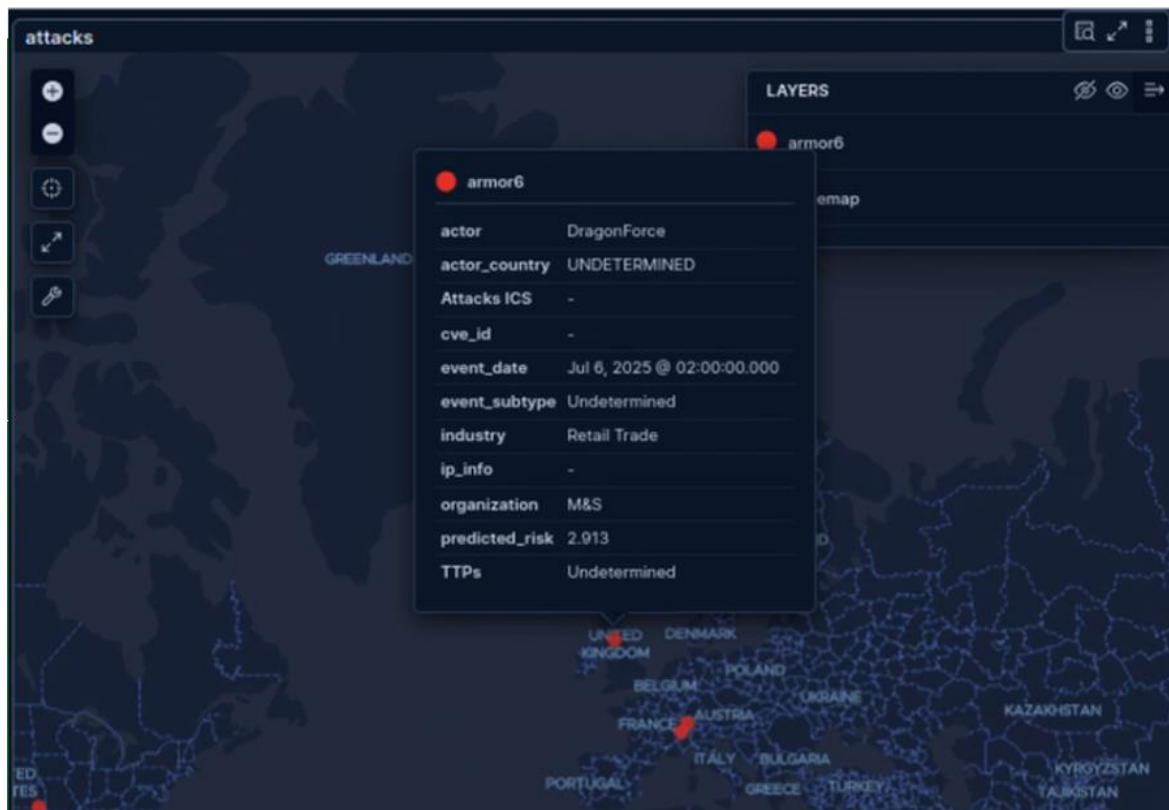


Ilustración 9. Visualización de datos en mapa de Kibana

Capítulo 3. ESTADO DE LA CUESTIÓN

3.1 CTI (CYBER THREAT INTELLIGENCE)

La inteligencia de ciberamenazas, CTI por sus siglas en inglés (*Cyber Threat Intelligence*), se refiere a la recopilación, procesamiento e interpretación de datos con el objetivo de transformarlos en información o elementos accionables. Esta información está relacionada con las capacidades, acciones e intenciones de adversarios en el dominio cibernético, y es útil para los responsables de la toma de decisiones en el ámbito de ciberseguridad.

3.1.1 FUENTES DE CTI

La CTI proviene de fuentes internas y externas, ambas de gran importancia. Las fuentes externas incluyen informes de amenazas informáticas (públicos o según contratos comerciales), reportes de amenazas en comunidades de SOC's nacionales o internacionales, y fuentes abiertas publicadas en repositorios públicos. Por otro lado, las fuentes internas incluyen a los analistas de seguridad que analizan información sobre adversarios con datos del inventariado de activos o sondas de seguridad.

Los datos se consideran CTI solo si están claramente relacionados con información contextual del comportamiento del adversario o atacante que describen. Por ejemplo, una dirección IP por sí sola no es CTI. Dicha IP debe “localizarse” en el contexto, determinando si se han realizado acciones sospechosas o si se ha intentado explotar alguna vulnerabilidad desde ella [2].

Los elementos de CTI deben cumplir las siguientes características:

- *Accionable*: que se puede traducir en una acción por parte del SOC
- *Oportuno*: significa que es información reciente, aplicable en el presente

- *Relevante*: significa que es aplicable en el contexto de la organización
- *Preciso*: significa que describe precisa y correctamente lo que ocurre

Algunos ejemplos de elementos CTI y otros que no se consideran CTI (si no se asocian a un contexto de atacante o adversario) se muestran en la siguiente tabla:

<i>Ejemplos de CTI</i>	<i>No CTI</i>
Informes de amenazas no estructurados Informes de amenazas estructurados Open-Source Intelligence (OSINT) Informes y comentarios de suscriptores	Direcciones IP Nombres de dominio Direcciones de email Muestras de malware Firmas de virus Logs Redes sociales

Tabla 1. Ejemplos de CTI [2]

Los informes de amenazas no estructurados incluyen informes de inteligencia y análisis, a veces muy extensos, que describen adversarios, observables y análisis de contexto. Por ejemplo, informes de CTI de varias páginas producidos por muchos SOC's, a menudo distribuidos en formato PDF y HTML. Los informes de amenazas estructurados incluyen tácticas, técnicas y procedimientos (TTPs) contextualizados, como MITRE ATT&CK. Por otro lado, OSINT se refiere al conocimiento obtenido de fuentes de acceso público. Por último, los informes y comentarios de suscriptores o clientes contienen información de amenazas anonimizada de estos mismos. Por ejemplo, Mandiant Advantage Free.

Los indicadores de compromiso (IoCs), que incluyen direcciones IP, nombres de dominio, direcciones de email, etc., no se consideran CTI si no se relacionan con un contexto específico. Además, cabe destacar que los IoCs tienen un carácter dinámico, puesto que permiten a los atacantes que los usan cambiarlos fácilmente. Esto se refleja claramente en el concepto de la Pirámide del Dolor, en la cual los IoCs se encuentran en la parte más inferior.

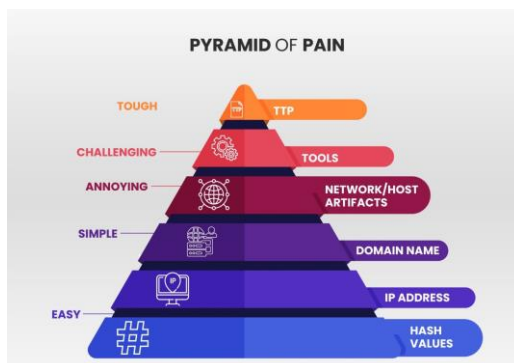


Ilustración 10. Pirámide del dolor [5]

3.1.2 RECOGIDA DE CTI

Para que un SOC pueda determinar acciones efectivas a partir de CTI, es fundamental contar con los datos adecuados. Esto implica considerar tres aspectos clave: el entorno técnico, la información sobre los adversarios y la relevancia.

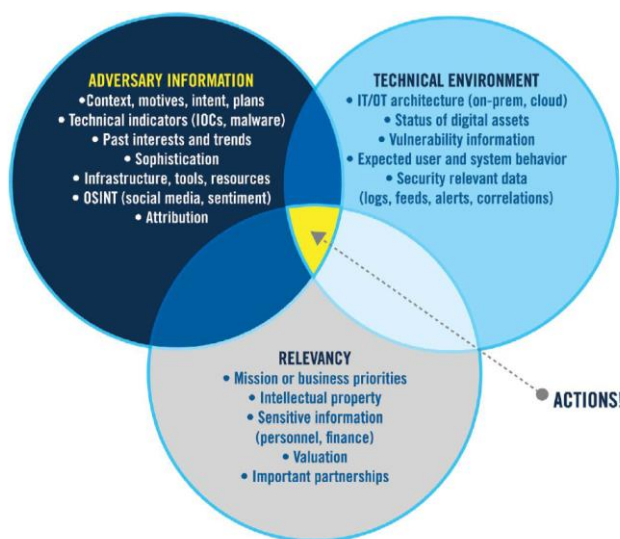


Ilustración 11. Aspectos a considerar al recopilar CTI [2]

El **entorno técnico** abarca los datos sobre las tecnologías a proteger, incluyendo activos IT/OT y sus estados, información sobre vulnerabilidades, registros o logs, comportamiento de usuarios, etc. La información sobre los **adversarios** engloba todo lo que se sabe sobre ellos y cómo se manifiestan en la tecnología. Esto incluye información de contexto,

indicadores técnicos, tendencias pasadas, OSINT, etc. La **relevancia** se refiere al conocimiento de las prioridades del negocio y la información importante de la organización.

Estos tres aspectos implican una gran cantidad de datos, que se incluyen en detalle en secciones posteriores de este documento. Con tantos datos entre los que elegir, es necesario identificar aquellos que son más útiles para los objetivos de la empresa, para evitar perder información relevante. La gráfica siguiente muestra cómo esta pérdida de información puede ocurrir tanto por una recolección deficiente de datos como por una recolección excesiva, que hace que las herramientas y/o analistas estén sobrecargados.

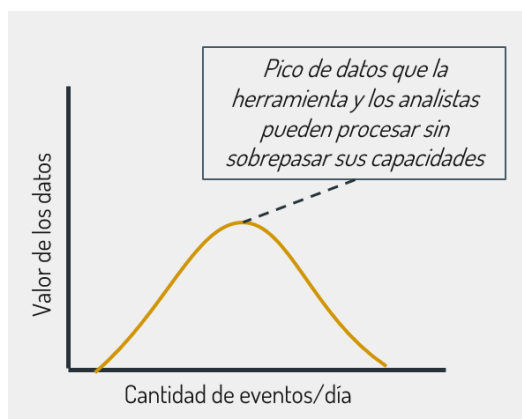


Ilustración 12. Recolección de datos [2]

Esta idea se aplica en este trabajo a la hora de elegir las fuentes a tratar para el desarrollo de la herramienta.

3.1.3 STIX

STIX (*Structured Threat Information eXpression*) es un lenguaje estandarizado diseñado para expresar información o inteligencia sobre ciberamenazas de manera estructurada. Permite compartir, almacenar y analizar datos de amenazas en un formato consistente, facilitando la automatización y la colaboración entre diferentes entidades de seguridad. STIX utiliza un léxico basado en JSON, por lo que es fácilmente legible para diferentes aplicaciones y herramientas de ciberseguridad. Por otro lado, TAXII (*Trusted Automated eXchange of Intelligence Information*) es un protocolo de transporte que soporta la

transferencia de STIX sobre HTTPS (*Hyper Text Transfer Protocol Secure*). STIX y TAXII son estándares independientes: STIX no se basa en un mecanismo de transporte específico, y TAXII se puede utilizar para transferir datos que no sean STIX.

En GitHub, la organización MITRE ATT&CK proporciona una serie de repositorios que contienen datos expresados en STIX. Se trata de ficheros JSON que representan (y no se limitan a) tácticas, técnicas y procedimientos de adversarios o atacantes observados en el mundo real. Los datos en formato STIX permiten a los analistas de seguridad y desarrolladores integrar y utilizar la información de MITRE ATT&CK en sus propias herramientas y sistemas, mejorando la capacidad de respuesta ante amenazas cibernéticas.

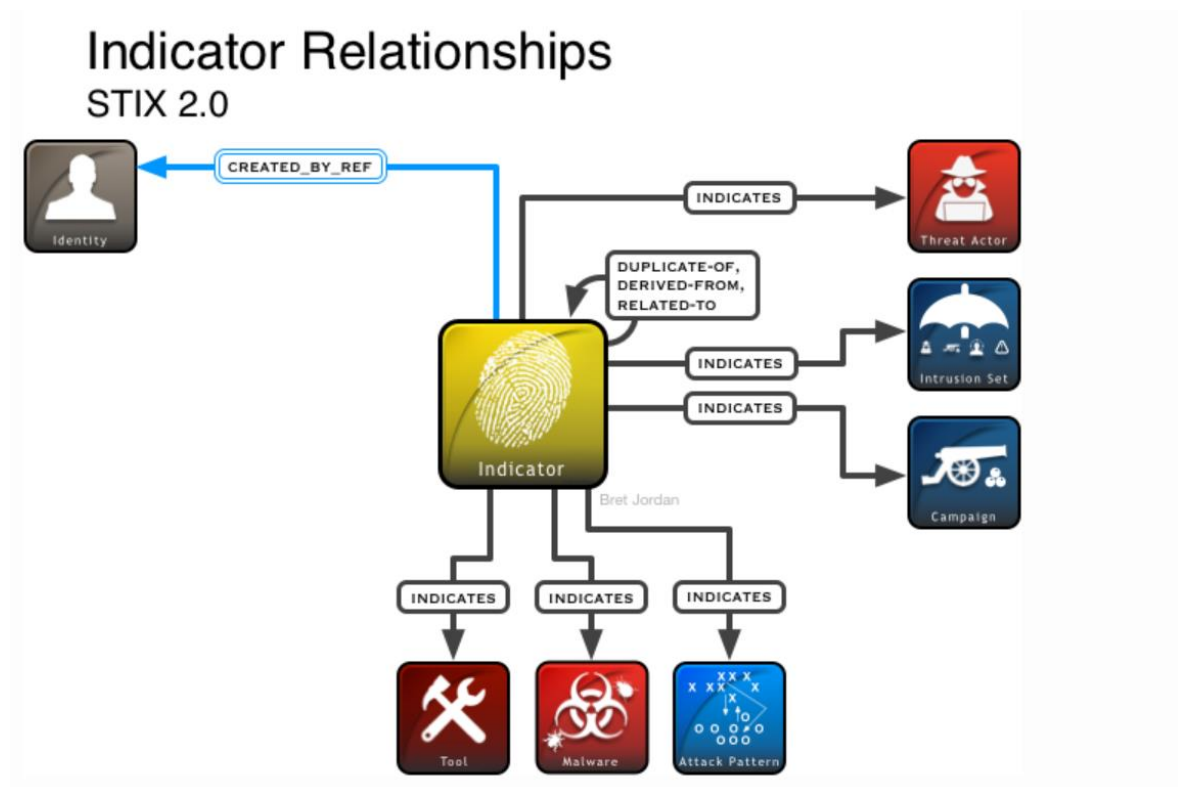


Ilustración 13. Indicadores STIX

3.2 SOLUCIONES TECNOLÓGICAS Y TRABAJOS DE INVESTIGACIÓN

En esta sección se describen diferentes soluciones tecnológicas existentes que abordan el ámbito de la ciberinteligencia y la inteligencia de amenazas.

3.2.1 SOLUCIONES TECNOLÓGICAS EXISTENTES

Intel471. Proveedor líder de CTI o inteligencia de ciberamenazas. Una de sus principales herramientas es una plataforma SaaS (*Software as a Service* o “software como servicio”) llamada TITAN que proporciona visibilidad de los actores de ciberataques y las amenazas que estos implican para las organizaciones. También permite, entre otras cosas, monitorizar eventos geopolíticos que puedan afectar al riesgo de ciberseguridad de una empresa, enriqueciendo así la respuesta a incidentes.

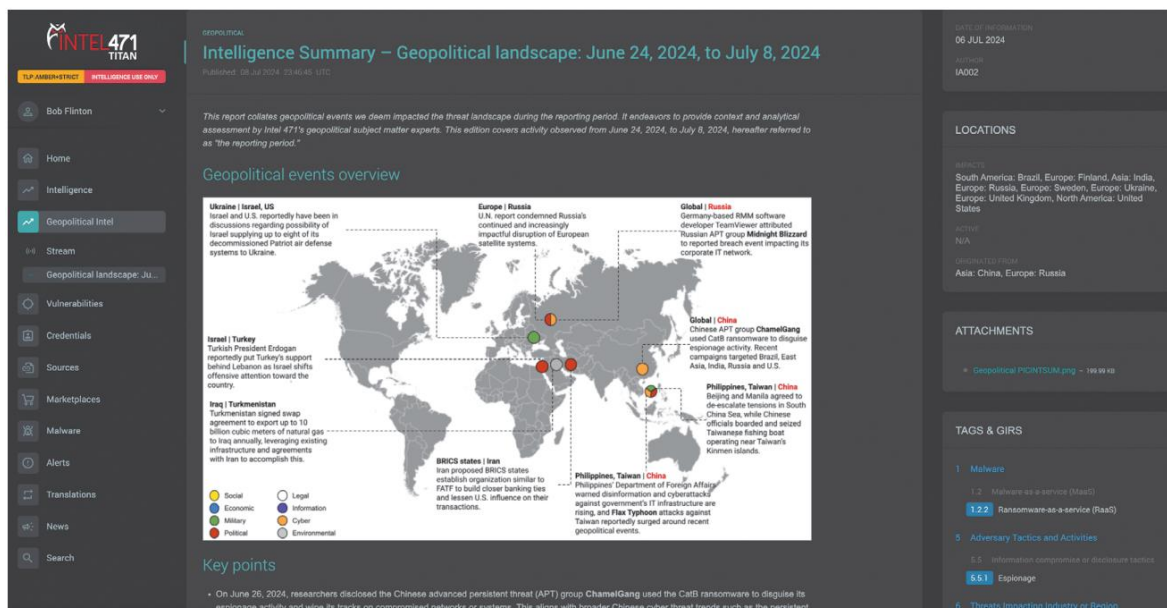


Ilustración 14. Interfaz gráfica de Intel741 [6]

Nozomi Networks OT & IoT Threat Intelligence. Uno de los productos que ofrece la empresa Nozomi Networks es una solución avanzada de inteligencia de amenazas enfocada en la seguridad de tecnologías operativas (OT) y el Internet de las Cosas (IoT). Su plataforma

permite la detección de anomalías, la monitorización de tráfico en redes industriales y la identificación de vulnerabilidades en los activos de estas.

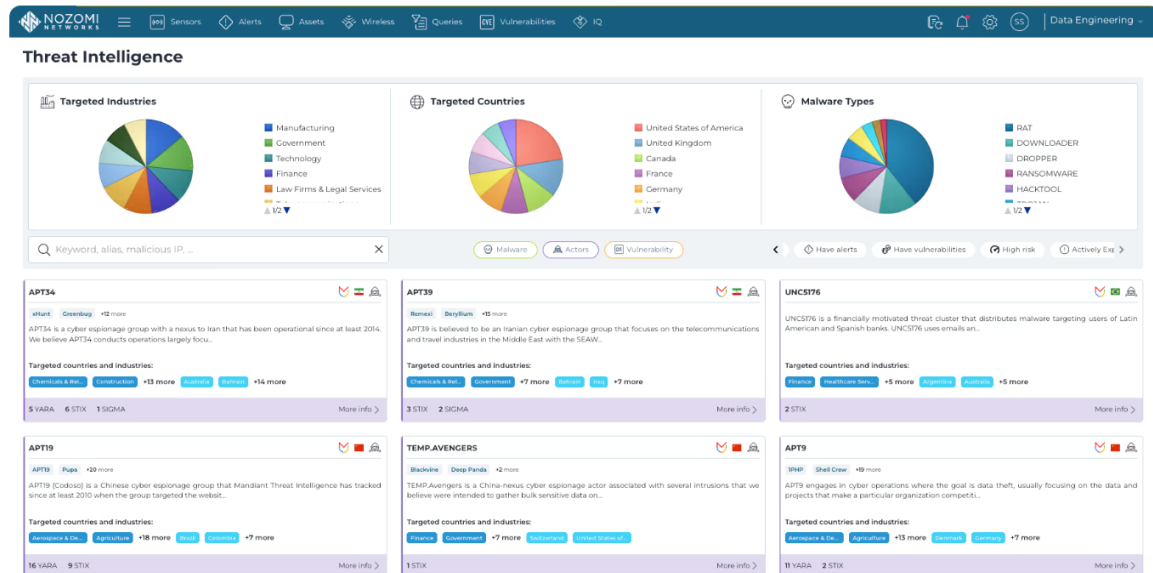


Ilustración 15. Interfaz gráfica de Nozomi Networks OT & IoT Threat Intelligence [7]

Fortiguard Labs Threat Intelligence Platform. Es una plataforma de inteligencia de amenazas que combina diferentes tecnologías, como *machine learning*, análisis de comportamiento y detección de malware, para identificar y prevenir ciberataques. Su solución está diseñada para integrarse con el ecosistema de seguridad de Fortinet, proporcionando actualizaciones en tiempo real sobre nuevas amenazas, protección contra malware y estrategias de mitigación basadas en análisis globales de seguridad. Permite consultar inteligencia de amenazas relativa a una alerta específica, como se muestra en la Ilustración 16. Fortiguard

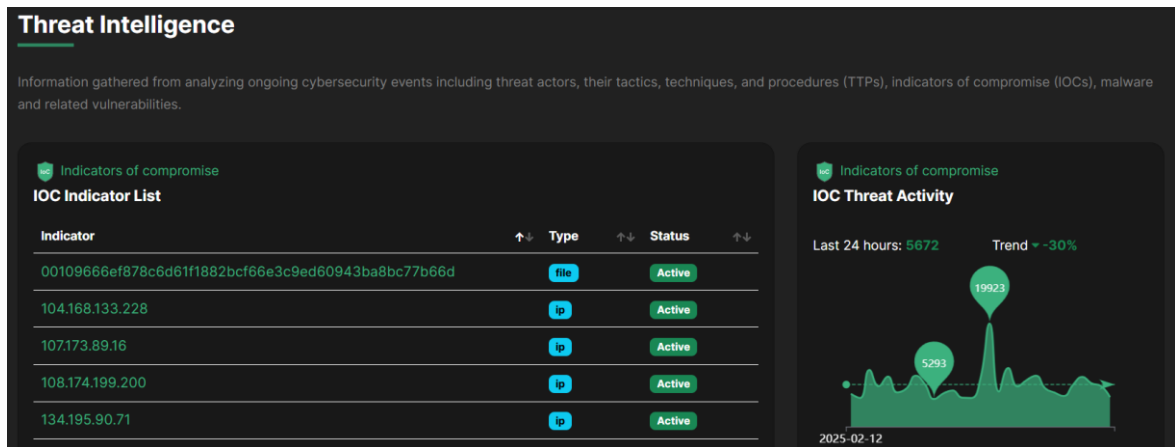


Ilustración 16. Fortiguard [8]

3.2.2 TRABAJOS DE INVESTIGACIÓN RELACIONADOS

La inteligencia de ciberamenazas o CTI es el foco de muchos trabajos de investigación. Aunque existen estándares específicos para su compartición, como STIX, la diversidad de sus fuentes supone un reto a la hora de obtenerla y filtrarla. Las herramientas de Inteligencia Artificial (IA) han supuesto un gran avance en esto, como las relacionadas con *Machine Learning* (ML) y LLMs (*Large Language Models* o modelos de lenguaje de gran tamaño).

Un ejemplo destacado es el trabajo "*Enhancing Cyber Threat Identification in Open-Source Intelligence Feeds Through an Improved Semi-Supervised Generative Adversarial Learning Approach With Contrastive Learning*" de la Universidad de Leeds (Reino Unido). Este estudio se centra en la identificación automatizada de amenazas utilizando ML, combinando datos etiquetados y no etiquetados para mejorar la identificación de amenazas, a partir de una versión avanzada del modelo GAN-BERT (*Generative Adversarial Learning – Bidirectional Encoder Representations from Transformers*).

Otro estudio relevante es "*The Use of Large Language Models (LLM) for Cyber Threat Intelligence (CTI) in Cybercrime Forums*", que evalúa el rendimiento de un sistema LLM basado en el modelo GPT-3.5-turbo para extraer información de CTI de foros de ciberdelincuencia. Los resultados muestran una alta precisión y relevancia en la extracción

de información crítica sobre amenazas emergentes.

Por otro lado, el artículo "*Using LLMs to Automate Threat Intelligence Analysis Workflows in Security Operation Centers*" explora cómo los LLMs pueden automatizar tareas repetitivas en los SOC, mejorando la eficiencia y reduciendo la intervención humana (aunque esta es siempre necesaria).

El uso de LLMs se ha convertido en algo esencial en todo lo relativo a la extracción de inteligencia de ciberamenazas de diferentes fuentes abiertas.

Capítulo 4. DEFINICIÓN DEL TRABAJO

4.1 JUSTIFICACIÓN

4.1.1 NECESIDAD ESTRATÉGICA

En el contexto actual de creciente sofisticación de las amenazas de ciberseguridad, especialmente dirigidas a infraestructuras críticas en el sector industrial energético, es imprescindible reforzar las capacidades de detección y respuesta del departamento de ciberseguridad de Iberdrola. La digitalización de los entornos OT, tradicionalmente aislados, ha incrementado su exposición a vectores de ataque avanzados, lo que claramente exige una evolución en las herramientas de monitorización e inteligencia. En este sentido, el desarrollo de una solución específica de *Cyber Threat Intelligence* para OT permitirá identificar, correlacionar y anticipar amenazas dirigidas a los sistemas industriales, mejorando significativamente la visibilidad y la capacidad de reacción ante incidentes.

Este proyecto se enmarca en la estrategia de Iberdrola para fortalecer la resiliencia de sus operaciones críticas, alineándose con los estándares internacionales de ciberseguridad industrial y las recomendaciones de organismos reguladores. La herramienta propuesta integrará fuentes de inteligencia específicas del ámbito OT y capacidades de análisis contextualizado, lo que permitirá al equipo de ciberseguridad tomar decisiones informadas y proactivas. De este modo, se contribuye no solo a la protección de los activos industriales, sino también a la continuidad del servicio y la confianza de los clientes y otros grupos de interés de la empresa.

4.1.2 PROYECTO MODULAR

El proyecto desarrollado se ha concebido como un sistema modular y ampliable, diseñado para adaptarse a las necesidades actuales y futuras de Iberdrola. Esta modularidad permitirá que el sistema crezca en cuanto a las fuentes de datos que lo alimentan, facilitando su ajuste conforme evolucionen las amenazas y los requisitos de seguridad de la empresa en el futuro.

Por ejemplo, para la identificación inicial de vulnerabilidades (CVE), se ha utilizado la base de datos de vulnerabilidades del *National Vulnerability Database* (NVD). Esta elección se ha basado en la fiabilidad y la exhaustividad de la información proporcionada por NVD, pero el diseño modular del proyecto permite la integración de otras fuentes de información, como FIRST (*Forum of Incident Response and Security Teams*), una organización global que proporciona información sobre vulnerabilidades y amenazas.

También se ha usado GreyNoise, una plataforma que analiza el tráfico de Internet para identificar escaneos y ataques automatizados, pero se podrían incluir otras fuentes como Criminal IP, un motor de búsqueda de inteligencia sobre amenazas que utiliza tecnología de aprendizaje automático para supervisar puertos abiertos y proporcionar informes detallados con una puntuación de riesgo. Se optó por GreyNoise porque se obtuvo una licencia de estudiantes para utilizar su API.

La elección de las fuentes de información y las herramientas utilizadas en este proyecto se ha realizado intentando cubrir diversas tipologías de fuentes de datos, siempre dentro de las necesidades actuales de Iberdrola. No obstante, la arquitectura modular del sistema permite que este se pueda ampliar con otras fuentes de datos complementarias en el futuro. A medida que surjan nuevas tecnologías o cambien las necesidades de seguridad, el sistema puede adaptarse fácilmente para incorporar nuevas fuentes de datos y herramientas de análisis.

4.2 OBJETIVOS

El proyecto persigue cuatro objetivos principales:

1. **Entrenamiento de un modelo de *Machine Learning*** que tome como entrada una serie de datos sobre ciberataques (fecha, empresa u organización afectada, tipo de ataque, atacante, países involucrados, situación geopolítica en dichos países, etc.) y que tenga como variable objetivo un valor del riesgo que dicho evento supone para la empresa.

2. **Desarrollo de una herramienta de ingesta y contextualización de información de Threat Intelligence en entornos OT.** Esta herramienta debe utilizar el modelo de ML mencionado, el contexto geopolítico y otros datos sobre amenazas en tiempo real para evaluar el riesgo que un evento supone para la empresa.
3. **Integración con el SOC.** El resultado de la combinación de fuentes se debe presentar en un formato adecuado para ser comunicado al *Security Operations Center* y que este pueda tomar decisiones informadas en base a ello.
4. **Interfaz gráfica o *dashboard*.** Para visualización de inteligencia sobre OT, se requiere de una interfaz gráfica que proporcione conciencia situacional de forma clara y en tiempo real. Para este objetivo, se contempla el uso de herramientas como Kibana.

4.3 METODOLOGÍA

La metodología a seguir en este proyecto se ilustra en el siguiente esquema:

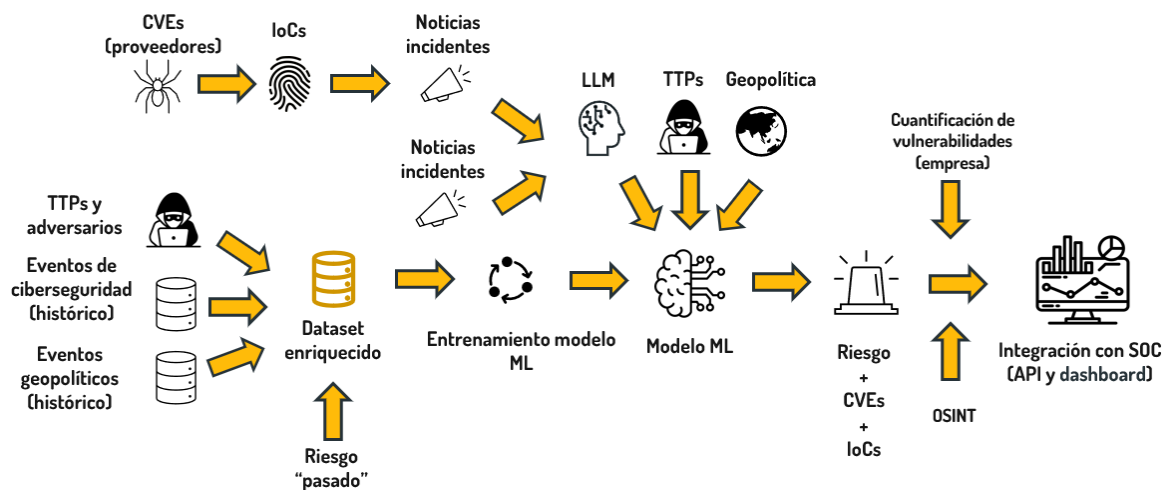


Ilustración 17. Esquema de la herramienta

4.4 PLANIFICACIÓN Y ESTIMACIÓN ECONÓMICA

4.4.1 PLANIFICACIÓN

A continuación, se muestra un cronograma de tipo Gantt que ilustra el plan de trabajo:

		Feb 3 - Mar 2				Mar 3 - 30				Mar 31 - Abr 27				Abr 28 - May 25				May26-Jun15		
		S1	S2	S3	S4	S1	S2	S3	S4	S1	S2	S3	S4	S1	S2	S3	S4	S1	S2	S3
FASE 1: Investigación y análisis de fuentes	1.1: Recopilación de datos y fuentes																			
	1.2: Pruebas con fuentes en Python																			
FASE 2: Desarrollo de la herramienta de contextualización	2.1: Definición datos y variable objetivo																			
	2.2: Entrenamiento de modelo ML																			
	2.3: Definición datos en tiempo real																			
	2.4: Aplicación modelo tiempo real																			
FASE 3: Integración con SOC y ajustes técnicos	3.1: Definición modelo comunicación																			
	3.2: Implantación de comunicación																			
FASE 4: Desarrollo de interfaz gráfica	4.1: Diseño de interfaz gráfica																			
	4.2: Desarrollo de interfaz gráfica																			

Ilustración 18. Gantt del proyecto

El proyecto consta de cuatro fases principales:

Fase 1: Investigación y análisis de fuentes. En la primera fase del proyecto se seleccionan las fuentes de información a utilizar. Esta fase es esencial para asegurar que la herramienta se basa en fuentes fiables. Se hacen dos selecciones:

- Fuentes “históricas” o de eventos pasados, para construir un conjunto de datos enriquecido como se muestra en la Ilustración 17. A partir de estas fuentes (y una serie de pruebas de calibración) se definirá un nivel de riesgo de cada evento y su contexto para la empresa energética. El riesgo se basará en una combinación de factores clave, que se explica más adelante (ver Definición del riesgo), procedente de diferentes fuentes. Estas fuentes son:
 - Eventos de ciberseguridad históricos. Registros de incidentes previos, en entornos IT y OT. Se obtienen del repositorio de eventos de ciberseguridad de la Universidad de Maryland.
 - Eventos geopolíticos históricos. Información relevante sobre conflictos que pueden aumentar el riesgo de ciberataques. Se utiliza para esto la base de datos de ACLED.

- Tácticas, técnicas y procedimientos (TTPs). Se hace foco en las técnicas utilizadas por actores maliciosos en ataques (a sistemas industriales y también a sistemas de información tradicionales).
- Fuentes en tiempo real, para obtener información en el momento y poder compararla con los históricos.
 - GoogleNews. Para obtener noticias (de eventos de ciberseguridad) en tiempo real (y posteriormente filtrarlas y normalizarlas con un LLM).
 - Eventos geopolíticos actuales. Para obtener datos geopolíticos (de conflictos en países involucrados en los eventos de ciberseguridad) en tiempo real se usa la API de ACLED.
 - CVEs (*Common Vulnerabilities and Exposures*). Vulnerabilidades conocidas que podrían ser explotables. Se hace foco en CVEs de proveedores de Iberdrola que puedan afectar a la empresa.
 - KEV (*Known Exploited Vulnerabilities*). Vulnerabilidades (CVEs) que han sido explotadas o tienen un exploit conocido público.
 - EPSS (*Exploit Prediction Scoring System*). Modelo basado en datos que ayuda a predecir la probabilidad de que una vulnerabilidad de software sea explotada.
 - Indicadores de compromiso y otra información de contexto. Para ello se utilizan herramientas y motores de búsqueda de direcciones IP y otros datos vinculados a CVEs, como GreyNoise o Shodan.

El esquema específico, con las herramientas usadas, queda:

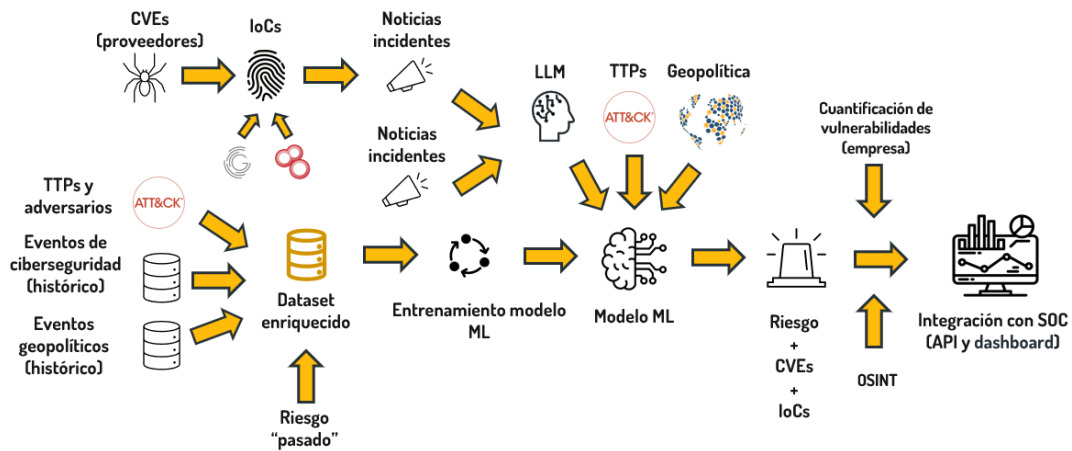


Ilustración 19. Esquema completo de la herramienta

Fase 2. Desarrollo de la herramienta de contextualización. Se definen dos etapas:

- Entrenamiento de un modelo de Machine Learning con el dataset enriquecido. Se harán las limpiezas y normalizaciones necesarias sobre el dataset completo para poder entrenar el modelo con él, y se compararán diferentes modelos.
- Aplicación y enriquecimiento del modelo en tiempo real. Se utilizará el modelo anterior sobre los datos en tiempo real y se combinará con las fuentes adicionales de la fase 1.

Fase 3. Integración con SOC y ajustes técnicos. Se definirá e implantará un modelo de comunicación de los resultados de la fase anterior con el SOC. Para esto se considerarán diferentes alternativas (comunicaciones continuas, comunicaciones periódicas, informes o mensajes breves, uso de STIX o no, API, etc.).

Fase 4. Desarrollo de interfaz gráfica. Se diseñará y se desarrollará una interfaz gráfica para la herramienta que complemente a las comunicaciones con el SOC. El propósito de esta herramienta será la visualización de los datos de inteligencia de amenazas de ciberseguridad obtenidos de todo el proceso anterior. Para esto se considera usar plataformas de visualización como Kibana.

4.4.2 ESTIMACIÓN ECONÓMICA

Se ha realizado una estimación económica del coste del proyecto teniendo en cuenta que se trata de un proyecto individual llevado a cabo durante un contrato de prácticas. Cabe destacar que muchas de las herramientas utilizadas se usaron bajo una licencia de estudiante (gratuita), por lo que no se han incluido algunos precios de licencias empresariales que no han sido necesarias durante el desarrollo del proyecto.

Coste de recursos humanos

<i>Concepto</i>	<i>Precio / Hora</i>	<i>Horas totales</i>	<i>Coste total</i>
Becaria	20€	360	7200€

Tabla 2. Coste de recursos humanos

Coste de licencias y herramientas

<i>Concepto</i>	<i>Precio / Mes</i>	<i>Usuarios</i>	<i>Meses</i>	<i>Coste total</i>
Microsoft 365	12€	1	5	60€
OpenRouter	1€	1	5	5€
Total	-	-	-	65€

Tabla 3. Coste de licencias y herramientas

Coste de hardware

<i>Concepto</i>	<i>Precio</i>
Portátil DELL	569€

Tabla 4. Coste de hardware

Capítulo 5. EL SISTEMA/MODELO DESARROLLADO

En este capítulo se describe el proceso seguido para el diseño y el desarrollo de la herramienta. El funcionamiento de la herramienta completa se refleja en el siguiente esquema:

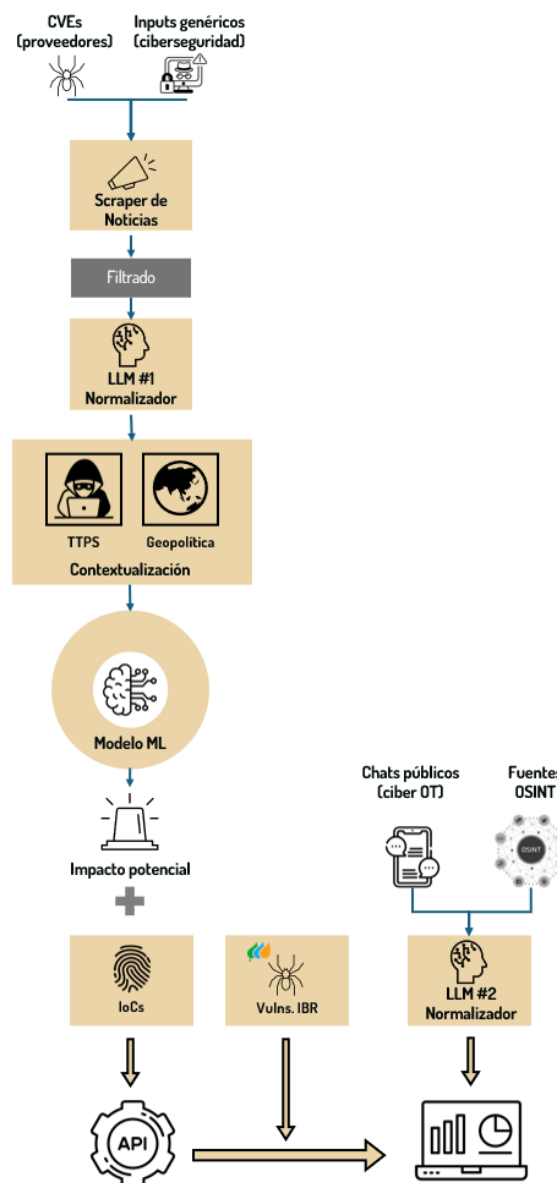


Ilustración 20. Funcionamiento de la herramienta

El primer elemento desarrollado fue el modelo de ML, por lo que será el primero que se explicará en este apartado.

5.1 MODELO DE MACHINE LEARNING PARA PREDICCIÓN DEL RIESGO

Se ha entrenado un modelo de red neuronal con los datos de ciberataques históricos (desde 2014 a 2024) y su contexto (geopolítico, MITRE, etc.). Previamente al entrenamiento, se define la variable objetivo del modelo: el “riesgo” o “impacto potencial” de cada ciberataque para Iberdrola.

5.1.1 DEFINICIÓN DEL RIESGO

El riesgo se define como el impacto potencial que un ciberataque externo puede tener sobre Iberdrola, considerando factores como la ubicación, el tipo de empresa afectada, el contexto geopolítico y las características de los atacantes, entre otros. La definición tradicional del riesgo como “probabilidad por impacto” se refleja en este proyecto como se explica a continuación:

- **Probabilidad.** Se ha entrenado el modelo con valores de “riesgo” basados, entre otros factores, en el impacto que tuvieron ciertos ciberataques reales sobre la empresa, por lo que la ocurrencia de ataques similares a otros con alto impacto supondría una mayor probabilidad de impacto que otros.
- **Impacto.** Se define como la magnitud de las consecuencias que un ciberataque puede generar sobre la organización, considerando tanto los efectos directos como los indirectos. En este proyecto, el impacto no se limita a una métrica técnica o económica aislada, sino que se construye como una combinación de múltiples factores, incluyendo el tipo de infraestructura afectada, el perfil y las capacidades del actor atacante, etc.

Dado que no existen datasets públicos que cuantifiquen el riesgo real asociado a cada ciberataque para una empresa concreta del sector energético como Iberdrola, se opta por

construir una métrica heurística de riesgo. Esta integra factores relacionados con la ubicación, tipo de industria, contexto geopolítico y técnicas utilizadas, todos ellos extraídos de fuentes como ACLED y MITRE ATT&CK. La métrica se diseña para aproximarse al nivel de impacto potencial que un ciberataque puede representar para un entorno OT en el sector energético, con foco en la generación eléctrica de Iberdrola España. No obstante, el cálculo de esta métrica puede adaptarse fácilmente a otros entornos OT generales mediante la parametrización de las variables clave según el contexto específico de cada organización, permitiendo así su aplicación en distintos sectores industriales con infraestructuras críticas.

Para establecer una base razonable sobre la que entrenar modelos predictivos y comparar riesgos entre eventos, y para definir una cuantificación inicial, se hacen diversas iteraciones sobre el dataset de ciberataques históricos para calibrar correctamente la cuantificación. La metodología seguida para la calibración se explica en el siguiente apartado.

La cuantificación del riesgo se basa en las ponderaciones de una serie de características de los ciberataques del dataset. En concreto:

- Ubicación (país en el que ocurre el ciberataque)
- Organización atacada
- Contexto geopolítico
- Contexto de técnicas y adversarios
- Tipo de ataque

Después de todas las iteraciones (ver Desarrollo de las pruebas), la cuantificación queda:

<i>Factor</i>	<i>Ponderación</i>	<i>Detalle de riesgo</i>
Ubicación (país en el que ocurre el ciberataque)	25%	10 – Si ocurre en España
		8 – Si ocurre en otro país donde opera Iberdrola
		5 – Si ocurre en un aliado de España

			3 – Si ocurre en la UE o en varios países a la vez (“MULTIPLE”) 0 – Si ocurre en un país sin relevancia geopolítica
Organización atacada	30%		10 – Si es Iberdrola 9.5 – Si es un proveedor clave 8 – Si es una <i>Utility</i> 5 – Si es otra empresa industrial 0 – Si es un sector irrelevante
Contexto geopolítico (*)	10%	1.0	Número de conflictos en país atacado en los 90 días previos al ciberataque Batallas (+2) • Protestas (+0.5) • Disturbios (+1) • Explosiones/violencia remota (+2) • Desarrollos estratégicos (+0.5) • Violencia contra civiles (+1.5)
		0.8	Número de conflictos en país del atacante en los 90 días previos al ciberataque • Batallas (+2) • Protestas (+0.5) • Disturbios (+1) • Explosiones/violencia remota (+2) • Desarrollos estratégicos (+0.5) • Violencia contra civiles (+1.5)

	1.5	Número de conflictos entre el país atacado y el país del atacante en los 90 días previos al ciberataque • Batallas (+2) • Protestas (+0.5) • Disturbios (+1) • Explosiones/violencia remota (+2) • Desarrollos estratégicos (+0.5) • Violencia contra civiles (+1.5)
	1.2	Número de conflictos entre aliados del país atacado y el país del atacante en los 90 días previos al ciberataque • Batallas (+2) • Protestas (+0.5) • Disturbios (+1) • Explosiones/violencia remota (+2) • Desarrollos estratégicos (+0.5) • Violencia contra civiles (+1.5)
	1.2	Número de conflictos entre el país atacado y aliados del país del atacante en los 90 días previos al ciberataque • Batallas (+2) • Protestas (+0.5) • Disturbios (+1) • Explosiones/violencia remota (+2) • Desarrollos estratégicos (+0.5)

		• Violencia contra civiles (+1.5)
Contexto de técnicas y adversarios	15%	10 – Si ataca a ICS 8 – Si el adversario es un actor de riesgo (de Rusia, Irán, Corea del Norte o China) 5 – Si no ataca a ICS pero usa TTPs específicas (**) 0 – Si no cumple ninguna de las anteriores
Tipo de ataque (***)	20%	10 – Ataque físico a infraestructura OT (“Physical Attack” en event_subtype) 6 - 'Message Manipulation', 'External Denial of Services', 'Internal Denial of Services', 'Data Attack' 0 - 'Exploitation of Sensors', 'Exploitation of End Host', 'Exploitation of Network Infrastructure', 'Exploitation of Application Server', 'Exploitation of Data in Transit'

Tabla 5. Cuantificación final del riesgo

(*) Las métricas de la función de cuantificación del riesgo parcial asociado al contexto geopolítico se han ajustado en función de las siguientes justificaciones:

<i>Contexto</i>	<i>Multiplicador</i>	<i>Justificación</i>
Número de conflictos en país atacado en los 90 días previos al ciberataque	1.0	Directamente en el país atacado. Riesgo directo.
Número de conflictos en país del atacante en los 90 días previos al ciberataque	0.8	En el país del actor atacante. Influye en su motivación o agresividad.
Número de conflictos entre el país atacado y el país del atacante en los 90 días previos al ciberataque	1.5	Conflicto bilateral entre ambos países. Tensión muy alta.
Número de conflictos entre aliados del país atacado y el país del atacante en los 90 días previos al ciberataque	1.2	Si el país atacado está en conflicto con aliados del actor. También genera riesgo indirecto.
Número de conflictos entre el país atacado y aliados del país del atacante en los 90 días previos al ciberataque	1.2	Si aliados del país atacado están en conflicto con el país del actor. Puede implicar ciberataques de represalia.

Tabla 6. Justificación de las métricas de riesgo geopolítico

<i>Tipo de conflicto</i>	<i>Definición</i>	<i>Riesgo base</i>
Protesta	Manifestaciones públicas en las que los participantes no recurren a la violencia,	0.5

	aunque puede usarse violencia contra ellos.	
Disturbios	Eventos violentos en los que manifestantes o multitudes participan en actos disruptivos.	1.0
Batallas	Interacciones violentas entre dos grupos armados organizados políticamente.	2.0
Explosiones (violencia remota)	Eventos violentos unilaterales en los que la herramienta utilizada para el conflicto crea una asimetría al eliminar la capacidad de respuesta del objetivo.	2.0
Violencia contra civiles	Eventos violentos en los que un grupo armado organizado inflige deliberadamente violencia a personas no combatientes y desarmadas.	1.5
Desarrollos estratégicos	Información contextualmente importante sobre las actividades de grupos violentos que no se registra como violencia política en sí misma, pero que puede desencadenar eventos futuros o contribuir a la dinámica política dentro de los estados o entre ellos.	0.5

Tabla 7. Detalle de métricas de riesgo geopolítico [9]

La asignación de riesgos base en la función responde a la gravedad y el impacto potencial de cada tipo de evento en la estabilidad geopolítica. Las batallas y explosiones tienen los valores más altos (2.0) porque representan violencia organizada o ataques con consecuencias

devastadoras, tanto humanas como estructurales. La violencia general se sitúa en un nivel intermedio (1.5) ya que, aunque no siempre es tan letal como una batalla, puede generar un clima de inseguridad persistente. Los disturbios (1.0) reflejan descontento social que puede escalar, pero no siempre lo hace, mientras que las protestas (0.5) suelen ser expresiones pacíficas de malestar, con menor impacto directo. Finalmente, los eventos estratégicos también tienen un riesgo bajo (0.5) porque, aunque pueden ser señales de tensión, no implican violencia directa ni consecuencias inmediatas. Esta jerarquía permite modelar el riesgo de forma proporcional a la amenaza real que representa cada tipo de evento.

(**) Técnicas de MITRE Enterprise Attack que suponen más riesgo para OT:

- *External Remote Services* (T1133). Los atacantes pueden aprovechar servicios remotos externos como VPNs y otros mecanismos de acceso para obtener acceso inicial y/o persistir dentro de una red. Esto es crítico para empresas industriales con accesos remotos.
- *Exploitation of Remote Services* (T1210). La explotación de vulnerabilidades en servicios remotos puede permitir a los atacantes acceder a sistemas internos.
- *Windows Management Instrumentation* (T1047). Los adversarios pueden abusar de *Windows Management Instrumentation* (WMI) para ejecutar comandos y cargas útiles maliciosas. WMI está diseñado para programadores y es la infraestructura para la gestión de datos y operaciones en sistemas Windows. WMI es una característica de administración que proporciona un entorno uniforme para acceder a los componentes del sistema Windows. Se usó en el ataque a la red eléctrica ucraniana en 2016.
- *OS Credential Dumping* (T1003). Obtención de credenciales almacenadas en el sistema operativo, como las de LSASS. Usado contra infraestructuras críticas en Europa por atacantes como EmberBear.
- *Masquerading* (T1036). Los atacantes pueden disfrazar archivos y procesos maliciosos como legítimos para evitar la detección
- *Lateral Tool Transfer* (T1570). Transferencia de herramientas maliciosas entre sistemas comprometidos para facilitar el movimiento lateral.

(***) Evaluación de los tipos de ataque. Para un entorno industrial o de energía, además de un ataque físico (*Physical Attack*) que supone un alto riesgo, se deben considerar también los siguientes eventos como de riesgo medio:

- *Data Attack*. La manipulación, destrucción o cifrado de datos puede tener un impacto significativo en la operación de sistemas críticos.
- *Internal Denial of Services*. La degradación o denegación de acceso a partes de la red IT desde dentro de la organización puede paralizar operaciones críticas.
- *External Denial of Services*. Aunque ejecutado desde fuera, puede impedir la comunicación con sistemas externos esenciales para la operación.

La expresión que calcula el riesgo de impacto final es:

$$Riesgo = W_1 \cdot R_{país} + W_2 \cdot R_{empresa} + W_3 \cdot R_{conflictos} + W_4 \cdot R_{MITRE} + W_5 \cdot R_{ataque}$$

Esta expresión se utiliza antes del entrenamiento del modelo para definir la variable objetivo del mismo. Se aplica sobre los datos iniciales de ciberataques con contexto, y luego ese dataset completo (con la variable de riesgo) se usa para entrenar el modelo.

La función de cuantificación final se define en Python así:

```
def cuantificar_riesgo(row, paises_criticos, proveedores, aliados,
paises_ue):

    #aliados = get_allies(row["country"],row["year"])

    if row["organization"] == "Iberdrola":
        riesgo_total = 10
    else:

        #riesgo = 0

        # 1. Ubicación
        ubi_ponderacion = 0.25
        ubi_riesgo = 0
        if row["country"] == "SPAIN":
            ubi_riesgo = 10
```

```

elif row["country"] in paises_criticos: # Países en los que
opera IBR
    ubi_riesgo = 8
elif row["country"] in aliados: # Países aliados de España
    ubi_riesgo = 5
elif row["country"] in paises_ue: # Países de la UE
    ubi_riesgo = 3
elif row["country"] == "MULTIPLE": # Varios países afectados
    ubi_riesgo = 3
ubi_sumando = ubi_ponderacion * ubi_riesgo

# 2. Organización atacada
org_ponderacion = 0.3
org_riesgo = 0
if row["organization"] == "Iberdrola":
    org_riesgo = 10
elif row["organization"] in proveedores:
    org_riesgo = 9.5
elif row["industry"] == "Utilities":
    org_riesgo = 8
elif row["industry"] == "Mining, Quarrying, and Oil and Gas
Extraction":
    org_riesgo = 5
org_sumando = org_ponderacion * org_riesgo

# 3. Técnicas y adversarios
ttps_ponderacion = 0.15
ttps_riesgo = 0
high_risk_actors = {"RUSSIA", "CHINA", "IRAN", "NORTH KOREA"}
if row["Attacks ICS"] == "Yes":
    ttps_riesgo = 10
elif row["actor_country"] in high_risk_actors:
    ttps_riesgo = 8
elif row["contains_risky_ttps"] == 1 or row["actor_country"] ==
"UNDETERMINED":
    ttps_riesgo = 5
ttps_sumando = ttps_ponderacion * ttps_riesgo

# 4. Geopolítica
geopol_ponderacion = 0.1
geopol_riesgo = 0
geopol_riesgo = calcular_riesgo_geopolitico(row)
geopol_sumando = geopol_ponderacion * geopol_riesgo

# 5. Tipo de ataque
ataque_ponderacion = 0.2
ataque_riesgo = 0
if row["event_subtype"] == "Physical Attack" or "Physical Attack"
in row["event_subtype"]:
    ataque_riesgo = 10
elif "Message Manipulation" in row["event_subtype"]:
    ataque_riesgo = 6
elif "External Denial of Services" in row["event_subtype"]:
    ataque_riesgo = 6
elif "Internal Denial of Services" in row["event_subtype"]:

```

```
        ataque_riesgo = 6
    elif "Data Attack" in row["event_subtype"]:
        ataque_riesgo = 6

    ataque_sumando = ataque_ponderacion * ataque_riesgo

    # Sumar todo
    riesgo_total = ubi_sumando + org_sumando + ttps_sumando +
geopol_sumando + ataque_sumando
    return round(riesgo_total, 4)
```

Fragmento de código 1. Función de cuantificación del riesgo final

Donde la función `calcular_riesgo_geopolitico` es:

```
def calcular_riesgo_geopolitico(row):
    # Riesgo base por tipo de evento (0.5 a 2.0)
    riesgos_base = {
        "protests": 0.5,
        "riots": 1.0,
        "battles": 2.0,
        "explosions": 2.0,
        "violence": 1.5,
        "strategic": 0.5
    }

    # Multiplicadores según el contexto geopolítico
    contextos = {
        "_country": 1.0,                # País atacado
        "_actor_country": 0.8,          # País del actor
        "_between": 1.5,                # Conflictos entre ambos
        "_vs_allies_of_actor": 1.2,     # Atacado contra aliados
        del actor
        "_allies_of_country_vs_actor": 1.2 # Aliados del atacado
        contra actor
    }

    riesgo_total = 0

    for tipo, base in riesgos_base.items():
        for sufixo, mult in contextos.items():
            col = f"{tipo}{sufijo}"
            if col in row and not pd.isna(row[col]):
                riesgo_total += row[col] * base * mult

    # Riesgo entre 0 y 10
    riesgo_geopolitico = min(riesgo_total, 10)

    return riesgo_geopolitico
```

Fragmento de código 2. Función de cálculo del riesgo parcial asociado al contexto geopolítico

5.1.2 PRUEBAS DE CALIBRACIÓN

Para definir la cuantificación del riesgo de cada ciberataque y su contexto se realizaron una serie de pruebas de calibración en varias iteraciones. Cada iteración consta de unas pruebas con casos sintéticos y otras con casos reales. En total, se realizaron **9 iteraciones** (numeradas de 0 a 8). En este apartado se describen principalmente los resultados finales (tras las 9 iteraciones), puesto que los cambios realizados a partir de los resultados de los casos reales afectaron a la siguiente iteración de los casos sintéticos primero.

5.1.2.1 Casos sintéticos

La primera fase de una iteración es una evaluación de **casos sintéticos**: ciberataques no reales que representan casuísticas específicas. Una vez definida la cuantificación del riesgo, se crean seis casos sintéticos:

1. Peor caso #1
2. Peor caso #2
3. Mejor caso #1
4. Mejor caso #2
5. Caso intermedio #1
6. Caso intermedio #2

Se evalúa si la expresión matemática definida previamente clasifica estos casos como tal; es decir, si el peor caso se clasifica con un riesgo 10 o cercano a 10 (al menos superior a 8.5), si el mejor caso se clasifica con un riesgo cercano a 0 (al menos inferior a 1.5), y si el caso intermedio se clasifica con un riesgo alrededor de 5 (entre 4.5 y 5.5).

Peor caso #1

El primer “worst case” o peor caso sintético representa un ataque a Iberdrola en España de tipo físico, por parte de un atacante ruso que emplea técnicas incluidas en el marco de MITRE para sistemas de control industrial. Se define como se muestra a continuación:

```
import pandas as pd
from datetime import datetime
```

```
# Diccionario con valores del worst case en el orden exacto solicitado
worst_case_dict = {
    'Unnamed: 0': 0,
    'event_date': datetime(2025, 5, 1),
    'year': 2025,
    'actor': 'APT28',
    'actor_type': 'State-sponsored',
    'organization': 'Iberdrola',
    'industry_code': '2211',
    'industry': 'Utilities',
    'motive': 'Sabotage critical infrastructure',
    'event_type': 'Disruptive',
    'event_subtype': 'Physical Attack',
    'description': 'A major physical sabotage affected Iberdrola's OT
systems.',
    'source_url': 'https://fakenews.org/worst-case-attack',
    'country': 'SPAIN',
    'actor_country': 'RUSSIA',
    'cve_assoc': 'CVE-2025-9999',
    'protests_country': 10,
    'violence_country': 10,
    'strategic_country': 10,
    'riots_country': 10,
    'battles_country': 10,
    'explosions_country': 10,
    'protests_actor_country': 10,
    'violence_actor_country': 10,
    'strategic_actor_country': 10,
    'riots_actor_country': 10,
    'battles_actor_country': 10,
    'explosions_actor_country': 10,
    'TTPs': 'OS Credential Dumping, Masquerading, Lateral Tool Transfer',
    'Attacks ICS': 'Yes',
    'protests_between': 10,
    'violence_between': 10,
    'strategic_between': 10,
    'riots_between': 10,
    'battles_between': 10,
    'explosions_between': 10,
    'protests_allies_of_country_vs_actor': 10,
    'protests_country_vs_allies_of_actor': 10,
    'violence_allies_of_country_vs_actor': 10,
    'violence_country_vs_allies_of_actor': 10,
    'strategic_allies_of_country_vs_actor': 10,
    'strategic_country_vs_allies_of_actor': 10,
    'riots_allies_of_country_vs_actor': 10,
    'riots_country_vs_allies_of_actor': 10,
    'battles_allies_of_country_vs_actor': 10,
    'battles_country_vs_allies_of_actor': 10,
    'explosions_allies_of_country_vs_actor': 10,
    'explosions_country_vs_allies_of_actor': 10,
    'contains_risky_ttps': 1
}
```

```
# Crear DataFrame
df_worst_case = pd.DataFrame([worst_case_dict])

df_worst_case.head()
```

Fragmento de código 3. Definición del “worst case 1”

La función de cuantificación asigna un riesgo de valor 10 a este caso:

```
cuantificar_riesgo(df_worst_case.iloc[0], paises_criticos, proveedores, aliados, paises_ue)
```

10

Ilustración 21. Cuantificación del riesgo del peor caso #1

Lo cual es de esperar porque al ser un ataque a Iberdrola el riesgo es 10 independientemente de los otros factores.

En la primera iteración (iteración 0), no se había incluido esta condición, por lo que los ataques reales a Iberdrola (del dataset histórico) aparecían con un riesgo menor (en concreto de 6.70 o inferior) en las pruebas con casos reales. Esto se debe a que los ataques ocurridos en el pasado incluidos en el dataset fueron ataques de tipo “Data Attack” que no tuvieron un impacto físico sobre los activos OT de la empresa. A pesar de ello, se ha considerado que cualquier ciberataque a la empresa en conjunto debe suponer un riesgo de impacto elevado independientemente de si tiene repercusiones físicas o no.

Peor caso #2

El segundo “worst case” o peor caso sintético representa un ciberataque con consecuencias físicas a un proveedor crítico de Iberdrola (Siemens Energy) por parte de un atacante norcoreano (Lazarus Group) que aparece en el MITRE ATT&CK ICS y utiliza técnicas definidas como peligrosas. En este caso se simulan conflictos geopolíticos intensos entre España (país atacado) y Corea del Norte (país del atacante), representados por altos valores (todos en 10) en categorías como protestas, violencia, batallas, etc. Además, se reflejan tensiones indirectas mediante conflictos entre aliados de España y Corea del Norte, y viceversa, lo que amplifica el riesgo geopolítico global del evento.

```
# Definimos un proveedor crítico ficticio
proveedor_critico = "Siemens Energy"
```

```
# Worst case con proveedor crítico en Utilities
worst_case_proveedor_dict = {
    'Unnamed: 0': 1,
    'event_date': datetime(2025, 5, 1),
    'year': 2025,
    'actor': 'Lazarus Group',
    'actor_type': 'State-sponsored',
    'organization': proveedor_critico,
    'industry_code': '2211',
    'industry': 'Utilities',
    'motive': 'Destabilize national energy systems',
    'event_type': 'Exploitative',
    'event_subtype': 'Physical Attack, Data Attack',
    'description': 'The Lazarus Group launched a coordinated physical and
cyber attack on Siemens Energy control systems.',
    'source_url': 'https://fakenews.org/high-impact-siemens-attack',
    'country': 'SPAIN',
    'actor_country': 'NORTH KOREA',
    'cve_assoc': 'CVE-2025-2222',
    'protests_country': 10,
    'violence_country': 10,
    'strategic_country': 10,
    'riots_country': 10,
    'battles_country': 10,
    'explosions_country': 10,
    'protests_actor_country': 10,
    'violence_actor_country': 10,
    'strategic_actor_country': 10,
    'riots_actor_country': 10,
    'battles_actor_country': 10,
    'explosions_actor_country': 10,
    'TTPs': 'OS Credential Dumping, External Remote Services, Lateral
Tool Transfer',
    'Attacks ICS': 'Yes',
    'protests_between': 10,
    'violence_between': 10,
    'strategic_between': 10,
    'riots_between': 10,
    'battles_between': 10,
    'explosions_between': 10,
    'protests_allies_of_country_vs_actor': 10,
    'protests_country_vs_allies_of_actor': 10,
    'violence_allies_of_country_vs_actor': 10,
    'violence_country_vs_allies_of_actor': 10,
    'strategic_allies_of_country_vs_actor': 10,
    'strategic_country_vs_allies_of_actor': 10,
    'riots_allies_of_country_vs_actor': 10,
    'riots_country_vs_allies_of_actor': 10,
    'battles_allies_of_country_vs_actor': 10,
    'battles_country_vs_allies_of_actor': 10,
    'explosions_allies_of_country_vs_actor': 10,
    'explosions_country_vs_allies_of_actor': 10,
    'contains_risky_ttps': 1
}
```

Fragmento de código 4. Definición del "worst case 2"

La función de cuantificación asigna un riesgo de valor 9.05 a este caso:

```
cuantificar_riesgo(df_worst_case_proveedor.iloc[0], paises_criticos, proveedores, aliados, paises_ue)

9.05
```

Ilustración 22. Cuantificación del riesgo del peor caso #2

Que se considera como válido al ser superior a 8.5.

Mejor caso #1

El primer mejor caso representa un ciberataque de tipo “Message Manipulation” (robo o suplantación de una cuenta de redes sociales o suplantación de dominios web) a una ONG pequeña de Costa Rica, por parte de un hacktivista desconocido que no aparece en el marco de MITRE para sistemas de control industrial y tampoco emplea técnicas consideradas peligrosas para OT.

```
best_case_dict = {
    'Unnamed: 0': 2,
    'event_date': datetime(2025, 5, 1),
    'year': 2025,
    'actor': 'Anonymous Hacktivist',
    'actor_type': 'Hacktivist',
    'organization': 'Local NGO',
    'industry_code': '8139',
    'industry': 'Other Services (except Public Administration)',
    'motive': 'Political protest',
    'event_type': 'Disruptive',
    'event_subtype': 'Message Manipulation',
    'description': 'Defacement of a small NGO website in Costa Rica by a
hacktivist group.',
    'source_url': 'https://example.com/ngo-defacement',
    'country': 'COSTA RICA',
    'actor_country': 'UNDETERMINED',
    'cve_assoc': '',
    'protests_country': 0,
    'violence_country': 0,
    'strategic_country': 0,
    'riots_country': 0,
    'battles_country': 0,
    'explosions_country': 0,
    'protests_actor_country': 0,
    'violence_actor_country': 0,
    'strategic_actor_country': 0,
    'riots_actor_country': 0,
```

```
'battles_actor_country': 0,  
'explosions_actor_country': 0,  
'TTPs': '',  
'Attacks ICS': 'No',  
'protests_between': 0,  
'violence_between': 0,  
'strategic_between': 0,  
'riots_between': 0,  
'battles_between': 0,  
'explosions_between': 0,  
'protests_allies_of_country_vs_actor': 0,  
'protests_country_vs_allies_of_actor': 0,  
'violence_allies_of_country_vs_actor': 0,  
'violence_country_vs_allies_of_actor': 0,  
'strategic_allies_of_country_vs_actor': 0,  
'strategic_country_vs_allies_of_actor': 0,  
'riots_allies_of_country_vs_actor': 0,  
'riots_country_vs_allies_of_actor': 0,  
'battles_allies_of_country_vs_actor': 0,  
'battles_country_vs_allies_of_actor': 0,  
'explosions_allies_of_country_vs_actor': 0,  
'explosions_country_vs_allies_of_actor': 0,  
'contains_risky_ttps': 0  
}
```

Fragmento de código 4. Definición del "best case 1"

La función de cuantificación asigna un riesgo de valor 1.2 a este caso:

```
cuantificar_riesgo(df_best_case.iloc[0], paises_criticos, proveedores, aliados, paises_ue)  
np.float64(1.2)
```

Ilustración 23. Cuantificación del riesgo del mejor caso 1

Que se considera como válido para el mejor caso, al ser inferior a 1.5.

Mejor caso #2

El segundo mejor caso representa un ciberataque de tipo “Data Attack” a una universidad de Nueva Zelanda por parte de un actor desconocido de tipo “Hobbyist” que tenía intención de modificar datos relativos a las calificaciones de los alumnos.

```
best_case_dict_2 = {  
    'Unnamed: 0': 3,  
    'event_date': datetime(2025, 5, 15),  
    'year': 2025,  
    'actor': 'Anonymous Hobbyist',  
    'actor_type': 'Hobbyist',  
    'organization': 'Local University',  
}
```

```
'industry_code': '61',
'industry': 'Educational Services',
'motive': 'Protest',
'event_type': 'Disruptive',
'event_subtype': 'Data Attack',
'description': 'Manipulation of grades data in a university website
by an unknown group.',
'source_url': 'https://example.com/university-defacement',
'country': 'NEW ZEALAND',
'actor_country': 'UNDETERMINED',
'cve_assoc': '',
'protests_country': 2,
'violence_country': 0,
'strategic_country': 0,
'riots_country': 0,
'battles_country': 0,
'explosions_country': 0,
'protests_actor_country': 0,
'violence_actor_country': 0,
'strategic_actor_country': 0,
'riots_actor_country': 0,
'battles_actor_country': 0,
'explosions_actor_country': 0,
'TTPs': '',
'Attacks ICS': 'No',
'protests_between': 0,
'violence_between': 0,
'strategic_between': 0,
'riots_between': 0,
'battles_between': 0,
'explosions_between': 0,
'protests_allies_of_country_vs_actor': 0,
'protests_country_vs_allies_of_actor': 0,
'violence_allies_of_country_vs_actor': 0,
'violence_country_vs_allies_of_actor': 0,
'strategic_allies_of_country_vs_actor': 0,
'strategic_country_vs_allies_of_actor': 0,
'riots_allies_of_country_vs_actor': 0,
'riots_country_vs_allies_of_actor': 0,
'battles_allies_of_country_vs_actor': 0,
'battles_country_vs_allies_of_actor': 0,
'explosions_allies_of_country_vs_actor': 0,
'explosions_country_vs_allies_of_actor': 0,
'contains_risky_ttps': 0
}
```

Fragmento de código 5. Definición del "best case 2"

La función de cuantificación asigna un valor de riesgo de 1.3 a este caso, que se considera válido como “best case”, al ser inferior a 1.5.

```
cuantificar_riesgo(df_best_case.iloc[0], paises_criticos, proveedores, aliados, paises_ue)
```

1.3

Ilustración 24. Cuantificación del riesgo del mejor caso 2

Caso intermedio #1

El primer caso intermedio representa un ciberataque de denegación de servicio a la parte IT de una empresa de Portugal del sector petroquímico, por parte de un actor de China (uno de los países definidos como actores de riesgo) que no aparece en el marco de MITRE para sistemas de control industrial.

```
intermediate_case_dict = {
    'Unnamed: 0': 3,
    'event_date': datetime(2025, 5, 10),
    'year': 2025,
    'actor': 'Cyber Criminal Group',
    'actor_type': 'Criminal',
    'organization': 'PetroReg Co.',
    'industry_code': '2211',
    'industry': 'Mining, Quarrying, and Oil and Gas Extraction',
    'motive': 'Financial gain',
    'event_type': 'Disruptive',
    'event_subtype': 'External Denial of Services',
    'description': 'DDoS attack on a regional energy provider in
Portugal.',
    'source_url': 'https://example.com/ddos-portugal',
    'country': 'PORTUGAL',
    'actor_country': 'CHINA',
    'cve_assoc': '',
    'protests_country': 0,
    'violence_country': 0,
    'strategic_country': 0,
    'riots_country': 0,
    'battles_country': 0,
    'explosions_country': 0,
    'protests_actor_country': 0,
    'violence_actor_country': 0,
    'strategic_actor_country': 0,
    'riots_actor_country': 0,
    'battles_actor_country': 0,
    'explosions_actor_country': 0,
    'TTPs': '',
    'Attacks ICS': 'No',
    'protests_between': 0,
    'violence_between': 0,
```

```
{
  'strategic_between': 0,
  'riots_between': 0,
  'battles_between': 0,
  'explosions_between': 0,
  'protests_allies_of_country_vs_actor': 0,
  'protests_country_vs_allies_of_actor': 0,
  'violence_allies_of_country_vs_actor': 0,
  'violence_country_vs_allies_of_actor': 0,
  'strategic_allies_of_country_vs_actor': 0,
  'strategic_country_vs_allies_of_actor': 0,
  'riots_allies_of_country_vs_actor': 0,
  'riots_country_vs_allies_of_actor': 0,
  'battles_allies_of_country_vs_actor': 0,
  'battles_country_vs_allies_of_actor': 0,
  'explosions_allies_of_country_vs_actor': 0,
  'explosions_country_vs_allies_of_actor': 0,
  'contains_risky_ttps': 0
}
```

Fragmento de código 6. Definición del “caso intermedio #1”

La función de cuantificación asigna un riesgo de valor 5.75 a este caso:

```
cuantificar_riesgo(df_intermediate_case.iloc[0], paises_criticos, proveedores, aliados, paises_ue)
```

5.75

Ilustración 25. Cuantificación del riesgo del caso intermedio #1

Que es ligeramente superior al valor de 5.5 que se definió, pero se acepta como válido como caso intermedio.

Caso intermedio #2

El segundo caso intermedio representa una brecha de datos de una empresa eléctrica de Colombia por parte de un adversario iraní.

```
intermediate_case_2_dict = {
  'Unnamed: 0': 4,
  'event_date': datetime(2025, 5, 15),
  'year': 2025,
  'actor': 'State-sponsored Group',
  'actor_type': 'State',
  'organization': 'Electrica Regional',
  'industry_code': '21',
  'industry': 'Utilities',
  'motive': 'Espionage',
  'event_type': 'Disruptive',
  'event_subtype': 'Data Attack',
}
```

```
'description': 'Data exfiltration from a regional electricity
provider in Colombia by a suspected state-sponsored group.',
'source_url': 'https://example.com/data-exfiltration-colombia',
'country': 'COLOMBIA',
'actor_country': 'IRAN',
'cve_assoc': '',
'protests_country': 1,
'violence_country': 0,
'strategic_country': 0,
'riots_country': 0,
'battles_country': 0,
'explosions_country': 0,
'protests_actor_country': 0,
'violence_actor_country': 0,
'strategic_actor_country': 0,
'riots_actor_country': 0,
'battles_actor_country': 0,
'explosions_actor_country': 0,
'TTPs': '',
'Attacks ICS': 'Yes',
'protests_between': 0,
'violence_between': 0,
'strategic_between': 0,
'riots_between': 0,
'battles_between': 0,
'explosions_between': 0,
'protests_allies_of_country_vs_actor': 0,
'protests_country_vs_allies_of_actor': 0,
'violence_allies_of_country_vs_actor': 0,
'violence_country_vs_allies_of_actor': 0,
'strategic_allies_of_country_vs_actor': 0,
'strategic_country_vs_allies_of_actor': 0,
'riots_allies_of_country_vs_actor': 0,
'riots_country_vs_allies_of_actor': 0,
'battles_allies_of_country_vs_actor': 0,
'battles_country_vs_allies_of_actor': 0,
'explosions_allies_of_country_vs_actor': 0,
'explosions_country_vs_allies_of_actor': 0,
'contains_risky_ttps': 1
}
```

Fragmento de código 7. Definición del “caso intermedio #2”

Para validar este caso se hizo un cambio (ya reflejado en los demás casos) en la tercera iteración (iteración 2), porque se había asignado un valor de riesgo parcial de 9 si el país del actor era de alto riesgo (Rusia, Irán, China, Corea del Norte), y se obtenía un valor superior a 6 para este caso intermedio. Se bajó el riesgo parcial indicado a 7, y así se obtiene un valor de 5.15.

```
cuantificar_riesgo(df_intermediate_case_2.iloc[0], paises_criticos, proveedores, aliados, paises_ue)
```

5.15

Ilustración 26. Cuantificación del riesgo del caso intermedio #2

5.1.2.2 Casos reales

La segunda fase es una evaluación de **casos reales**. Se tomaron como referencia varios ciberataques relevantes para la empresa y se analizaron los valores de riesgo que tenían en cada prueba.

- **Ciberataque a Vestas Wind Systems (2021).** En 2021, Vestas Wind Systems sufrió un ataque de ransomware que comprometió sus sistemas informáticos y resultó en una filtración de datos confidenciales. El impacto principal fue la interrupción de sus operaciones internas y la exposición de información sensible de empleados y socios comerciales, lo que afectó tanto su reputación como su cadena de suministro en el sector de energía renovable. Tuvo un impacto sobre las operaciones de parques eólicos de Iberdrola con elementos provistos por Vestas porque el proveedor contaba con accesos remotos a algunos de estos parques y se tuvieron que cortar esas comunicaciones.
- **Ciberataques a Schneider Electric (2018, 2023, 2024)**
 - Malware en drives USB de dispositivos Schneider Electric para sistemas de energía solar (2018). En 2018, se descubrió que algunas unidades USB distribuidas con productos de Schneider Electric contenían malware, lo que generó un riesgo crítico para los sistemas de energía solar donde fueron utilizadas. Este incidente puso en evidencia las vulnerabilidades en la cadena de suministro y planteó una amenaza directa para infraestructuras energéticas que utilizaban estos dispositivos.
 - Brecha de datos a Schneider Electric a través de ransomware Clon (2023). En 2023, Schneider Electric fue víctima del grupo de ransomware Clon, que explotó una vulnerabilidad en el software MOVEit Transfer. El ataque resultó en la filtración de datos corporativos y personales, afectando a empleados y

clientes. La brecha tuvo un impacto reputacional y operativo, además de exponer riesgos legales por la posible violación de regulaciones de protección de datos.

- Brecha de datos a Schneider Electric a través de ransomware Cactus (2024).

En 2024, Schneider sufrió un ataque de ransomware Cactus en su plataforma de desarrollo. Esto supuso el robo de 40 GB de datos del servidor JIRA de la empresa

- **Ciberataque WannaCry a Iberdrola (2017).** En mayo de 2017, Iberdrola fue una de las muchas empresas afectadas por el ataque de ransomware WannaCry. Este ataque explotó una vulnerabilidad en el protocolo SMB de Windows, conocida como EternalBlue, para propagarse rápidamente a través de las redes. WannaCry cifró los archivos de los sistemas infectados y exigió un rescate en Bitcoin para su liberación. En Iberdrola, el ataque causó interrupciones significativas en sus operaciones, afectando la disponibilidad de sistemas críticos y obligando a la empresa a implementar medidas de contención y recuperación de emergencia.
- **Industroyer (2016).** Industroyer, también conocido como CrashOverride, es un malware diseñado específicamente para atacar sistemas de control industrial (ICS). En diciembre de 2016, Industroyer fue utilizado para causar un apagón en Ucrania, atacando la red eléctrica de Kiev. El malware estaba diseñado para interactuar con protocolos industriales como IEC 60870-5-104 u OPC DA, utilizado en sistemas de control de subestaciones eléctricas. Este ataque demostró la capacidad de los actores maliciosos para causar interrupciones físicas significativas mediante el uso de malware dirigido a infraestructuras críticas.
- **Industroyer (2022).** En abril de 2022, una variante del malware Industroyer, conocida como Industroyer2, fue utilizada en un ataque contra la red eléctrica de Ucrania. Este nuevo ataque empleó una versión más sofisticada del malware original, modificada para atacar dispositivos específicos dentro de la infraestructura de control industrial. Industroyer2 permitía a los atacantes definir objetivos más personalizables y específicos dentro de la red atacada.

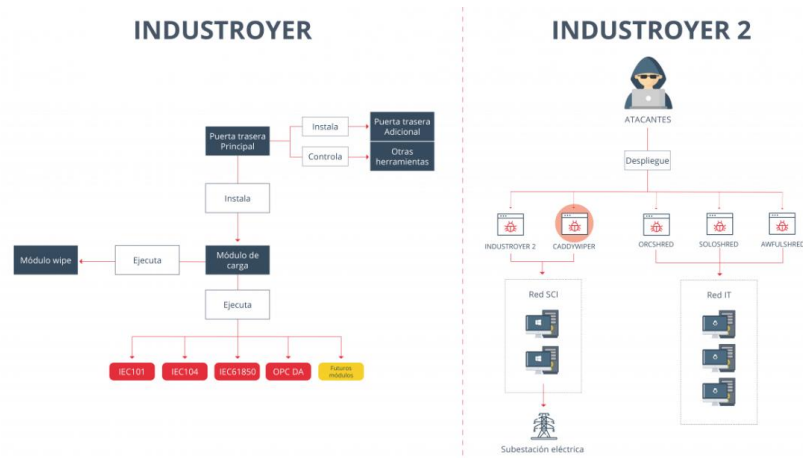


Ilustración 27. Comparativa Industroyer e Industroyer2 [9]

Se espera que todos aquellos casos reales que tuvieran un impacto sobre Iberdrola (tanto los directos a la empresa como aquellos a proveedores que impactaran sus operaciones) tengan un valor de riesgo de 8 o superior.

5.1.2.3 Desarrollo de las pruebas

A continuación, se detallan las diferentes pruebas realizadas (entre las cuales varían las ponderaciones de cada factor a tener en cuenta para el riesgo) y sus resultados.

Como ya se ha explicado, el modelo de cuantificación del riesgo se basa en:

<i>Factor</i>	<i>Ponderación</i>	<i>Valor (0-10) según:</i>
Ubicación	X%	España, país Iberdrola, UE, etc.
Organización atacada	X%	Utilities, Mining, etc.
Geopolítica	X%	Tipos de conflicto (en país atacado, atacante, aliados)
Técnicas y adversarios	X%	Atacante en MITRE ATT&CK ICS, técnicas de Enterprise aplicables, etc.
Tipo de ataque	X%	Físico, de datos, etc.

Tabla 8. Resumen de la cuantificación del riesgo

La ponderación (y los valores de 0-10 asignados para cada factor) se elige durante una serie de pruebas de calibración iterativas. En este apartado se describen los cambios realizados en las 9 iteraciones sobre las ponderaciones y los valores de cada factor.

Iteración 0

Inicialmente, se parte del siguiente código de cuantificación:

```
def cuantificar_riesgo(row, paises_criticos, proveedores, aliados,
paises_ue):
    riesgo = 0

    # 1. Ubicación
    ubi_ponderacion = 0.2
    ubi_riesgo = 0
    if row["country"] == "SPAIN":
        ubi_riesgo = 10
    elif row["country"] in paises_criticos: # Países en los que opera
IBR
        ubi_riesgo = 8
    elif row["country"] in aliados: # Países aliados de España
        ubi_riesgo = 5
    elif row["country"] in paises_ue: # Países de la UE
        ubi_riesgo = 3
    ubi_sumando = ubi_ponderacion * ubi_riesgo

    # 2. Organizaciónn atacada
    org_ponderacion = 0.2
    org_riesgo = 0
    if row["organization"] == "Tberdrola":
        org_riesgo = 10
    elif row["organization"] in proveedores:
        org_riesgo = 9
    elif row["industry"] == "Utilities":
        org_riesgo = 8
    elif row["industry"] == "Mining, Quarrying, and Oil and Gas
Extraction":
        org_riesgo = 5
    org_sumando = org_ponderacion * org_riesgo

    # 3. Técnicas y adversarios
    ttps_ponderacion = 0.2
    ttps_riesgo = 0
    if row["Attacks ICS"] == "Yes":
        ttps_riesgo = 10
    elif row["contains_risky_ttps"] == 1:
        ttps_riesgo = 5
    ttps_sumando = ttps_ponderacion * ttps_riesgo

    # 4. Geopolítica
```

```
geopol_ponderacion = 0.2
geopol_riesgo = 0
geopol_riesgo = calcular_riesgo_geopolitico(row)
geopol_sumando = geopol_ponderacion * geopol_riesgo

# 5. Tipo de ataque
ataque_ponderacion = 0.2
ataque_riesgo = 0
if row["event_subtype"] == "Physical Attack":
    ataque_riesgo = 10
elif "Message Manipulation" in row["event_subtype"]:
    ataque_riesgo = 6
elif "External Denial of Services" in row["event_subtype"]:
    ataque_riesgo = 6
elif "Internal Denial of Services" in row["event_subtype"]:
    ataque_riesgo = 6
elif "Data Attack" in row["event_subtype"]:
    ataque_riesgo = 6

ataque_sumando = ataque_ponderacion * ataque_riesgo

# Sumar todo
riesgo_total = ubi_sumando + org_sumando + ttps_sumando +
geopol_sumando + ataque_sumando
return round(riesgo_total, 4)
```

Fragmento de código 8. Código de cuantificación (iteración 0)

Con esta cuantificación, el caso sintético “peor caso #1”, en el cual la organización afectada por un ciberataque industrial era Iberdrola, se puntuaba con un valor de riesgo de 6.5 / 10, y el caso real de WannaCry que afectó a la empresa se puntuaba con un valor de 6.4 / 10. Se decide entonces añadir a la cuantificación una condición a los valores relativos a la organización atacada: si es Iberdrola, el factor ahora vale 10 / 10. Se valida así el “peor caso #1” sintético (y los casos reales con impacto en Iberdrola).

Iteración 1

Se parte del código de cuantificación inicial con la nueva condición de Iberdrola mencionada.

```
def cuantificar_riesgo(row, paises_criticos, proveedores, aliados,
paises_ue):

    #aliados = get_allies(row["country"],row["year"])

    if row["organization"] == "Iberdrola":
        riesgo_total = 10
```

```
else:
    # código inicial
```

Fragmento de código 9. Cambios en código de cuantificación (iteración 1)

Ahora todo ataque a Iberdrola aparece con cuantificación riesgo 10 / 10. Por otro lado, el “caso intermedio #1” sintético queda con un riesgo asignado de 4.2 / 10, que no cumple con los umbrales definidos para los casos intermedios (entre 4.5 y 5.5 aproximadamente). Este primer caso intermedio definía un ataque por parte de un actor “de riesgo” (China, de entre Rusia, Irán, China y Corea del Norte), por lo que esto debería reflejarse en la puntuación. Se añade una nueva condición en los valores de “técnicas y adversarios”: si el atacante es Rusia, Corea del Norte, China o Irán, el factor vale 7 (antes no se valoraba esto). Se valida así el caso intermedio sintético.

Iteración 2

Se parte del código con el cambio anterior, y ahora el riesgo de todo ataque realizado por uno de esos países se refleja en la puntuación como es debido. En este momento las ponderaciones son las siguientes:

<i>Factor</i>	<i>Ponderación</i>
Ubicación	20%
Organización atacada	20%
Geopolítica	20%
Técnicas y adversarios	20%
Tipo de ataque	20%

Tabla 9. Ponderaciones en iteración 2

Se observa que los casos reales elegidos en los que las organizaciones atacadas eran proveedores de Iberdrola obtienen un riesgo entre 6 y 7 sobre 10, y el riesgo que obtienen otros ataques ajenos a Iberdrola y sus proveedores es mayor. Se decide hacer una modificación en las ponderaciones para asignar mayor peso a la organización atacada.

Iteración 3

Las nuevas ponderaciones son:

<i>Factor</i>	<i>Ponderación</i>
Ubicación	20%
Organización atacada	25%
Geopolítica	15%
Técnicas y adversarios	20%
Tipo de ataque	20%

Tabla 10. Ponderaciones en iteración 3

Se quita un 5% de peso a la parte geopolítica y se añade a la parte relativa a la organización (que es donde se define un mayor valor de riesgo si la organización atacada es un proveedor de Iberdrola). Se decide quitar peso a la parte geopolítica porque muchos de los ataques que quedaban con un riesgo alto ocurrían en países con grandes conflictos internos que no afectaban a España ni a otros países en los que opera Iberdrola.

Se obtienen valores superiores para los casos reales de Vestas y Schneider, pero siguen siendo inferiores a 8 en su mayoría. En concreto:

- Los ataques a Schneider tienen asignados unos valores de riesgo de 7.70, 8.10 y 7.05, pero el primero debería tener un riesgo mayor al tratarse de un caso de malware en USBs para sistemas de energía solar.
- El ataque a Vestas que afectó a la operación de Iberdrola (porque hubo que hacer ciertas paradas por precaución), tiene asignado un valor de solo 7.05.

Iteración 4

<i>Factor</i>	<i>Ponderación</i>

Ubicación	25%
Organización atacada	30%
Geopolítica	10%
Técnicas y adversarios	15%
Tipo de ataque	20%

Tabla 11. Ponderaciones en iteración 4

Se opta por dar más peso a la organización atacada para poder validar los casos reales de ataques a proveedores, y también se da más peso a la ubicación, para poder reflejar más fácilmente que los riesgos de un ataque a un país aliado de España son mayores e implican más peso que los conflictos geopolíticos internos en países que no son tan relevantes para la operación de Iberdrola (como algunos en África y Asia).

Con estos cambios, los casos reales de Schneider y Vestas quedan:

- Los ataques a Schneider tienen puntuaciones 7.80, 8.10, 7.80.
- El ataque a Vestas ahora tiene asignado un valor de 7.80.

Iteración 5

Se añade un mayor valor dentro de “técnicas y adversarios” si el atacante es indefinido (“UNDETERMINED”), puesto que hay muchos casos de ataques con atacantes anónimos o desconocidos que tienen o pueden tener un gran impacto. Esto es el caso del ataque a Schneider con dispositivos USB, que al ser un ataque puramente físico debería reflejar un impacto mayor que los otros dos al mismo fabricante.

```
# 3. Técnicas y adversarios
ttps_ponderacion = 0.15
ttps_riesgo = 0
high_risk_actors = {"RUSSIA", "CHINA", "IRAN", "NORTH KOREA"}
if row["Attacks ICS"] == "Yes":
    ttps_riesgo = 10
elif row["actor_country"] in high_risk_actors:
    ttps_riesgo = 8
```

```
elif row["contains_risky_ttps"] == 1 or row["actor_country"] ==  
"UNDETERMINED":  
    ttps_riesgo = 5  
    ttps_sumando = ttps_ponderacion * ttps_riesgo
```

Fragmento de código 10. Cambios en código de cuantificación (iteración 5)

- Los ataques a Schneider tienen puntuaciones 8.45, 8.10, 7.80. Esto es lógico al ser el primero el ataque relacionado con dispositivos USB, que obviamente pudo suponer una afección mayor sobre los sistemas OT.
- El ataque a Vestas ahora tiene asignado un valor de 7.80.

Iteración 6

Hasta ahora el modelo que se entrena sobre los datos históricos etiquetados reproduce patrones pasados correctamente, pero no siempre captura la relevancia contextual o geoestratégica de ciertos eventos. Algunos ciberataques etiquetados con la fórmula heurística con un riesgo medio han demostrado tener consecuencias catastróficas (como cortes de suministro eléctrico masivos), lo cual no se refleja en su puntuación actual. Por tanto, es legítimo introducir ajustes cualitativos justificados para que el modelo aprenda a valorar adecuadamente estos eventos.

Se consideran ataques candidatos a ajuste aquellos que han provocado interrupciones físicas del suministro eléctrico mediante acción directa sobre sistemas de control industrial (ICS/SCADA). Esto incluye, además de Industroyer 1 (2016) y Industroyer 2 (2022), otros como BlackEnergy (2015, Ucrania), que implicó un corte de luz que afectó a aproximadamente 230,000 personas, y TRITON/Trisis (2017), que comprometió sistemas de seguridad de una planta energética en Arabia Saudí que, aunque sin corte explícito, estuvo cerca de causar un accidente físico.

El incremento de +0.2 puntos sobre una escala de 0 a 10 representa una variación del 2%. Este valor es suficientemente pequeño como para no distorsionar el resto de la distribución, pero suficientemente grande como para que el modelo aprenda a diferenciar eventos con consecuencias físicas severas. Se está aplicando el ajuste a un número muy reducido de casos

(<5% del dataset), por lo que el impacto global es controlado. Este ajuste no busca sesgar el modelo, sino compensar una infraestimación sistemática de ciertos eventos críticos.

La decisión de aplicar este ajuste se basa en el consenso de expertos en ciberseguridad industrial de distintas geografías dentro de la organización, así como colaboradores externos especializados en entornos OT. Este enfoque garantiza que el ajuste refleje una valoración experta y contextualizada del impacto real de estos incidentes. Entre los incidentes ajustados quedan:

- Industroyer 1 (2016): de 5.6 a 5.8
- Industroyer 2 (2022): de 5.8 a 5.6
- Triton (2017): de 5.7 a 5.9

Iteración 7

A raíz de algunas pruebas con datos en tiempo real asociados a vulnerabilidades de equipos de fabricantes (ver Ingesta de noticias de ciberataques en tiempo real), se observó que muchas explotaciones de CVEs ocurrían en diferentes países que no se revelaban al público. El modelo trataba entonces de cuantificar su riesgo basándose en el nombre de país “MULTIPLE”, que no se reflejaba en la cuantificación heurística inicial (aunque algunos datos ya aparecían etiquetados así). Se asigna así una puntuación igual a la que obtienen los países de la UE:

```
# 1. Ubicación
ubi_ponderacion = 0.25
ubi_riesgo = 0
if row["country"] == "SPAIN":
    ubi_riesgo = 10
elif row["country"] in paises_criticos: # Países en los que opera IBR
    ubi_riesgo = 8
elif row["country"] in aliados: # Países aliados de España
    ubi_riesgo = 5
elif row["country"] in paises_ue: # Países de la UE
    ubi_riesgo = 3
elif row["country"] == "MULTIPLE": # Varios países afectados
    ubi_riesgo = 3
ubi_sumando = ubi_ponderacion * ubi_riesgo
```

Fragmento de código 11. Cambios en código de cuantificación (iteración 7)

Iteración 8

La última iteración incorpora todos los cambios realizados en las anteriores. El resultado final es el mostrado en la sección de Definición del riesgo. Con este código de cuantificación, las puntuaciones que reciben los casos evaluados son las siguientes:

<i>Tipo de caso</i>	<i>Caso</i>	<i>Riesgo</i>
Sintético	Peor caso #1	10
Sintético	Peor caso #2	9.05
Sintético	Mejor caso #1	1.20
Sintético	Mejor caso #2	1.30
Sintético	Caso intermedio #1	5.75
Sintético	Caso intermedio #2	5.15
Real	Vestas 2021	7.80
Real	Schneider 2018	8.45
Real	Schneider 2023	8.25
Real	Schneider 2024	7.80
Real	Iberdrola Wannacry 2017	10
Real	Industroyer 1 (2016)	5.80
Real	Industroyer 2 (2022)	6
Real	TRITON (2017)	5.90

Tabla 12. Resultados de las pruebas de calibración

5.1.3 ESPECIFICACIÓN DEL MODELO Y ENTRENAMIENTO

Para la predicción del riesgo de los ciberataques se ha entrenado con los datos históricos un modelo de regresión (adecuado para predecir el valor continuo del riesgo, que va de 0 a 10) basado en una arquitectura de red neuronal utilizando la API de Keras de TensorFlow.

Aunque la cuantificación del riesgo que se ha explicado previamente podría parecer suficiente para cuantificar el riesgo de ciberataques en tiempo real, los modelos de ML son capaces de detectar patrones complejos y relaciones no lineales entre variables que podrían no ser evidentes a través de reglas predefinidas. Usar este modelo de red neuronal permite obtener valores más precisos para la predicción del riesgo.

El dataset de ciberataques históricos tiene algunas columnas categóricas y otras numéricas. Las numéricas son todas las relativas a los conflictos geopolíticos entre los países involucrados en el evento de ciberseguridad y sus aliados (se pueden ver en los ejemplos definidos en la sección Casos sintéticos).

Dentro de las categóricas hay algunas columnas que tienen muchos valores posibles (por ejemplo, la columna “organization”, que indica el nombre de organización afectada), y otras que tienen menos (por ejemplo, “event_type”, que tiene solo cuatro valores posibles que clasifican el ciberataque a alto nivel). Por ello, para codificar estas entradas al modelo se han utilizado dos técnicas diferentes:

1) One-Hot encoding. Este tipo de codificación transforma las categorías en vectores binarios, permitiendo que el modelo las procese como entradas numéricas. Esto implica que una columna con cuatro valores posibles se convierte en cuatro columnas diferentes. Las columnas categóricas siguientes se codifican utilizando One-Hot Encoding.

- “actor_type”: indica la naturaleza del actor responsable del evento, que puede ser *Criminal* (organización que accede ilícitamente a redes con fines financieros), *Nation-State* (agencia gubernamental, militar o afiliada), *Terrorist* (actor no estatal)

que busca influir en condiciones políticas o militares atacando a civiles), *Hacktivist* (individuo o grupo motivado por activismo social o político), *Hobbyist* (individuo motivado por curiosidad o prestigio).

- “event_type”: indica si los efectos principales del evento fueron *Disruptive* (interrumpen operaciones), *Exploitive* (robo de información sensible) o *Mixed* (ambos, como en un ataque de ransomware).
- “Attacks ICS”: indica si el actor está incluido en el marco de MITRE ATT&CK ICS.
- “contains_risky_ttps”: indica si entre las técnicas asociadas al actor están algunas de las consideradas de riesgo (ver Definición del riesgo).

2) Embedding. Un embedding es una representación vectorial de datos categóricos en un espacio de menor dimensión. En lugar de representar categorías como vectores one-hot (que pueden ser muy grandes y dispersos), los embeddings permiten representar cada categoría como un vector denso de números reales. Las columnas categóricas “actor”, “organization”, “industry”, “event_subtype”, “country”, “actor_country”, y “TTPs” se codifican como enteros y luego se transforman mediante capas de embedding. Cada columna categórica tiene una capa de embedding que transforma las categorías en vectores densos.

- “actor”: el actor responsable del ataque.
- “organization”: la organización, entidad o empresa atacada.
- “industry”: la industria a la que pertenece la organización afectada por el ataque, clasificada en función del sistema NAICS (*North American Industry Classification System*).
- “event_subtype”: clasifica el evento según la parte de la infraestructura de TI más impactada, como manipulación de mensajes, denegación de servicios, ataques a datos, componentes físicos o robo de datos desde sensores, hosts, redes, servidores o en tránsito.
- “country”: el país en el que ha ocurrido el ciberataque.
- “actor_country”: el país al que pertenece o se asocia el actor.
- “TTPs”: el listado de técnicas asociadas al actor responsable del ataque.

```
n_categories = df[col].nunique()
embedding_dim = min(50, (n_categories // 2) + 1) # regla común para dim
```

Fragmento de código 12. Definición de embeddings

- *n_categories* calcula el número de categorías únicas en la columna.
- *embedding_dim* determina la dimensión del vector de embedding. La dimensión se calcula como el menor valor entre 50 y la mitad del número de categorías más uno. Esta es una regla común para establecer la dimensión del embedding.

Una vez codificados los datos, se concatenan los embeddings y las entradas codificadas con One-Hot en un solo tensor. Este tensor se pasa por las capas de la red neuronal, o capas densas, que son capas totalmente conectadas en las que cada neurona está conectada a todas las de la capa anterior.

```
x = Concatenate()(embedding_layers + [onehot_input, numerical_input])
x = Dense(256, activation='relu')(x)
x = BatchNormalization()(x)
x = Dropout(0.3)(x)
x = Dense(128, activation='relu')(x)
x = Dropout(0.2)(x)
x = Dense(64, activation='relu')(x)
x = Dense(32, activation='relu')(x)
x = Dense(16, activation='relu')(x)
x = Dense(12, activation='relu')(x)
x = Dense(10, activation='relu')(x)
x = Dense(8, activation='relu')(x)
output = Dense(1)(x)
```

Fragmento de código 13. Definición de la red neuronal

Cada capa `Dense(N, activation='relu')` tiene N neuronas y utiliza la función de activación ReLU, que introduce no linealidades y permite capturar relaciones complejas. Se utilizan capas de `Dropout()` para prevenir el sobreajuste y también `BatchNormalization()` para estabilizar y acelerar el entrenamiento.

El modelo se crea utilizando la API funcional de Keras. Se definen sus entradas y salidas, y se compilan con el optimizador Adam (*Adaptive Moment Estimation*, que ajusta automáticamente la tasa de aprendizaje para cada peso de la red.), una tasa de aprendizaje de 0.001, la función de pérdida MSE (error cuadrático medio) y la métrica MAE (error absoluto medio).

```
model = Model(inputs=inputs, outputs=output)
model.compile(optimizer=Adam(learning_rate=0.001), loss='mse',
metrics=['mae'])
```

Fragmento de código 14. Modelo definido

Por último, los datos se dividen en conjuntos de entrenamiento y prueba, con el 80% de los datos destinados al entrenamiento y el 20% restante para la validación del modelo. Los resultados que se obtienen con los datos de prueba son los siguientes:

<i>Métrica</i>	<i>Valor</i>	<i>Interpretación</i>
MSE	0.12	Las predicciones se acercan a los valores reales.
MAE	0.25	Las predicciones se acercan a los valores reales.
R2-Score	0.91	El modelo explica la mayoría de la variabilidad en los datos.

Fragmento de código 15. Interpretación de las métricas obtenidas

Layer (type)	Output Shape	Param #	Connected to
actor_input (InputLayer)	(None, 1)	0	-
organization_input (InputLayer)	(None, 1)	0	-
industry_input (InputLayer)	(None, 1)	0	-
event_subtype_input (InputLayer)	(None, 1)	0	-
country_input (InputLayer)	(None, 1)	0	-
actor_country_input (InputLayer)	(None, 1)	0	-
TTPs_input (InputLayer)	(None, 1)	0	-
embedding_84 (Embedding)	(None, 1, 50)	61,400	actor_input[0][0]
embedding_85 (Embedding)	(None, 1, 50)	653,100	organization_input[0][0]
embedding_86 (Embedding)	(None, 1, 11)	242	industry_input[0][0]
embedding_87 (Embedding)	(None, 1, 48)	4,560	event_subtype_input[0][0]
embedding_88 (Embedding)	(None, 1, 50)	8,300	country_input[0][0]
embedding_89 (Embedding)	(None, 1, 41)	3,362	actor_country_input[0][0]
embedding_90 (Embedding)	(None, 1, 21)	861	TTPs_input[0][0]
flatten_84 (Flatten)	(None, 50)	0	embedding_84[0][0]
flatten_85 (Flatten)	(None, 50)	0	embedding_85[0][0]
flatten_86 (Flatten)	(None, 11)	0	embedding_86[0][0]
flatten_87 (Flatten)	(None, 48)	0	embedding_87[0][0]
flatten_88 (Flatten)	(None, 50)	0	embedding_88[0][0]
flatten_89 (Flatten)	(None, 41)	0	embedding_89[0][0]
flatten_90 (Flatten)	(None, 21)	0	embedding_90[0][0]
onehot_input (InputLayer)	(None, 15)	0	-
numerical_input (InputLayer)	(None, 30)	0	-
concatenate_12 (Concatenate)	(None, 316)	0	flatten_84[0][0], flatten_85[0][0], flatten_86[0][0], flatten_87[0][0], flatten_88[0][0], flatten_89[0][0], flatten_90[0][0], onehot_input[0][0], numerical_input[0][0]

Ilustración 28. Arquitectura del modelo Keras

Los casos sintéticos desarrollados previamente también se prueban con el modelo ya entrenado y se obtienen las cuantificaciones pertinentes. Por ejemplo, si se prueba el peor caso 1:

```
1/1 ————— 0s 166ms/step
Predicción de riesgo para el caso worst-case-1:
      actor organization country predicted_risk
0 Undetermined Iberdrola SPAIN 9.019754
```

Ilustración 29. Ejemplo de prueba válida del modelo con caso sintético

Se obtiene un riesgo superior a 9 (aunque no es 10 a pesar de ser Iberdrola). Se considera válido puesto que el R2 es 0.91.

5.1.4 PRUEBAS CON CIBERATAQUES RECIENTES

Los datos históricos incluyen ciberataques desde 2014 hasta 2024. Se realizan una serie de pruebas con ciberataques recientes para verificar el funcionamiento adecuado del modelo. En este apartado se muestra como ejemplo el ciberataque a Endesa ocurrido en marzo de 2025: una brecha de datos, por parte de un actor ruso llamado AgencyInt.

Al aplicar el modelo sobre este ataque se obtiene un valor de 4.9: se trata de un riesgo intermedio (o algo inferior), que efectivamente no tuvo ninguna implicación posterior para Iberdrola.

5.2 *INGESTA DE NOTICIAS DE CIBERATAQUES EN TIEMPO REAL*

El proceso de ingesta de datos de ciberataques en tiempo real se basa en un pipeline bifurcado que integra fuentes abiertas, enriquecimiento contextual y análisis avanzado para la evaluación de riesgos. En este proceso se utilizan librerías y herramientas como GoogleNews y Article, y un modelo de lenguaje (LLM) para normalización.

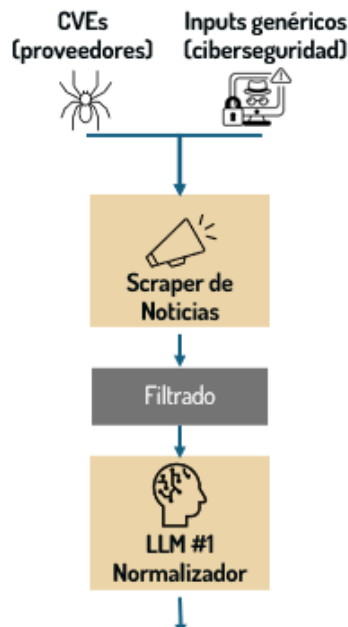


Ilustración 30. Proceso de ingestión de noticias de ciberataques en tiempo real

5.2.1 BÚSQUEDA AUTOMATIZADA DE NOTICIAS

Las dos líneas en las que se bifurca este proceso son la obtención de noticias de ciberataques generales y la obtención de noticias de ciberataques relacionados con CVEs de proveedores.

Ciberataques generales. Esta línea comienza con la búsqueda automatizada de noticias sobre ciberataques mediante la librería GoogleNews de Python.

```
# --- Obtención de noticias de ciberataques recientes
def buscar_noticias():
    googlenews = GoogleNews(lang='en', period='7d')
    googlenews.search('cyber attack')
    results = googlenews.results()[:5] # Solo las primeras 5
    noticias_procesadas = []
    for r in results:
        parsed_date = dateparser.parse(r['date'])
        noticias_procesadas.append({
            'event_date': parsed_date,
            'title': r.get('title', ''),
            'link': r.get('link', ''),
            'description': r.get('desc', '')
        })

    df_noticias = pd.DataFrame(noticias_procesadas)
```

```
return df_noticias
```

Fragmento de código 16. Obtención de noticias de ciberataques generales

Ciberataques relacionados con CVEs de proveedores. En paralelo, la segunda línea del pipeline se enfoca en CVEs relevantes para los proveedores de la empresa. Primero, se identifican los CVEs más recientes que afecten a software de proveedores críticos. Luego, se buscan noticias específicas relacionadas a esos CVEs usando GoogleNews, siguiendo el mismo proceso.

5.2.2 SCRAPING Y FILTRADO

Una vez obtenidas las noticias y sus enlaces o URLs, se hace *web scraping* del cuerpo completo de estas usando la clase *Article* de la librería *newspaper3k* de Python.

```
article = Article(url)
article.download()
article.parse()
full_text = article.text.strip()
```

Fragmento de código 17. Scraping de las noticias

Una vez extraído el contenido de cada noticia, se hace un filtrado de “fake news” en el que se valida el TLD (*Top-Level Domain*) de cada una (para verificar que es uno de los que indica el Fragmento de código 18. TLD confiables) y se detectan ciertas palabras de índole sensacionalista que puedan indicar que es una noticia falsa o exagerada.

```
tld_confiables = (".com", ".org", ".net", ".es", ".co.uk", ".io")
```

Fragmento de código 18. TLD confiables

Las noticias de ambas líneas de obtención se agrupan en un único Pandas DataFrame que contiene los campos: title, full_text, url, date, cve, proveedor.

5.2.3 NORMALIZACIÓN CON LLM

El contenido extraído se normaliza usando un modelo LLM, que estructura la información clave: organizaciones e industrias afectadas, actores involucrados, tipos de ataque, impactos, fechas, etc. Para esto se ha utilizado OpenRouter, y en concreto se ha empleado el modelo llama-3.3-70b-instruct:free, desarrollado por Meta.

La consulta que se envía al LLM es la siguiente (se puede ver completa en el anexo):

```
# Construcción de prompt
def construir_prompt(title, full_text, url, date, cve, proveedor):
    return f"""You are a cybersecurity analyst.
Given the following news article:

Title: {title}
Full text: {full_text}
URL: {url}
Date: {date}
CVE: {cve}
Proveedor: {proveedor}

Extract the following fields if possible:
- event_date en este formato: 20/01/2025 0:00:00 (can't be Undetermined,
if you know year and month write day 1 of that month or the date of the
article; you must refer to Date)
- actor (author of the attack; if not known, write "Undetermined")
- organization (attacked; write content of Proveedor indicated at the
start of this prompt if not "None", or write "Undisclosed" if not
mentioned, or find it yourself in Full text)
- industry ('Agriculture, Forestry, Fishing and Hunting',
'Mining, Quarrying, and Oil and Gas Extraction', 'Utilities',
'Construction', 'Manufacturing', 'Retail Trade',
'Professional, Scientific, and Technical Services',
[...],
'Other Services (except Public Administration)',
'Public Administration', 'Undetermined')
- event_type (Disruptive, Exploitative, Mixed)
- event_subtype ('Message Manipulation', 'Exploitation of Sensors',
'Exploitation of End Host', 'Physical Attack', [...],
'External Denial of Service',
'Internal Denial of Service',
'Exploitation of End User', 'Data Attack', 'Exploitation
of Sensors', 'Exploitation of Infrastructure', 'Undetermined'),
can be several with commas between them
- country (of the attack, in upper case)
- actor_country (if known write in upper case, if not write
"UNDETERMINED")
- cve_id: CVE (indicated at the start of this prompt)

"country" and "actor_country", if not UNDETERMINED, must be within:
'THAILAND', 'GERMANY', 'SUDAN', 'PANAMA', 'JORDAN', 'BOLIVIA',
'MOROCCO', [...], 'TOGO'

Return them as a JSON object.
"""
```

Fragmento de código 19. Consulta o prompt al primer LLM

Esta consulta devuelve un JSON embebido en un fragmento de texto del que se puede obtener el objeto JSON fácilmente, y a partir de él se genera otro DataFrame que contiene

los mismos campos (relativos al ciberataque) que el DataFrame utilizado para entrenar el modelo (ver Modelo de Machine Learning para predicción del riesgo).

5.3 INGESTA DE INFORMACIÓN DE CONTEXTO EN TIEMPO REAL

Una vez generado el DataFrame parcial que contiene únicamente información general sobre el ciberataque, se añaden los campos que faltan relativos a la información de los grupos de adversarios y sus técnicas, y al contexto geopolítico de los países involucrados en el evento de ciberseguridad.

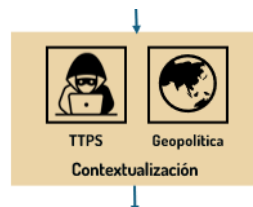


Ilustración 31. Ingesta de información de contexto en tiempo real

5.3.1 INFORMACIÓN DE CONTEXTO GEOPOLÍTICO

Para enriquecer el DataFrame parcial con las variables numéricas que representan el número de conflictos en los 90 días previos a cada ciberataque en el país atacado, en el país del actor, entre ellos, entre sus aliados, etc., se utiliza la API de ACLED. Se consulta dicha API de forma general:

```
# URL base
url = "https://api.acleddata.com/acled/read"

# Filtros y formato
params = {
    "key": api_key_acled,
    "email": email,
    "format": "json", # No usar CSV porque el contenido real no lo es
    "year": "2025",
    "limit": 10000 # máximo permitido por llamada
}
```

Fragmento de código 20. Consulta a la API de ACLED

Y se hace el filtrado y agrupación de datos después, para aumentar la eficiencia del proceso.

5.3.2 INFORMACIÓN DE ADVERSARIOS Y TTPs

Para completar los campos del DataFrame relativos a la información de los adversarios, se utiliza la librería mitre-attack de Python, que permite la integración con el marco de MITRE ATT&CK. Se buscan los nombres de los actores en el marco de MITRE Enterprise y en el marco de MITRE ICS (igual que se hizo durante la recopilación inicial de datos de entrenamiento), y posteriormente se asocia a cada uno sus técnicas con un código como el siguiente:

```
# Función para obtener las técnicas asociadas a un actor
def obtener_ttps(actor, data):
    ttp_objs = []

    # Buscar relaciones 'uses' para encontrar las técnicas asociadas
    for relationship in data['objects']:
        if relationship.get('type') == 'relationship' and
relationship.get('source_ref') == actor['id'] and
relationship.get('relationship_type') == 'uses':
            target_ref = relationship.get('target_ref')
            # Buscar el objeto de la técnica asociada (attack-pattern)
            related_obj = next((obj for obj in data['objects'] if
obj['id'] == target_ref), None)
            if related_obj and related_obj.get('type') == 'attack-
pattern':
                ttp_objs.append(related_obj)

    return ttp_objs
```

Fragmento de código 21. Obtención de técnicas asociadas a un actor

5.4 APLICACIÓN DEL MODELO ENTRENADO

Una vez llevado a cabo el entrenamiento, se guardan los codificadores/normalizadores de datos con Pickle, un formato de serialización de objetos de Python, y el modelo entrenado con Keras. Así, se obtienen los archivos `numerical_scaler.pkl`, `onehot_encoder.pkl`, `embedding_encoders.pkl` y `modelo_riesgo.keras`, permitiendo así su reutilización para aplicar el modelo entrenado a datos en bruto como los del DataFrame completo obtenido en tiempo real.

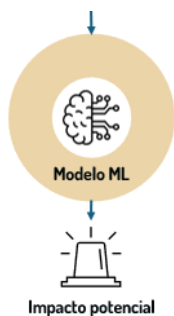


Ilustración 32. Aplicación del modelo entrenado

El modelo predice el riesgo o impacto potencial de los ataques y los añade a una columna al final del DataFrame con el siguiente formato:

```
Tipos de entrada: [dtype('float32'), dtype('float32'), dtype('float32'), dtype('float32'), dtype('float32'), dtype('float32'), dtype('float32'), dtype('float32'), dtype('float32'), dtype('float32')]
1/1 ----- 0s 186ms/st ----- 0s 211ms/step
predicted_risk
0      2.233811
1      2.233811
2      2.184066
3      2.184066
4      1.801864
5      2.263243
6      1.842065
7      2.184066
8      2.251808
9      2.745708
10     0.775901
```

Ilustración 33. Predicciones del modelo en tiempo real

5.5 ENRIQUECIMIENTO CON INDICADORES DE COMPROMISO

Aunque el modelo entrenado predice el riesgo o impacto potencial de los ataques evaluados a partir de los tres bloques de información ya mencionados (datos generales del ataque, datos de adversarios y técnicas, datos de geopolítica), la información de los ataques asociados a un CVE puede enriquecerse aún más con elementos como indicadores de compromiso (IoCs). Existe una gran variedad de herramientas y fuentes abiertas que permiten mapear CVEs a IoCs. En el desarrollo de este proyecto se han utilizado solo algunas de ellas, explicadas a continuación, pero la herramienta podría ampliarse a futuro para incluir otras diferentes.



Ilustración 34. Enriquecimiento con IoCs

5.5.1 IPs MALICIOSAS: GREYNOISE

La plataforma GreyNoise (ver GreyNoise) permite hacer búsquedas de CVEs para obtener información en tiempo real sobre ellos. Una de sus funcionalidades, que se ha utilizado en este proyecto, es el mapeo de IPs asociadas a la explotación de un CVE. Se añade al DataFrame anterior una nueva columna llamada “ips” que contiene un listado de aquellas IPs asociadas con actividad maliciosa y explotación de cada CVE. GreyNoise también clasifica las IPs (unknown o desconocidas, malicious o maliciosas, suspicious o sospechosas, benign o benignas); para este caso de uso se han añadido a los datos solo las maliciosas y sospechosas. Además, para cada listado de IPs, GreyNoise permite identificar cuáles de ellas tienen tráfico saliente hacia países determinados. Se añade otro elemento que indica si las IPs de la anterior tienen tráfico hacia España.

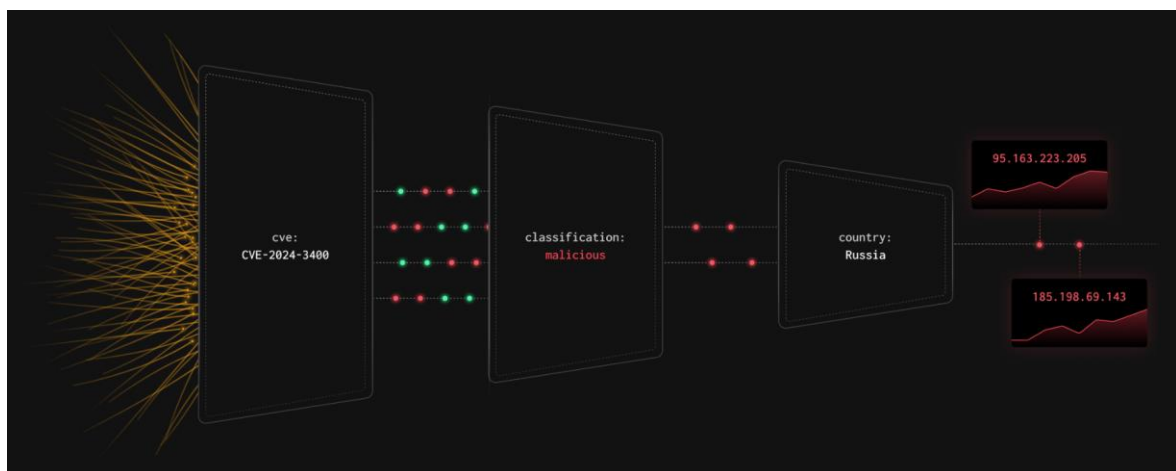


Ilustración 35. Funcionamiento de la clasificación de IPs de GreyNoise

Esto permite a los analistas valorar qué CVEs priorizar en una pasada inicial.

5.5.2 ANÁLISIS EN PROFUNDIDAD DE IPS MALICIOSAS: SHODAN

Se ha utilizado Shodan, un motor de búsqueda especializado en dispositivos conectados a Internet, para ampliar la información sobre las direcciones IP identificadas previamente por GreyNoise como asociadas a la explotación activa de vulnerabilidades (CVEs) relevantes. A partir de estas IPs, Shodan obtiene (en algunos casos) detalles adicionales como el proveedor de servicios de Internet (ISP), el nombre de la organización (ORG) propietaria del activo con dicha IP, su país de origen, etc. Esta información es fundamental para evaluar la procedencia y el posible riesgo asociado a la infraestructura maliciosa detectada.

Shodan se utiliza mucho para buscar activos expuestos a Internet, como parte de una dinámica de prevención en las empresas, pero en este caso se utiliza como proveedor adicional de datos para inteligencia de amenazas.

5.6 INGESTA DE DATOS DE INTELIGENCIA OT

La herramienta completa tiene otra línea paralela a la obtención de datos de ciberataques y su contextualización y cuantificación. Esta otra línea se centra en recopilar información más genérica de inteligencia de ciberseguridad en el entorno OT. Para ello recopila datos de dos fuentes principales: grupos públicos de Telegram de ciberseguridad, y OSINT (*Open Source Intelligence* o inteligencia de fuentes abiertas).

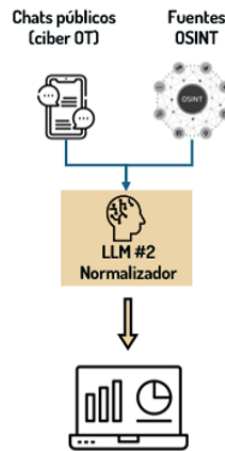


Ilustración 36. Ingesta de datos de inteligencia OT

5.6.1 TELEGRAM

Telegram es una plataforma de mensajería instantánea con un gran foco en la seguridad y la privacidad. Esta plataforma ofrece además una API que, combinada con herramientas como Telethon en Python, permite acceder de forma automatizada a mensajes de grupos o chats públicos buscando por palabras clave. Telethon es una librería que actúa como cliente asíncrono de Telegram. Se ha utilizado en este proyecto para monitorizar las conversaciones de comunidades abiertas en el contexto de ciberseguridad. Algunos de los grupos monitorizados son: *offensiveOT*, *TheHackerNews*, *ThreatIntelligenceSharing*.

Por otro lado, las palabras clave utilizadas para filtrar son:

```
keywords = ['SCADA', 'ICS', 'ENERGY', 'OT', 'NERC', 'PLC', 'CRITICAL  
INFRASTRUCTURE']
```

Ilustración 37. Palabras clave para filtrado de mensajes de Telegram

Una vez obtenidos los mensajes relevantes en base a las palabras clave mencionadas, se pasan por un motor LLM que los analiza y los “normaliza” a un formato de “inteligencia contextual” común para facilitar su integración en la API y en la herramienta de visualización.

5.6.2 FEEDS DE OSINT

Para obtener inteligencia de fuentes abiertas se han utilizado algunas fuentes como:

- CISA Alerts": https://www.cisa.gov/sites/default/files/cisa_alerts.xml
- Dragos Blog": <https://www.dragos.com/feed/>
- "Nozomi Networks Blog": <https://www.nozominetworks.com/feed/>
- "ICS-CERT": <https://www.us-cert.gov/ncas/alerts.xml>

En este caso se ha filtrado con feedparser y se han buscado las palabras clave:

```
KEYWORDS = ["industrial", "ot", "substation", "grid", "electric", "iec  
104", "cip", "plc", "energy", "iberdrola"]
```

Ilustración 38. Palabras clave para OSINT

En este caso también se pasa toda la información por un LLM que la normaliza.

5.6.3 NORMALIZACIÓN CON LLM

La consulta que se le hace al LLM para cada entrada (contenido de mensajes de Telegram y OSINT):

```
# Prompt plantilla para el LLM OpenRouter  
prompt_template = """  
Eres un analista de ciberseguridad especializado en amenazas al sector  
industrial eléctrico.  
A partir del siguiente contexto extraído de fuentes OSINT, extrae:  
  
1. Nombres de actores o grupos de amenazas recientes mencionados.  
2. Técnicas, tácticas y procedimientos (TTPs) usados en esos ataques.  
3. Indicadores de compromiso (IoCs) relevantes (IPs, hashes, URLs,  
dominios).  
  
Contexto:  
{context}  
  
Por favor, responde con la información de forma clara y estructurada,  
usando listas cuando sea posible.  
"""
```

Ilustración 39. Prompt al LLM que normaliza la inteligencia OT

Posteriormente se hace otra consulta al mismo LLM para que simplemente resuma el contenido, y todo ello se convierte en un JSON con un formato como:

```
{
  "fecha": "2025-06-04",
  "tipo": "Alerta de Vulnerabilidades en Schneider y Mitsubishi",
  "descripcion": "Se mencionan unas vulnerabilidades reportadas en equipos Wiser Avataron 6K Freelocate, Wiser Cuadro H 5P Socket de marcas como Schneider Electric y Mitsubishi Electric, donde la explotación permite saltar los controles de autenticación o inyectar código a través de un buffer overflow, reportados en ICSA-25-153-01 y en CVE-2023-4041.",
  "enlaces": [
    "https://www.cisa.gov/news-events/ics-advisories/icsa-25-153-01",
    "https://www.incibe.es/incibe-cert/alerta-temprana/avisos-sci/divulgacion-de-informacion-y-denegacion-de-servicio-dos-en-productos-mitsubishi",
    "https://www.cisa.gov/news-events/ics-advisories/icsa-25-153-03",
    "https://www.incibe.es/incibe-cert/alerta-temprana/avisos-sci/error-en-la-gestion-de-sesiones-en-eibport-lan-de-abb",
    "https://search.abb.com/library/Download.aspx?ID%20de%20documento=9AKK108471A1621&Código%20de%20idioma=en&Id.%20de%20parte%20del%20documento=pdf&Acción=Iniciar"
  ]
}
```

Ilustración 40. Ejemplo de entrada de JSON de inteligencia OT

Se trata de formato JSON que se puede mostrar fácilmente en una interfaz web, o compartir vía API. En el alcance del proyecto, se ha mostrado simplemente en una interfaz web.

5.7 CUANTIFICACIÓN DE VULNERABILIDADES: MODELO EMPRESARIAL

Como ya se ha mencionado, la herramienta desarrollada tiene una arquitectura modular que permite integrar otras fuentes de datos e inteligencia fácilmente. Debido a su uso de CVEs, se ha decidido integrar la herramienta que utiliza Iberdrola para cuantificar las vulnerabilidades expresadas como CVEs.

Aunque los CVEs pueden estar clasificados con puntuaciones CVSS (*Common Vulnerability Scoring System*) críticas o altas, estas evaluaciones no siempre reflejan con precisión el riesgo real en un entorno industrial específico. En particular, en contextos industriales como plantas de generación energética, es fundamental aplicar cuantificaciones de vulnerabilidades que estén correctamente contextualizadas.

Un mismo CVE puede tener un impacto muy distinto según las características del activo afectado y su entorno. Por ejemplo, un activo que está conectado a una red externa o a Internet suele ser mucho más vulnerable y representativo de un riesgo inmediato que uno que se encuentra en una red aislada, sin acceso directo a otras infraestructuras o sistemas críticos, como suele ser habitual en muchas instalaciones industriales.

Por esta razón, se ha integrado en la herramienta la solución utilizada por Iberdrola para la cuantificación de vulnerabilidades expresadas como CVEs, la cual permite una evaluación de riesgo mucho más alineada con el contexto operativo y de negocio. De lo contrario, si se utilizaran las métricas de CVSS para evaluar las vulnerabilidades de los equipos en entornos industriales, sería inviable gestionar todas las vulnerabilidades que tienen estos dispositivos, puesto que muchos de ellos son *legacy*. Esta herramienta ayuda a los usuarios a evaluar el riesgo para el negocio (*business risk*) de cada vulnerabilidad, y a definir una recomendación sobre la acción a tomar, basada en respuestas a una serie de preguntas acerca del activo y su entorno, tales como:

- ¿El activo está conectado a una red externa?
- ¿La vulnerabilidad afecta a más de un activo?
- ¿La materialización de la vulnerabilidad puede implicar riesgos operativos o de seguridad física?

En base a estas respuestas, y a información de CVEs, KEV y EPSS, la herramienta calcula el riesgo para el negocio y sugiere una decisión recomendada entre las siguientes opciones:

- *Defer & Track*: posponer la mitigación inmediata, pero mantener un seguimiento continuo, si el riesgo es bajo o controlado.
- *Mitigate*: aplicar acciones de mitigación para reducir el riesgo cuando la vulnerabilidad presenta un impacto moderado y se puede controlar adecuadamente.
- *Remediate*: corregir la vulnerabilidad de forma prioritaria cuando el riesgo es alto y puede comprometer la operación o seguridad.

- *Accept*: aceptar el riesgo de forma consciente cuando el coste o impacto de la mitigación supera el beneficio, siempre con documentación y justificación.

Esta metodología permite que las decisiones de seguridad no se basen únicamente en puntuaciones CVSS genéricas, sino en un análisis contextualizado, optimizando así la gestión de vulnerabilidades y la protección efectiva de los activos críticos.

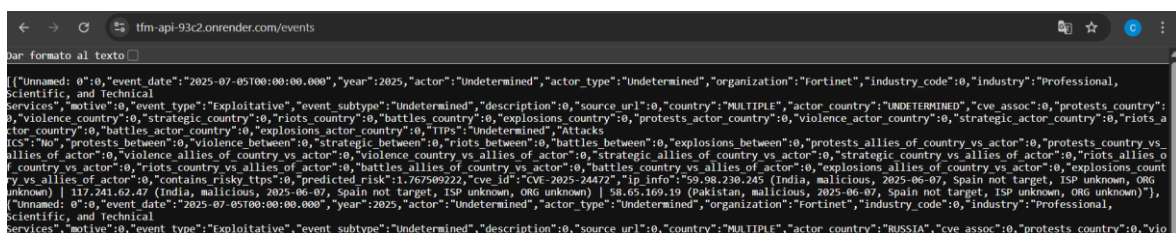
Esta funcionalidad se ha integrado en la interfaz gráfica de la herramienta.

5.8 DISEÑO DE LA COMUNICACIÓN CON EL SOC

Una vez estructurada en un único conjunto de datos toda la información obtenida a partir de fuentes en tiempo real, como noticias recientes, CVEs explotados e IoCs, se despliega una API (*Application Programmable Interface*) utilizando la plataforma gratuita Render. La dirección de la API para hacer consultas es: <https://tfm-api-93c2.onrender.com/>



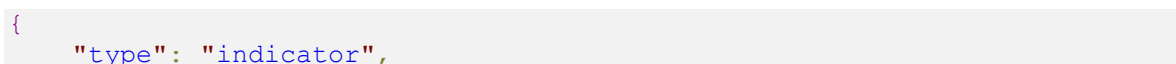
Ilustración 41. Inicio de la API



Fragmento de código 22. API creada

Esta API es pública y permite la comunicación con el SOC.

Aunque los datos que publica API están en formato JSON básico, también es posible publicarlos en formato STIX (ver STIX):



```
{
  "spec_version": "2.1",
  "id": "indicator--20250606095535",
  "created": "2025-06-06T09:55:35.977261",
  "modified": "2025-06-06T09:55:35.977267",
  "name": "Exploitative Event",
  "description": "Exploitative",
  "pattern": "[ipv4-addr:value = '168.0.228.61']",
  "valid_from": "2025-01-01T00:00:00.000",
  "labels": [
    "Exploitative"
  ],
  "confidence": 50,
  "lang": "en",
  "external_references": [
    {
      "source_name": "cve",
      "external_id": "CVE-2025-24472"
    }
  ],
  "object_marking_refs": [
    "marking-definition--613f2e26-407d-48c7-9eca-b8e91df99dc9"
  ],
  "granular_markings": [],
  "defanged": false,
  "revoked": false
}
```

Fragmento de código 23. Publicación de datos en formato STIX

Aunque se propone como un formato ideal para el contexto de CTI, para este proyecto se ha optado por mantener el formato JSON habitual.

5.9 INTERFAZ GRÁFICA

Para la interfaz gráfica se ha optado por desarrollar una página web en HTML (*HyperText Markup Language*) que actúa como punto de acceso visual para el sistema. En esta página se embebe un mapa interactivo de ciberataques, generado a partir de los datos recogidos en tiempo real de la API y almacenados en una base de datos gestionada por Elasticsearch. La visualización se realiza mediante Kibana, una herramienta de análisis y representación gráfica integrada con Elasticsearch, que permite crear dashboards y mapas dinámicos. Este mapa muestra información clave como la ubicación de los atacantes, las técnicas de ataque utilizadas y los indicadores de compromiso, facilitando así una visión centralizada, geolocalizada y en tiempo real del panorama de amenazas.

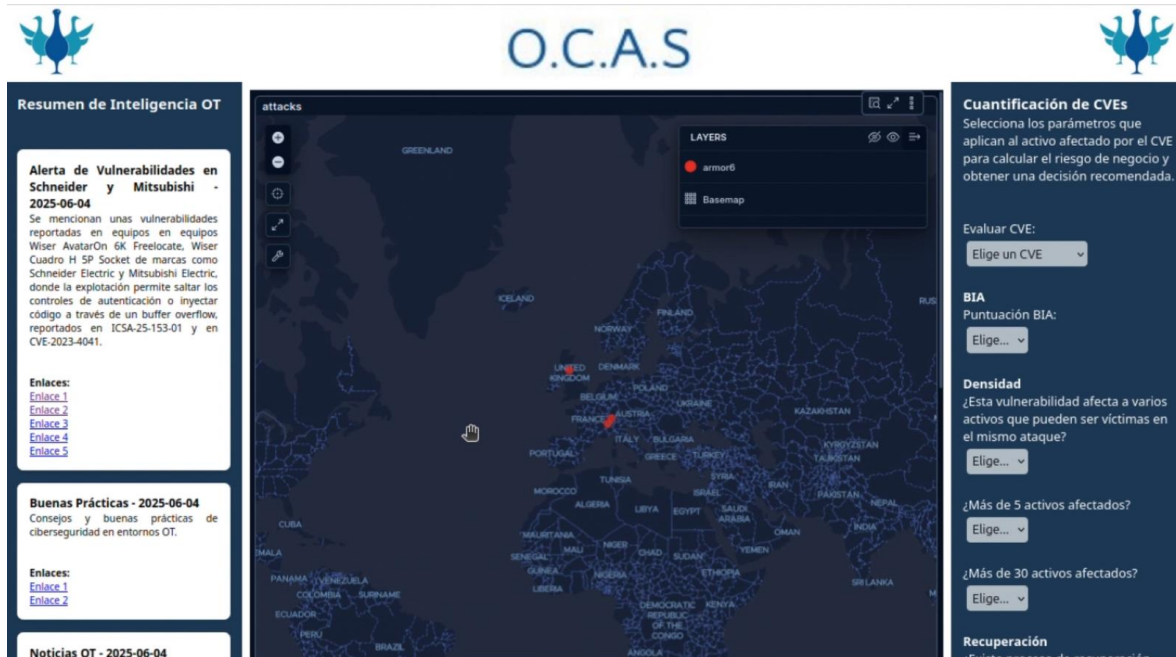


Ilustración 42. Interfaz gráfica

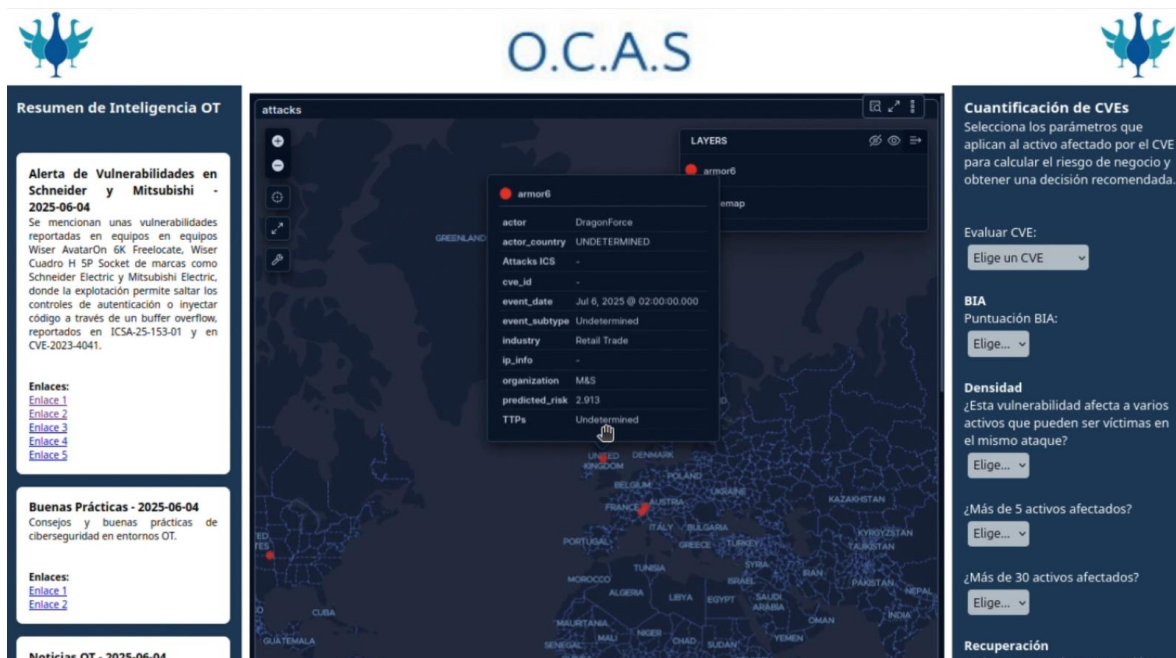
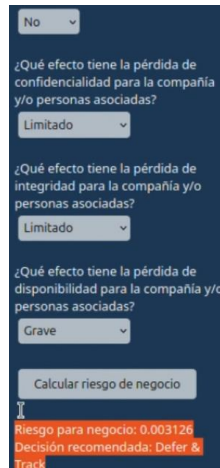


Ilustración 43. Interfaz gráfica (funcionalidades)

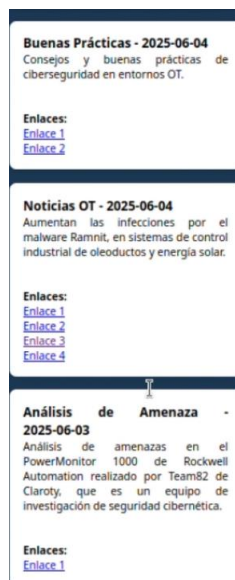
El panel derecho de la página HTML contiene el formulario de cuantificación de CVEs en contexto (ver Cuantificación de vulnerabilidades: modelo empresarial).



The screenshot shows a web form for quantifying CVEs. It has a dark blue background. At the top, there is a dropdown menu with 'No' selected. Below it, there are three questions, each with a dropdown menu: '¿Qué efecto tiene la pérdida de confidencialidad para la compañía y/o personas asociadas?' (Limited), '¿Qué efecto tiene la pérdida de integridad para la compañía y/o personas asociadas?' (Limited), and '¿Qué efecto tiene la pérdida de disponibilidad para la compañía y/o personas asociadas?' (Grave). At the bottom, there is a button labeled 'Calcular riesgo de negocio'. Below the button, the results are displayed: 'Riesgo para negocio: 0.003126', 'Decisión recomendada: Defer &', and 'Track'.

Ilustración 44. Formulario de cuantificación de CVEs en contexto

El panel izquierdo contiene una serie de tarjetas con la información de inteligencia OT recopilada (ver Ingesta de datos de inteligencia OT).



The screenshot shows a panel of OT intelligence cards. The first card is titled 'Buenas Prácticas - 2025-06-04' and contains the text 'Consejos y buenas prácticas de ciberseguridad en entornos OT.' and two links: 'Enlace 1' and 'Enlace 2'. The second card is titled 'Noticias OT - 2025-06-04' and contains the text 'Aumentan las infecciones por el malware Ramnit, en sistemas de control industrial de oleoductos y energía solar.' and four links: 'Enlace 1', 'Enlace 2', 'Enlace 3', and 'Enlace 4'. The third card is titled 'Análisis de Amenaza - 2025-06-03' and contains the text 'Análisis de amenazas en el PowerMonitor 1000 de Rockwell Automation realizado por Team82 de Claroty, que es un equipo de investigación de seguridad cibernética.' and one link: 'Enlace 1'.

Ilustración 45. Panel de inteligencia OT

Cada tarjeta contiene una serie de enlaces de referencia que el usuario puede consultar también.

Capítulo 6. ANÁLISIS DE RESULTADOS

En este apartado se analiza el producto final del proyecto realizado. El flujo final de la herramienta, ya mostrado previamente, es el siguiente:

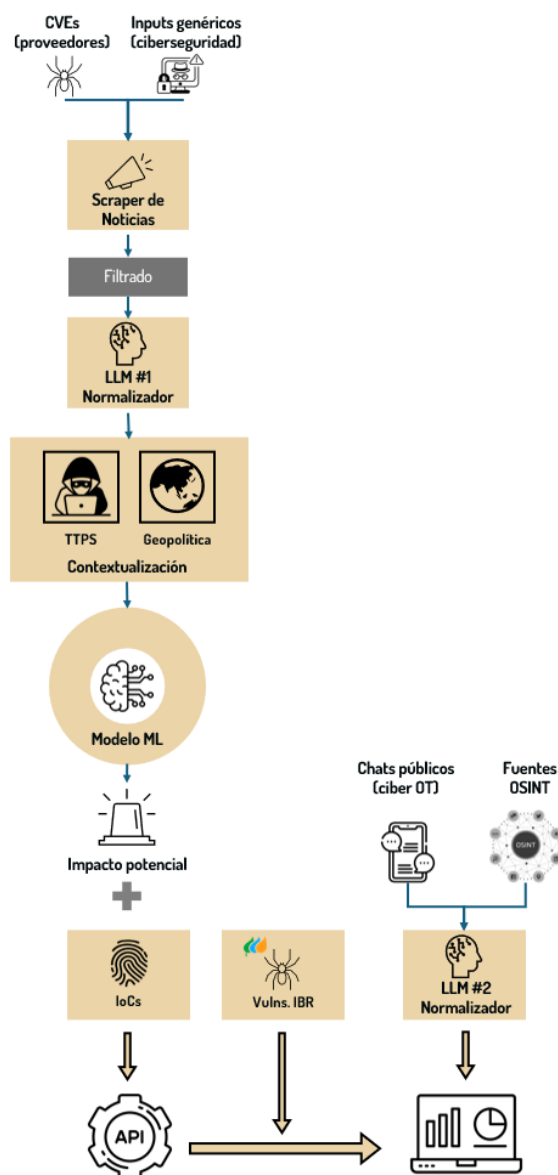


Ilustración 46. Flujo final de la herramienta

Cabe destacar que la interfaz gráfica (ver Interfaz gráfica), aunque aparezca al final de todo el flujo como producto resultante, no es más que una posibilidad de presentación de los resultados. La API de datos de ciberataques enriquecidos y contextualizados, que la propia interfaz gráfica utiliza, puede ser accedida por muchas otras aplicaciones y sistemas, entre ellos el SOC.

6.1 RESULTADOS DEL MODELO DE MACHINE LEARNING

La incorporación de un modelo de ML en el flujo de análisis resulta esencial para la predicción precisa del riesgo o impacto potencial asociado a los ciberataques, debido a las limitaciones inherentes a las fórmulas heurísticas tradicionales. Estas fórmulas, aunque útiles como aproximación inicial, no capturan la complejidad ni la no linealidad de los patrones presentes en eventos reales. Además, debido a la necesidad de realizar ajustes basados en las evaluaciones de expertos sobre ciertos casos críticos (ver Pruebas de calibración), las heurísticas por sí solas pierden consistencia, por lo que el uso del modelo es esencial.

El modelo de ML, de tipo regresión, ha demostrado una excelente capacidad predictiva, alcanzando un coeficiente de determinación R^2 de 0.9111, lo que indica que puede explicar más del 91% de la varianza en la variable de riesgo. Esta precisión se debe a su capacidad para detectar patrones complejos y relaciones multivariables entre los distintos atributos extraídos del ataque (actor, organización afectada, técnicas usadas, geolocalización, etc.), que serían difíciles de capturar con métodos heurísticos.

Aunque el modelo de regresión muestra un excelente rendimiento general, es importante destacar una característica del conjunto de entrenamiento que puede afectar a su precisión en ciertos contextos específicos. De los aproximadamente 14.000 ciberataques históricos recopilados, menos de 1.000 correspondían a sistemas de control industrial (SCI). Esta menor representación de incidentes en entornos industriales refleja una realidad observable: históricamente, los ciberataques a SCI han sido menos frecuentes que en otros sectores.

Este desequilibrio en la distribución se manifiesta en los histogramas de la variable de riesgo, donde los ataques a ICS aparecen con menor densidad. Esto podría hacer que, en algunos casos concretos, los resultados del modelo no converjan de manera inequívoca en escenarios OT. Sin embargo, esta limitación es temporal: las fuentes de datos sobre amenazas a sistemas industriales están en aumento, tanto en volumen como en calidad, lo que permitirá que el modelo gane progresivamente en precisión y capacidad de generalización para este tipo de entornos. Además, el modelo ha sido diseñado con una arquitectura modular y adaptable, lo que facilita su ajuste o reentrenamiento a medida que se incorporen nuevos datos.

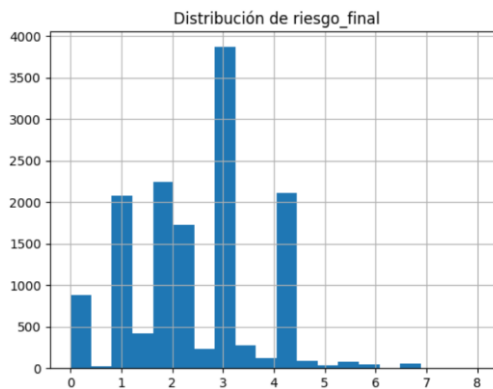


Ilustración 47. Distribución de la variable de riesgo

6.2 RESULTADOS DE LA INTEGRACIÓN DE FUENTES EN TIEMPO REAL

La integración de las diferentes fuentes mencionadas para la ingesta de datos en tiempo real ha sido uno de los aspectos más relevantes y enriquecedores del proyecto, ya que ha permitido construir una visión amplia, dinámica y contextualizada del panorama de amenazas de ciberseguridad. La variedad de fuentes utilizadas (desde buscadores de noticias como GoogleNews, hasta bases de datos de amenazas como GreyNoise, Shodan, MITRE ATT&CK y ACLED) ha permitido combinar indicadores técnicos, información táctica y contexto geopolítico. Todo ello es esencial para una evaluación del riesgo realmente informada y contextualizada.

No obstante, también se han identificado algunas limitaciones en esta fase. Una de las principales ha sido el uso limitado de ciertos modelos de lenguaje (LLM), tanto en tiempo como en recursos, debido a que se trata de un proyecto aún en fase de prueba de concepto. Esta restricción ha impedido realizar un fine-tuning personalizado sobre los modelos utilizados.

Aunque los modelos empleados han ofrecido buenos resultados en tareas de normalización y extracción semántica, su rendimiento podría mejorar si se ajustaran específicamente a textos técnicos o a patrones lingüísticos propios del ámbito de la ciberseguridad. Está previsto que, en fases futuras del proyecto, se destinen recursos adicionales para llevar a cabo un fine-tuning más profundo, lo que permitirá optimizar el rendimiento del sistema y adaptarlo aún más a las necesidades específicas del dominio OT.

Por otro lado, algunas fuentes que se habían considerado inicialmente no pudieron integrarse como se había previsto. Por ejemplo, GDELT, que en su momento ofrecía un gran potencial para análisis globales de eventos, parece haber quedado en desuso o sin soporte actualizado, por lo que no se pudo utilizar (y por ello se optó por GoogleNews). También se evaluó la incorporación de AlienVault OTX, pero en este caso se observó que su uso para buscar IoCs asociados a CVEs no implicaba una ventaja diferencial significativa frente a otras fuentes como GreyNoise, que además proporcionaba una mayor contextualización sobre actividad maliciosa en tiempo real. Se plantea incorporar AlienVault OTX en el futuro, para aprovechar otras funcionalidades que ofrece, como su capacidad de agrupar IoCs y pulsos [10] sobre amenazas específicas.

Por último, el uso de Shodan ha supuesto un reto en términos de rendimiento, ya que el elevado número de peticiones a su API puede generar respuestas lentas que afectan al flujo general de la herramienta, incluyendo la API desarrollada y la interfaz gráfica. Shodan fue seleccionada frente a otras alternativas como Censys debido a su mayor cobertura en dispositivos expuestos, lo que resulta especialmente relevante para el enfoque del proyecto. Se considerará optimizar su integración en el futuro y valorar el uso de otras fuentes similares si fuera necesario.

Los resultados de la integración de fuentes en tiempo real son muy positivos, aunque dejan margen para seguir optimizando tanto la calidad como la cobertura de datos a medida que evolucionan las herramientas disponibles y se ajustan los recursos técnicos del sistema.

6.3 RESULTADOS DE LA API

La API desplegada mediante la plataforma Render ha resultado ser una herramienta útil y versátil para facilitar la integración del sistema con otros componentes, como los dashboards de la interfaz web. Su diseño RESTful permite realizar consultas estructuradas sobre eventos, riesgos e indicadores enriquecidos, lo que favorece la interoperabilidad con otros sistemas de ciberinteligencia o plataformas de monitorización en tiempo real.

El principal problema que se ha encontrado con la API ha sido la lentitud en las respuestas, debida en parte a las operaciones que implican el uso del modelo de lenguaje (LLM) para la normalización de textos o el análisis semántico. Esto se debe al procesamiento secuencial del pipeline (que contiene algunas consultas pesadas) y a que el modelo no está desplegado de forma persistente en un servidor propio, sino que se activa bajo demanda. Además, se utilizó la versión gratuita de Render, que tiene un límite de memoria RAM a consumir de 512 MB. Esto implica que, cuando las noticias obtenidas o la información de IoCs son muy largas, el servicio de la API se cae. Por ello, se ha establecido ya como línea de mejora el buscar formas alternativas de implementación de la API.

Capítulo 7. CONCLUSIONES Y TRABAJOS FUTUROS

En este apartado se comentan las conclusiones del proyecto, los objetivos cumplidos y las posibles líneas de trabajo futuras a considerar.

7.1 CONCLUSIONES

La herramienta desarrollada, O.C.A.S. (*OT Cyberrisk Assessment System*), es una solución eficaz en su propósito de proporcionar inteligencia contextualizada a sistemas OT, concretamente en el sector energético, un ámbito cada vez más expuesto a ciberataques complejos con motivaciones geopolíticas. En un contexto donde los sistemas de control industrial (SCI) se han convertido en objetivos prioritarios de actores estatales y no estatales, esta plataforma proporciona una capacidad de análisis fundada en la inteligencia de amenazas contextualizada, y ayuda a responder con conciencia de situación adaptada al entorno en el que ocurre.

O.C.A.S. ha sido capaz de integrar fuentes de datos en tiempo real, aplicar técnicas de procesamiento de lenguaje natural para extraer conocimiento estructurado, y utilizar modelos de machine learning para predecir de forma automatizada el riesgo asociado a un evento en una instalación, negocio y/o empresa con infraestructura crítica OT. Todo ello permite una evaluación de amenazas más precisa, no solo en función de los aspectos técnicos del ataque, sino también teniendo en cuenta el contexto global en el que se produce.

Además, puede afirmarse que los cuatro objetivos principales del proyecto han sido plenamente alcanzados. Se ha entrenado un modelo de machine learning para la predicción de riesgo con datos históricos de ciberataques; se ha desarrollado una herramienta completa de ingesta, análisis y contextualización de eventos de ciberinteligencia para entornos OT; se ha integrado el sistema con el SOC a través de una API que permite consumir los resultados del análisis en tiempo real; y se ha creado una interfaz gráfica funcional, que incluye un

dashboard en Kibana conectado a los datos procesados, facilitando la visualización dinámica de amenazas relevantes.

A pesar de algunos desafíos, como la lentitud puntual del modelo de lenguaje debido a su ejecución en entornos de desarrollo o las limitaciones para integrar algunas fuentes inicialmente consideradas (como GDELT o AlienVault), la herramienta ha demostrado ser funcional, adaptable y útil para la toma de decisiones estratégicas en ciberseguridad industrial. Estas dificultades técnicas han sido compensadas con soluciones alternativas y decisiones de diseño que priorizaron la eficiencia del sistema.

En definitiva, O.C.A.S. cumple con su propósito y con los objetivos definidos al inicio del proyecto, y los resultados obtenidos han sido tan positivos que ya se están considerando diversas líneas de evolución del sistema, como el despliegue en servidores dedicados con mayor capacidad de procesamiento y la integración completa con las plataformas ya disponibles en el CSIRT (*Computer Security Incident Response Team*) corporativo, como Mandiant, VirusTotal y otras. Esto permitirá enriquecer aún más el sistema, aumentar su precisión y consolidarlo como una herramienta estratégica para la ciberseguridad industrial.

7.2 TRABAJOS FUTUROS

Algunas líneas de trabajo futuras a considerar son:

- Integración de más fuentes de inteligencia. A medida que se vaya evolucionando la herramienta modular y, en caso de que cambien los objetivos de la empresa o el panorama de amenazas de ciberseguridad, se irán añadiendo nuevas fuentes. Por ejemplo, se podría valorar agregar otras fuentes como IP Address Locator o Censys para enriquecer los datos de IoCs, además de otras como EDRs/XDRs (CrowdStrike, FortiEDR, etc.).
- Integración con otras herramientas empresariales. Un trabajo futuro será integrar la herramienta desarrollada con la CMDB (*Configuration Management DataBase*) de activos, lo que permitirá contextualizar aún más la valoración de vulnerabilidades.

- Mejoras en la API. Para mejorar su rendimiento, se plantea como posible solución la automatización del proceso de carga de datos o el despliegue en un servidor dedicado con capacidad de proceso solvente, lo que permitiría mantener el LLM activo y reducir significativamente los tiempos de respuesta. Otra posibilidad que se valorará es utilizar la versión empresarial de Render en lugar de la versión de prueba que se ha utilizado en el proyecto, para asegurar que la memoria a consumir puede ser mayor.
- Integración de datos de competencias entre empresas a la información de contexto geopolítico.
- Fine-tuning de los LLMs empleados para mejorar su rendimiento.

Capítulo 8. BIBLIOGRAFÍA

[1]	MITRE, «ATT&CK,» MITRE, 2025. [En línea]. Available: https://attack.mitre.org/ .
[2]	K. Knerler, I. Parker y C. Zimmerman, «11 Strategies of a World-Class Cybersecurity Operations Center,» MITRE, 2022.
[3]	Instituto Español de Geopolítica, «#OpSpain: Saltan las alarmas entre las autoridades y la prensa españolas ante los ciberataques de los hackers rusos por su apoyo a Zelesnky,» 2025. [En línea]. Available: https://geoestrategia.eu/noticia/44264/ultimas-noticias/opspain:-saltan-las-alarmas-entre-las-autoridades-y-la-prensa-espanolas-ante-los-ciberataques-de-los-hackers-rusos-por-su-apoyo-a-zelesnky.html .
[4]	GreyNoise, «Using the GreyNoise Query Language (GNQL),» 2025. [En línea]. Available: https://docs.greynoise.io/docs/using-the-greynoise-query-language-gnql .
[5]	Terrabyte, «The Pyramid of Pain in Cybersecurity: A Strategic Approach to Threat Hunting,» 2024. [En línea]. Available: https://www.terrabytegroup.com/pyramid-of-pain-in-cyber-security/ .
[6]	Intel471, «Cyber Geopolitical Intelligence,» 2025. [En línea]. Available: https://intel471.com/titan/cyber-geopolitical-intelligence .
[7]	Nozomi Networks, «Nozomi Threat Intelligence,» 2025. [En línea]. Available: https://es.nozominetworks.com/products/threat-intelligence .

[8]	Fortinet, «FortiGuard Labs Outbreak Alerts,» 2025. [En línea]. Available: https://www.fortinet.com/fortiguard/labs .
[9]	INCIBE, «Industroyer2, el amperio contraataca,» INCIBE, 2023. [En línea]. Available: https://www.incibe.es/incibe-cert/blog/industroyer2-el-amperio-contraataca .
[10]	S. Delgado, «Schneider Electric confirma una brecha de datos más de 40 GB robados de su plataforma de desarrollo,» BitLife Media, 2024. [En línea]. Available: https://bitlifemedia.com/2024/11/schneider-electric-brecha-de-datos-gb-robados-plataforma/ .
[11]	INCIBE, «El grupo criminal Cl0p responsable de la oleada de ciberataques a MOVEit,» INCIBE, 2023. [En línea]. Available: https://www.incibe.es/incibe-cert/publicaciones/bitacora-de-seguridad/el-grupo-criminal-cl0p-responsable-de-la-oleada-de-ciberataques-moveit .
[12]	U. Amir, «Schneider Electric Shipped USB Drives Loaded with Malware,» HackRead, 2018. [En línea]. Available: https://hackread.com/schneider-electric-shipped-usb-drives-loaded-with-malware/ .
[13]	Reuters, «Vestas data 'compromised' by cyber attack,» Reuters, 2021. [En línea]. Available: https://www.reuters.com/markets/europe/vestas-data-compromised-by-cyber-attack-2021-11-22/ .
[14]	Iberdrola, «Descubriendo el potencial de las energías renovables en España,» Iberdrola, 2025. [En línea]. Available: https://www.iberdrolaespana.com/sostenibilidad/energias-renovables .
[15]	Kaspersky, «¿Qué es el ransomware WannaCry?,» Kaspersky, 2019. [En línea]. Available: https://www.kaspersky.es/resource-center/threats/ransomware-

	wannacry?srsltid=AfmBOorkt7x-CwwwIh_L2nWRCgmmw1Uu8_EcRlIzXu7DhAzrcuvhi01K.
[16]	Oasis Open, «Comparing STIX 1.X/CybOX 2.X with STIX 2,» 2024. [En línea]. Available: https://oasis-open.github.io/cti-documentation/stix/compare.html .
[17]	J. Deshpande y S. Shinde, «Wannacry Ransomware Attack: Lessons from a,» <i>International Journal of Novel Research and Development</i> , vol. 10, nº 2, pp. 5-7, 2025.
[18]	INCIBE, «CrashOverride, el malware que sabotó el suministro eléctrico,» INCIBE, 2017. [En línea]. Available: https://www.incibe.es/incibe-cert/publicaciones/bitacora-de-seguridad/crashoverride-el-malware-saboteo-el-suministro-electrico .
[19]	FortiGuard Labs, «2025 Global Threat Landscape Report,» 2025.
[20]	Pacto Mundial, «17 Objetivos de Desarrollo Sostenible (ODS),» Pacto Mundial, 2025. [En línea]. Available: https://www.pactomundial.org/que-puedes-hacer-tu/ods/?gad_source=1&gad_campaignid=21296951996&gclid=EAIaIQobChMI67Ls6qPSjQMVjTsGAB2lKi0aEAAYASAAEgI1BPD_BwE .
[21]	Streamlit, «Streamlit documentation,» 2025. [En línea]. Available: https://docs.streamlit.io/ .
[22]	IndustrialCyberCo, «Growing convergence of geopolitics and cyber warfare continue to threaten OT and ICS environments in 2024,» 2024. [En línea]. Available: https://industrialcyber.co/features/growing-convergence-of-geopolitics-and-cyber-warfare-continue-to-threaten-ot-and-ics-environments-in-2024/ .

[23]	Elastic, «Kibana,» Elastic, 2025. [En línea]. Available: https://www.elastic.co/es/kibana .
[24]	C. Harry y N. Gallagher, «Classifying Cyber Events: A proposed taxonomy,» <i>CISSM Working Paper</i> , pp. 5-9, 2018.
[25]	ACLED, «Tipos de eventos y subtipos de eventos geopolíticos,» ACLED, 2025. [En línea]. Available: https://acleddata.com/knowledge-base/codebook/#event-types-and-sub-event-types .
[26]	A. G, «Los gansos que salvaron Roma de los galos,» National Geographic, 2021. [En línea]. Available: https://historia.nationalgeographic.com.es/a/gansos-salvaron-roma-galos_15359 .
[27]	INCIBE, «Triton, un nuevo malware que afecta a infraestructuras industriales,» INCIBE, 2017. [En línea]. Available: https://www.incibe.es/incibe-cert/publicaciones/bitacora-de-seguridad/triton-nuevo-malware-afecta-infraestructuras-industriales .
[28]	Newspaper, «Newspaper3k: Article scraping & curation,» Newspaper, 2025. [En línea]. Available: https://newspaper.readthedocs.io/en/latest/ .
[29]	United States Government, «North American Industry Classification System,» United States Census Bureau, 2025. [En línea]. Available: https://www.census.gov/naics/ .
[30]	GreyNoise, «GreyNoise for Threat Hunting Teams,» GreyNoise, 2025. [En línea]. Available: https://www.greynoise.io/products/threat-hunting-teams .
[31]	CaseIQ, «101+ OSINT Resources for Investigators,» CaseIQ, 2025. [En línea]. Available: https://www.caseiq.com/resources/101-osint-resources-for-investigators/ .

[32]	Cloudflare, «¿Qué es STIX/TAXII?,» 2025. [En línea]. Available: https://www.cloudflare.com/es-es/learning/security/what-is-stix-and-taxii/ .
[33]	University of Maryland, «Cyber Events Database,» 2024. [En línea]. Available: https://cissm.umd.edu/cyber-events-database .
[34]	PyPI, «GoogleNews 1.6.15,» 2025. [En línea]. Available: https://pypi.org/project/GoogleNews/ .
[35]	NIST, «National Vulnerability Database,» 2025. [En línea]. Available: https://nvd.nist.gov/ .
[36]	IBM, «What is an API?,» 2023. [En línea]. Available: https://www.ibm.com/topics/api .
[37]	UCDP, «World Incidents and Conflicts 1989-2024,» 2024. [En línea]. Available: https://www.kaggle.com/datasets/keithtutwiler/world-incidents-and-conflicts-1989-2024 .
[38]	IBM Security, «X-Force Threat Intelligence Index,» IBM, 2024.
[39]	V. Clairoux-Trépanier, I.-M. Beauchamp, E. Ruellan, M. Paquet-Clouston, S.-O. Paquette y E. Clay, «The Use of Large Language Models (LLM) for Cyber Threat,» 2024.
[40]	OTX AlienVault, «Open Threat Exchange™ FAQ,» OTX AlienVault, 2025. [En línea]. Available: https://success.alienvault.com/s/question/0D53q0000BpKk4bCQC/what-are-pulses-in-alienvault .

ANEXO A. ALINEACIÓN CON OBJETIVOS DE DESARROLLO SOSTENIBLE (ODS)

Los Objetivos de Desarrollo Sostenible (ODS) forman parte de la Agenda 2030 para el Desarrollo Sostenible, un plan de acción global aprobado en 2015 por la Asamblea General de Naciones Unidas para dar respuesta a los grandes desafíos mundiales: la pobreza, el hambre, la corrupción, el cambio climático, etc. Los ODS son en total 17 objetivos:



Ilustración 48. Objetivos de Desarrollo Sostenible

Alinear sus proyectos con los ODS permite a las empresas fortalecer su compromiso con la sociedad y el medio ambiente y tener un impacto más positivo en su entorno. Es esencial integrar estos objetivos en cualquier iniciativa o proyecto de desarrollo.

El proyecto desarrollado se alinea con varios de los ODS, entre ellos:

ODS 7. Energía asequible y no contaminante.

El Objetivo de Desarrollo Sostenible 7 busca, en el corto plazo, asegurar que todas las personas tengan acceso a una energía que sea asequible, segura, sostenible y moderna, con el fin de mejorar su calidad de vida. A largo plazo, también tiene como meta fomentar el uso de fuentes de energía renovables, para reducir progresivamente la dependencia de los combustibles fósiles. Al proteger infraestructuras críticas del sector energético, en concreto de Iberdrola y sus instalaciones de generación de energía renovable, la herramienta desarrollada ayuda a garantizar un suministro energético seguro, fiable y resiliente, lo cual es esencial para el desarrollo sostenible.

ODS 9. Industria, innovación e infraestructura.

El Objetivo de Desarrollo Sostenible 9 consiste en construir infraestructuras resilientes, promover la industrialización inclusiva y sostenible, y fomentar la innovación. En el contexto del sector energético, la digitalización y automatización de infraestructuras críticas requiere soluciones avanzadas de ciberseguridad. Este proyecto impulsa la innovación tecnológica en ciberinteligencia industrial, fortaleciendo la resiliencia de los sistemas de control y contribuyendo a una infraestructura energética más segura.

ODS 13. Acción por el clima.

El Objetivo de Desarrollo Sostenible 13 promueve la adopción de medidas urgentes para combatir el cambio climático y sus efectos. Las infraestructuras energéticas, especialmente las relacionadas con energías renovables, son fundamentales para la transición hacia un modelo bajo en emisiones. Al proteger estas infraestructuras frente a amenazas de ciberseguridad, la herramienta contribuye a la continuidad operativa de sistemas que apoyan la descarbonización, ayudando así a mitigar los impactos del cambio climático.

ODS 17. Alianzas para lograr los objetivos.

El Objetivo de Desarrollo Sostenible 17 busca revitalizar la Alianza Mundial para el Desarrollo Sostenible, fomentando la cooperación entre gobiernos, sector privado y sociedad civil. La herramienta desarrollada promueve el uso de fuentes de inteligencia (algunas de ellas abiertas) y, al ser un proyecto modular, permite la integración de otras fuentes fácilmente. A futuro, podría utilizarse también para compartir inteligencia con otras empresas del sector energético.

ANEXO B. CÓDIGOS COMPLETOS

Se incluyen unicamente los códigos que engloban el flujo completo. El resto de códigos se puede consultar en el repositorio.

Código de la API

```
from fastapi import FastAPI, Query
from fastapi.responses import JSONResponse, Response
from threading import Thread
import pandas as pd
import time

from pipeline import run_pipeline

app = FastAPI(title="Threat Intel API")

# Cargar el DataFrame al iniciar
#df = pd.read_excel("DATA_API3.xlsx")

# Variable cacheada
cached_df = pd.DataFrame()

# Función que actualiza el cache cada 30 minutos
def actualizar_cache():
    global cached_df
    while True:
        try:
            print("Actualizando datos...")
            cached_df = run_pipeline()
            print("Cache actualizado correctamente.")
        except Exception as e:
            print(f"Error al actualizar cache: {e}")
            time.sleep(1800) # Espera 30 minutos

# Iniciar el hilo de actualización en segundo plano
@app.on_event("startup")
def start_background_tasks():
    Thread(target=actualizar_cache, daemon=True).start()

@app.get("/")
def read_root():
    return {"message": "Welcome to the O.C.A.S. API"}

@app.get("/events")
def get_all_events():
    #json_text = df.to_json(orient="records", date_format="iso")
    #return Response(content=json_text, media_type="application/json")
```

```
json_text = cached_df.to_json(orient="records", date_format="iso")
return Response(content=json_text, media_type="application/json")

@app.get("/search")
def search(query: str = Query(..., description="Busca por CVE o IP")):
    #filtered = df[df["cve_id"].astype(str).str.contains(query, na=False,
case=False) |
    #
df["ip_info"].astype(str).str.contains(query, na=False,
case=False)]
    #return JSONResponse(content=filtered.to_dict(orient="records"))
    filtered =
cached_df[cached_df["cve_id"].astype(str).str.contains(query, na=False,
case=False) |
cached_df["ip_info"].astype(str).str.contains(query, na=False,
case=False)]
    return JSONResponse(content=filtered.to_dict(orient="records"))
```

Código del pipeline completo en tiempo real

```
# pipeline.py

import requests
import time
import pandas as pd
import dateparser
import re
import json
import feedparser
from urllib.parse import urlparse
from datetime import datetime
from datetime import timedelta

from GoogleNews import GoogleNews
from greynoise import GreyNoise
from mitreattack.stix20 import MitreAttackData
import shodan
from tqdm import tqdm

import tempfile

import joblib
import numpy as np
from tensorflow.keras.models import load_model
from sklearn.metrics import mean_squared_error, mean_absolute_error,
r2_score
from sklearn.preprocessing import OneHotEncoder, LabelEncoder,
StandardScaler
from sklearn.model_selection import train_test_split
from tensorflow.keras.models import Model
from tensorflow.keras.layers import Input, Embedding, Flatten,
Concatenate, Dense, Dropout, BatchNormalization
from tensorflow.keras.optimizers import Adam
```

```
# NVD
api_key_nvd = "302bba5e-a251-4880-b590-7dc767aa4f58"
# GreyNoise
api_key_gn
="epcwbdMMsQOwMsIa2jRBrblcSrRtql3kurjv7zPlKRyEo9wHUj52veubchRoCU9Y"
# LLM OpenRouter
#api_key_open = "sk-or-v1-
2c3c02273fe6f7e87d39d75db37b7a2f48c13c63e648a2c038c64de8f6c3e1f8"
api_key_open = "sk-or-v1-
721d1e0bfdf56c9d8e84d144e563b201830d4ce50bee6f720280716e4edf7d71"
# ACLED
api_key_acled = "DxFtYBdua9YS1fVxqVAv"
email = "201905597@alu.comillas.edu"
# SHODAN
api_key_shodan = "guAgZYlDUy9Uga27UptB3VmzeTlTxnxw"

# CVEs de proveedores
# CVEs de NVD

headers_nvd = {"apiKey": api_key_nvd}

# Lista de proveedores
proveedores = [
    'Forti', 'ABB', 'Siemens', 'Siemens Energy', 'Vestas Wind Systems',
    'Cisco',
    'Palo Alto', 'Schneider Electric', 'General Electric', 'Atos',
    'Nokia',
    'Windows', 'ARC Informatique', 'Yokogawa', 'Westinghouse Electric
Company', 'Amper'
]

# Función para buscar CVEs por proveedor
def buscar_cves_por_proveedor(proveedor):
    url = "https://services.nvd.nist.gov/rest/json/cves/2.0"
    params = {"keywordSearch": proveedor, "resultsPerPage": 1000}
    r = requests.get(url, params=params, headers=headers_nvd)
    cves = []
    if r.status_code == 200:
        data = r.json()
        for item in data.get("vulnerabilities", []):
            cve_data = item["cve"]
            cve_id = cve_data["id"]
            if cve_id.startswith("CVE-2025"):
                description = cve_data["descriptions"][0]["value"]
                metrics = cve_data.get("metrics", {})
                base_score = None

                # Determinar score de CVSS v3 si existe
                for key in metrics:
                    if key.startswith("cvssMetricV3"):
                        base_score =
metrics[key][0]["cvssData"]["baseScore"]

                cves.append({
```

```
        "proveedor": proveedor,
        "cve_id": cve_id,
        "description": description,
        "score": base_score,
        "url": f"https://nvd.nist.gov/vuln/detail/{cve_id}"
    })
    return cves

# Noticias CVEs
# Buscar noticias asociadas a los CVEs
def buscar_noticias_cves(cves, days=7):
    googlenews = GoogleNews(lang='en', period=f'{days}d')
    noticias = []

    for cve in cves:
        cve_id = cve['cve_id']
        proveedor = cve.get("proveedor", "")
        googlenews.search(f"cyber attack {cve_id}")
        results = googlenews.results()
        for r in results[:1]: # Limitar a 1 noticia por CVE
            parsed_date = dateparser.parse(r['date'])
            noticias.append({
                "event_date": parsed_date,
                "organization": proveedor,
                "title": r['title'],
                "description": r['desc'],
                "link": r['link'],
                "cve_id": cve_id,
            })
    return pd.DataFrame(noticias)

# --- Obtención de noticias de ciberataques recientes generales
def buscar_noticias():
    googlenews = GoogleNews(lang='en', period='7d')
    googlenews.search('cyber attack')
    results = googlenews.results()[:5] # Solo las primeras 5
    noticias_procesadas = []
    for r in results:
        parsed_date = dateparser.parse(r['date'])
        noticias_procesadas.append({
            'event_date': parsed_date,
            'title': r.get('title', ''),
            'link': r.get('link', ''),
            'description': r.get('desc', '')
        })

    df_noticias = pd.DataFrame(noticias_procesadas)
    return df_noticias

# UNIFICAR NOTICIAS
def unificar_noticias(noticias_cves_df, noticias_df):
    noticias_df["organization"] = "None"
```

```

    noticias_df["cve_id"] = "None"
    df_unificado = pd.concat([noticias_cves_df, noticias_df],
ignore_index=True)
    return df_unificado

# VERIFICACIÓN DE NOTICIAS CONFIABLES
def es_noticia_confiable(row):

    url = row['link']
    titulo = row.get('title', '')
    descripcion = row.get('description', '')

    # Dominio con TLD confiable
    dominio = urlparse(url).netloc
    tld_confiables = (".com", ".org", ".net", ".es", ".co.uk", ".io")
    tld_ok = dominio.endswith(tld_confiables)

    # Alertas
    señales_fake = [
        r"BREAKING",          # En mayúsculas, típico de titulares
sensacionalistas
        r"URGENT",            # Urgencia falsa
        r"SHOCKING",          # Sensacionalismo
        r"EXCLUSIVE",         # Clickbait
        r"YOU WON'T BELIEVE", # Engaño clásico
        r"MUST SEE",          # Exageración
        r"THIS CHANGES EVERYTHING", # Hipérbole
        r"EVERYONE IS TALKING ABOUT", # Falacia de popularidad
        r"HIDDEN TRUTH",      # Conspiración
        r"SECRET [A-Z\s]+"    # Sensacionalista
    ]
    contiene_alertas = any(re.search(p, titulo + " " + descripcion,
re.IGNORECASE) for p in señales_fake)
    return not contiene_alertas and tld_ok

# FILTRADO NOTICIAS FAKE
def filtrar_noticias_confiables(df_noticias):
    df_filtrado = df_noticias[df_noticias.apply(es_noticia_confiable,
axis=1)].reset_index(drop=True)
    return df_filtrado

# SCRAPING + LLM
# Construcción de prompt
def construir_prompt(title, full_text, url, date, cve, proveedor):
    return f"""You are a cybersecurity analyst.
Given the following news article:

Title: {title}
Full text: {full_text}
URL: {url}
Date: {date}
CVE: {cve}
Proveedor: {proveedor}

```

Extract the following fields if possible:

- event_date en este formato: 20/01/2025 0:00:00 (can't be Undetermined, if you know year and month write day 1 of that month or the date of the article; you must refer to Date)
- actor (author of the attack; if not known, write "Undetermined")
- organization (attacked; write content of Proveedor indicated at the start of this prompt if not "None", or write "Undisclosed" if not mentioned, or find it yourself in Full text)
- industry ('Agriculture, Forestry, Fishing and Hunting',
'Mining, Quarrying, and Oil and Gas Extraction', 'Utilities',
'Construction', 'Manufacturing', 'Wholesale Trade', 'Retail
Trade',
'Transportation and Warehousing', 'Information',
'Finance and Insurance', 'Real Estate and Rental and Leasing',
'Professional, Scientific, and Technical Services',
'Management of Companies and Enterprises',
'Administrative and Support and Waste Management and Remediation
Services',
'Educational Services', 'Health Care and Social Assistance',
'Arts, Entertainment, and Recreation',
'Accommodation and Food Services',
'Other Services (except Public Administration)',
'Public Administration', 'Undetermined')
- event_type (Disruptive, Exploitative, Mixed)
- event_subtype ('Message Manipulation', 'Exploitation of Sensors',
'Exploitation of End Host', 'Physical Attack',
'Exploitation of Network Server', 'Exploitation of Data
in Transit', 'Exploitation of Network Infrastructure',
'External Denial of Service', 'Exploitation of
Application Server',
'Internal Denial of Service', 'Exploitation of End
Hosts',
'Exploitation of End User', 'Data Attack', 'Exploitation
of Sensors', 'Exploitation of Infrastructure', 'Undetermined'),
can be several with commas between them
- country (of the attack, in upper case)
- actor_country (if known write in upper case, if not write
"UNDETERMINED")
- cve_id: CVE (indicated at the start of this prompt)

"country" and "actor_country", if not UNDETERMINED, must be within:

'THAILAND', 'GERMANY', 'SUDAN', 'PANAMA', 'JORDAN', 'BOLIVIA',
'MOROCCO', 'POLAND', 'ROMANIA', 'MEXICO', 'SOUTH KOREA',
'LESOTHO',
'BONAIRE, SINT EUSTATIUS AND SABA', 'ITALY', 'CROATIA', 'QATAR',
'LAOS', 'NAURU', 'LIBYA', 'NORTH KOREA', 'TURKEY', 'ZAMBIA',
'CHILE', 'IRAN', 'MONACO', 'ISRAEL', 'INDIA', 'BULGARIA', 'MALI',
'AFGHANISTAN', 'BERMUDA', 'HUNGARY', 'TAIWAN', 'SLOVENIA',
'NORTH MACEDONIA', 'TRINIDAD AND TOBAGO', 'VATICAN CITY',
'GUATEMALA', 'LIECHTENSTEIN', 'MONTENEGRO', 'PARAGUAY',
'GUADELOUPE', 'HONG KONG', 'CHINA', 'NEPAL', 'GEORGIA', 'VIETNAM',
'TAJIKISTAN', 'VANUATU', 'SINGAPORE', 'CYPRUS', 'NORWAY',
'UNITED KINGDOM', 'AZERBAIJAN', 'AUSTRIA', 'CABO VERDE', 'UGANDA',
'SENEGAL', 'AUSTRALIA', 'RWANDA', 'ICELAND', 'CZECH REPUBLIC',

```
'INDONESIA', 'CAYMAN ISLANDS', 'NIGERIA', 'MALAYSIA', 'CUBA',
'TANZANIA', 'CZECHIA', 'SAINT VINCENT AND THE GRENADINES',
'SRI LANKA', 'ECUADOR', 'MALDIVES', 'NEW ZEALAND', 'ALGERIA',
'VENEZUELA', 'TUNISIA', 'BOSNIA AND HERZEGOVINA', 'IRAQ', 'OMAN',
'GHANA', 'PHILIPPINES', 'PAKISTAN', 'DENMARK', 'LITHUANIA',
'CANADA', 'AMERICAN SAMOA', 'MULTIPLE', 'NAMIBIA', 'LATVIA',
'URUGUAY', 'FIJI', 'GREECE', 'SERBIA', 'ESTONIA', 'RUSSIA',
'CAMBODIA', 'TURKMENISTAN', 'BAHRAIN', 'MONGOLIA', 'MALAWI',
'JAPAN', 'SYRIA', nan, 'KAZAKHSTAN', 'FRANCE', 'MALTA', 'SPAIN',
'BAHAMAS', 'UNDETERMINED', 'SIERRA LEONE', 'BELGIUM', 'COSTA
RICA',
'ANGOLA', 'MACAU', 'SAINT THOMAS', 'GUAM', 'LUXEMBOURG',
'EUROPEAN UNION', 'ETHIOPIA', 'BRAZIL', 'YEMEN', 'ARGENTINA',
'EGYPT', 'UZBEKISTAN', 'EL SALVADOR', 'KOSOVO', 'PUERTO RICO',
'ALBANIA', 'SAUDI ARABIA', 'KIRIBATI', 'UKRAINE', 'ZIMBABWE',
'SOUTH AFRICA', 'ISLE OF MAN', 'KUWAIT', 'TONGA', 'PORTUGAL',
'SWEDEN', 'BANGLADESH', 'COLOMBIA', 'PERU', 'SWITZERLAND',
'KENYA',
'BARBADOS', 'UNITED ARAB EMIRATES', 'SLOVAKIA', 'NETHERLANDS',
'FINLAND', 'GREENLAND', 'NICARAGUA', 'DOMINICAN REPUBLIC',
'ANDORRA', 'PAPUA NEW GUINEA', 'ARMENIA', 'SEYCHELLES', 'MYANMAR',
'JAMAICA', 'PALESTINE', 'GIBRALTAR', 'MOLDOVA', 'BELARUS',
'IRELAND', 'PALAU', 'UNITED STATES OF AMERICA', 'LEBANON', 'TOGO'
```

Return them as a JSON object. Remember we are in 2025 so all the news collected are from 2025.

```
"""

def extraer_info_llm(prompt):
    headers = {
        "Authorization": f"Bearer {api_key_open}", # tu token
        "Content-Type": "application/json"
    }
    data = {
        "model": "meta-llama/llama-3.3-70b-instruct:free",
        #"model": "google/gemma-3n-e4b-it:free",
        "messages": [{"role": "user", "content": prompt}]
    }

    response =
requests.post("https://openrouter.ai/api/v1/chat/completions", json=data,
headers=headers)
    if response.status_code == 200:
        return response.json()[0]['choices'][0]['message']['content']
    else:
        print(f"Error en LLM: {response.status_code}")
        return None

def scrap_and_extract(noticias_df):
    llm_outputs = []
    delay_seconds = 2

    for _, row in noticias_df.iterrows():
        url = row['link']
```

```
try:
    article = Article(url)
    article.download()
    article.parse()
    full_text = article.text.strip()

    if not full_text:
        full_text = row['description'] # fallback
except:
    full_text = row['description'] # fallback si hay error en
scraping

# Convertir fecha al formato requerido: mm/dd/yyyy 0:00:00
raw_date = row['event_date']
try:
    parsed_date = pd.to_datetime(raw_date)
    formatted_date = parsed_date.strftime("%d/%m/%Y 0:00:00")
except:
    formatted_date = "01/06/2025 0:00:00" # fallback si falla el
parsing

prompt = construir_prompt(row['title'], full_text, url,
formatted_date, row['cve_id'], row['organization'])
# Sugerencia explícita al LLM de usar esta fecha si no puede
extraer una
prompt += f'\n\nIf the date of the attack is not clearly stated
in the article, use this publication date instead: {formatted_date}'

info = extraer_info_llm(prompt)
llm_outputs.append(info)
#time.sleep(delay_seconds)
return llm_outputs

# EXPANDIR FILAS DF
def expand_row(row, countries_complete_list):
    orgs = [o.strip() for o in str(row.get('organization',
'')).split(',') if o.strip()]
    countries = [c.strip() for c in str(row.get('country',
'')).split(',') if c.strip()]
    actor_countries = [c.strip() for c in str(row.get('actor_country',
'')).split(',') if c.strip()]

    countries = [c for c in countries if c in countries_complete_list]
    actor_countries = [c for c in actor_countries if c in
countries_complete_list]

    orgs = orgs or ["UNDETERMINED"]
    countries = countries or ["UNDETERMINED"]
    actor_countries = actor_countries or ["UNDETERMINED"]

    return [
        {
```

```
        **row,
        'organization': o,
        'country': c,
        'actor_country': a
    }
    for o in orgs
    for c in countries
    for a in actor_countries
]

# Paso 1 - conflictos en país del ciberataque
def update_event_columns(ataque_new, datos_geo):
    event_types = {
        'Protests': 'protests_country',
        'Violence against civilians': 'violence_country',
        'Strategic developments': 'strategic_country',
        'Riots': 'riots_country',
        'Battles': 'battles_country',
        'Explosions/Remote violence': 'explosions_country',
    }

    # Preconteo de eventos por país y tipo
    event_counts = datos_geo.groupby(['country',
    'event_type']).size().unstack(fill_value=0)

    for index, row in ataque_new.iterrows():
        country = row['country']
        for event_type, column_name in event_types.items():
            if country in event_counts.index and event_type in
event_counts.columns:
                ataque_new.at[index, column_name] =
event_counts.at[country, event_type]
            else:
                ataque_new.at[index, column_name] = 0 # Valor por
defecto si no hay eventos

# Paso 2 - conflictos en país del atacante
def update_event_columns(ataque_new, datos_geo):
    event_types = {
        'Protests': 'protests_actor_country',
        'Violence against civilians': 'violence_actor_country',
        'Strategic developments': 'strategic_actor_country',
        'Riots': 'riots_actor_country',
        'Battles': 'battles_actor_country',
        'Explosions/Remote violence': 'explosions_actor_country',
    }

    # Preconteo de eventos por país y tipo
    event_counts = datos_geo.groupby(['country',
    'event_type']).size().unstack(fill_value=0)

    for index, row in ataque_new.iterrows():
        country = row['actor_country']
```

```
        for event_type, column_name in event_types.items():
            if country in event_counts.index and event_type in
event_counts.columns:
                ataque_new.at[index, column_name] =
event_counts.at[country, event_type]
            else:
                ataque_new.at[index, column_name] = 0 # Valor por
defecto si no hay eventos

# Función para descargar los datos JSON desde las URLs de MITRE
def load_attack_data_from_url(url):
    response = requests.get(url)
    response.raise_for_status() # Verificar que la solicitud fue exitosa
    return response.json() # Devolver los datos JSON

# Filtrar intrusion sets (actores) de los datos STIX
def get_intrusion_sets(data):
    return [obj for obj in data['objects'] if obj.get('type') ==
'intrusion-set']

# Función para buscar actor por nombre o alias
def buscar_actor(nombre, intrusion_sets):
    for actor in intrusion_sets:
        nombre = str(nombre).lower()
        if nombre.lower() in actor['name'].lower() or any(nombre.lower()
in alias.lower() for alias in actor.get('aliases', [])):
            return actor
    return None

# Función para obtener las técnicas asociadas a un actor
def obtener_ttps(actor, data):
    ttp_objs = []

    # Buscar relaciones 'uses' para encontrar las técnicas asociadas
    for relationship in data['objects']:
        if relationship.get('type') == 'relationship' and
relationship.get('source_ref') == actor['id'] and
relationship.get('relationship_type') == 'uses':
            target_ref = relationship.get('target_ref')
            # Buscar el objeto de la técnica asociada (attack-pattern)
            related_obj = next((obj for obj in data['objects'] if
obj['id'] == target_ref), None)
            if related_obj and related_obj.get('type') == 'attack-
pattern':
                ttp_objs.append(related_obj)

    return ttp_objs

# Función para obtener los datos desde el diccionario
def get_actor_data(actor_name, actor_info):
    actor_name = str(actor_name).lower().strip()
    #actor_name = actor_name.lower()
```

```

    return actor_info.get(actor_name, {"TTPs": None, "Attacks ICS":
None})

# Paso 5 - Conflictos entre aliados de "country" y "actor_country" y
Conflictos entre "country" y aliados de "actor_country"
def get_allies(country: str, year: int) -> list:
    alliances = {
        'NATO': [
            ("ALBANIA", 2009), ("BELGIUM", 1949), ("BULGARIA", 2004),
            ("CANADA", 1949),
            ("CROATIA", 2009), ("CZECH REPUBLIC", 1999), ("DENMARK",
1949), ("ESTONIA", 2004),
            ("FINLAND", 2023), ("FRANCE", 1949), ("GERMANY", 1955),
            ("GREECE", 1952),
            ("HUNGARY", 1999), ("ICELAND", 1949), ("ITALY", 1949),
            ("LATVIA", 2004),
            ("LITHUANIA", 2004), ("LUXEMBOURG", 1949), ("MONTENEGRO",
2017),
            ("NETHERLANDS", 1949), ("NORTH MACEDONIA", 2020), ("NORWAY",
1949),
            ("POLAND", 1999), ("PORTUGAL", 1949), ("ROMANIA", 2004),
            ("SLOVAKIA", 2004),
            ("SLOVENIA", 2004), ("SPAIN", 1982), ("SWEDEN", 2024),
            ("TURKEY", 1952),
            ("UNITED KINGDOM", 1949), ("UNITED STATES OF AMERICA", 1949)
        ],
        'CSTO': [
            ("RUSSIA", 1992), ("BELARUS", 1992), ("ARMENIA", 1992),
            ("KAZAKHSTAN", 1992), ("KYRGYZSTAN", 1992), ("TAJIKISTAN",
1992)
        ],
        'EU': [
            ("BELGIUM", 1958), ("FRANCE", 1958), ("GERMANY", 1958),
            ("ITALY", 1958),
            ("NETHERLANDS", 1958), ("LUXEMBOURG", 1958), ("DENMARK",
1973), ("IRELAND", 1973),
            ("GREECE", 1981), ("SPAIN", 1986), ("PORTUGAL", 1986),
            ("AUSTRIA", 1995),
            ("SWEDEN", 1995), ("FINLAND", 1995), ("CZECH REPUBLIC",
2004), ("ESTONIA", 2004),
            ("HUNGARY", 2004), ("LATVIA", 2004), ("LITHUANIA", 2004),
            ("POLAND", 2004),
            ("SLOVAKIA", 2004), ("SLOVENIA", 2004), ("BULGARIA", 2007),
            ("ROMANIA", 2007),
            ("CROATIA", 2013)
        ],
        'Five_Eyes': [
            ("UNITED STATES OF AMERICA", 1946), ("UNITED KINGDOM", 1946),
            ("CANADA", 1946), ("AUSTRALIA", 1946), ("NEW ZEALAND", 1946)
        ],
        'SCO': [
            ("CHINA", 2001), ("KAZAKHSTAN", 2001), ("KYRGYZSTAN", 2001),
            ("RUSSIA", 2001),
            ("TAJIKISTAN", 2001), ("UZBEKISTAN", 2001), ("INDIA", 2017),
            ("PAKISTAN", 2017),

```

```

        ("IRAN", 2023), ("BELARUS", 2024)
    ],
    'FPDA': [
        ("UNITED KINGDOM", 1971), ("AUSTRALIA", 1971),
        ("NEW ZEALAND", 1971), ("MALAYSIA", 1971), ("SINGAPORE",
1971)
    ],
    'NARC': [
        ("ALGERIA", 2004), ("LIBYA", 2004), ("EGYPT", 2004),
        ("TUNISIA", 2004), ("MAURITANIA", 2004), ("WESTERN SAHARA",
2004)
    ],
    'G5_Sahel': [
        ("BURKINA FASO", 2014), ("CHAD", 2014), ("MALI", 2014),
        ("MAURITANIA", 2014), ("NIGER", 2014)
    ],
    'AES': [
        ("MALI", 2023), ("BURKINA FASO", 2023), ("NIGER", 2023)
    ]
}

# Normaliza el nombre del país
country = country.upper()
allies = set()

for bloc, members in alliances.items():
    for member, entry_year in members:
        if year >= entry_year and country == member:
            # Añade a todos los demás miembros activos en ese año
            for other, other_year in members:
                if other != member and year >= other_year:
                    allies.add(other)

return list(allies)

def extraer_info_ips_por_cve(cve_id):
    try:
        if not cve_id or cve_id == "None":
            return []

        query = f"cve:{cve_id}"
        response = gn.query(query)
        resultados = []

        for entry in response.get("data", []):
            ip = entry.get("ip")
            country = entry.get("metadata", {}).get("country", "")
            classification = entry.get("classification", "")
            last_seen = entry.get("last_seen", "")
            destinos = entry.get("destination_countries", [])

            if ip:
                resultados.append({
                    "ip": ip,
                    "country": country,

```

```
        "classification": classification,
        "last_seen": last_seen,
        "targets_spain": "Spain" in destinos
    })

    return resultados

except Exception as e:
    print(f"Error con GreyNoise para {cve_id}: {e}")
    return []

# EJECUCIÓN DEL PIPELINE ENTERO
def run_pipeline():

    # Noticias CVEs
    todos_los_cves_2025 = []
    for proveedor in proveedores:
        cves = buscar_cves_por_proveedor(proveedor)
        todos_los_cves_2025.extend(cves)
        time.sleep(0.3) # No saturar la API
    cves_activos = [cve for cve in todos_los_cves_2025 if cve["score"]
and cve["score"] >= 7.0]
    noticias_cves_df = buscar_noticias_cves(cves_activos)

    # Búsqueda de noticias de ciberataques generales
    noticias_df = buscar_noticias()

    # Unificar noticias
    df_noticias_completas = unificar_noticias(noticias_cves_df,
noticias_df)

    # Noticias verificadas
    df_noticias_completas_filtradas =
filtrar_noticias_confiabiles(df_noticias_completas)

    # Noticias normalizadas
    noticias_normalizadas =
scrap_and_extract(df_noticias_completas_filtradas)

    # ...
    countries = pd.read_excel("countries_complete_list.xlsx")
    countries_complete_list = countries["valor"].unique()

    json_data = []
```

```
for out in noticias_normalizadas:
    # Buscar bloques dentro de cualquier tipo de triple comillas o
    sin prefijo 'json'
    matches = re.findall(r'```(?:json)?\n(?:.*?)\n```', out, re.DOTALL)

    # Si no se encontró ningún bloque con backticks, intenta con
    comillas simples o dobles
    if not matches:
        matches = re.findall(r'{.*?}', out, re.DOTALL)

    for match in matches:
        try:
            parsed = json.loads(match)
            if isinstance(parsed, list):
                json_data.extend(parsed)
            else:
                json_data.append(parsed)
        except json.JSONDecodeError:
            continue # Silenciosamente ignora los bloques inválidos

# Convertir a DataFrame
df_json = pd.DataFrame(json_data)

# Lista de países válida que debes definir
countries_complete_list = countries_complete_list

# Expandir filas
expanded_rows = []
for _, row in df_json.iterrows():
    expanded_rows.extend(expand_row(row, countries_complete_list))

df_expanded = pd.DataFrame(expanded_rows)

# TTPs + ACLED
plantilla = pd.read_excel("prueba_ataque_reciente.xlsx")
plantilla = plantilla[0:0]
# Asegurarse de que solo se usen columnas que están en ambos
DataFrames
common_cols = plantilla.columns.intersection(df_expanded.columns)
# Añadir las filas (solo las columnas en común)
plantilla = pd.concat([plantilla, df_expanded[common_cols]],
ignore_index=True)
plantilla.fillna(0, inplace=True)
df_cyber_events = plantilla.copy()

# Parámetros de tu cuenta
api_key_acled = "DxFtYBdua9YS1fVxqVAv"
email = "201905597@alu.comillas.edu"
# URL base
url = "https://api.acleddata.com/acled/read"
# Filtros y formato
params = {
    "key": api_key_acled,
```

```
"email": email,
"format": "json", # No usar CSV porque el contenido real no lo
es
"year": "2025",
"limit": 10000 # máximo permitido por llamada
}

# Hacer la petición
response = requests.get(url, params=params)

# Procesar el JSON
if response.status_code == 200:
    data = response.json()
    df_acled = pd.DataFrame(data["data"]) # accedemos a la clave
"data"
    #print(df_acled.head())
else:
    print(f"Error {response.status_code}: {response.text}")

# Creamos una versión más ligera de df_acled, solo con lo que
necesitamos
df_acled = df_acled[['event_date', 'event_type', 'actor1',
'actor2', 'country']].copy()

acled_datos = df_acled.copy()

# Actualizar las columnas en df_cyber_events
update_event_columns(df_cyber_events, acled_datos)

# Actualizar las columnas en df_cyber_events
update_event_columns(df_cyber_events, acled_datos)

# URLs de los archivos JSON de MITRE ATT&CK
enterprise_url = "https://raw.githubusercontent.com/mitre-
attack/attack-stix-data/master/enterprise-attack/enterprise-attack.json"
ics_url = "https://raw.githubusercontent.com/mitre-attack/attack-
stix-data/master/ics-attack/ics-attack.json"

# Descargar los datos de MITRE ATT&CK Enterprise e ICS
enterprise_attack_data = load_attack_data_from_url(enterprise_url)
ics_attack_data = load_attack_data_from_url(ics_url)

# Obtener los intrusion sets (actores)
enterprise_intrusion_sets =
get_intrusion_sets(enterprise_attack_data)
ics_intrusion_sets = get_intrusion_sets(ics_attack_data)

mis_actores = df_cyber_events["actor"].unique()
resultado = []
for nombre in mis_actores:
    actor = buscar_actor(nombre, enterprise_intrusion_sets)
    if actor:
        ttp_objs = obtener_ttps(actor, enterprise_attack_data)
        ttps = [t["name"] for t in ttp_objs]
```

```
# Verificar si el país está presente
country = actor.get("country")
country = country if country else "Not specified" # Si no
tiene país, usar "Not specified"

resultado.append({
    "Name": actor["name"],
    "Aliases": ", ".join(actor.get("aliases", [])),
    "Country": country,
    "TTPs": ", ".join(ttps),
    "Attacks ICS": "Yes" if buscar_actor(nombre,
ics_intrusion_sets) else "No"
})

else:
    resultado.append({
        "Name": nombre,
        "Aliases": "Not found",
        "Country": "Not found",
        "TTPs": "",
        "Attacks ICS": "No"
    })

# Resultado
df_mitre = pd.DataFrame(resultado)
df_mitre = df_mitre.drop(["Country"],axis=1)

df_c = df_mitre.copy()

# Diccionario principal
actor_info = {}

for _, row in df_c.iterrows():
    main_actor = row["Name"]
    ttps = row["TTPs"]
    attacks_ics = row["Attacks ICS"]

    # Añadir el actor principal
    actor_info[main_actor.lower()] = {"TTPs": ttps, "Attacks ICS":
attacks_ics}

    # Añadir los aliases si existen
    if pd.notna(row["Aliases"]):
        aliases = [alias.strip().lower() for alias in
row["Aliases"].split(",")]
        for alias in aliases:
            actor_info[alias] = {"TTPs": ttps, "Attacks ICS":
attacks_ics}

# Aplicar al DataFrame original
df_cyber_events[["TTPs", "Attacks ICS"]] =
df_cyber_events["actor"].apply(
    lambda x: pd.Series(get_actor_data(x,actor_info))
)

# Paso 4 - conflictos between pais victima y pais atacante
```

```
countries_complete_list_df =
pd.read_excel("countries_complete_list.xlsx")
countries_complete_list =
countries_complete_list_df["valor"].tolist()

# Limpiar la lista de países: convertir todo a string y quitar NaNs
countries_cleaned = [str(c).strip().upper() for c in
countries_complete_list if pd.notnull(c)]

# Función para extraer el país del actor
def extract_country_f(actor):
    if pd.isnull(actor):
        return 'Unknown'
    actor_upper = str(actor).upper()
    for country in countries_cleaned:
        if str(country).upper() in actor_upper:
            return country.title().upper()
    return 'Unknown'

# Usamos la función anterior para extraer algunos de los actores
df_acled["actor1_country"] =
df_acled["actor1"].apply(extract_country_f)
df_acled["actor2_country"] =
df_acled["actor2"].apply(extract_country_f)

# Excel que mapea otros actores a países
df_unk_f = pd.read_excel("unknown_groups_final.xlsx")

# Creamos un diccionario para búsqueda rápida de actor -> país
actor_to_country = dict(zip(df_unk_f["actor"], df_unk_f["country"]))

# Rellenamos actor1_country donde es 'Unknown' y actor1 coincide
mask1 = df_acled["actor1_country"] == "Unknown"
df_acled.loc[mask1, "actor1_country"] = df_acled.loc[mask1,
"actor1"].map(actor_to_country).fillna(df_acled.loc[mask1,
"actor1_country"])

# Rellenamos actor2_country donde es 'Unknown' y actor2 coincide
mask2 = df_acled["actor2_country"] == "Unknown"
df_acled.loc[mask2, "actor2_country"] = df_acled.loc[mask2,
"actor2"].map(actor_to_country).fillna(df_acled.loc[mask2,
"actor2_country"])

# Selecciona filas donde actor1 es el grupo indicado
masks_varios = [df_acled["actor1"] == "JNIM: Group for Support of
Islam and Muslims", df_acled["actor1"] == "Unidentified Gang and/or
Police Militia",
df_acled["actor1"] == "Unidentified Gang and/or
Colectivo",
df_acled["actor1"] == "JNIM: Group for Support of Islam
and Muslims and/or Islamic State Sahel Province (ISSP)",
df_acled["actor1"] == "Islamic State West Africa Province
(ISWAP)"]

for m in masks_varios:
```

```
# Copia el valor de 'country' a 'actor1_country' en esas filas
df_acled.loc[m, "actor1_country"] = df_acled.loc[m, "country"]

actor1_list =
df_acled[df_acled["actor1_country"]=="Unknown"]["actor1"].unique()
# Creamos la máscara donde actor1 esté en actor_list
mask = df_acled["actor1"].isin(actor1_list)
# Asignamos el valor de 'country' a 'actor1_country' en esas filas
df_acled.loc[mask, "actor1_country"] = df_acled.loc[mask, "country"]

actor2_list =
df_acled[df_acled["actor2_country"]=="Unknown"]["actor2"].unique()
# Creamos la máscara donde actor1 esté en actor_list
mask = df_acled["actor2"].isin(actor2_list)
# Asignamos el valor de 'country' a 'actor1_country' en esas filas
df_acled.loc[mask, "actor2_country"] = df_acled.loc[mask, "country"]

cols_to_upper = ["country", "actor1_country", "actor2_country"]

for col in cols_to_upper:
    df_acled[col] = df_acled[col].astype(str).str.upper()

# Creamos una versión más ligera de df_acled, solo con lo que
necesitamos
df_acled_light = df_acled[['event_date', 'event_type',
'actor1_country', 'actor2_country']].copy()

# Fechas a datetime
df_cyber_events['event_date'] =
pd.to_datetime(df_cyber_events['event_date'])
df_acled_light['event_date'] =
pd.to_datetime(df_acled_light['event_date'])

# Tipos de conflicto a contar
event_types = [
    'Protests', 'Violence against civilians', 'Strategic
developments',
    'Riots', 'Battles', 'Explosions/Remote violence'
]

# Recorrer cada fila de df_cyber_events
for idx, row in df_cyber_events.iterrows():
    target_country = row['country']
    attacker_country = row['actor_country']
    event_date = row['event_date']

    # Saltar si el atacante es UNDETERMINED
    if attacker_country == "UNDETERMINED":
        continue

    # Filtrar eventos relevantes por fecha y países involucrados
    (ambas direcciones)
```

```

mask_time = (df_acled_light['event_date'] >= event_date -
pd.Timedelta(days=90)) & \
(df_acled_light['event_date'] <= event_date)

mask_country_pair = (
((df_acled_light['actor1_country'] == target_country) &
(df_acled_light['actor2_country'] == attacker_country)) |
((df_acled_light['actor1_country'] == attacker_country) &
(df_acled_light['actor2_country'] == target_country))
)

df_filtered = df_acled_light[mask_time & mask_country_pair]

# Contar eventos por tipo
for evt_type in event_types:
    count = (df_filtered['event_type'] == evt_type).sum()
    col_name = evt_type.lower().replace(' ', '_').replace('/',
' ').replace('violence_against_civilians', 'violence_between') +
"_between"
    df_cyber_events.at[idx, col_name] = count

# columnas en formato datetime
df_cyber_events["event_date"] =
pd.to_datetime(df_cyber_events["event_date"])
df_acled_light["event_date"] =
pd.to_datetime(df_acled_light["event_date"])

# Tipos de eventos de interés
event_types = [
    'Protests',
    'Riots',
    'Battles',
    'Strategic developments',
    'Explosions/Remote violence',
    'Violence against civilians'
]

def add_ally_conflict_columns(df_cyber, df_acled):
    for event_type in event_types:
        df_cyber[f'conflict_{event_type}_allies_of_country_vs_actor']
= 0
        df_cyber[f'conflict_{event_type}_country_vs_allies_of_actor']
= 0

    for idx, row in df_cyber.iterrows():
        attack_date = row["event_date"]
        year = attack_date.year
        country = row["country"]
        actor = row["actor_country"]

        if pd.isna(country) or pd.isna(actor) or actor ==
"UNDETERMINED":
            continue

    # Obtener aliados

```

```
allies_country = get_allies(country, year)
allies_actor = get_allies(actor, year)

# Filtrar eventos en ventana de 90 días antes del ataque
time_filter = (df_acled["event_date"] >= attack_date -
timedelta(days=90)) & (df_acled["event_date"] <= attack_date)
df_window = df_acled[time_filter]

for event_type in event_types:
    df_event = df_window[df_window["event_type"] ==
event_type]

    # Conflictos: aliados del país vs actor
    condition1 = (
        (df_event["actor1_country"].isin(allies_country) &
(df_event["actor2_country"] == actor)) |
        (df_event["actor2_country"].isin(allies_country) &
(df_event["actor1_country"] == actor))
    )

    # Conflictos: país vs aliados del actor
    condition2 = (
        ((df_event["actor1_country"] == country) &
df_event["actor2_country"].isin(allies_actor)) |
        ((df_event["actor2_country"] == country) &
df_event["actor1_country"].isin(allies_actor))
    )

    df_cyber.at[idx,
f'conflict_{event_type}_allies_of_country_vs_actor'] = condition1.sum()
    df_cyber.at[idx,
f'conflict_{event_type}_country_vs_allies_of_actor'] = condition2.sum()

    return df_cyber

df_cyber_events =
add_ally_conflict_columns(df_cyber_events, df_acled_light)

# Lista de TTPs a buscar
risky_ttps = [
    "External Remote Services",
    "Exploitation of Remote Services",
    "Windows Management Instrumentation",
    "OS Credential Dumping",
    "Masquerading",
    "Lateral Tool Transfer"
]

# Crear una expresión regular que busque cualquiera de esos términos
pattern = '|'.join([re.escape(ttp) for ttp in risky_ttps])

# Asegurarse de que 'TTPs' sea string (evitar errores si hay NaNs o
números)
df_cyber_events["TTPs"] = df_cyber_events["TTPs"].astype(str)
```

```
# Crear columna booleana basada en coincidencias
df_cyber_events["contains_risky_ttps"] =
df_cyber_events["TTPs"].str.contains(pattern, case=False,
na=False).astype(int)

df_cyber_aux = df_cyber_events.copy()
df_cyber_aux =
df_cyber_aux.rename(columns={'violence_between_between':
'violence_between'})

df_cyber_aux =
df_cyber_aux.rename(columns={'strategic_developments_between':
'strategic_between'})
df_cyber_aux =
df_cyber_aux.rename(columns={'explosions_remote_violence_between':
'explosions_between'})

df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Protests_allies_of_country_vs_acto
r': 'protests_allies_of_country_vs_actor'})
df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Protests_country_vs_allies_of_acto
r': 'protests_country_vs_allies_of_actor'})

df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Riots_country_vs_allies_of_actor':
'riots_country_vs_allies_of_actor'})
df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Riots_allies_of_country_vs_actor':
'riots_allies_of_country_vs_actor'})

df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Battles_allies_of_country_vs_actor
': 'battles_allies_of_country_vs_actor'})
df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Battles_country_vs_allies_of_actor
': 'battles_country_vs_allies_of_actor'})

df_cyber_aux = df_cyber_aux.rename(columns={'conflict_Strategic
developments_allies_of_country_vs_actor':
'strategic_allies_of_country_vs_actor'})
df_cyber_aux = df_cyber_aux.rename(columns={'conflict_Strategic
developments_country_vs_allies_of_actor':
'strategic_country_vs_allies_of_actor'})

df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Explosions/Remote
violence_allies_of_country_vs_actor':
'explosions_allies_of_country_vs_actor'})
df_cyber_aux =
df_cyber_aux.rename(columns={'conflict_Explosions/Remote
violence_country_vs_allies_of_actor':
'explosions_country_vs_allies_of_actor'})
```

```
df_cyber_aux = df_cyber_aux.rename(columns={'conflict_Violence
against_civilians_allies_of_country_vs_actor':
'violence_allies_of_country_vs_actor'})
df_cyber_aux = df_cyber_aux.rename(columns={'conflict_Violence
against_civilians_country_vs_allies_of_actor':
'violence_country_vs_allies_of_actor'})

df_cyber_events = df_cyber_aux.copy()
#df_cyber_events = df_cyber_events.iloc[:, :-12]
df_cyber_events = df_cyber_events.loc[:,
~df_cyber_events.columns.duplicated()]
#df_cyber_events.head()

# Cargar modelo y transformadores
#model = load_model("modelo_riesgo.keras", compile=False)
#onehot_encoder = joblib.load("onehot_encoder.pkl")
#embedding_encoders = joblib.load("embedding_encoders.pkl")
#scaler = joblib.load("numerical_scaler.pkl")

# Cargar y preparar los datos
df_1 = df_cyber_events.copy()

# Columnas utilizadas
one_hot_columns = [
    'actor_type', 'event_type', 'Attacks ICS', 'contains_risky_ttps'
]

embedding_columns = [
    'actor', 'organization', 'industry', 'event_subtype',
    'country', 'actor_country', 'TTPs'
]

numerical_columns = [
    'protests_country', 'violence_country', 'strategic_country',
    'riots_country',
    'battles_country', 'explosions_country',
    'protests_actor_country',
    'violence_actor_country', 'strategic_actor_country',
    'riots_actor_country',
    'battles_actor_country', 'explosions_actor_country',
    'protests_between',
    'violence_between', 'strategic_between',
    'riots_between', 'battles_between', 'explosions_between',
    'protests_allies_of_country_vs_actor',
    'protests_country_vs_allies_of_actor',
    'violence_allies_of_country_vs_actor',
    'violence_country_vs_allies_of_actor',
    'strategic_allies_of_country_vs_actor',
    'strategic_country_vs_allies_of_actor',
    'riots_allies_of_country_vs_actor',
    'riots_country_vs_allies_of_actor',
    'battles_allies_of_country_vs_actor',
    'battles_country_vs_allies_of_actor',
    'explosions_allies_of_country_vs_actor',
    'explosions_country_vs_allies_of_actor'
]
```

```
]

# Rellenar NaNs
df_1[one_hot_columns + embedding_columns] = df_1[one_hot_columns +
embedding_columns].fillna("Undetermined")
#df_1[numerical_columns] = df_1[numerical_columns].fillna(0)

# Convertir a tipo string explícitamente (para evitar problemas con
tipos)
df_1[one_hot_columns] = df_1[one_hot_columns].astype(str)

# Cargar modelo y transformadores
model = load_model("modelo_riesgo.keras", compile=False)
onehot_encoder = joblib.load("onehot_encoder.pkl")
embedding_encoders = joblib.load("embedding_encoders.pkl")
scaler = joblib.load("numerical_scaler.pkl")

# One-hot transform
X_onehot =
onehot_encoder.transform(df_1[one_hot_columns]).astype(np.float32)

# Embedding transform
X_embed = []
for col in embedding_columns:
    le = embedding_encoders[col]

    # Asegurar string antes de comprobar pertenencia
    df_1[col] = df_1[col].astype(str)

    # Sustituir desconocidos por "Undetermined"
    df_1[col] = df_1[col].apply(lambda x: x if x in le.classes_ else
'Undetermined')

    # Añadir "Undetermined" si no está en las clases
    if 'Undetermined' not in le.classes_:
        le.classes_ = np.append(le.classes_, 'Undetermined')

    # Transformar y agregar a X_embed
    X_embed.append(le.transform(df_1[col]).reshape(-1,
1).astype(np.float32))

# Finalmente transformar
#df_1[col] = le.transform(df_1[col])

# Escalar numéricos
# Verifica las columnas que espera el scaler
#print("Esperado por el scaler:", list(scaler.feature_names_in_))
#print("Columnas actuales en df_1:",
list(df_1[numerical_columns].columns))
X_numerical =
scaler.transform(df_1[numerical_columns]).astype(np.float32)

# Preparar entradas para todas las filas
X_input = X_embed + [X_onehot, X_numerical]
```

```
# Verificar los tipos de entrada antes de la predicción
#print(f"Tipos de entrada: {[x.dtype for x in X_input]}") # Verifica
que todas las entradas sean de tipo numérico

# Predicción para todas las filas
y_pred = model.predict(X_input).flatten()
del model # libera memoria

# Mostrar las predicciones de riesgo para cada fila
#df_1['predicted_risk'] = y_pred
#print(df_1[['predicted_risk']])

# GREYNOISE
df_1["cve_id"] = df_expanded["cve_id"].values
gn = GreyNoise(api_key=api_key_gn)
# Añade una columna 'ip_info' que contiene listas de dicts con info
por IP
df_1["ip_info"] = df_1["cve_id"].apply(extraer_info_ips_por_cve)

# DEVUELVE EL DATASET COMPLETO
return df_1
```