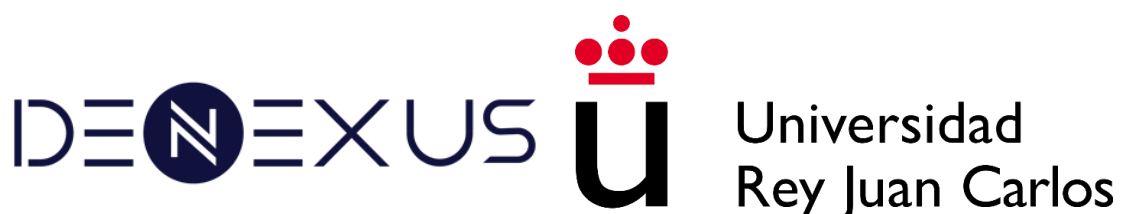


DICYME:

Dynamic Industrial CYberrisk Modelling based on Evidence

Iniciativa Conjunta de:



ENTREGABLE 2.1:

Primer prototipo del módulo de medida/estimación de probabilidad e impacto

Coordinadores:

Romy R. Ravines (DeNexus Tech)

Isaac Martín de Diego (Universidad Rey Juan Carlos)

Alberto Fernández Isabel (Universidad Rey Juan Carlos)



Contenido

1	Introducción	3
2	Modelos de Medida/Estimación de Probabilidades e Impactos.	3
3	ATR2ATK: Atractivo de una organización a ciberataques.....	5
3.1	Antecedentes	5
3.2	Metodología	5
3.3	Indicadores incluidos en la base de datos DICYME.....	6
3.4	Resultados.....	7
4	CVE2ATK: Vulnerabilidades y las Técnicas de las matrices MITRE	9
4.1	Antecedentes	9
4.2	Metodología	10
4.3	Resultados.....	12
4.4	Indicadores incluidos en la base de datos DICYME.....	12
5	THRACT: Actores de Amenazas y Víctimas.....	12
6	Base de Datos y Código.....	14
7	Conclusiones y Sigüientes Pasos	15

1 Introducción

Este documento resume los aspectos más importantes de diseño e implementación del entregable 2.1 del Proyecto, es decir, del primer prototipo de módulo de medida/estimación de probabilidad e impacto. Este entregable es fruto del trabajo realizado en las tareas 2.1, 2.2, 2.3 y 2.4 del Proyecto, dentro de la Actividad 2 “Módulos de medida/estimación de probabilidad e impacto”. Como puede observarse en la **Ilustración 1. Actividades en el Plan de Trabajo de DICYME.**, esta actividad se desarrolla entre los meses 6 y 30 del Proyecto.

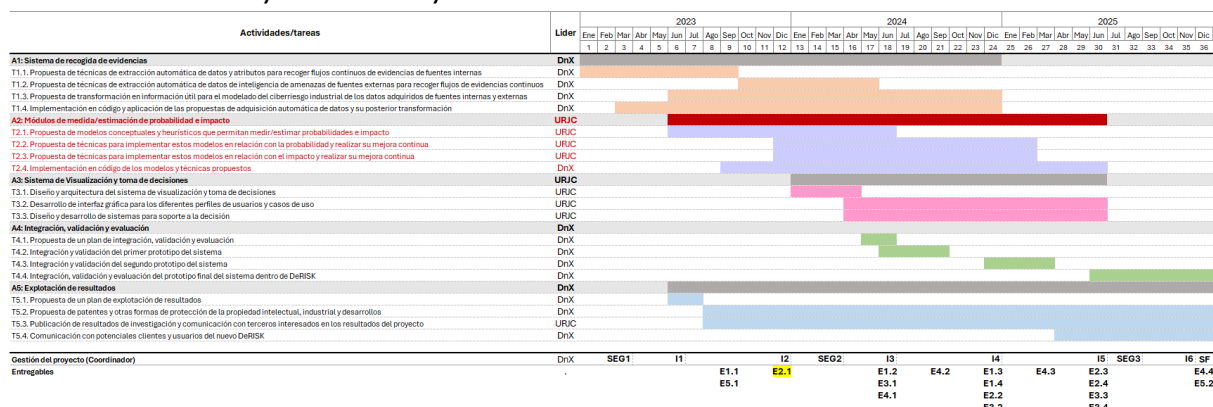


Ilustración 1. Actividades en el Plan de Trabajo de DICYME.

En el resto de secciones de este documento se abordan los siguientes temas:

- Los modelos desarrollados, su relación con el ciber riesgo y su explotación o integración en la plataforma de CRQM de DeNexus
- El tipo de datos que se utilizarán en esta primera versión y cómo se transformarán para obtener información útil que se empleará como entrada para el cálculo de probabilidades o impactos.
- Los flujos de evidencias de ciber riesgo provenientes de diferentes fuentes externas relacionadas con las organizaciones y el contexto de ciber amenazas que se han tenido en cuenta para este primer prototipo.
- Los datos y código que se encuentran disponibles en este enlace ([deriskGroup / DICyME Project · GitLab](#)) y cómo pueden emplearse.

2 Modelos de Medida/Estimación de Probabilidades e Impactos.

La cuantificación y gestión del riesgo asociado a ciberseguridad (CRQM) es el proceso de evaluación y mitigación del potencial impacto financiero que las amenazas cibernéticas pueden causar en una organización.

Existen varios métodos de cuantificación del ciber riesgo. Uno de los más conocidos es FAIR, siglas en ingles de *Análisis Factorial del Riesgo de la Información*. FAIR es un marco

de trabajo y metodología que ayuda a las organizaciones a identificar, analizar y comunicar los factores que afectan el riesgo de la información. FAIR divide el riesgo en dos componentes: frecuencia de eventos de pérdida (LEF) y magnitud de pérdida (LM), que es una fórmula bastante común para el sector de seguros. En el sector de ciberseguridad, esa misma descomposición se expresa Probabilidad e Impacto. Si tomamos como ejemplo la [Ilustración 2. Fórmula del ciberriesgo](#). **Fuente: Balbix.com**, se aprecia como relacionar factores importantes como vulnerabilidades, amenazas, exposición, etc. con ambos componentes.

$$\text{risk} = \text{likelihood} \times \text{impact}$$

$\text{likelihood} = f(\text{vulnerabilities, exposure, threats, mitigating controls})$
 $\text{impact} = g(\text{business criticality})$

Ilustración 2. Fórmula del ciberriesgo.
Fuente: Balbix.com

En DeRISK (<https://www.denexus.io/products/derisk/cyber-risk-quantification>), producto de DeNexus para la cuantificación y gestión del ciber riesgo en instalaciones industriales, se adopta el enfoque mostrado en la [Ilustración 3. Cuantificación del ciberriesgo en DeRISK](#). En DeRISK, las pérdidas financieras se estiman como el producto (o convolución de distribuciones) de Frecuencia y Severidad, donde la Frecuencia se divide en dos componentes: Número de Intentos de Ataque y Probabilidad de éxito de cada ataque.

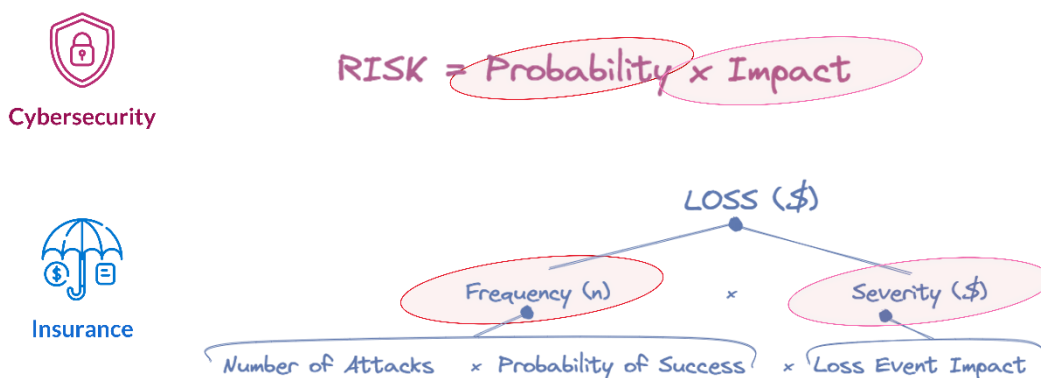


Ilustración 3. Cuantificación del ciberriesgo en DeRISK.

En este documento describimos tres iniciativas de modelización llevadas a cabo en DICYME, que proporcionan datos y estimaciones relevantes sobre los factores que afectan al ciber riesgo. Estas iniciativas son sistemas de modelos individuales que pueden usar de manera independiente (*stand-alone systems*) para análisis descriptivo de dichos factores y/o pueden usarse en combinación con un sistema más completo como DeRISK. Los sistemas son:

1. **ATR2ATK**: Estimación del atractivo de una organización a ciber ataques basada en información pública, accesible por cualquier individuo o actor de amenazas. Índice

que muestra cómo evoluciona el atractivo de una empresa a ser víctima de un ciber ataque.

2. **CVE2ATK**: Vulnerabilidades y las Técnicas de las matrices MITRE ATT&CK. Identificación de las técnicas que los actores de amenazas podrían explotar con mayor facilidad debido a la existencia de vulnerabilidades.
3. **THRACT**: Los actores de amenazas por organización. Índice que muestra cómo evoluciona un actor de amenazas respecto a un grupo objetivo y/o índice de cómo cambia o evoluciona el conjunto de actores de amenazas respecto a un potencial objetivo.

3 ATR2ATK: Atractivo de una organización a ciberataques

Este término se define como la posesión de atributos que aumentan el interés de ser víctima de un adversario. A continuación, se detallan los motivos y objetivos establecidos en la creación de este indicador, así como resultados obtenidos.

3.1 Antecedentes

Las medidas actuales para estimar la probabilidad de que ocurran incidentes de ciberseguridad se basan en un análisis de riesgo detallado que incluye evidencias externas sobre la compañía en cuestión. Estas evidencias pueden comprender datos sobre incidentes anteriores, información de inteligencia de amenazas, vulnerabilidades conocidas y configuraciones de seguridad. Metodologías como el análisis de riesgo cuantitativo, que utiliza modelos estadísticos para predecir la frecuencia de los incidentes, y el análisis cualitativo, que evalúa la severidad potencial de un incidente sin asignar probabilidades numéricas, son comunes. La integración de datos firmográficos y de amenazas enriquece estos modelos permitiendo ajustar las probabilidades basadas en características específicas de la empresa y su entorno operativo.

En este primer prototipo se ha trabajado únicamente con datos firmográficos, es decir, aquellos que definen una compañía por cómo es: localización de su sede principal, sector de operaciones, si cotizan en bolsa, si son lucrativas, tamaño (cantidad de empleados), y facturación y ganancias anuales.

3.2 Metodología

Para estimar la probabilidad de que se produzcan ciber incidentes en distintos entornos, se propone el uso de modelos probabilísticos basados en técnicas de aprendizaje automático y análisis estadístico. Para ello, el modelo propuesto por DICYME integra un enfoque de minería de reglas mediante el algoritmo *FP-Growth*, seguido por la clasificación con un árbol de decisión. Este enfoque aprovecha las fortalezas complementarias de ambos métodos para una identificación más precisa de riesgos potenciales de la forma que se detalla a continuación:

1. **Minería de reglas con *FP-Growth***: Este modelo se centra inicialmente en la extracción de patrones frecuentes de datos sobre incidentes de ciberseguridad. *FP-Growth* facilita la identificación de conjuntos de elementos frecuentes sin generar

candidatos explícitamente, lo que lo hace más rápido y escalable que otros métodos alternativos como, por ejemplo, A priori. La salida de este proceso son reglas que describen las condiciones bajo las cuales los incidentes de ciberseguridad han ocurrido con más frecuencia. Se siguen tres etapas:

- 1.1. **Identificación de conjuntos frecuentes:** los atributos que coexisten con alta frecuencia, como la relación entre ganancias y tamaño de la empresa. Por ejemplo, un patrón frecuente es unos beneficios anuales desconocidos con una cantidad baja de empleados.
 - 1.2. **Generación de reglas asociativas:** se crean reglas del tipo "Si X, entonces Y", donde X (denominado antecedente) representa un conjunto de condiciones e Y (denominado consecuente) el resultado asociado.
 - 1.3. **Evaluación de reglas:** mediante las métricas de soporte, confianza y *lift*. Al ejecutar el algoritmo *FP-Growth*, se establecen unos umbrales mínimos de soporte y confianza igual a 0,01, para descartar aquellas reglas que no tienen un nivel suficiente de interés o relevancia.
2. **Árbol de Decisión:** Posteriormente, se utiliza un árbol de decisión para clasificar nuevos casos en base a las reglas derivadas del *FP-Growth*. El árbol de decisión permite visualizar y entender cómo se toman las decisiones según las reglas, ya que representa una serie de decisiones binarias que gradualmente clasifican los datos. Este método es útil por su capacidad para manejar datos no lineales y por proporcionar modelos altamente comprensibles y explicables.

3.3 Indicadores incluidos en la base de datos DICYME

Empleando el trabajo que se está desarrollando en el marco de la actividad 1 del proyecto, y especialmente el trabajo descrito en el entregables E1.1, la base de datos DICYME contiene los indicadores descritos en la [Tabla 1. Indicadores de la base de datos DICYME usados en el modelo de atractivo](#), que sirven como entradas del modelo de atractivo.

Nombre del atributo	Expresión para su cálculo
Categoría del incidente	4 posibles categorías de incidente y 1 de no incidente.
País	Nombre del país al que pertenece la compañía.
Categoría de la compañía	Nombre del sector al que pertenece la compañía.
Beneficios anuales	Cantidad de dinero facturado como beneficios anualmente por la compañía.
Ingresos anuales	Cantidad de dinero facturado anualmente por la compañía.
Cotización en bolsa	Valor de True o False en función de si la compañía cotiza en bolsa o no.
Número de empleados	Cantidad de empleados contratados en el momento del incidente.

Lucrativo	Valor de True o False en función de si la compañía busca obtener beneficios o no.
-----------	---

Tabla 1. Indicadores de la base de datos DICYME usados en el modelo de atractivo

3.4 Resultados

Este modelo sigue en constante mejora. Los principales resultados se han resumido y compartido en un artículo enviado a una revista de ciberseguridad. Destacan las reglas descubiertas por la minería de reglas (véase la [Ilustración 4. Primeras 10 reglas obtenidas del algoritmo FP-Growth.](#)) y el alto poder discriminatorio entre las empresas que han sido víctima de ciberataques de las que no lo han sido (véanse la [Ilustración 5. Métricas sobre el poder de clasificación del modelo.](#) y la [Ilustración 6. Distribución de incidentes según el atractivo.](#)).

Rule	Sup.	Conf.	Lift
Earnings=NaN → Employees=0.0	0.845	0.927	0.996
Employees=0.0 → Earnings=NaN	0.845	0.909	0.996
Earnings=NaN → Publicly traded=False	0.871	0.955	1.096
Publicly traded=False → Earnings=NaN	0.871	1.000	1.096
Publicly traded=False → Employees=0.0	0.807	0.926	0.995
Employees=0.0 → Publicly traded=False	0.807	0.867	0.995
Earnings=NaN, Publicly traded=False → Employees=0.0	0.807	0.926	0.995
Earnings=NaN, Employees=0.0 → Publicly traded=False	0.807	0.954	1.095
Publicly traded=False, Employees=0.0 → Earnings=NaN	0.807	1.000	1.096
Earnings=NaN → Publicly traded=False, Employees=0.0	0.807	0.884	1.096

Ilustración 4. Primeras 10 reglas obtenidas del algoritmo FP-Growth.

Metric	Training Data	Testing Data
Accuracy	0.931	0.923
Kappa Statistic	0.847	0.835
Recall	0.875	0.929
Precision	0.925	0.867

Ilustración 5. Métricas sobre el poder de clasificación del modelo.

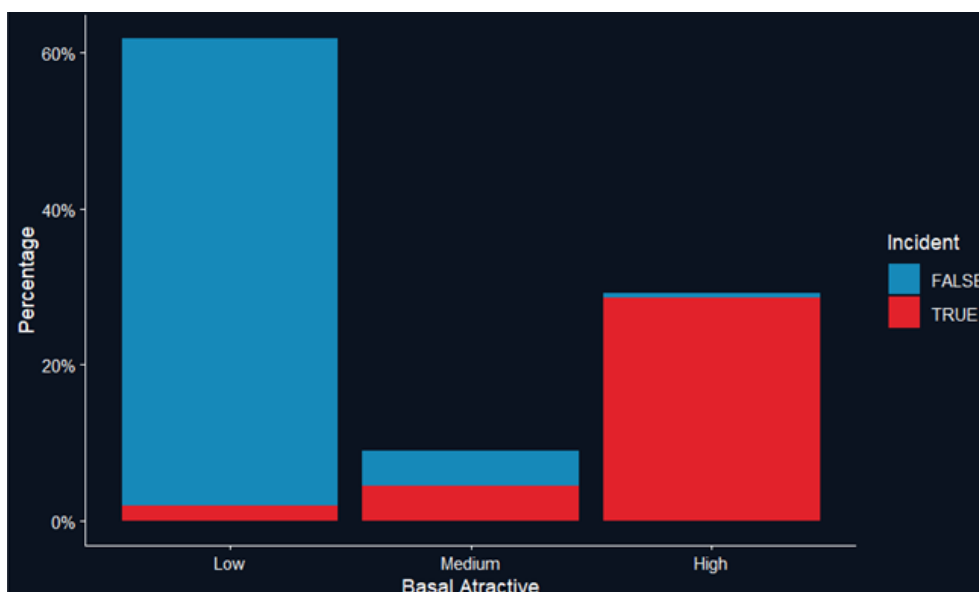


Ilustración 6. Distribución de incidentes según el atractivo.

Al hacer el estudio de densidad de las reglas obtenidas de los datos extraídos se obtiene lo representado en la [Ilustración 7. Estudio de densidad de las reglas](#). (DF1 contiene las empresas que no han sufrido un incidente, mientras que DF2 aquellas que sí han sufrido un incidente, denominadas no incidentes e incidentes, respectivamente, por simplicidad y claridad):

- **Densidad de reglas (rules_count):** se observa una diferencia clara en la distribución del número de reglas generadas entre los conjuntos de datos de incidentes y no incidentes. Los no incidentes presentan una mayor densidad en valores más bajos de reglas, mientras que los incidentes tienen una densidad más uniforme y extendida hacia valores altos, indicando que el conjunto de incidentes genera reglas más numerosas o diversificadas.
- **Correlación entre número de reglas y soporte total (rules_count vs. total_support):** existe una correlación positiva fuerte entre la cantidad de reglas generadas y el soporte total. Esto indica que un mayor número de reglas tiende a estar asociado con combinaciones de atributos que ocurren con más frecuencia en los datos.
- **Densidad de soporte total (total_support):** la distribución de soporte total muestra diferencias significativas entre los conjuntos. Los no incidentes tienen valores de soporte más bajos en general, mientras que los incidentes presentan una distribución más amplia hacia valores medios y altos. Esto podría implicar que los datos de no incidentes contienen patrones más comunes o recurrentes.
- **Relación entre soporte y confianza totales (total_support vs. total_confidence):** se observa una relación positiva entre el soporte y la confianza, lo que es esperado en algoritmos de minería de reglas. Sin embargo, algunos puntos de incidentes destacan con alta confianza relativa para valores bajos de soporte, lo cual podría ser interesante para identificar reglas específicas con menor frecuencia, pero alta fiabilidad.
- **Densidad de confianza total (total_confidence):** la confianza total en incidentes se distribuye más homogéneamente en rangos altos, mientras que los no incidentes tienen mayor densidad en valores bajos. Esto sugiere que las reglas generadas en incidentes tienden a ser más confiables en general.
- **Relación entre soporte y lift (total_support vs. total_lift):** aunque hay una correlación general, se observa dispersión en los valores de *lift*, especialmente en incidentes. Esto indica que algunas reglas tienen mayor relevancia relativa (*lift* alto), incluso si su soporte es bajo, sugiriendo patrones más especializados en los datos.
- **Densidad de lift total (total_lift):** la distribución de *lift* en no incidentes tiene un pico más pronunciado en valores bajos, mientras que en incidentes se muestra

una distribución más uniforme hacia valores más altos, lo que puede implicar que los incidentes contienen relaciones más fuertes y específicas entre los atributos.

- **Correlación general entre métricas (pairplots):** las relaciones entre rules_count, total_support, total_confidence y total_lift muestran tendencias generales similares entre los conjuntos, pero con diferencias notables en las densidades y rangos alcanzados. Esto refuerza que los datos en incidentes generan reglas más variadas y posiblemente más útiles para análisis posteriores.

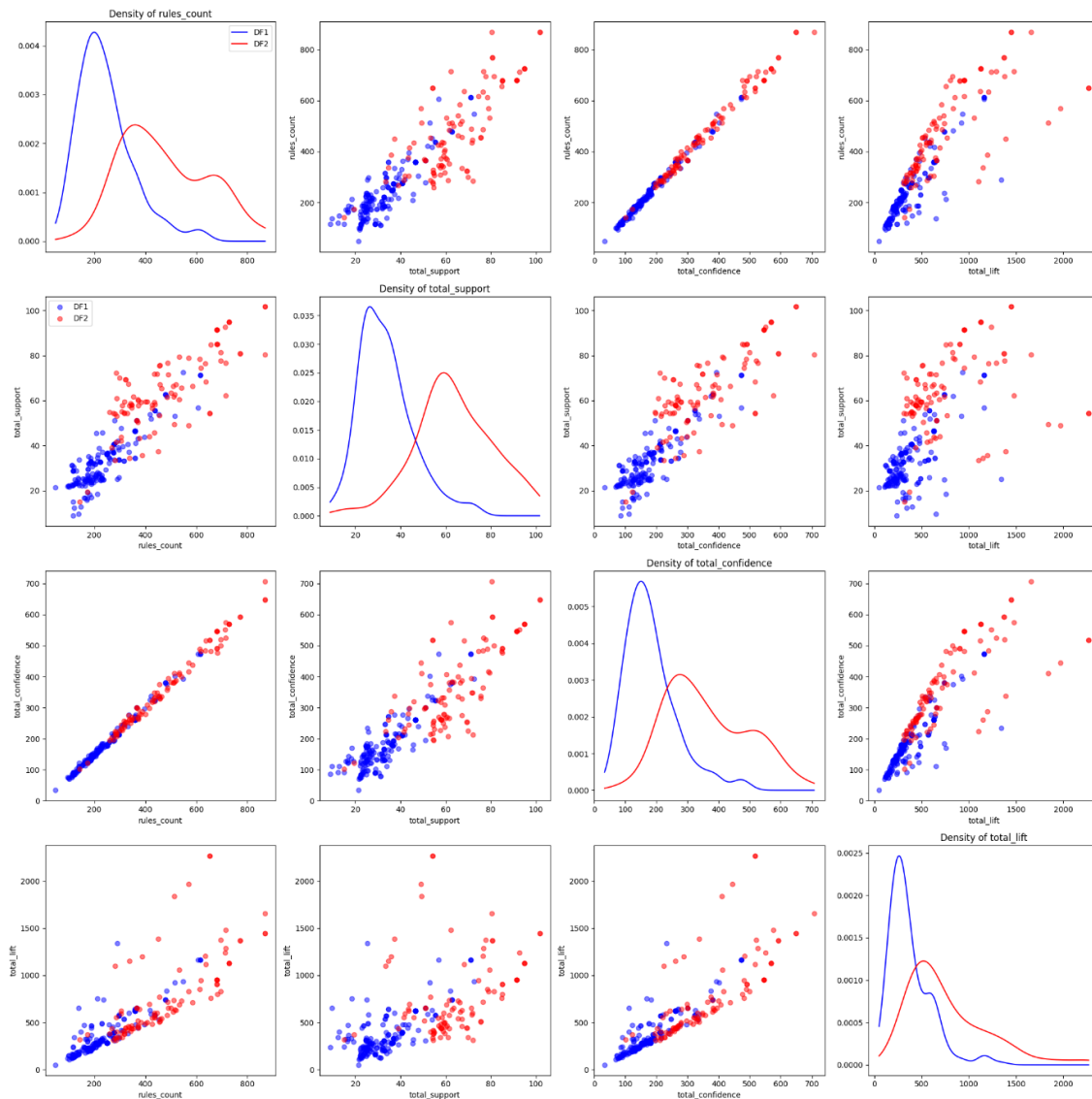


Ilustración 7. Estudio de densidad de las reglas.

4 CVE2ATK: Vulnerabilidades y las Técnicas de las matrices MITRE

4.1 Antecedentes

Un CVE (por sus siglas en inglés de *Common Vulnerabilities and Exposures*) es un identificador único asignado a una vulnerabilidad de seguridad conocida. Cada CVE proporciona una base de datos estandarizada para que las organizaciones mejoren su seguridad. Un TTP (por sus siglas en inglés de *Tactics, Techniques, and Procedures*) describe cómo los actores de amenazas operan durante un ataque. La relación entre CVEs y TTPs es crucial, ya que permite a las organizaciones entender cómo se puede explotar una vulnerabilidad específica (CVE) mediante técnicas específicas (TTPs), y así anticipar y mitigar el impacto potencial de los ataques.

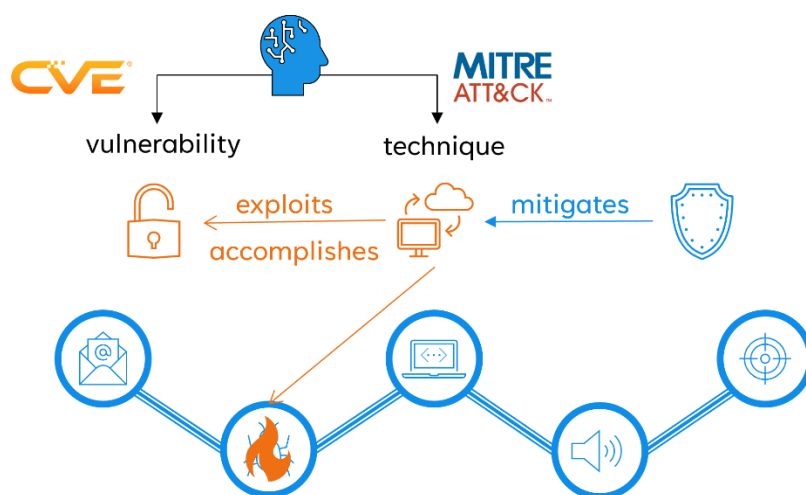


Ilustración 8. Objetivo del sistema CVE2TTs

Por su parte, la conexión CVEs y TTPs se ha explorado mediante avanzadas técnicas de Procesamiento del Lenguaje Natural (NLP). Los CVEs generalmente incluyen un campo de texto libre que describe la vulnerabilidad y su contexto. Utilizando este texto y su relación con los TTPs, se establece un modelo de asociación empleando NLP (véase la [Ilustración 8. Objetivo del sistema CVE2TTs](#)). De manera específica, esta tarea se realiza principalmente mediante modelos semánticos de redes neuronales de tipo *transformer* como, por ejemplo, BERT y similares.

4.2 Metodología

Este desafío presenta una naturaleza multiclase y multietiqueta, dado que cada CVE puede vincularse con múltiples TTPs. Para abordarlo, se adopta una estrategia de fine-tuning sobre un modelo base BERT estándar ya pre-entrenado. Esto implica añadir una nueva capa al modelo existente, con dimensiones correspondientes al número de TTPs que buscamos mapear, y luego someterlo a un proceso de re-entrenamiento. De esta manera, se logra modelar la asociación entre CVEs y TTPs. La [Ilustración 9. CVE2TTs: Datos de Entrenamiento](#) muestra el proceso del modelo a alto nivel.

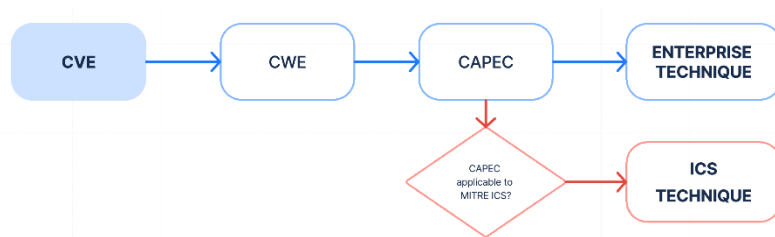


Ilustración 9. CVE2TTs: Datos de Entrenamiento

Tenemos dos formas de predecir las técnicas dada una vulnerabilidad CVE (véase la [Ilustración 10. Alternativas para predecir técnicas de un CVE](#)):

- Con tipo de vulnerabilidad: predecir las técnicas directamente a partir del modelo entrenado, donde las entradas son la descripción y el tipo de vulnerabilidad.
- Sin tipo de vulnerabilidad: primero predecir el tipo de vulnerabilidad a partir de la descripción y luego usar la descripción y el tipo de vulnerabilidad predicho para predecir las técnicas.

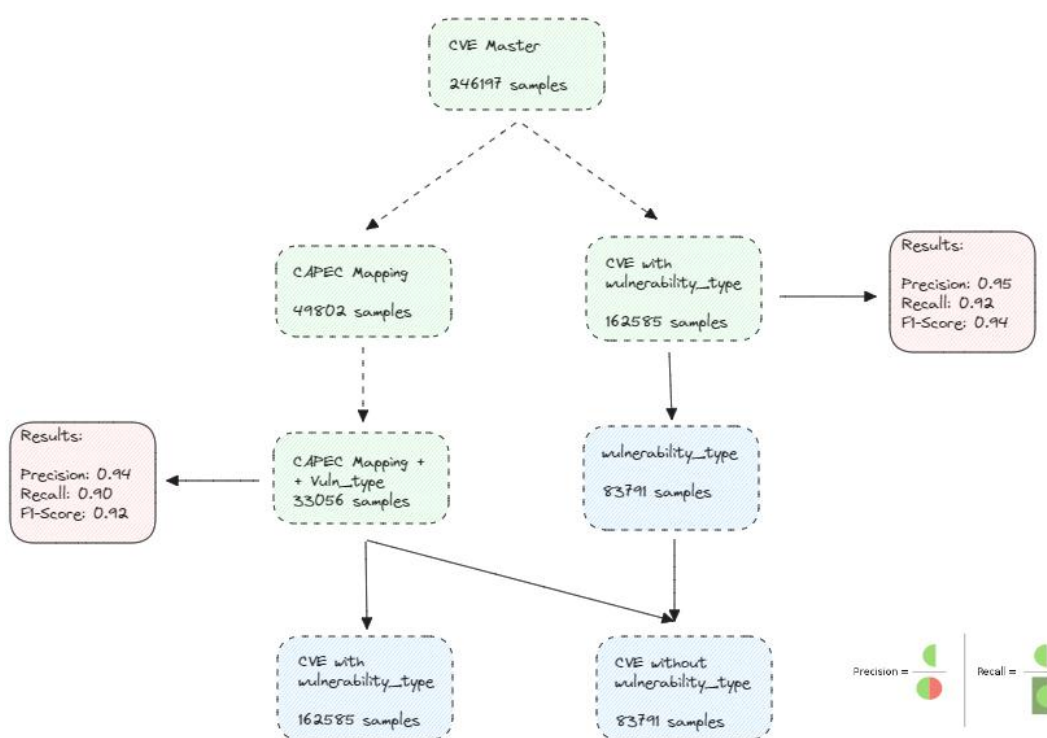


Ilustración 10. Alternativas para predecir técnicas de un CVE

El modelo está configurado de la manera mostrada en la [Ilustración 11. CVE2TTs: Modelo de Predicción](#).

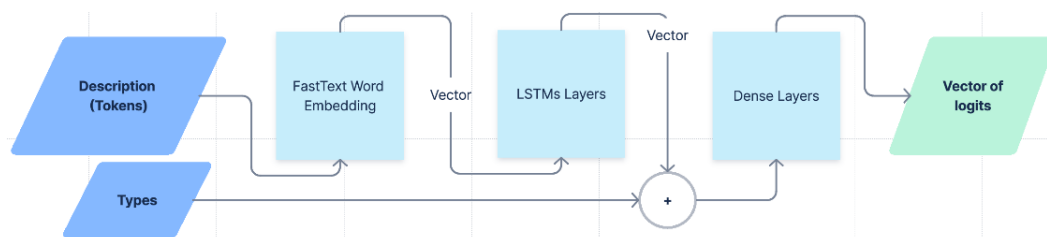


Ilustración 11. CVE2TTs: Modelo de Predicción

4.3 Resultados

La [Ilustración 12. Model loss para la matriz Enterprise](#) muestra los resultados obtenidos para el modelo de la matriz Enterprise de MITRE ATT&CK, mientras que la [Ilustración 13. Model loss para la matriz ICS](#) muestra los resultados para la matriz ICS.

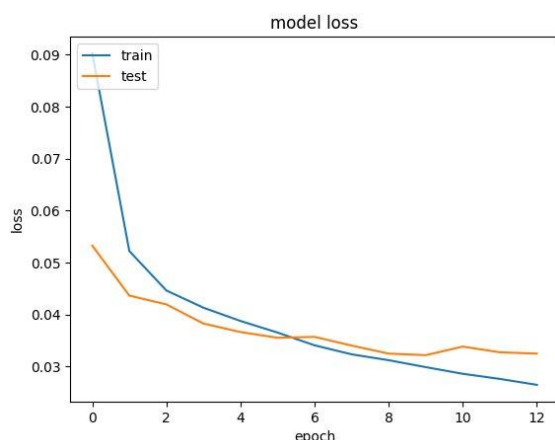


Ilustración 12. Model loss para la matriz Enterprise

- Precision: 0.94
- Recall: 0.90
- F1-Score: 0.92
- TPP: 0.84
- FPP: 0.086
- NPP: 0.045

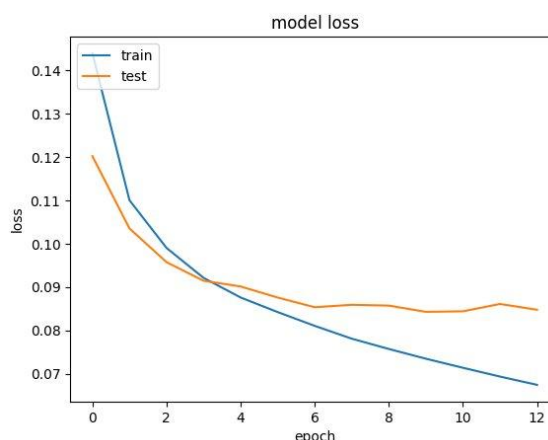


Ilustración 13. Model loss para la matriz ICS

- Precision: 0.91
- Recall: 0.87
- F1-Score: 0.89
- TPP: 0.90
- FPP: 0.098
- NPP: 0.00056

4.4 Indicadores incluidos en la base de datos DICYME

La [Tabla 2. Indicadores de la base de datos DICYME usados en el modelo CVE2TTs](#), contiene los indicadores incluidos en la base de datos DICYME que sirven de entrada para el modelo CVE2TTs.

Nombre del atributo	Expresión para su cálculo
Texto CVE	Texto libre.
Attack vector	Campos booleanos.
Tipo de vulnerabilidad	Campos booleanos.

Tabla 2. Indicadores de la base de datos DICYME usados en el modelo CVE2TTs

5 THRACT: Actores de Amenazas y Víctimas

La información sobre actores de amenazas es muy útil para analizar el contexto y panorama actual del cibercrimen y también para enriquecer la información sobre ciber incidentes hechos públicos (en los que se haya hecho atribución y se pueda relacionar la información por actor). La fuente de datos utilizada para obtener información sobre los actores es la enciclopedia de actores de amenazas desarrollada y mantenida por la *Electronic Transactions Development Agency*, disponible públicamente en [Threat Group Cards: A Threat Actor Encyclopedia](#). Esta fuente de datos proporciona información

como los diferentes nombres asociados a los actores (alias), los países en los que están establecidos, su motivación, la fecha de su primera aparición registrada, una descripción, los sectores y países a los que han atacado, las diferentes herramientas y software que emplean, así como diversas campañas conocidas atribuidas a los mismos.

El índice de Actores de Amenazas (Threat Actor Index) es una métrica que depende de datos de las potenciales víctimas, se puede calcular para todos los actores activos en el mundo y se puede monitorizar a lo largo del tiempo.

La métrica final propuesta para *Threat Actor Index* se compone de tres métricas o puntuaciones parciales:

- Puntuación de Actividad (*Activity Score*)
- Puntuación de Capacidad (*Capacity Score*)
- Puntuación de Objetivo (*Target Score*)

La fórmula de la métrica *Threat Actor Index* es:

- **Puntuación de Actividad:** se calcula sólo para los actores de amenazas activos y depende de la información disponible sobre la última vez que estuvo activo

$$\text{last_seen_score} = \begin{cases} 1.0, & \text{if last_activity_date} \leq 30 \text{ days} \\ 0.8, & \text{if } 30 < \text{last_activity_date} \leq 90 \text{ days} \\ 0.6, & \text{if } 90 < \text{last_activity_date} \leq 365 \text{ days} \\ 0.3, & \text{if last_activity_date} > 365 \text{ days} \end{cases}$$

Luego, $\text{Activity Score} = \text{last_seen_score}$

- **Puntuación de Capacidad:** es cualitativa. En particular, se dispone de una variable categórica (*capability_values*) que toma valores: {'Below average', 'Average', 'Above average'}. Para calcular la puntuación de capacidad se aplica una escala entre 0 y 1 tal que:

$$\text{sophistication_score} = \begin{cases} 0.3, & \text{if capability_values is equal to 'Below average'} \\ 0.5, & \text{if capability_values is equal to 'Average'} \\ 0.8, & \text{if capability_values is equal to 'Above average'} \end{cases}$$

Luego, $\text{Capacity Score} = \text{sophistication_score}$

- **Puntuación de Objetivo:** representa si el ámbito de acción del actor coincide con el país e industria de una organización. Para calcular la puntuación de objetivo (*Target Score*), se utilizan tres datos disponibles: Países donde ha actuado el actor (*target_countries_values*), Industrias donde ha actuado el actor (*target_industries_values*) y motivación del actor (*motivations_values*).
- Paso 1. País donde está localizada la infraestructura (*target_country*) y sus países fronterizos (*adjacent_countries*).

- Paso 2. Industria a la que pertenece la instalación (target_industry).
- Paso 3. Motivación del actor.
- Paso 4. Cálculo de la puntuación de objetivo:

$$\text{Target Score} = \frac{w_{\text{country_target_score}} * \text{country_target_score} + w_{\text{industry_target_score}} * \text{industry_target_score} + w_{\text{motivations_score}} * \text{motivations_score}}{w_{\text{country_target_score}} + w_{\text{industry_target_score}} + w_{\text{motivations_score}}}$$

La **Ilustración 14. Ejemplo de resultados de THRACT** muestra algunos ejemplos de casos de uso de la métrica para actores de amenazas. Los datos y código se encuentran en este enlace: [deriskGroup / DICyME Project · GitLab](#), dentro del directorio *E.1.2 First prototype of automatic external data extraction module*. La carpeta *Etda_cyber_threat_actors* contiene los algoritmos de extracción de datos sobre actores de amenazas, así como la información resultante.

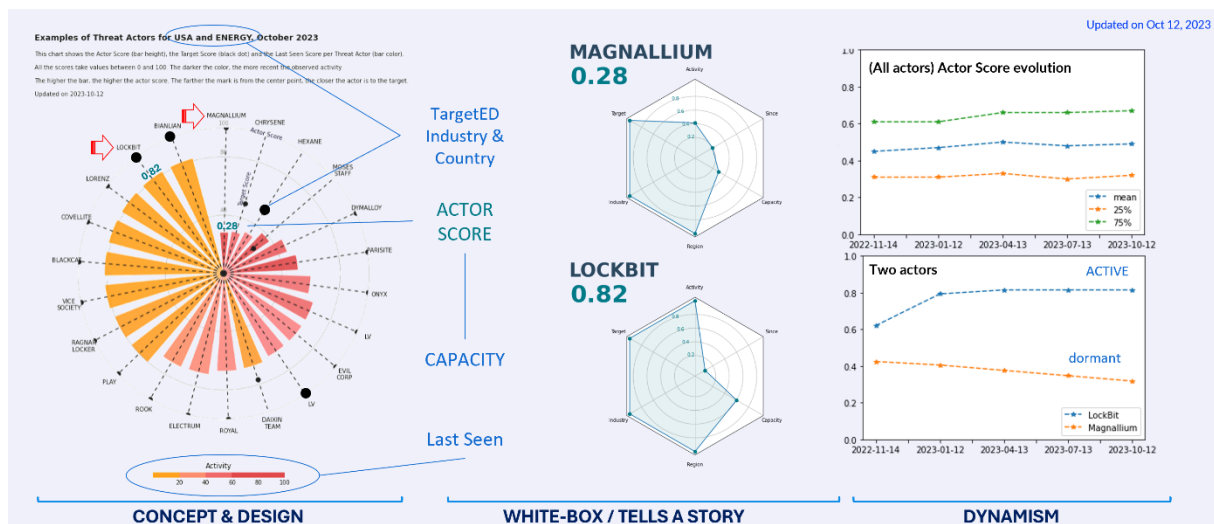


Ilustración 14. Ejemplo de resultados de THRACT

6 Base de Datos y Código

En el enlace [deriskGroup / DICyME Project · GitLab](#) se pueden encontrar los siguientes elementos, dentro del directorio *E.2.1 First prototype of measurement and estimation of probability and impact of a cyberincident*:

1. Basal_Atractiveness ([Repositorio en GitLab](#)):

Cuaderno de Python diseñado para realizar análisis sobre datos firmográficos de compañías con la finalidad de extraer un indicador sobre cómo de atractiva es una compañía dado un estado de entrada.

- **Code:** contiene los algoritmos de exploración de datos, entrenamiento de los modelos escogidos y validación de los resultados.
- **Output:** contiene los resultados en formato de matrices de confusión de los modelos entrenados y validados.

2. CVE_to_TTPs ([Repositorio en GitLab](#)):

- **CVE2TTP.ipynb:** Cuaderno de Python donde se crea y valida un modelo de aprendizaje automático para el mapeo de CVEs a TTPs.

- *CVE2TTP.keras*: Modelo exportado en formato keras. Este modelo se cargará en el entorno de producción para realizar el mapeo directamente.
- *CVE2TTP_metrics.csv*: Métricas obtenidas con el mejor modelo.

Cabe aclarar algunos aspectos sobre el entorno de desarrollo empleado para generar el prototipo. Se han utilizado:

- *Python 3.11.8*
- *Pandas 2.2.1*
- *Mlxtend 0.23.1*
- *Numpy 1.26.4*
- *Scikit-learn 1.4.2*
- *Matplotlib 3.8.3*
- *Scipy 1.12.0*
- *Keras-nlp 0.12.1*
- *Keras 3.3.2*
- *Torch 2.3.1*

Para poder ejecutar este código sin errores se recomienda usar el mismo entorno que el de desarrollo, arriba señalado.

7 Conclusiones y Sigüientes Pasos

En DICYME se están desarrollando modelos de medida o estimación de factores relacionados con el ciber riesgo. Por ahora el prototipo no integra los diferentes modelos, eso se hará cuando la *web app* esté funcionando y sea posible usar un modelo simplificado de cuantificación del ciber riesgo.

El actual estado del sistema es un conjunto de bases de datos y modelos que se explotan para proporcionar indicadores con información relacionada al ciber riesgo, cada uno de ellos tiene valor y significado en si mismo, pero se podrán traducir a lenguaje de negocio e información útil para los tomadores de decisiones cuando se vinculen a DeRISK y o a un motor de cuantificación del ciber riesgo.

Cabe destacar que ésta es una primera entrega y como primer prototipo se requiere más trabajo en esta rama. En concreto, los equipos van a valorar fuentes de datos diferentes, que aporten más información especialmente con los datos relacionado con el ciber incidente a nivel técnico, que permitan completar y matizar más aún los datos extraídos en este primer prototipo. Por otra parte, continuarán mejorando los modelos y metodologías utilizadas, mejorando las técnicas de análisis y aprendizaje automático para incrementar la precisión y explicabilidad de los modelos.