



Premios Isdefe I+D+i

Los diez artículos finalistas del DESEi+d 2020

Madrid, noviembre de 2020



MINISTERIO DE DEFENSA



Premios Isdefe I+D+i

Los diez artículos finalistas del DESEi+d 2020

Madrid, noviembre de 2020



MINISTERIO DE DEFENSA



Catálogo de Publicaciones de Defensa
<https://publicaciones.defensa.gob.es>



Catálogo de Publicaciones de la Administración General del Estado
<https://cpage.mpr.gob.es>

publicaciones.defensa.gob.es
cpage.mpr.gob.es

Edita:



Paseo de la Castellana 109, 28046 Madrid

© Autores y editor, 2022

NIPO 083-21-234-1 (edición impresa)

NIPO 083-21-235-7 (edición en línea)

Depósito legal M 34467-2021

Fecha de edición: febrero de 2022

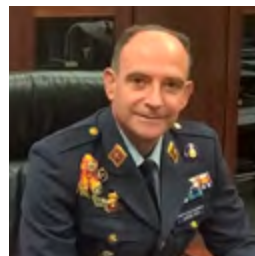
Maqueta e imprime: Imprenta Ministerio de Defensa

Las opiniones emitidas en esta publicación son exclusiva responsabilidad de los autores de la misma.

Los derechos de explotación de esta obra están amparados por la Ley de Propiedad Intelectual. Ninguna de las partes de la misma puede ser reproducida, almacenada ni transmitida en ninguna forma ni por medio alguno, electrónico, mecánico o de grabación, incluido fotocopias, o por cualquier otra forma, sin permiso previo, expreso y por escrito de los titulares del copyright ©.

En esta edición se ha utilizado papel 100% libre de cloro procedente de bosques gestionados de forma sostenible.

Prólogo DIGEREM



Desde 2013, DESEI se ha convertido en el foro y punto de encuentro de todos los actores relacionados con la I+D en el ámbito de la Defensa y la Seguridad, en el que los centros universitarios de la Defensa participan destacadamente, tanto en los comités Director, Organizador y Científico como en las ponencias y comunicaciones presentadas.

Este octavo congreso, organizado en esta ocasión por la Subdirección General de Enseñanza Militar, ha logrado alcanzar un año más, sus objetivos a través de un intercambio de experiencias, ideas y proyectos de investigación en las diferentes áreas de conocimiento tratadas, entre las que cabe destacar la exploración del espacio, desde la dimensión de la investigación, el desarrollo e innovación y la generación del talento.

La situación de emergencia sanitaria bajo la que se ha desarrollado el congreso y nuestro compromiso en gestionar de la mejor manera posible esta situación, han sido la clave para el planteamiento, por primera vez, de un congreso virtual, minimizando los actos presenciales, y donde la presentación de comunicaciones se ha realizado por videoconferencia.

Esta edición nos ha dado la oportunidad de mostrar la dimensión investigadora de los centros universitarios de la Defensa, enfocada a intercambiar conocimientos, identificar futuras líneas de estudio, concebir proyectos de investigación en áreas relacionadas con la seguridad y defensa, y vislumbrar posibles líneas de colaboración entre los diferentes actores del Ministerio de Defensa, universidades y empresas del sector. Coincide con la iniciativa de la subsecretaría de Defensa para impulsar la investigación a través de los centros universitarios de la Defensa que ha supuesto la publicación de la Orden DEF/375/2021, de 20 de abril, por la que se aprueban las directrices generales para la ordenación de la investigación y la transferencia del conocimiento en los centros universitarios de la defensa.

Por último, quiero agradecer el trabajo llevado a cabo por los miembros del Comité Científico en el proceso de revisión de las 144 comunicaciones recibidas, que ha dado como resultado la selección de las 94 comunicaciones presentadas durante el congreso. Las diez comunicaciones recogidas en este libro han sido las mejor valoradas por el jurado del «Premio ISDEFE I+D+i - Antonio Torres», y son un reconocimiento a la gran calidad de todas las comunicaciones recibidas.



Prólogo Consejero Delegado de ISDEFE

La octava edición del Congreso Nacional de I+D en Defensa y Seguridad (DESEi+d 2020) ha sido capaz de adaptarse a las restricciones impuestas por la pandemia de la covid-19, realizándose por primera vez de forma virtual, lo que ha supuesto para la organización un gran esfuerzo de transformación y colaboración.

La Subdirección General de Enseñanza Militar de la Dirección General de Reclutamiento y Enseñanza Militar, la Dirección General de Armamento y Material, los centros universitarios de Defensa e Isdefe han conseguido sumar fuerzas para transformar la crisis en oportunidad, permitiendo que la comunidad científica nacional de seguridad y defensa se haya reunido a lo largo de tres días para compartir su conocimiento y experiencia. Como en pasadas ediciones, investigadores y expertos de Fuerzas Armadas, cuerpos de seguridad, universidades y empresas del sector han presentado y difundido los resultados de las últimas investigaciones y trabajos realizados.

Isdefe apoya al Ministerio de Defensa en la organización y celebración del DESEi+d desde sus orígenes, lanzando en 2020 la quinta convocatoria del Premio Isdefe I+D+i «Antonio Torres» para seguir contribuyendo al objetivo de promover la cultura de participación y compartición del conocimiento en el ámbito de la defensa y seguridad.

En esta edición del congreso, el Comité Científico ha evaluado casi 150 comunicaciones de gran calidad y profundidad, pese a la adversidad del momento y la reducción de los plazos respecto a anteriores ediciones. En este libro figuran los diez trabajos finalistas al Premio Isdefe I+D+i «Antonio Torres», incluyendo el del ganador, tras la valoración del jurado, a cuyos miembros les agradezco su labor de evaluación.

No puedo dejar de resaltar el impulso que las direcciones generales de Armamento y Material y de Reclutamiento y Enseñanza Militar del Ministerio de Defensa proporcionan a esta iniciativa, además de la hospitalidad recibida por parte de la Academia Básica del Aire de León, sede seleccionada para celebrar el congreso y que finalmente no pudo ceder su espacio físico para esta gran iniciativa.

Por último, me gustaría agradecer a todos los ponentes, autores y asistentes del DESEi+d su participación y su gran flexibilidad para adecuarse al nuevo fomento del evento. A todos, os animo a participar en el próximo DESEi+d, sea cual sea su formato, esperando que esta nueva edición sea igual o más fructífera que las pasadas.

Francisco Quereda Rubio
Consejero Delegado de Isdefe

Fallo Jurado 5ª edición Premio Isdefe I+D+i "Antonio Torres"

En el marco del VIII Congreso Nacional de I+D en Defensa y Seguridad (DESEI+d 2020) a celebrarse de forma híbrida, en la Academia Básica del Aire (ABA) de León y de forma virtual, durante los días 24, 25 y 26 de noviembre de 2020, bajo la coordinación de la Dirección General de Reclutamiento y Enseñanza Militar en colaboración con la Dirección General de Armamento y Material, los Centros Universitarios de Defensa e Isdefe (Ingeniería de Sistemas para la Defensa de España), se convocó la Quinta Edición del Premio Isdefe I+D+i "Antonio Torres".

Las 141 comunicaciones recibidas para el congreso han sido analizadas por los miembros del Comité Científico Permanente, valorando los siguientes aspectos:

- Los **resultados** obtenidos por los trabajos presentados.
- El grado de **innovación** de las comunicaciones presentadas.
- La **contribución** de las mismas al I+D Nacional.
- La **proyección** de los trabajos en el ámbito de Defensa y Seguridad.

Según lo establecido en las bases del premio, para la evaluación final se conformó un Jurado Sancionador¹. Este jurado ha revisado y valorado las comunicaciones nominadas², de las que se propone como ganador del premio a:

D. David Novoa Paradela, del Centro de Investigación en Tecnologías de la Información y las Comunicaciones (CITIC) de la Universidad de A Coruña, por su artículo:

Mantenimiento predictivo de motores de buques mediante aprendizaje automático

Para que conste, se firma y traslada esta acta al Comité Organizador Permanente del Congreso.



Belinda Misiego Tejeda
Secretaria Jurado 5ª edición Premio Isdefe I+D+i "Antonio Torres"

¹ **Miembros del Jurado:** Secretaria: Dña. Belinda Misiego Tejeda – Isdefe; Vocal 1: Dña. Silvia Vicente Oliva – CUD de Zaragoza, CUD-AGM; Vocal 2: D. Miguel Ángel Urbiztondo Castro – CUD de Zaragoza, CUD-AGM; Vocal 3: TCOL (EA) Juan Manuel González Del Campo Martínez – Unidad de Prospectiva y Estrategia Tecnológica, SDGPLATIN; Vocal 4: D. Natalio García Hondurilla – CUD de Madrid, CUD-ACD; Vocal 5: D. Miguel Ángel Álvarez Feijoo – CUD de Marín, CUD-ENM; Vocal 6: D. Germán Rodríguez Bermúdez – CUD de San Javier, CUD-AGA; Vocal 7: Dña. María Teresa Martínez Inglés – CUD San Javier, CUD-AGA y Vocal 8: Dña. Belinda Misiego Tejeda – Isdefe.

² **Lista de 10 artículos finalistas:** #7 Mantenimiento predictivo de motores de buques mediante aprendizaje automático. 14 puntos; #100 Identificación de agentes de guerra biológica mediante inmunobiosensado de exoproteínas dependientes del Sistema de Secreción bacteriano de Tipo 2. 9 puntos; #88 Resultados de la primera misión española de observación IOD de prestaciones submétricas en la Estación Espacial Internacional. 6 puntos; #133 Análisis de sentimiento en las redes sociales producido por las Fuerzas Armadas Españolas. 6 puntos; #61 Mejora de la gestión en las bases de operaciones avanzadas por medio de sistemas IoT. 5 puntos; #59 Modelos epidemiológicos: Estimadores de estado basados en información incompleta. 4 puntos; #74 Navegación visual eficiente aplicada a sistemas aéreos no tripulados en entornos de GNSS denegado 4 puntos; #113 Formación física y eficacia operativa en unidades paracaidistas. 3 puntos; #116 Modelización de cargas explosivas para el diseño de estructuras blindadas mediante simulación. 2 puntos y #107 Influencia de la emisividad superficial en la medida de la temperatura de transición vítrea en el análisis dinámico de propulsores sólidos. 1 punto.

ÍNDICE

Prólogo DIGEREM	I
Prólogo Consejero Delegado de ISDEFE	III
Mantenimiento predictivo de motores de buques mediante aprendizaje automático	1
Modelos epidemiológicos: estimadores de estado basados en información incompleta	13
Mejora de la gestión en las bases de operaciones avanzadas por medio de sistemas IoT	23
Navegación visual eficiente aplicada a sistemas aéreos no tripulados en entornos de GNSS denegado	35
Resultados de la primera misión española de observación IOD de prestaciones submétricas en la Estación Espacial Internacional	45
Identificación de agentes de guerra biológica mediante inmunobiosensado de exoproteínas dependientes del sistema de secreción bacteriano de tipo 2	55
Influencia de la emisividad superficial en la medida de la temperatura de transición vítrea en el análisis dinamomecánico de propulsores sólidos	65
Formación física y eficacia operativa en unidades paracaidistas	75
Modelización de cargas explosivas para el diseño de estructuras blindadas mediante simulación	85
Análisis de sentimiento en las redes sociales producido por las Fuerzas Armadas españolas	95

Mantenimiento predictivo de motores de buques mediante aprendizaje automático

**Novoa Paradela, David^{1,*}; Eiras Franco, Carlos¹; Fontenla Romero, Óscar¹;
Lamas López, Francisco²; Sanz Muñoz, David³**

¹ CITIC, Universidad de A Coruña. Campus de Elviña s/n, 15071, A Coruña. david.novoa@udc.es, carlos.eiras.franco@udc.es, oscar.fontenla@udc.es

² CESADAR, Armada Española. Arsenal de Cartagena, 30290, Cartagena. flamlop@mde.es

³ DigitalLabs / IA, INDRA. Avenida de Bruselas 35, 28108, Madrid. dsmunoz@indra.es

* Autor principal

Resumen: la aparición de anomalías durante el funcionamiento de los sistemas industriales puede apuntar a la presencia de degradaciones y fallos, que deriven con el tiempo en comportamientos indeseados, pérdidas de operatividad y la rotura final del sistema. Las técnicas de mantenimiento predictivo se encargan de monitorizar el estado de los sistemas para llevar a cabo la detección de estas anomalías en fases incipientes, permitiendo programar las tareas de mantenimiento de forma óptima. En este trabajo se presenta una solución de mantenimiento predictivo para motores de buques basada en técnicas de inteligencia artificial. Para ello se hace uso de la información de los sensores (temperaturas, presiones, etc.) recogida en tiempo real por los buques y transmitidos diariamente al centro de control. El sistema desarrollado es capaz de predecir la aparición de modos de fallos a partir de un histórico de datos del motor. Además, para que su uso sea escalable a grandes flotas, se ha implementado la solución en el entorno distribuido Spark.

Palabras clave: aprendizaje automático, mantenimiento predictivo, Spark.

1. Introducción

El coste de mantenimiento representa una parte importante de los costes operativos de la industria. En algunos casos, como en la industria metalúrgica, estos costes pueden llegar a suponer el 15 %-60 % de los costes totales de producción. Además, de estos, una tercera parte de la inversión se desperdicia como resultado de actividades innecesarias o incorrectas [1]. Sin embargo, el mantenimiento es crucial puesto que el fallo de un sistema puede llegar a suponer enormes costes económicos.

En el pasado, la imposibilidad de manejar los grandes y continuos flujos de datos ha llevado a emplear, en muchos casos, técnicas estadísticas. El mantenimiento predictivo actual, sin embargo, sigue una filosofía más avanzada: en lugar de depender de estas estadísticas industriales (p. ej., tiempo medio entre fallos) para programar actividades de mantenimiento, se lleva a cabo una monitorización en tiempo real del sistema para determinar su estado. La capacidad de cómputo actual permite tanto procesar cantidades de datos mayores, así como la utilización de técnicas más sofisticadas para llevar a cabo predicciones, detección de estados anormales y un diagnóstico del sistema. Por tanto, el mantenimiento predictivo se puede entender como un mantenimiento preventivo [2] condicionado al estado actual del sistema y a las predicciones a futuro realizadas a partir de un histórico de operación.

En este trabajo de investigación se presenta el desarrollo de un sistema de mantenimiento predictivo enmarcado dentro del proyecto SOPRENE en su aplicación a los motores de buques de la Armada. El sistema planteado ha analizado y empleado las técnicas de aprendizaje automático en entornos distribuidos. En este sentido, las metodologías consideradas se pueden dividir según lo expuesto por Ran *et al.* [3]:

1.1 Categorías en función del propósito

En base al criterio de optimización seguido, podemos distinguir entre:

- Minimización de costes. La métrica de coste empleada es habitualmente el tiempo de vida útil (RUL, *Remaining Useful Life*) del sistema, aunque también es posible definir un modelo de coste *ad hoc* [4].
- Maximización de la fiabilidad y la disponibilidad. Las métricas son la probabilidad de un sistema de encontrarse en un estado de funcionamiento normal dado un intervalo de tiempo [5] y la probabilidad de que el sistema se encuentre operativo [6].
- Optimización multiobjetivo. Se busca optimizar múltiples métricas simultáneamente para lograr un mejor equilibrio entre los objetivos. Además de las antes mencionadas, emplean métricas como el riesgo, la seguridad o la viabilidad. Generalmente, es imposible obtener los valores óptimos para todos los obje-

tivos al mismo tiempo, por lo que se han desarrollado una gran variedad de modelos multiobjetivo [7]-[9].

1.2 Categorías en función de la aproximación

En base al tipo de aproximación empleada, podemos distinguir entre:

- Aproximaciones basadas en conocimiento. Se utiliza conocimiento experto y procesos de razonamiento deductivo. Existen aproximaciones basadas en ontologías [10], en reglas [11] o en modelos que tratan de vincular los procesos físicos de un sistema con modelos matemáticos, como modelos Gaussianos [12], modelos de sistemas lineales [13] o modelos de Markov [14].
- Aproximaciones basadas en técnicas clásicas de aprendizaje automático (*machine learning*, ML). Se han utilizado redes de neuronas artificiales [15], los árboles de decisión [16] (incluyendo el algoritmo Random Forest [17]), así como máquinas de soporte de vectores (SVM), tanto de forma supervisada [18] como no supervisada [19]. Por último, la técnica de vecinos más cercanos (k-NN) es una de las más comunes para la clasificación de fallos [20], para la predicción de tiempo de vida útil (RUL) [21] y la detección temprana [22].
- Aproximaciones basadas en aprendizaje profundo (*deep learning*). Una de las más empleadas son las redes neuronales Autoencoder, cuya capa de salida busca reproducir los datos presentados en su capa de entrada después de haber pasado por una fase de compresión dimensional, permitiendo crear modelos robustos ante el ruido [23]. También se han utilizado redes neuronales recurrentes (RNN), basadas en celdas Long Short Term Memory Networks (LSTM) [24], que pueden aprender dependencias a más largo plazo. Este tipo de redes son muy potentes para el análisis de secuencias [25].

2. Contextualización del problema

La aplicación de las técnicas de mantenimiento predictivo a los motores de los buques de la Armada está definida por las características de los mismos: los buques considerados registran cada 10 s los valores de cerca de 300 variables asociadas al motor, representando una amplia variedad de aspectos físicos del mismo (ej., temperatura del gas, presión en los filtros, etc.), y se dispone de 4 años de históricos de operación de 4 buques. Sin embargo, los datos registrados presentan una calidad heterogénea debido a fallos en los sensores (valores erróneos o perdidos) o problemas en las comunicaciones (frecuencias de muestreo no uniformes). Por otro lado, existen tres modos de funcionamiento del motor en función de las revoluciones por minuto (RPM) a las que trabaja: motor apagado (RPM cercanas a cero), motor a ralentí (RPM aproximadas a un umbral μ) y motor en operación (RPM superiores a un umbral μ), teniendo que caracterizar y filtrar los estados de interés.

En cuanto a los fallos de funcionamiento en el histórico de datos, debido a la extensión de los archivos, su identificación y etiquetado manual por parte de un experto es inviable. De este modo, no se dispone de un conjunto de modos de fallo localizados y, por tanto, la detección de anomalías debe ser desarrollada de forma no supervisada. Para suplir esta carencia se utiliza un análisis FMECA (*Failure Mode, Effects, and Criticality Analysis*) que describe de forma teórica los modos de fallo que se pueden producir durante el funcionamiento del motor (las variables o elementos que intervienen en dichos modos y sus valores nominales, máximos y mínimos).

Finalmente, debido a que la solución propuesta no debe analizar un único buque, sino una flota, la gestión y control se lleva a cabo de manera centralizada. Cada buque envía los datos registrados por sus sensores a un nodo central donde se almacenan y procesan. Ya que los conjuntos de datos con los que se va a tratar pueden llegar ser de gran tamaño, y para que el sistema sea escalable a un número elevado de buques, el sistema ha de ser distribuido: se ha utilizado el sistema de ficheros distribuido HDFS para almacenar los datos y el entorno Apache Spark para entrenar y ejecutar los modelos de forma distribuida.

3. Solución propuesta

La solución diseñada combina las soluciones a varios subproblemas. Así, la arquitectura general se compone de cuatro bloques o tareas claramente delimitadas (véase figura 1):

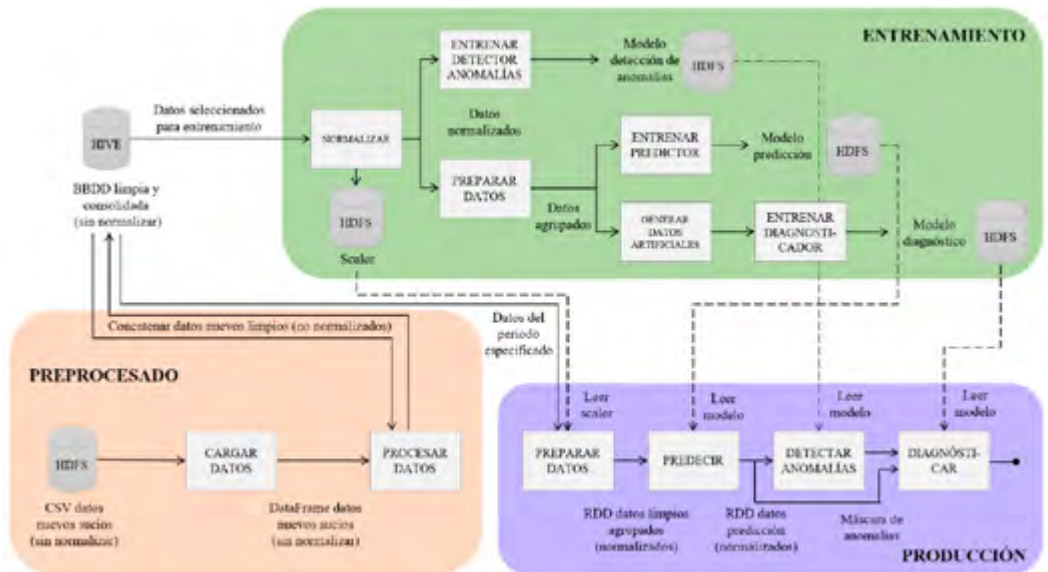


Figura 1. Arquitectura de la solución.

3.1 Módulo de preprocesamiento

Para entrenar los modelos de ML se necesita disponer de un conjunto de datos robusto y representativo de los diferentes estados del motor. Los valores medidos por los sensores del sistema se encuentran almacenados en ficheros CSV en HDFS, pero, como se ha detallado en la sección 3, estos datos precisan ser preprocesados. Así, en este bloque se comprueba la existencia de valores perdidos, se unifica la frecuencia de muestreo y lleva a cabo un proceso de selección de variables:

- Cargar datos. El flujo comienza mediante la lectura del conjunto de datos que se desea procesar y que se encuentran almacenados en HDFS.
- Procesar datos. Para solventar el problema de la baja calidad de los datos originales (sección 3) se llevan a cabo dos acciones:
 1. Selección de variables. De todas las variables se emplearán aquellas que presentan suficiente variabilidad en sus valores. Esta selección se lleva a cabo tanto automáticamente (desechando variables que presentan valores constantes) y manualmente, eliminando las seleccionadas por el usuario.
 2. Estandarización de las frecuencias. Para evitar el segundo problema es necesario homogeneizar la frecuencia de muestreo. Así, se crea, a partir de los datos iniciales, un conjunto en el que exista una medida simultánea para todas las variables cada 60 s. Para ello se rellena la ausencia de valores de las variables que presenten una frecuencia inferior con el último valor disponible.

Los datos procesados son almacenados en una base de datos consolidada de Hive.

- Normalizar datos. Debido a que los valores recogidos por los sensores oscilan en rangos muy diferentes entre sí para cada variable se ejecuta un proceso de normalización de forma individual. El normalizador entrenado se almacena en HDFS para ser utilizado en producción con la llegada de nuevos datos.
- Preparar datos. Para dotar de flexibilidad al sistema, el usuario puede establecer la unidad del horizonte con el que llevar a cabo la predicción (horas, días, semanas o meses), de forma que los datos normalizados se agrupen temporalmente (ej. si se van a realizar predicciones en días, los datos normalizados se agrupan en un único dato por día). Además, antes de agrupar los datos, estos se filtran en función del valor de la variable RPM para eliminar aquellos que se correspondan con instantes del motor apagado.

3.2 Módulo de predicción

El objetivo principal es conocer el estado del motor en un instante de tiempo futuro. De esta manera, para llevar a cabo una predicción desde un instante de tiempo t_i , para

un instante futuro que dista a horizonte unidades de tiempo de uso del motor ($t_i + \text{horizonte}$), el sistema deberá recibir la información recogida por los sensores durante las últimas ventana unidades de tiempo anteriores a t_i ($t_i - \text{ventana}$), definida junto al horizonte por el usuario. En la figura 2 se muestra una representación gráfica de estos conceptos.

Los métodos de predicción disponibles son: regresión lineal (en su versión distribuida), y de redes LSTM y Elastic Net (de forma no distribuida). Tras el entrenamiento de los modelos, estos son almacenados en HDFS para ser usados en producción del mismo modo que los normalizadores.

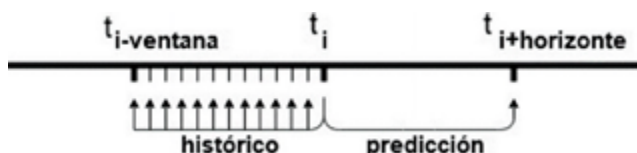


Figura 2. Utilización de una ventana de datos previos para llevar a cabo una predicción.

3.3 Módulo de detección de anomalías

Una vez realizada la predicción del estado del motor, es necesario determinar si dicho estado se corresponde con un valor normal o no. Al no disponer de anomalías etiquetadas, el proceso de detección se realiza de forma no supervisada utilizando una red Autoencoder. Mediante el error de reconstrucción podemos discernir entre datos normales (bajo error de reconstrucción) y datos anómalos (errores altos). Este valor puede presentarse como el error cuadrático medio de todas las variables de entrada (un único valor) o su descomposición, el error de cada una de las variables o nodos de entrada de la red. El objetivo, dado un conjunto de registros, es determinar cuáles son anómalos y qué variables provocan dichas anomalías, lo que se hace en tres subfases secuenciales:

1. Detectar anomalías. Para determinar qué registros son anómalos se realiza un primer filtrado mediante el error cuadrático medio. A partir de un error umbral precalculado se clasifican los datos que lo superen como anómalos y el resto como normales. El usuario escoge si el cálculo de este error umbral se realiza mediante el rango intercuartílico o estableciendo un porcentaje de datos anómalos en el conjunto.
2. Independizar contribuciones. Para determinar qué variables han sido las causantes de la aparición de la anomalía en los datos clasificados como anómalos, se emplea la descomposición del error de reconstrucción. Así, se pasa de un único error global a tantos como variables conformen el registro, pudiendo

ordenar las variables por su error de reconstrucción y determinar de forma automática la contribución de cada variable utilizando el método Elbow [26].

3. Construir una máscara de anomalías. A partir de la selección de la subfase anterior, se construye una matriz o máscara de salida de dimensiones $m \times n$ (siendo m el número de filas o registros y n el número de columnas o variables) en la que las variables anómalas se marcan con un uno y las normales con un cero.

3.4 Módulo de diagnóstico

El bloque de diagnóstico es el encargado de, a partir de la predicción y la máscara de detección de anomalías, determinar qué modos de fallo pueden producirse y su probabilidad. Para determinar la probabilidad de cada modo de fallo se ha decidido utilizar un modelo de clasificación supervisado basado en redes neuronales. Debido a que no se dispone de conjuntos de datos etiquetados de todos los posibles modos de fallo, se ha implementado un generador de datos artificiales que los produzca a partir de las características teóricas de los modos de fallo recogidas en el documento FMECA (variables implicadas, valores nominales, valores toques de rango, etc.). Así, el modelo es entrenado considerando a cada uno de los modos de fallo como una clase de salida, teniendo en su capa de entrada tantas neuronas como variables del motor, y en su capa de salida tantas softmax como modos de fallo.

El clasificador de modos de fallo se emplea en combinación con la máscara de salida del bloque de detección de anomalías para llevar a cabo el diagnóstico del motor, utilizada para acotar los modos de fallo posibles, omitiendo los modos de fallo que todas sus variables implicadas han sido consideradas normales. La salida de este módulo (y final del sistema) contiene el identificador del modo de fallo en el FMECA, la fecha prevista en la que se produzca el modo de fallo y la probabilidad/certidumbre asociada a la misma.

4. Resultados

Dado que no se dispone de un conjunto de *targets* o valores de referencia, no ha sido posible realizar una evaluación cuantitativa del sistema. El desarrollo de las distintas soluciones parciales que conforman la solución final ha sido supeditado a una evaluación mayormente cualitativa por los expertos de la organización implicada:

- Predicción. El filtrado de los datos por RPM y su agrupación reduce considerablemente el número de ejemplos disponibles. Esto incapacita a modelos como las redes LSTM que precisan grandes cantidades de datos. En estos escenarios se ha visto que los modelos más sencillos reportan mejores resultados, siendo posible realizar predicciones hasta 10 días/semanas/meses con una calidad razonable.

- Detección de anomalías. Los resultados del detector de anomalías se han comparado con las alarmas registradas en los buques a lo largo de los cuatro años. Tal y como se observa en la figura 3, el modelo fue capaz de detectar la mayoría de estas anomalías e incluso anticipar la ocurrencia de algunas de ellas.

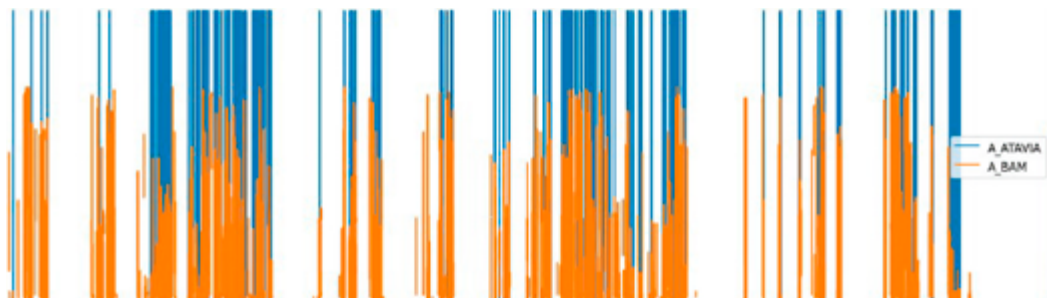


Figura 3. Estados anómalos etiquetados manualmente por un experto humano (azul) frente a los detectados por el sistema (naranja).

- Diagnóstico. La generación de conjuntos de datos artificiales basados en el FMECA ha permitido construir modelos de clasificación que determinen qué modos de fallo se están produciendo. Sin embargo, este comportamiento teórico del motor no tiene por qué corresponderse siempre con la realidad, ya que su funcionamiento puede variar con el uso, reemplazamiento o la reparación de piezas.

5. Conclusión

La solución desarrollada permite predecir la aparición de los diferentes modos de fallo descritos en el FMECA del motor de combustión de un buque. Las tareas de predicción y detección de anomalías son totalmente independientes, por lo que esta última se puede llevar a cabo tanto para momentos futuros (datos provenientes de la predicción), presentes (tiempo real) o pasados (análisis a posteriori). La mayor responsabilidad recae en la tarea de predicción ya que las operaciones posteriores parten de la salida de este módulo. Para dotarlo de flexibilidad se proporciona un amplio abanico de métodos, fácilmente extensible si es preciso, a la vez que se permite al usuario configurar los parámetros de predicción para obtener los resultados más adecuados en cada situación. El sistema es altamente configurable y su uso se puede extrapolar a otros buques que presenten características similares. Su ejecución es distribuida, por lo que los tiempos de predicción, detección y diagnóstico son bajos, siendo el mayor cuello de botella la tarea de preprocesado, operación que no se lleva a cabo de manera distribuida.

Agradecimientos

Queremos agradecer a la Armada Española la colaboración en el proyecto SOPRENE, así como a la Consejería de Educación, Universidad y Formación Profesional de la Xunta de Galicia, por la financiación ofrecida al CITIC a través del Fondo Europeo de Desarrollo Regional (FEDER) y de la Secretaría General de Universidades (Ref. ED431G 2019/01).

Referencias

1. R. K. Mobley, *An Introduction to Predictive Maintenance*, 2nd ed., Elsevier, 1990.
2. O. Motaghare, A. S. Pillai, and K. I. Ramachandran, «Predictive maintenance architecture», *IEEE International Conference on Computational Intelligence and Computing Research (ICCIC)*, 2018, pp. 1–4.
3. Y. Ran, X. Zhou, P. Lin, Y. Wen, and R. Deng, «A survey of predictive maintenance: Systems, purposes and approaches», [arXiv:1912.07383](https://arxiv.org/abs/1912.07383), 2019.
4. Y. He, X. Han, C. Gu, and Z. Chen, «Cost-oriented predictive maintenance based on mission reliability state for cyber manufacturing systems», *Advances in Mechanical Engineering*, 2018.
5. S. Song, D. W. Coit, and Q. Feng, «Reliability analysis of multiple-component series systems subject to hard and soft failures with dependent shock effects», *IIE Transactions*, vol. 48, no. 8, pp.720–735, 2016.
6. M. A. Gravette and K. Barker, «Achieved availability importance measure for enhancing reliability-centered maintenance decisions», *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, vol. 229, no. 1, pp. 62–72, 2015.
7. L. Lin, B. Luo, and S. Zhong, «Multi-objective decision-making model based on CBM for an aircraft fleet with reliability constraint», *International Journal of Production Research*, vol. 56, no. 14, pp. 4831–4848, 2018.
8. J. Zhao and L. Yang, «A bi-objective model for vessel emergency maintenance under a condition-based maintenance strategy», *Simulation*, vol. 94, no. 7, pp. 609–624, 2018.
9. Y. Xiang, D. Coit, and Z. Zhu, «A multi-objective joint burn-in and imperfect condition-based maintenance model for degradation-based heterogeneous populations», *Quality and Reliability Engineering International*, vol. 32, no. 8, pp. 2739–2750, 2016.
10. A. Konys, «An ontology-based knowledge modelling for a sustainability assessment domain», *Sustainability*, vol. 10, 300, 2018.
11. M. D. Ying Peng and M. J. Zuo, «Current status of machine prognostics in condition-based maintenance: a review», *The International Journal of Advanced Manufacturing Technology*, vol. 50, pp. 297–313, 2010.

12. N. Chen, Z.-S. Ye, Y. Xiang, and L. Zhang, «Condition-based maintenance using the inverse Gaussian degradation model assessment domain», *European Journal of Operational Research*, vol. 243, no. 1, pp. 190–199, 2015.
13. A. Lucifredi, C. Mazzieri, and M. Rossi, «Application of multiregressive linear models, dynamic kriging models and neural network models to predictive maintenance of hydroelectric powers ystems», *Mechanical Systems and Signal Processing*, vol. 14, no. 3, pp. 471–494, 2000.
14. M. Calder and M. Sevegnani, «Stochastic model checking for predicting component failures and service availability», *IEEE Transactions on Dependable and Secure Computing*, vol. 16, no. 1, pp. 174–187, 2019.
15. B. Samanta and K. Al-Balushi, «Artificial neural network based fault diagnostics of rolling element bearings using time-domain features», *Mechanical Systems and Signal Processing*, vol. 17, pp. 317–328, 2003.
16. J. R. Quinlan, «Improved use of continuous attributes in C4.5», *CoRR*, vol. cs.AI/9603103, 1996. [Online]. Available: <https://arxiv.org/abs/cs/9603103>.
17. J. Shi, N. Niu, and X. Zhu, «A fault diagnosis method for multi-condition system based on random forest», *Prognostics and System Health Management Conference (PHM-Paris)*, pp. 350–355, 2019.
18. Y. Chen, H. Yigang, Z. Li, L. Chen, and C. Zhang, «Remaining useful life prediction and state of health diagnosis of lithium-ion battery based on second-order central difference particle filter», *IEEE Access*, vol. 8, pp. 37305–37313, 2020.
19. G. Ratsch, S. Mika, B. Scholkopf, and K. Muller, «Constructing boosting algorithms from SVMs: an application to one-class classification», *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 24, no. 9, pp. 1184–1199, 2002.
20. G. A. Susto, A. Schirru, S. Pampuri, S. McLoone, and A. Beghi, «Machine learning for predictive maintenance: A multiple classifier approach», *IEEE Transactions on Industrial Informatics*, vol. 11, no. 3, pp. 812–820, 2015.
21. Z. Liu, W. Mei, X. Zeng, C. Yang, and X. Zhou, «Remaining useful life estimation of insulated gate bipolar transistors (IGBTs) based on a novel volterra k-nearest neighbor optimally pruned extreme learning machine (VKOPP) model using degradation data», *Sensors*, vol. 17, no. 11: 2524, 2017.
22. X. Long Chen, P. hong Wang, Y. Sheng Hao, and M. Zhao, «Evidential KNN-based condition monitoring and early warning method with applications in power plant», *Neurocomputing*, vol. 315, pp. 18 – 32, 2018.
23. M. Ma, C. Sun, and X. Chen, «Deep coupling autoencoder for fault diagnosis with multimodal sensory data», *IEEE Transactions on Industrial Informatics*, vol. 14, pp. 1137–1145, 2018.

24. J. Yuan and Y. Tian, «An intelligent fault diagnosis method using GRU neural network towards sequential data in dynamic processes», *Processes*, vol. 7, no. 3: 152, 2019.
25. R. Yang, M. Huang, Q. Lu, and M. Zhong, «Rotating machinery fault diagnosis using long-short-term memory recurrent neural network», *IFAC-Papers OnLine*, vol. 51, no. 24, pp. 228–232, 2018, 10th IFAC Symposium on Fault Detection, Supervision and Safety for Technical Processes (SAFEPROCES).
26. P. Bholowalia and A. Kumar, «Article: Ebk-means: A clustering technique based on elbow method and k-means in WSN», *International Journal of Computer Applications*, vol. 105, no. 9, pp. 17–24, 2014.

Modelos epidemiológicos: estimadores de estado basados en información incompleta

Gómez López, Miguel Ángel^{1,*}; Solera Rico, Alberto²; Gómez Pérez, Ignacio³;
Valero Sánchez, Eusebio⁴

1 Universidad Politécnica de Madrid (UPM). Calle de las Carretas, 13, 28012, Madrid.
angel.glopez@alumnos.upm.es

1 Instituto Nacional de Técnica Aeroespacial (INTA). Ctra. M-301, Km 10,500, 28330, San Martín de la Vega (Madrid).

2 Universidad Politécnica de Madrid (UPM). Calle de las Carretas, 13, 28012, Madrid.
a.solera@alumnos.upm.es

2 Instituto Nacional de Técnica Aeroespacial (INTA). Ctra. M-301, Km 10,500, 28330, San Martín de la Vega (Madrid). gomezlma@inta.es

3 Universidad Politécnica de Madrid (UPM). Calle de las Carretas, 13, 28012, Madrid.
ignacio.gomez@upm.es

4 Universidad Politécnica de Madrid (UPM). Calle de las Carretas, 13, 28012, Madrid.
eusebio.valero@upm.es

* Autor principal

Resumen: el modelado y análisis de los procesos de difusión ha sido un área de investigación intensa y que hace uso a menudo de muchos campos diferentes de la ciencia, desde la biología matemática, la física o la informática y economía, hasta la ingeniería. Uno de los primeros modelos epidémicos pertenece a Daniel Bernoulli en 1760, motivado por el estudio de la propagación de la viruela. Además de Bernoulli, había muchos investigadores diferentes que también trabajaban en modelos matemáticos de epidemias en esta época. Y aunque estos modelos iniciales fueron bastante simplistas, marcaron el inicio y las vías de desarrollo y de dichos modelos, con el objetivo de proporcionar información sobre cómo se pueden propagar diversas enfermedades a través de una población. En los últimos años, ha habido un resurgimiento del interés en estos problemas a medida que el concepto de «redes» gana terreno sobre estos modelos primitivos y se vuelve cada vez más frecuente en el modelado de muchos aspectos de este y otros diferentes del mundo actual.

Sin embargo, el talón de Aquiles de la modelización de estos procesos continúa siendo la elección o calibración de estos parámetros. Pese a un mayor número de datos

disponibles, la heterogeneidad entre países o regiones a la hora de recoger dichos datos hace que el proceso de ajuste del modelo sea complejo y a menudo infructuoso sin los datos adecuados.

Palabras clave: control moderno, epidemiología, observabilidad, Kalman.

1. Introducción

El proyecto tiene como objetivo estimar, de manera óptima, los casos de infectados en una población en una enfermedad de rápida y fácil transmisión basándose en la teoría de control moderno. Concretamente, el proyecto se centra en el estudio de estimadores de estado óptimos (filtro de Kalman extendido) y los conceptos de observabilidad y controlabilidad, aplicados a la dinámica de poblaciones de distintos modelos epidemiológicos, tanto deterministas como estocásticos, vistos como si de sistemas de control se tratara.

La propagación de la enfermedad a través de las poblaciones humanas es compleja. Las características de la propagación de la enfermedad evolucionan con el tiempo, como resultado de una multitud de factores ambientales y geográficos. Esta múltiple interdependencia no estacionaria es uno de los factores que complican enormemente la predicción en tiempo real, y la toma de decisiones en momentos claves de expansión de un brote contagioso. Además, en general durante el brote las autoridades priorizan, como es lógico, la gestión de los recursos críticos y asistencia a los infectados frente a la toma de datos sobre el terreno. Esto da lugar a que la fuente de datos para realizar los análisis en tiempo real sea a menudo ruidosa e incompleta. Para describir correctamente la propagación de la enfermedad exploramos por tanto un enfoque flexible, basado en modelos estocásticos con variación temporal para la dinámica de la enfermedad.

Aquí es donde juega un papel importante en las últimas décadas la teoría moderna de control. Este campo de la matemática se ha usado con éxito en el análisis de procesos complejos, con fuertes interdependencias y un número grande de variables de interés (desde decenas a cientos). Una de las herramientas nacidas de esta teoría, el filtro de Kalman, es prácticamente un estándar en el análisis de datos y filtrado de los mismos, pero se ha desarrollado poco la teoría y su uso en este campo como «observador» hasta recientemente [1], frente al fuerte desarrollo que ha tenido en otros campos [2].

Este enfoque nos permite reconstruir la evolución temporal de parámetros incluso cuando no se pueden medir directamente, e incluso proporciona cotas y límites concretos respecto a qué se puede estimar y calcular, y que no, con los datos proporcionados. El dual de esta propiedad de los sistemas dinámicos, conocida como observabilidad, es la controlabilidad [3]. Así, esta misma herramienta tiene la capacidad de proporcionarnos una magnitud de en qué medida los parámetros –como la tasa de infección, el tiempo hospitalizados o el número de camas disponibles– pueden usarse para variar el curso de la epidemia, en qué instantes temporales son más efectivos, y cuánto puede moverse el «pico» de la infección o achatar su extensión [4].

2. Materiales y métodos

En primer lugar, desarrollaremos el uso de los algoritmos en un modelo sencillo, donde se conocen los parámetros y el proceso de observación, y se introducen incertidumbres estocásticas, así como sesgos, que el algoritmo tendrá que sortear, para a continuación, aplicarlo a conjuntos de datos reales, nuestro. Sugerimos una metodología que puede usarse como un primer paso hacia una mejor comprensión de un complejo epidemia, y toma de decisiones, en situaciones donde los datos son limitados y/o inciertos. La unidad de tiempo básica de nuestro modelo es el día, permitimos variar los parámetros del modelo de día a día. El modelo ha de ser suficientemente detallado como capturar la evolución de la enfermedad unos 6 días, ya que la aplicación de las políticas para evitar nuevos casos, tardan unos días en verse reflejadas y, por tanto, su eficacia.

2.1 Modelado sistema

- Modelo del sistema (paso de propagación del filtro), es decir, el sistema a controlar, usando los conceptos de «estado» e incertidumbres.
- Modelo epidemiológico compartimentado, el modelo SIR (susceptible, infectado y recuperado). La relación entre los compartimentos se detalla en la figura 1.
- El modelo tiene parámetros como:
 - La ratio de infección: $\beta I S$.
 - La ratio de paso de infectado a recuperado: γI .



Figura 1. Gráfica con el modelo compartimental.

2.2 Medidas

El modelo de medidas (paso de corrección del filtro), es decir, los observables con a su vez sus incertidumbres. La toma de decisiones durante el desarrollo de un brote de contagios depende en gran medida de la información que se tenga sobre la evolución del número de infectados reales cada día. Sin embargo, a menudo, solo se cuenta con el número de infectados detectados diariamente. El número de infectados detectados está relacionado con el número de infectados reales, pero numéricamente pueden llegar a ser muy diferentes. Razones múltiples:

- Una persona infectada un día puede no mostrar síntomas hasta unos días después.
- Aun mostrando síntomas puede tardar todavía en acudir a los servicios sanitarios.
- Aun acudiendo a los servicios sanitarios, puede tardar en ser diagnosticada como infectado.
 - Falta de disponibilidad no se le puedan aplicar las pruebas («test») que determinen su infección.
 - Resultados de estas pruebas tarden un tiempo en conocerse.

Es por tanto indispensable emplear un modelo de medida, o técnicas basadas en los estadísticos usados, que permiten estimar el número de infectados reales conociendo diversas distribuciones de datos sobre la epidemia, tanto sobre el número de infectados detectados o el número de muertes producidas, su desglose y cualquier otra información relevante (distribución por edad y por zonas, focos de infectados, etc.). Para hacer estimaciones fiables usando estas técnicas en un país o una región ayudará la información precisa que se tenga de lo ocurrido en países o regiones comparables (parecido número de habitantes, sistemas de salud similares, etc.). En este caso en el modelo se incluyen dos medidas, $Z(1)$ y $Z(2)$, que son respectivamente el número de individuos enfermos en ese momento, y el número de individuos que hasta la fecha han pasado la enfermedad o la están pasando (recuperados sumados a los enfermos). Dichas medidas tienen un ruido asociado, de distribución normal, media nula, y varianza desconocida a priori, pero supuesta acotada superiormente por 0,1.

2.3 Implementación

El simulador se ha implementado en MATLAB-Simulink por conveniencia, aunque sin hacer uso de las funciones de alto nivel incluidas en el *software*, por lo que sería reproducible en cualquier otro lenguaje de programación. La integración de las ecuaciones diferenciales se realiza con un esquema Runge-Kutta.

3. Resultados y discusión

3.1 Observabilidad

La observabilidad es una propiedad de los modelos de los sistemas dinámicos de la que se puede inferir si es posible conocer los estados internos de un sistema a partir de la medida de sus salidas. El concepto de observabilidad fue introducido por Rudolf E. Kalman para sistemas dinámicos lineales, a partir de él se pueden diseñar sistemas dinámicos para estimar el estado de un sistema a partir de las mediciones de las salidas.

Estos sistemas se denominan observadores o estimadores de estado para dicho sistema. Este modelo generalizado de sistema se puede escribir como:

$$x(t+1) = f(x(t), v)$$

$$y(t) = g(x(t), \mu)$$

Donde $x(t)$ es el estado, $y(t)$ son las medidas, y v y μ son los ruidos o incertidumbres asociadas a las mismas. Para saber si el modelo es observable con las medidas elegidas, se linealiza entorno al punto de trabajo en cada instante, obteniendo las matrices jacobianas de cambio de estado y de medida:

$$F(t) = \begin{pmatrix} 1 - x_2(t^*)\beta & -\beta & 0 \\ \beta & 1 - x_1(t^*)\beta - \gamma & 0 \\ 0 & \gamma & 1 \end{pmatrix}$$

$$H(t) = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \end{pmatrix}$$

Y con ellas se construye la matriz de observabilidad:

$$O(t) = \begin{pmatrix} H \\ HF(t) \\ HF(t)^2 \\ \dots \\ HF(t)^{n-1} \end{pmatrix}$$

Resultando de rango igual al número de estados en cada paso temporal y, por tanto, observable.

3.2 Ecuaciones del filtro de Kalman extendido

Una vez decidido el modelo epidemiológico a utilizar, las medidas que estarán disponibles, y las incertidumbres asociadas, se realizan las simulaciones

$$\mathbf{X}_k = F_k(\mathbf{X}_{k-1}, \mathbf{W}_k)$$

$$\mathbf{Y}_k = G_k(\mathbf{X}_k, \mathbf{V}_k)$$

Donde $\mathbf{W}_k$ y $\mathbf{V}_k$ son respectivamente los ruidos de proceso y de medida, con sus matrices de covarianza asociadas $\mathbf{Q}_k = E\{\mathbf{W}_k \mathbf{W}_k^T\} - \overline{\mathbf{W}}_k \overline{\mathbf{W}}_k^T$ y $\mathbf{R}_k = E\{\mathbf{V}_k \mathbf{V}_k^T\} - \overline{\mathbf{V}}_k \overline{\mathbf{V}}_k^T$.

The F_k representa precisamente el modelo dinámico de cambio del vector de estado (el SIR) y $G_k(\cdot)$ caracteriza la relación entre los estados y las medidas.

Las ecuaciones de predicción y corrección del filtro se recogen aquí por claridad, haciendo uso de la notación más o menos estandarizada, pero para más detalles acudir a [2]:

$$\begin{aligned}\hat{\mathbf{X}}_k^- &= F_k(\hat{\mathbf{X}}_{k-1}, \bar{\mathbf{W}}_k) \\ \mathbf{I}_k &= \mathbf{Y}_k - G_k(\hat{\mathbf{X}}_k^-, \bar{\mathbf{V}}_k) \\ P_k^- &= A_k P_{k-1} A_k^T + B_k Q_k B_k^T \\ K_k &= P_k^- C_k^T (C_k P_k^- C_k^T + D_k R_k D_k^T)^{-1} \\ P_k &= P_k^- - K_k C_k P_k^- \\ \hat{\mathbf{X}}_k &= \hat{\mathbf{X}}_k^- + K_k \mathbf{I}_k\end{aligned}$$

donde $\hat{\mathbf{X}}_k^- = E\{\mathbf{X}_k | \mathbf{Y}_0, \dots, \mathbf{Y}_{k-1}\}$ es la estimación *a priori* de $\mathbf{X}_k$ que usa todas las observaciones hasta el instante $k-1$. $\mathbf{I}_k$ es el vector de innovaciones, $P_k^- = E\{(\mathbf{X}_k - \hat{\mathbf{X}}_k^-)(\mathbf{X}_k - \hat{\mathbf{X}}_k^-)^T\}$ y $P_k = E\{(\mathbf{X}_k - \hat{\mathbf{X}}_k)(\mathbf{X}_k - \hat{\mathbf{X}}_k)^T\}$ son respectivamente las estimaciones de la covarianza *a priori* y *a posteriori*. Finalmente, $\hat{\mathbf{X}}_k = E\{\mathbf{X}_k | \mathbf{Y}_0, \dots, \mathbf{Y}_k\}$ es la estimación final *a posteriori* de $\hat{\mathbf{X}}_k$.

3.3 Resultados

Una vez decidido el modelo epidemiológico a utilizar, las medidas que estarán disponibles, y las incertidumbres asociadas, se realizan las simulaciones.

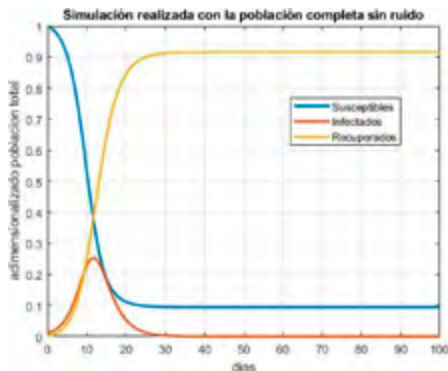


Figura 2. a) Sistema sobre el que se realizarán estimaciones (planta) supuestos conocidos todos los parámetros, simulación tipo.



Figura 2. b) Medidas realizadas durante la pandemia, corrompidas con ruido de distribución normal, varianza de un 10% sobre el total de la población, y media nula

En primer lugar, se muestra el proceso de la infección de la enfermedad propagada por la población, de manera ideal, esto se puede ver en la figura 2a. A continuación, se usa dicho modelo ideal para generar las medidas que, corrompidas con un ruido blanco de varianza desconocida pero acotada, alimentarán al observador en la siguiente etapa de la simulación (figura 2b). Se ha elegido un paso temporal para dichas simulaciones de 1 día, dentro de un horizonte temporal de 100 días, o 3 meses.

Finalmente se aplica el observador de estado, usando como propagador el modelo SIR propuesto anteriormente, en su versión no lineal, calculando los jacobianos a cada paso, para los cálculos de la ganancia de Kalman (matriz K) y de la covarianza de la innovación (matriz S).

Como las medidas, por su parte, si son lineales se puede entonces usar la matriz jacobiana H directamente para el cálculo de la innovación (diferencia entre las medidas predichas por el modelo en este paso temporal y las medidas tomadas en el mismo).

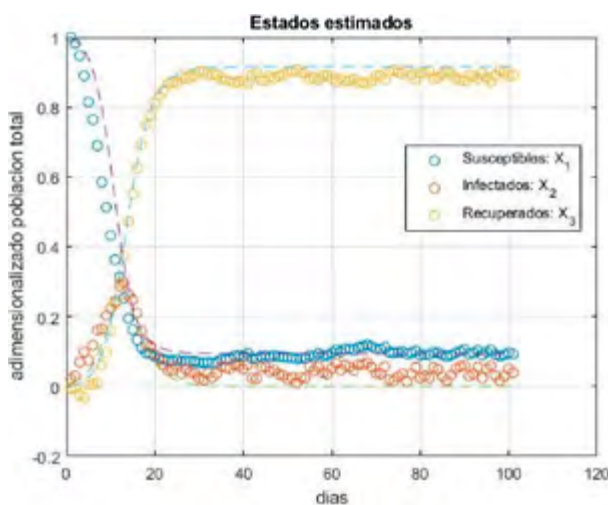


Figura 3. Resultado del modelo aplicado a unas medidas corrompidas y filtradas mediante EKF. En línea de trazos los resultados del modelo teóricos, mientras en puntos se observan los estados estimados a partir de las medidas.

En la figura 3 se pueden ver los resultados en la estimación del estado completo de dimensión 3 tras haber aplicado el filtro a las medidas, frente a en línea discontinua la propagación real que se obtendría si se pudiera acceder sin ruido a los tres estados reales.

4. Conclusiones

Los resultados del apartado anterior, aunque modestos, muestran que es posible diseñar y construir un estimador de estado completo tipo filtro de Kalman extendido, basado en modelos compartimentales, para conocer los estados internos no medibles

directamente aún en presencia de ruido. En este caso dichos estados han sido el número de individuos susceptibles de contraer la enfermedad y el número de individuos que han pasado y se han recuperado de dicha enfermedad a lo largo del tiempo.

El procedimiento mostrado se puede aplicar a modelos epidemiológicos compartimentales más complicados, que incluyan retardos temporales, parámetros que cambian con el tiempo y más comportamientos que reflejen el proceso que siguen los individuos a través del sistema sanitario de un país.

Agradecimientos

Se agradece la colaboración para este estudio al Instituto Nacional de Técnica Aeroespacial.

Referencias

1. R. Li, S. Pei, B. Chen, Y. Song, T. Zhang, W. Yang y J. Shaman, *Substantial undocumented infection facilitates the rapid dissemination of novel coronavirus SARS-CoV-2*, Science, ISSN: 0036-8075. doi: 10.1126/science.abb3221.
2. M. S. Grewal y A. P. Andrews, *Applications of Kalman Filtering in Aerospace 1960 to the Present [Historical Perspectives]*, IEEE Control Systems Magazine, 2010.
3. E. G. GILBERT, *Controllability and Observability in Multivariate Control Systems*, J.S.I.A.M. CONTROL, 1963.
4. J. S. Weitz, S. J. Beckett, A. R. Coenen, D. Demory, M. Dominguez-Mirazo, J. Dusho, C.-Y. Leung, G. Li, A. Malie, S. W. Park, R. Rodríguez-González, S. Shivam y C. Y. Zhao, *Modeling shield immunity to reduce COVID-19 epidemics spread*, Nature Medicine, 2020. doi: 10.1038/s41591-020-0895-3.
5. N. G. Davies, A. J. Kucharski, R. M. Eggo, A. Gimma, C. C.-1. working group y W. J. Edmunds, *The effect of non-pharmaceutical interventions on COVID-19 cases, deaths and demand for hospital services in the UK: a modelling study*, 2020.

Mejora de la gestión en las bases de operaciones avanzadas por medio de sistemas IoT

**Pereñíguez García, Fernando^{1,*}; Martínez Inglés, María Teresa²; Alonso Morán, Cristian³;
Molina Cegarra, Jaime⁴; Bermúdez Rico, Jorge⁵**

¹ Departamento de Ciencias e Informática, Centro Universitario de la Defensa (CUD). C/ Coronel López Peña, s/n, 30720, San Javier, Murcia. fernando.pereniguez@ cud.upct.es

² Departamento de Ingeniería y Técnicas Aplicadas, Centro Universitario de la Defensa (CUD). C/ Coronel López Peña, s/n, 30720, San Javier, Murcia. mteresa.martinez@ cud.upct.es

³ EZAPAC. Avda. de Lorca, s/n, 30820, Alcantarilla, Murcia. calomor@mde.es

⁴ Ala 14 Albacete, Ejército del Aire. Carretera de Murcia km 3, 02071. Los Llanos, Albacete. jmolceg@mde.es

⁵ Grupo Central de Mando y Control (GRUCEMAC), Crt. Barcelona Km 22,800, 28850, Torrejón de Ardoz Madrid. jberric@mde.es

* Autor principal: fernando.pereniguez@ cud.upct.es

Resumen: gracias al desarrollo de nuevas tecnologías, hoy en día se dispone de sistemas capaces de obtener información en tiempo real y de forma remota. Una de estas tecnologías es el Internet de las cosas (IoT), que interconecta dispositivos capaces de recopilar información y compartirla de forma autónoma. Parte del éxito de los sistemas IoT reside en el uso de dispositivos que tienen un bajo consumo de energía. Estas características del IoT facilitan actividades donde la toma de decisiones en base a datos obtenidos en tiempo real sea clave. Por ello, su uso en escenarios de gestión como las bases de operaciones avanzadas (FOB) puede ser de gran utilidad.

En este trabajo se pretende mostrar los beneficios que la tecnología IoT puede aportar a las FOB. Para ello, primero se presentan los aspectos básicos de los sistemas IoT con el fin de entender cómo funcionan dichos sistemas. Después, se presentan casos de uso que podría tener el IoT para la gestión de dichas bases y, por último, se muestra la validación de la implementación de un prototipo experimental mostrando los resultados que se obtuvieron en un escenario de monitorización de parámetros ambientales.

Palabras clave: FOB, IoT, LoRaWAN.

1. Introducción

El Internet de la cosas (*Internet of Things* – IoT) se ha convertido en un término esencial en el ámbito de las tecnologías de la información y las comunicaciones. El IoT abarca los dominios relacionados con la sensorización, actuación, computación e, incluso, tratamiento inteligente de datos mediante técnicas basadas en inteligencia artificial o *big data* [1]. La combinación de todas estas disciplinas tiene un objetivo claro: permitir la realización de operaciones de manera remota y no supervisada con un coste ínfimo si lo comparamos con el uso de mecanismos tradicionales. Además, la tecnología IoT ha sido desarrollada empleando estándares abiertos, garantizando la compatibilidad con los actuales sistemas de comunicaciones. Por tanto, las redes IoT pueden ser desplegadas para usos públicos o privados sobre las actuales infraestructuras de red.

Como consecuencia, las soluciones basadas en la tecnología IoT han sido ampliamente adoptadas en el ámbito civil tanto por la industria como por organismos públicos para implementar diversos tipos de aplicaciones en sector de la salud, transporte, logística, ciudades inteligentes o la agricultura, entre otros [1]. Sin embargo, el IoT es reconocido como una tecnología de doble uso cuya aplicación no se limita al ámbito civil, sino que a nivel militar ofrece numerosas ventajas a la hora de abordar ciertos cometidos asignados a las organizaciones encargadas de la defensa de los países. La entrada del ciberespacio como el cuarto escenario en el campo de batalla, tras los ya conocidos terrestres, navales y aéreos, y las características de los conflictos actuales, que destacan por su asimetría y complejidad, hacen del IoT una tecnología cuya aplicación es irrenunciable para afrontar con garantías los próximos teatros de operaciones donde se desarrollen los conflictos.

De hecho, la conveniencia de incorporar sistemas IoT al ámbito militar ha sido reconocida en el contexto de la OTAN. El grupo de trabajo IST-147 finalizó recientemente su trabajo identificando potenciales áreas de aplicación para el mando y control, conciencia situacional (*situational awareness*) y monitorización médica en el campo de batalla [2]. En la actualidad, el grupo IST-176 continúa este trabajo centrando su atención en la interoperabilidad de los sistemas de mando y control (C2C) y los sistemas IoT [2]. Por otro lado, la aplicabilidad no se limita al campo de batalla, sino que también reporta beneficios al funcionamiento operacional y logístico de las bases militares. En este sentido, el Ejército del Aire ha lanzado recientemente el proyecto BACSI (Base Aérea Conectada Sostenible Inteligente) [3] proponiendo, entre otros, el desarrollo de soluciones IoT que permitan reducir riesgos, mejorar el rendimiento y gestión de las operaciones que se llevan a cabo en una base militar.

Teniendo en cuenta los antecedentes existentes, este trabajo de investigación centra su atención en la aplicabilidad del Internet de las cosas en un escenario particular

que, según el conocimiento de los autores, no ha sido explorado hasta el momento: la gestión de las bases de operaciones avanzadas (en inglés FOB, *Forward Operating Base*). Una FOB se establece de manera temporal o semi-permanente para ofrecer apoyo al desarrollo de operaciones tácticas en áreas localizadas. Aunque las FOB podrían albergar instalaciones como un hospital, pista de aterrizaje u otras dependencias, reciben el apoyo de la base principal de operaciones [4]. La figura 1 muestra una fotografía tomada en el año 2006 de una FOB desplegada por el ejército alemán en Afganistán.

Una FOB proporciona protección al personal y almacenan recursos materiales como sistemas de armas, armamento, equipamiento de apoyo, etc. Sin embargo, su despliegue debe afrontar diversas restricciones relacionadas, por ejemplo, con el acceso a suministros de energía, combustible y agua. El objetivo de este artículo consiste en demostrar que la tecnología IoT puede adaptarse a estas restricciones y ofrecer una mejora en la gestión de las operaciones. Concretamente, se propondrán diversos ejemplos de aplicación, se identificarán las tecnologías más recomendables y se presentará un prototipo experimental usando equipamiento de bajo coste.

Este trabajo se estructura de la siguiente manera: en el apartado 2 se revisan aspectos básicos de los sistemas IoT. El apartado 3 ejemplifica los beneficios que reportan los sistemas IoT para la gestión de las FOB mediante la propuesta de diversas aplicaciones. El apartado 4 muestra la implementación de un prototipo experimental junto a resultados de validación del mismo. Finalmente, se resume las aportaciones de este trabajo.



Figura 1. Ejemplo de una FOB [5].

2. Aspectos básicos de los sistemas basados en Internet de las cosas (IoT)

El objetivo del IoT es interconectar dispositivos inteligentes basados en sensores que, a la vez que recopilan datos y toman decisiones, sean capaces de estar conectados en una red con otros dispositivos que compartan también información de forma autónoma. Aunque existe un amplio abanico de tipos de dispositivos que pueden conformar una red IoT, todos deben poseer una característica común: autonomía. En otras palabras, deben requerir un bajo consumo energético que les permita operar un largo periodo de tiempo por medio de baterías [6].

Con respecto a la arquitectura de los sistemas IoT, el modelo más empleado es el basado en capas, con 4 niveles diferentes [7], tal y como se muestra en la figura 2. El nivel más bajo lo compone la capa de percepción, donde se encuentran los sensores captando información. Seguidamente encontramos la capa de transmisión, donde reside un dispositivo conocido como enrutador o *gateway*, encargado de recopilar los datos recibidos de diferentes sensores y transmitirlos a la capa de computación. En esta capa, un *software* analiza y procesa los datos para extraer conclusiones que serán mostradas al usuario por la capa de aplicación a través de aplicaciones *software* basadas en navegador web, smartphone, etc.



Figura 2. Arquitectura simplificada IoT.

Uno de los aspectos más importantes de los sistemas IoT es la forma en la que los sensores (ubicados en la capa de percepción) hacen llegar los datos que recopilan al *gateway* (que reside en la capa de transmisión). Aunque existen múltiples tecnologías de comunicación que permiten hacer esto, debemos tener en cuenta que el máximo nivel de aplicabilidad de los sistemas IoT se alcanza cuando no existen restricciones que limiten el entorno donde se pueden desplegar sensores. En otras palabras, optar por una tecnología de comunicación cableada (p.ej. la tecnología Ethernet comúnmente empleada en la oficina o para conectarse al router doméstico) supondría un limitante importante. Es por ello que se asume que la comunicación entre los sensores y el *gateway* debe ser inalámbrica. Para este fin, el uso de tecnologías de comunicación que implementan las conocidas redes de área extensa y bajo consumo (*Low Power Wide Area Networks* – LPWAN) ha alcanzado una destacada popularidad debido a que proporcionan escalabilidad, reducido coste económico, amplio rango de cobertura y elevada eficiencia en el consumo energético.

Una de las aproximaciones LPWAN que está recibiendo una notable atención es LoRaWAN, que define una solución de transmisión a nivel físico y de control de acceso al medio (MAC). Respecto al nivel físico, LoRaWAN define un esquema de modulación propietario llamado LoRa (*Long Range*) que ofrece transmisiones de baja energía, pero de largo alcance. LoRa emplea frecuencias de transmisión ISM, concretamente 868 MHz para Europa. Respecto al nivel MAC de LoRaWAN, este aumenta las capacidades de transmisión de LoRa gracias a la inclusión de esquemas de confirmación de recepción, seguridad en la información transmitida, etc.

Haciendo uso de todas estas características, LoRaWAN permite alcanzar áreas de cobertura mayores a 10 kilómetros en entornos urbanos y rurales, a la vez que requiere un consumo mínimo de energía que permite a los dispositivos operar durante 10 años usando baterías [8]. La limitación de LoRaWAN se encuentra en su baja velocidad de transmisión de datos (máximo 50 kb/s).

3. Mejora de gestión de una FOB por medio de sistemas IoT: casos de uso

Una FOB típicamente no dispone de acceso a una red eléctrica y hace uso de sistemas de generación autónoma. Por lo tanto, es recomendable el uso de equipos que no solo optimicen el consumo de energía, sino que también sean capaces de operar usando baterías durante un tiempo prolongado. Tal y como se ha descrito en la sección 2, los sistemas IoT basados en LoRaWAN satisfacen ambas restricciones. Además, el largo alcance de comunicación permite plantear aplicaciones donde la sensorización (capa de percepción) se realice fuera de los límites de la FOB. Por último, el uso de frecuencias de transmisión de banda libre que, por tanto, no requieren de licencia, son idóneas para el despliegue de redes IoT privadas de uso militar.

A continuación se muestran ejemplos de uso dentro de una FOB donde se demuestra que los sistemas IoT basados en LoRaWAN reportan una mejora en la gestión de las operaciones desarrolladas en una FOB.

3.1 Establecimiento de perímetro de seguridad

Es posible realizar la vigilancia de un perímetro de seguridad, de forma rápida y sencilla, mediante dispositivos equipados con los sensores adecuados. Concretamente, el uso de detectores de sonido y vibración son idóneos para establecer un perímetro de seguridad alrededor de las zonas a defender, incluso fuera de los límites fortificados de la FOB. Al tratarse de pequeños dispositivos desechables, con amplia autonomía y de fácil uso, estos pueden ser desplegados rápidamente por las tropas en las carreteras y zonas de paso más susceptibles de ser usadas por el enemigo. Gracias al amplio rango de cobertura de LoRaWAN, estos dispositivos dan la oportunidad al mando a cargo de la misión de ser consciente de las zonas que le rodean, permitiéndole saber en todo

momento el estado de las zonas de paso y conocer si han sido comprometidas por el enemigo. En caso de que alguna fuerza no amiga se aproxime al perímetro de seguridad, el mando podría tomar conocimiento de este hecho con suficiente antelación de manera que la toma de decisiones será mucho más fácil debido al control de variables que no podrían ser conocidas de otro modo.

3.2 Control de almacenamiento de municiones y explosivos

Según la legislación relativa al almacenamiento de municiones y explosivos, debido a las singulares características de estos consumibles, es necesario mantener unas condiciones ambientales muy específicas en lo relativo a temperatura, humedad o polvo [9]. En una base principal de operaciones se suele disponer de un polvorín, que es un edificio (en ocasiones subterráneo) construido para este fin y que está dotado de sistemas artificiales que aseguren las condiciones óptimas de conservación. Sin embargo, en una FOB no es posible disponer de una construcción específica con estas características, por lo que es esencial realizar una monitorización de la instalación donde se almacenan estos recursos.

La tecnología IoT puede servir de ayuda para llevar a cabo esta tarea de manera flexible y efectiva. Sería posible desarrollar dispositivos equipados con sensores de humedad, temperatura, luz solar o, incluso, capaces de medir la concentración de polvo. Estos dispositivos se instalarían en la dependencia donde se encuentra el armamento, de manera que enviarán de manera periódica datos de las variables ambientales que se deben controlar. La capa de computación procesaría todos los datos recibidos de manera que, en caso de detectar cualquier anomalía que ponga en peligro la conservación del armamento, se informaría al usuario a través de la capa de aplicación.

3.3 Monitorización del estado de las dependencias militares

Tal y como se ha mencionado con anterioridad, una FOB suele depender de equipos autónomos de generación de energía, por lo que el uso eficiente de la misma es un requisito básico. De esta necesidad surge otra aplicación interesante consistente en controlar, a través de un sistema IoT, la eficiencia energética de las dependencias existentes en una FOB. A través de sensores de parámetros ambientales como temperatura, humedad o luminancia, se pueden tener datos en tiempo real y tomar decisiones acerca de la idoneidad de las condiciones que se dan. Además, si combinamos estos datos junto con los obtenidos de otros tipos de sensores como detección de presencia y sonido, la capa de computación puede inferir si la dependencia militar se encuentra vacía y determinar si se hace un uso eficiente desde el punto de vista energético. En última instancia, la capa de aplicación puede mostrar al usuario de manera visual los datos capturados por los sensores en cada dependencia y alertar de situaciones anómalas.

4. Caso de estudio experimental: monitorización de dependencias militares en una FOB

De entre las propuestas descritas en la sección 3, a continuación se presenta el desarrollo de un prototipo experimental de sistema IoT para monitorizar el estado de dependencias militares.

4.1 Configuración del escenario de pruebas

La tabla 1 muestra de forma resumida el *hardware* y *software* empleado para implementar cada una de las 4 capas de las que consta el sistema IoT desarrollado. Como se puede observar, la capa de percepción se ha implementado mediante un microcontrolador Arduino (figura 3a), distribuido por la empresa Seeedstudio [10], que incluye un módulo de comunicaciones LoRaWAN. Este dispositivo se encuentra alimentado por una batería externa de gran capacidad (10.000 mAh) y ha sido equipado con distintos tipos de sensores que proporcionan información ambiental: temperatura, humedad y ruido. Asimismo, se han incluido dos sensores (distancia y presencia) cuyo valor combinado indicará si existe personal militar dentro de la dependencia.

CAPA	HARDWARE	SOFTWARE
Percepción	Arduino + Base Shield V2	
	Módulo comunicaciones LoRaWAN RHF76-052AM	
	Xiaomi Mi PowerBank (10000 mAh)	
	Sensor de luminancia (APDS-9002)	
	Sensor temperatura y humedad (DHT22)	
	Sensor de ruido (LM2904)	
	Sensor de distancia (VL53L0X)	Arduino IDE
	Sensor de presencia (AK9753)	
Transmisión	Raspberry Pi 3	PuTTY
	PRI 2 Bridge RHF4T002	LorIoT Gateway Software (rhf1257)
	Modulo Gateway RHF0M301-868	
Computación	Servidor red en la nube	LorIoT.io
Aplicación	Servidor de aplicaciones en la nube	AllThingsTalk

Tabla 1. Resumen del *hardware* y *software* empleado, desglosado por capas de la arquitectura.

En la capa de transmisión se ha implementado un *gateway* LoRaWAN de bajo coste haciendo uso de un dispositivo Raspberry Pi 3 (figura 3b) que incluye un módulo de comunicaciones LoRaWAN modelo RisingHF (RHF). Por último, las capas de computación y aplicación se han desarrollado por medio de conocidos proveedores de servicio en la nube en el ámbito del IoT: LorIoT.io y AllThingsTalk.

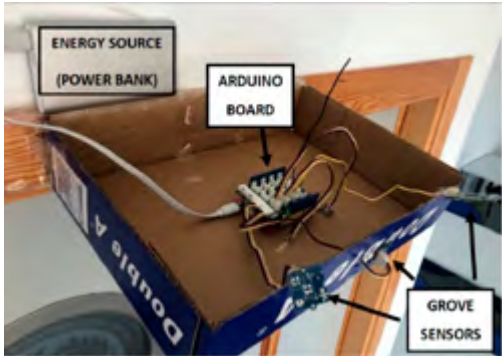


Figura 3. (a) Hardware capa percepción.



Figura 3. (b) Hardware capa transmisión.

Para conseguir un correcto funcionamiento, fue necesario realizar configuraciones *software* en todas las capas. En primer lugar, se desarrolló un programa (empleando Arduino IDE) con el objetivo de dotar de comportamiento a nuestro dispositivo sensorizado. De manera resumida, el programa inicia y configura sensores. Seguidamente se inicializa el módulo de comunicaciones LoRa, estableciendo una configuración conservadora que ofrece una transmisión de datos fiable pero lenta (entorno a 250 bps). Seguidamente, el programa comienza una ejecución iterativa (por defecto, cada 5 segundos) que toma muestra de datos de todos los sensores, crea un mensaje y lo envía por el canal inalámbrico LoRa.

En el *gateway* se desplegó la aplicación *LorIoT Gateway Software* encargada de recibir datos por la interfaz LoRa y reenviarlos al servidor ubicado en la capa de computación. Por simplicidad, en nuestro prototipo experimental no se realiza un procesamiento de datos en la capa de computación, sino que se reenvían directamente a la capa de aplicación. Por último, haciendo uso de las herramientas de configuración disponibles en AllThingsTalk, se desarrolló un panel de monitorización (ver figura 4) que muestra los datos recopilados por los sensores en un entorno amigable al usuario mediante gráficas, histogramas, etc.

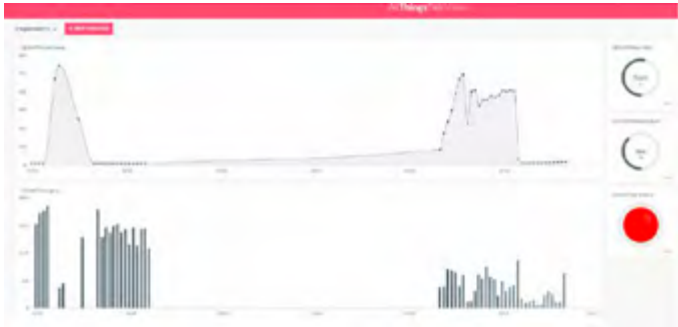


Figura 4. Panel de visualización AllThingsTalk (capa de aplicación).

Sin pérdida de generalidad, el sistema IoT ha sido testeado en un escenario que intenta replicar las condiciones que se darían en una dependencia militar ubicada en una FOB. Se ha desplegado en un aula docente ubicada en un pabellón de la Academia General del Aire. El uso que se hace de esta aula ofrece un entorno ideal para testear la capacidad de monitorización de nuestro sistema.

El dispositivo sensorizado (Arduino) se instaló en la zona de acceso al aula docente (ver figura 3a) ya que es la ubicación ideal para que los detectores de presencia operen correctamente. El *gateway* (Raspberry Pi) se instaló en otro edificio diferente ubicado a unos 600 metros aproximadamente.

4.2 Resultados

El sistema IoT estuvo operando durante una semana de manera autónoma. Además, los datos reportados por los sensores han demostrado ser lo suficientemente precisos como para detectar cambios en la actividad desarrollada dentro del aula. Como ejemplo, las figuras 5 a 7 muestran lecturas relevantes de los sensores ambientales. En la figura 5 se observa una clara disminución en la intensidad lumínica reportada por el sensor, lo que significa que no hay luces encendidas ni persianas abiertas, un hecho compatible con el final de una clase. Lo mismo sucede en la figura 6, donde el sensor de ruido reporta valores comprendidos en un rango (líneas horizontales rojas) que representan ausencia de silencio. Por último, la figura 7 muestra la evolución ascendente de los valores de temperatura y humedad, lo que puede interpretarse como el uso de equipos de climatización.

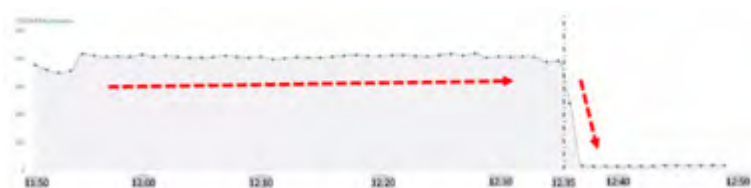


Figura 5. Lecturas proporcionadas por sensor de luminancia.



Figura 6. Lecturas proporcionadas por sensor de ruido.



Figura 7. Lecturas proporcionadas por sensor de temperatura y humedad.

5. Conclusiones

En este trabajo se ha presentado la tecnología IoT basada en LoRaWAN como opción para ayudar a la gestión de las FOB. Para ello, se han descrito las características fundamentales de IoT y su aplicabilidad en ejemplos de uso en una FOB. Además, se ha mostrado la implementación de un prototipo experimental para controlar parámetros ambientales y de acceso a instalaciones. La validación de este prototipo fue llevada a cabo en la Academia General del Aire, donde se monitorizó un aula a una distancia de 600 metros. Se han mostrado resultados de los parámetros de luminancia, ruido, temperatura y humedad que reflejan el correcto funcionamiento del sistema.

Agradecimientos

Los autores agradecen el apoyo económico recibido por el Centro Universitario de la Defensa de San Javier (CUD-AGA) para la realización de este trabajo de investigación.

Referencias

1. Al-Fuqaha, A., Guizani, M., Mohammadi, M., Aledhari, M., & Ayyasah, M. (2015). «Internet of Things: A Survey on Enabling Technologies, Protocols, and Applications». IEEE Communication Surveys & Tutorials, vol 17, no. 4, pp. 2347-2376. doi: 10.1109/COMST.2015.2444095.
2. Technical activities of the STO. (s. f.). North Atlantic Treaty Organization - Science and Technology Organization. Recuperado 24 de septiembre de 2020, de <https://www.sto.nato.int/Pages/activitieslisting.aspx>.
3. Base Aérea Conectada Sostenible Inteligente. (s. f.). Ejército del Aire. Recuperado 24 de septiembre de 2020, de <https://ejercitodelaire.defensa.gob.es/EA/bacsi/>.

4. Forward operating base. (s. f.). Wikipedia - The Free Encyclopedia. Recuperado 24 de septiembre de 2020, de https://en.wikipedia.org/wiki/Forward_operating_base.
5. Overview Camp Marmal. (s. f.). Wikimedia Commons. Recuperado 24 de septiembre de 2020, de <https://commons.wikimedia.org/w/index.php?curid=834164>.
6. G. Aceto, V. Perisco and A. Pescapé, «A Survey on Information and Communication Technologies for Industry 4.0: State-of-the-Art, Taxonomies, Perspectives, and Challenges», IEEE Communication Surveys & Tutorials, vol. 21, no. 4, pp. 3470-3483, 2019. doi: 10.1109/COMST.2019.2938259.
7. A. J. Trappey, C. V. Trappey, U. H. Govindarajan, A. C. Chuang and J. J. Sun, A review of essential standards and patent landscapes for the Internet of Things: A key enabler for Industry 4.0, Advanced Engineering Informatics, no. 33, pp. 208-229, 2017.
8. F. Adelantado, X. Vilajosana, P. Tuset-Peiro, B. Martínez, J. Melia-Segui and T. Watteyne, «Understanding the Limits of LoRaWAN», in IEEE Communications Magazine, vol. 55, no. 9, pp. 34-40, Sept. 2017. doi: 10.1109/MCOM.2017.1600613.
9. Real Decreto 130/2017, de 24 de febrero, por el que se aprueba el Reglamento de Explosivos. *Boletín Oficial del Estado*, 4 de marzo de 2017, núm. 54. Disponible: https://www.boe.es/diario_boe/txt.php?id=BOE-A-2017-2313.
10. LoRa/LoRaWAN Gateway Kit. (s. f.). Seeed Wiki. Recuperado 24 de septiembre de 2020, de https://wiki.seeedstudio.com/LoRa_LoRaWan_Gateway_Kit/.

Navegación visual eficiente aplicada a sistemas aéreos no tripulados en entornos de GNSS denegado

Echevarría Corrales, Fidel; Pulido Fernández, José María; de Frutos Carro, Miguel Ángel*

UAV NAVIGATION S.L. Av. Pirineos 7 B-11, San Sebastián de los Reyes, 28703.
www.uavnavigation.com

*(POC) Miguel Ángel de Frutos. mfrutos@uavnavigation.com

Resumen: la capacidad de estimar la posición que ocupa la plataforma, es esencial en los sistemas de navegación de aeronaves no tripuladas, para ejecutar de manera efectiva las tareas de guiado y navegación. El posicionamiento mediante sistemas no autónomos, como el GNSS, tiene algunas limitaciones en cuanto a disponibilidad y continuidad de servicio. UAV Navigation está desarrollando un sistema de posicionamiento, absoluto y autónomo, basado en el empleo de técnicas de visión artificial. El sistema utiliza una novedosa técnica de posicionamiento absoluto que podría utilizarse como alternativa al GNSS, o como elemento de supervisión y detección de información maliciosa, en vehículos aéreos no tripulados. La técnica propuesta en este trabajo necesita un mapa guardado en la memoria interna de la aeronave accesible durante el vuelo. Este mapa debe corresponder a una imagen aérea de la zona de vuelo, ya bien sea descargada de un servicio de imágenes satelitales u obtenidas por el mismo, u otras aeronaves, en misiones anteriores sobre la zona de operación. La técnica se basa en realizar una búsqueda sistemática de la imagen tomada por la cámara embarcada, que siempre apunta en la dirección de la vertical hacia el terreno, dentro del mapa interno guardado en memoria. Para conseguir una operación de búsqueda eficiente a lo largo del mismo, se emplea un filtro de partículas para delimitar la región de búsqueda más probable mediante la propagación del modelo dinámico de la aeronave.

Palabras clave: correlación cruzada, filtro de partículas, GNC, navegación visual, *template matching*, UAV.

1. Introducción

La determinación de la posición que ocupa la aeronave en el espacio, de una manera precisa y fiable, es clave para llevar a cabo múltiples lógicas, como la navegación, el guiado, o incluso los sistemas anti-colisiones, que permitirán una ejecución segura y eficiente de la misión. Por estos motivos es crucial disponer de la capacidad de estimar dicha posición mediante el empleo de múltiples sensores u observadores embarcados, cuyas medidas individuales se cruzan entre sí, para dar como resultado una medida, o estimación, más fiable y precisa en el tiempo. Una primera clasificación de estos sensores se puede hacer atendiendo a si estos necesitan de una infraestructura externa (e.g. radioayuda terrestres o mediante constelación satelital) para poder operar o, por el contrario, son sistemas autónomos capaces de entregar información precisa sin la ayuda de dichos elementos externos [1]. A su vez, podemos establecer nuevamente dos categorías atendiendo, esta vez, al tipo de información que son capaces de entregar [2]. En este caso distinguiremos entre métodos de posicionamientos relativos y métodos absolutos.

1.1 Métodos de posicionamiento relativo

Los métodos relativos proporcionan incrementos de posición respecto a una referencia conocida. Típicamente el uso de este tipo de técnicas tiene asociada una acumulación de error en cada iteración que provoca un error absoluto en posición, creciente con el tiempo. Las técnicas de posicionamiento relativo suelen ser eficientes, lo que permite su ejecución a una frecuencia elevada. Entre las técnicas de este tipo más comunes, de aplicabilidad al dominio del problema aquí tratado, se encuentran la propagación de la posición mediante la integración de medidas inerciales, integración de velocidades aerodinámicas verdaderas o la odometría visual a través de cámaras. Estas técnicas, si bien en el corto periodo son capaces de alcanzar niveles de precisión destacables, suelen terminar acumulando un error significativo en su estimación de posición absoluta como resultado de la acumulación iterativa de varias fuentes de error en el proceso.

1.2 Métodos de posicionamiento absoluto

Estos métodos proporcionan posiciones absolutas referenciadas a un sistema de coordenadas globales, con error constante y dependiente de las condiciones del entorno. Estas técnicas no suelen ser eficientes, por lo que es práctica común combinar técnicas de posicionamiento relativo a alta frecuencia con técnicas de posicionamiento absoluto a baja frecuencia, consiguiendo compensar el error acumulado por las relativas en el largo periodo [3]. Entre las técnicas de posicionamiento absoluto destaca el empleo de sistemas de navegación global por satélite, conocidos por sus siglas en inglés como GNSS.

1.3 Sistema de posicionamiento absoluto basado en visión

UAV Navigation está desarrollando un sistema de posicionamiento autónomo basado en visión. El sistema utiliza una novedosa técnica de posicionamiento absoluto que podría utilizarse como alternativa al GNSS, o como elemento de supervisión y detección de información maliciosa, en vehículos aéreos. Esta técnica de posicionamiento necesita un mapa guardado en la memoria interna de la aeronave. Este mapa puede corresponder a una imagen aérea de la zona de vuelo, bien descargada de un servicio de imágenes satelitales u obtenidas por el mismo u otro avión en misiones anteriores sobre la zona de operación. La técnica se basa en realizar una búsqueda de la imagen tomada por la cámara del avión, que siempre apunta en la dirección de la vertical hacia el terreno, dentro del mapa interno guardado en memoria. Este tipo de técnicas de «búsqueda de similitudes entre pares» se conoce como *template matching* [4]. Debido a que los paisajes y los patrones del suelo cambian con el tiempo, es necesario encontrar una técnica de búsqueda de plantillas que admita diferencias en las imágenes, y que sea robusta para poder detectar patrones con un ratio de falsos positivos bajo.

2. Materiales y métodos

La técnica de posicionamiento absoluto se ha desarrollado combinando una técnica de correlación cruzada normalizada entre dos imágenes con un filtro de partículas que tiene en cuenta un modelo dinámico genérico de aeronave. La correlación cruzada es una medida de la similitud entre dos series en función del desplazamiento relativo entre ellas. El *template matching* mediante correlación cruzada normalizada consiste en realizar una operación de convolución, deslizando la imagen plantilla a lo largo del mapa y aplicando una operación específica en cada una de estas posiciones [5]. De esta manera se genera una tercera imagen que corresponde al resultado de la convolución en cada una de las posiciones posibles durante el deslizamiento. La operación que se aplica en cada posición se define en la ecuación 1.

$$R(x, y) = \frac{\sum_{x', y'} (T(x', y') - I(x + x', y + y'))^2}{\sqrt{\sum_{x', y'} T(x', y')^2 \cdot \sum_{x', y'} I(x + x', y + y')^2}} \quad (eq. 1)$$

Donde las variables x e y corresponden a las coordenadas de los píxeles de cada convolución, x' e y' a las coordenadas de los píxeles de la operación anidada dentro de cada convolución, a los valores de los píxeles de la imagen de la cámara, e I a los valores de los píxeles de la imagen del mapa. Mediante esta operación es posible obtener correspondencia de patrones en el mapa [6]. Como ejemplo ilustrativo, se ha aplicado la operación para la búsqueda de la zona donde se ubica la Torre Eiffel dentro de un mapa aéreo de París. El ejemplo puede observarse en la figura 1.

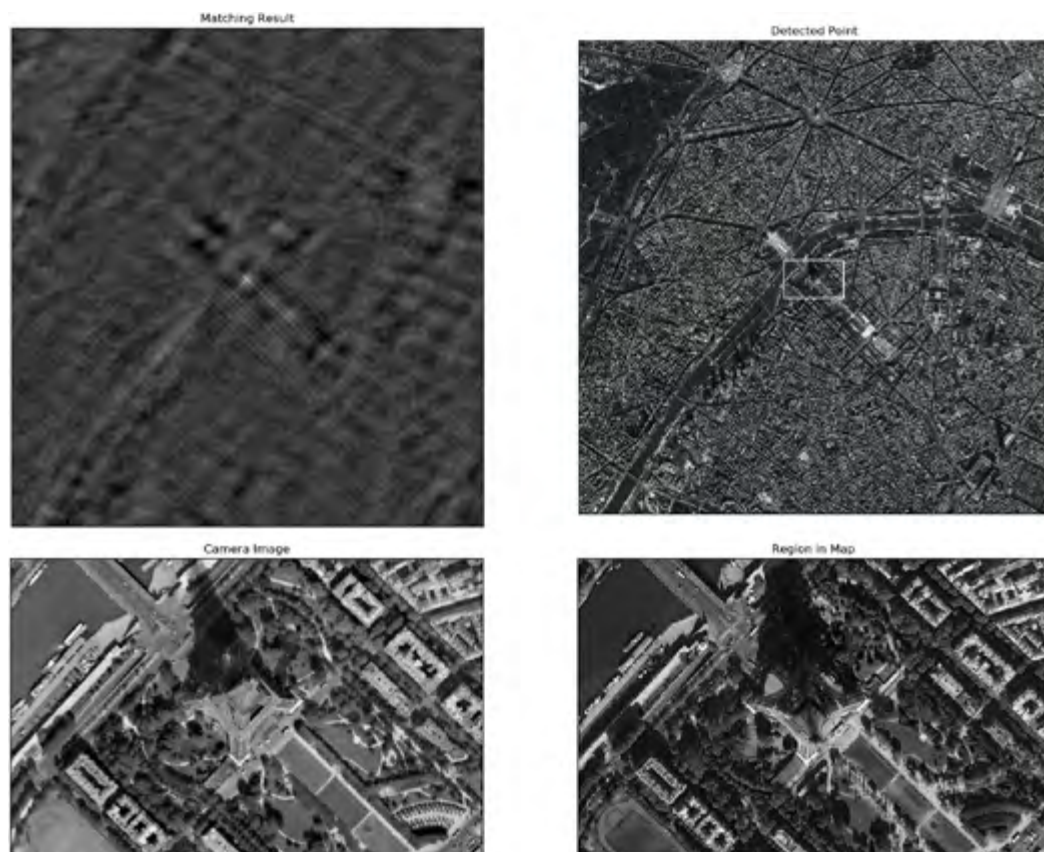


Figura 1. Correspondencia de patrones mediante correlación cruzada normalizada en un mapa de vista aérea de París. La imagen superior derecha corresponde con el mapa interno guardado en memoria. La imagen inferior izquierda corresponde a la imagen detectada por la cámara, que se pretende buscar en el mapa interno. La región asociada correspondiente en el mapa interno se observa en la imagen inferior derecha. Por último, la imagen superior izquierda corresponde al resultado de la operación de convolución mediante correlación cruzada normalizada. Las zonas más claras de esta imagen corresponden a las zonas donde la correspondencia tiene alta probabilidad.

En la imagen resultante de la operación puede observarse que se detecta el punto de correspondencia. También se observan algunas líneas claras diagonales que parten de dicho punto central. Estas líneas se originan debido a que la orientación de las calles hace que la correspondencia sea probable cuando la imagen se desliza en paralelo a ellas, ya que la intensidad de los píxeles en los edificios y en las calles tiende a ser uniforme.

Esta técnica demuestra ser robusta para la búsqueda de instancias. Sin embargo, existe la necesidad de que en la imagen haya presencia de patrones significativos que permitan el reconocimiento de la zona. Es decir, en un paisaje demasiado uniforme o carente de patrones únicos, como un paisaje desértico o marítimo, esta técnica no es

efectiva para el posicionamiento absoluto. Por otro lado, la operación de búsqueda global a lo largo del mapa no es una estrategia eficiente, y consume demasiados recursos, siendo esto una limitación tanto en tiempo de ejecución requerido como en capacidad computacional embarcada necesaria.

Debido a que tenemos una idea aproximada de la dinámica de la aeronave sobre la que el sistema va embarcado, es posible restringir la región de búsqueda a zonas cercanas a la propagación de dicho modelo dinámico, aumentando así la eficiencia de la operación de una manera notable. Esta técnica se ha implementado mediante un filtro de partículas.

El filtro de partículas es un estimador de estados muy utilizado en el área de procesamiento de señales [7]. Los problemas de filtrado donde suele utilizarse este tipo de filtros se caracterizan por tener observaciones parciales además de perturbaciones aleatorias en los sensores y en el propio sistema dinámico. El objetivo es calcular las distribuciones de probabilidad posterior de los estados de un proceso de Markov dadas las observaciones parciales y ruidosas. El filtro consta de un número de partículas que se propagan utilizando el modelo dinámico. Además, en cada propagación, se aplica el ruido inherente al sistema y a los sensores a cada partícula y medida. Las diferentes etapas de la lógica completa se ejecutan secuencialmente en un bucle. Estas se describen a continuación.

- Lógica de extracción de «picos»:
 - Se selecciona un área cuadrada alrededor de la posición estimada por la lógica. Se recorta una subimagen con esta geometría del mapa almacenado en memoria. El marco de trabajo se orienta utilizando el rumbo real estimado, de manera que sus lados permanezcan paralelos a los ejes e centrados en los ejes de referencia del avión.
 - La imagen capturada por la cámara situada en la parte inferior del avión, corresponde a un cuadrado recortado similar al del paso anterior, pero con lados más cortos.
 - Se aplica la técnica de convolución con correlación cruzada normalizada usando como plantilla la imagen capturada por la cámara embarcada y como mapa la imagen cuadrada obtenida del mapa interno. Se obtiene una tercera imagen resultante de la convolución.
 - En la imagen de correlación obtenida en el paso anterior se buscan los máximos locales, tras la aplicación de un factor de «suavizado» para no detectar máximos locales muy cercanos. Este suavizado se realiza mediante un filtro de desenfoque gaussiano [8]. Cada uno de los máximos locales obtenidos se considera un «pico».

- Lógica del filtro de partículas:
 - Durante la inicialización, todas las partículas están localizadas en la posición estimada de la aeronave.
 - Se propagan las partículas con los incrementos estimados de posición y de orientación para la iteración en curso. Estos incrementos pueden calcularse con una técnica de posicionamiento relativo, como odometría visual o propagación inercial.
 - Todas las partículas se «pesan» teniendo en cuenta su distancia al «pico» más cercano. Las partículas más cercanas a picos locales tendrán un factor de peso mayor que las lejanas.
 - En el siguiente paso se realiza un remuestreo de las partículas. Este proceso consiste en generar nuevas partículas a partir del remuestreo aleatorio de las anteriores. La probabilidad de que una partícula sea elegida es proporcional al peso de cada partícula, de manera que las partículas con mayor peso tienen mayor probabilidad de ser elegidas. Este proceso provoca una tendencia a que solo sobrevivan aquellas partículas que se encuentran cercanas a máximos locales.
 - La posición estimada se actualiza con la media de posición resultante de todas las partículas, de forma proporcional a su peso.

3. Resultados y discusión

Con el fin de evaluar el comportamiento de la lógica desarrollada en un caso realista, se ha simulado un vuelo ejecutando una trayectoria circular alrededor de un aeródromo de la Comunidad de Madrid. Se han obtenido dos imágenes satelitales de la zona utilizando la API de Microsoft Bing Maps. Las imágenes corresponden a diferentes épocas del año, por lo que existen numerosas diferencias estacionales en las características observadas en el terreno. Se ha realizado un vídeo a partir de una secuencia de fotogramas generados, que han sido diseñados para facilitar el análisis del comportamiento de la lógica durante el vuelo. Un ejemplo de fotograma puede observarse en la figura 2. Analizando la secuencia de imágenes, se aprecia que existen numerosos máximos locales que aparecen y desaparecen durante la trayectoria. Sin embargo, solo un máximo local tiende a permanecer estable, y este corresponde con la solución verdadera. Es por esto que las partículas que se acercan a máximos locales inestables tienden a desaparecer en sucesivas iteraciones, y únicamente las que se acercan al máximo local estable sobreviven, obteniendo una posición estimada robusta.

El algoritmo consigue una estimación perfecta a lo largo de casi toda la trayectoria. Únicamente existe una zona con escasez de patrones distinguibles, lo que ocasiona que

las partículas se alejen ligeramente de la posición real, pero pronto vuelven a converger a la solución correcta en cuanto se vuelven a detectar patrones distinguibles en el terreno. La velocidad de ejecución del algoritmo es proporcional al número de máximos locales detectados. Durante algunas secciones de la trayectoria, como en caso de sobrevolar un área residencial, se detecta un gran número de máximos locales debido a la estructura de rejilla resultante de las calles que conforman la zona urbanizada. Esto puede ralentizar el algoritmo, pero se puede solventar estableciendo un número máximo de picos para la detección, asegurando un tiempo máximo de ejecución por iteración en todos los casos.

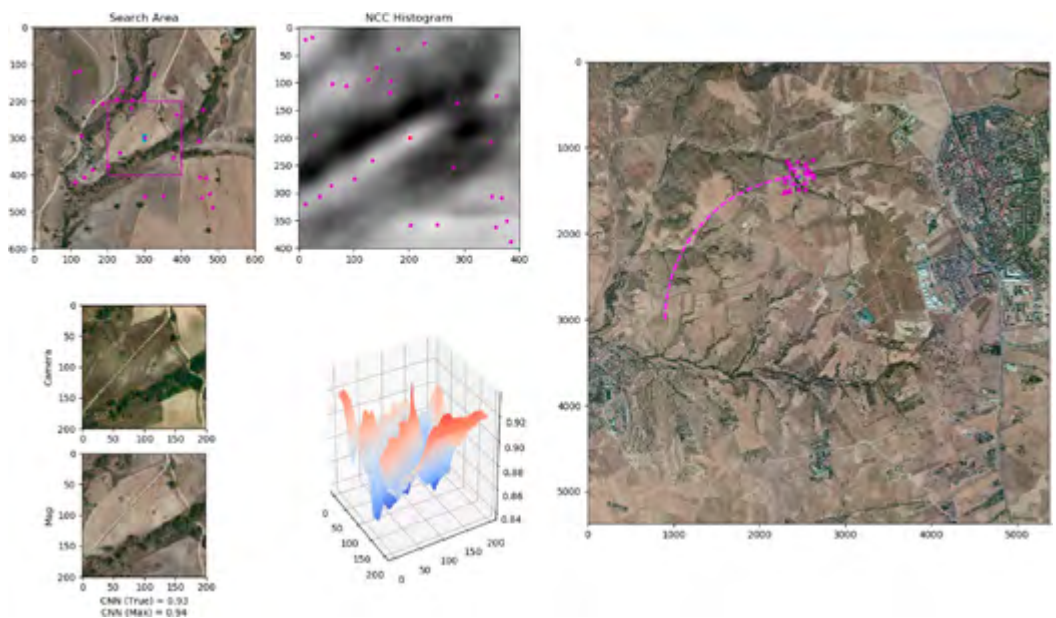


Figura 2. Fotograma resultante de la simulación. La imagen superior izquierda corresponde al cuadrado recortado que representa el mapa en esta iteración. A su derecha se observa la imagen resultante de la convolución con correlación cruzada normalizada. Las zonas claras representan las zonas con mayor probabilidad de correspondencia. El gráfico tridimensional corresponde a la misma imagen de correlación representada con una tercera dimensión en lugar de una intensidad de píxel. Las dos imágenes pequeñas situadas en la sección inferior izquierda corresponden a la imagen tomada por la cámara del avión (superior) y a la correspondiente zona en el mapa interno (inferior). Por último, la imagen grande de la sección derecha corresponde al mapa completo, en el que se puede observar la trayectoria real (magenta) y la estimada (cyan). En todas las imágenes los máximos locales detectados se representan con puntos de color magenta, mientras que las partículas se representan con puntos de color cyan. (Enlace al video con la ejecución en tiempo real: <https://www.youtube.com/watch?v=t3m3EdWZd88&feature=youtu.be>).

4. Conclusiones

En este trabajo se ha presentado una técnica de posicionamiento autónoma, porque solo depende de información embarcada, y absoluta, al ser capaz de proporcionar información suficiente para posicionar la aeronave en un marco de referencia ejes tierra,

basada en la búsqueda eficiente de patrones distinguibles en imágenes aéreas, de gran potencial para la navegación en entornos de GNSS denegado o como sistema supervisor de este. La innovación aportada mediante la utilización de un filtro de partículas, que tiene en cuenta un modelo dinámico de la aeronave, aumenta en gran medida la robustez y eficiencia de la técnica de búsqueda de similitudes por correlación cruzada normalizada, ya que por lo general no tiene porqué cumplirse que el máximo global en el histograma de convolución sea la posición verdadera. El filtro de partículas utiliza a su favor el hecho de que el máximo local que indica la solución siempre suele ser más estable en el tiempo que los demás máximos locales transitorios. Además, esta técnica permite utilizar otros métodos de posicionamiento relativo en el corto periodo, como propagación inercial, odometría visual, o la fusión de la dinámica de la plataforma con otros estados. Los requisitos de memoria necesarios para almacenar el mapa interno y la capacidad computacional necesaria para ejecutar el algoritmo a una frecuencia suficiente, no son excesivos, lo que permite su integración en sistemas embarcables para su operación real.

Agradecimientos

Los autores agradecen la confianza, apoyo y motivación para seguir investigando en el área de los Sistemas de Navegación Visual, recibido por el GRUPO OESÍA. En la actualidad estamos analizando las posibilidades de establecer planes de desarrollo conjuntos, basados en la complementariedad de nuestras capacidades tecnológicas, para integrarlas en soluciones técnicas para la navegación avanzada de sistemas aéreos en entornos difíciles, que sean de interés para la defensa.

Referencias

1. R. R. Murphy, «Designing a sensor suite», in *Introduction to AI Robotics*, 1st ed. Cambridge, MA, USA: The MIT Press, 2000, ch. 6, sec. 3, pp. 202-206.
2. J. González Jiménez and A. Ollero Baturone, «Estimación de la posición de un robot móvil», [position estimation in a mobile robot] Dpt. Systems and Automatic Engineering, Malaga and Seville Universities, Spain, Sept. 2020.
3. J. Borenstein and L. Feng, «Gyrodometry: A new method for combining data from gyros and odometry in mobile robots», in *Proc. 1996 IEEE Int. Conf. Robotics and Automation (ICRA)*, Minneapolis, USA, Apr. 22-28, 1996, pp. 423-428. [Online]. Available: <http://www-personal.umich.edu/~johannb/Papers/paper63.pdf>.
4. R. Szeliski, «Instance recognition», in *Computer Vision: Algorithms and Applications*, 1st ed. London, UK: Springer, 2011, ch. 14, sec. 3, pp. 602-611.
5. D. Buniatyan, T. Macrina, D. Ih, J. Zung and H. S. Seung, «Deep learning improves template matching by normalized cross correlation», Neuroscience Institute

- and Computer Science Dept., Princeton University, Princeton, NJ, May 2017, arXiv:1705.08593. [Online]. Available: <https://arxiv.org/pdf/1705.08593.pdf>.
6. M. B. Hisham, S. N. Yaakob, R. A. A. Raof, A. B. A. Nazren and N. M. Wafi, «Template matching using sum of squared difference and normalized cross correlation», in *2015 IEEE Stud. Conf. Res. and Devel. (SCORED)*, Kuala Lumpur, Malaysia, 2015, pp. 100-104. doi: 10.1109/SCORED.2015.7449303.
 7. D. Simon, «The particle filter», in *Optimal State Estimation: Kalman, H Infinity, and Nonlinear Approaches*, 1st ed. Hoboken, NJ, USA: Wiley-Interscience, 2006, ch. 15, pp. 461-485.
 8. R. Chandel and G. Gupta, «Image filtering algorithms and techniques: A review», *Int. J. Adv. Res. Comput. Sci. Softw. Eng.*, vol. 3, no. 10, pp. 198-202, Oct. 2013.

Resultados de la primera misión española de observación IOD de prestaciones submétricas en la Estación Espacial Internacional

Guzmán, Rafael^{1,*}; Conde, Aitor²; Ocerin, Eider³; Massimiani, Marta⁴; Davis, Stuart⁵; Buehler, David⁶

¹ SATLANTIS MICROSATS S.L. Edif. SEDE 2.º planta, Parque Científico, Campus UPV-EHU, 48940, Leioa (Vizcaya). guzman@satlantis.com

² SATLANTIS MICROSATS S.L. conde@satlantis.com

³ SATLANTIS MICROSATS S.L. ocerin@satlantis.com

⁴ SATLANTIS MICROSATS S.L. massimiani@satlantis.com

⁵ SATLANTIS MICROSATS S.L. davis@satlantis.com

⁶ SATLANTIS MICROSATS S.L. buhler@satlantis.com

* Rafael Guzmán; guzman@satlantis.com

Resumen: SATLANTIS es una pyme vasca que desarrolla cámaras ópticas de muy alta prestación para la observación de la Tierra. Su tecnología central – iSIM-170 – ha sido validada en órbita este año a través de una misión IOD a la Estación Espacial Internacional (ISS). Con esta misión, SATLANTIS ha sido la primera empresa no japonesa en instalar una carga útil en la plataforma externa i-SEEP de la ISS y ha conseguido demostrar en órbita sus capacidades submétricas por medio de las 60.000 imágenes adquiridas.

Esta validación en órbita posiciona SATLANTIS entre las mejores empresas de cámaras ópticas del mundo y demuestra el valor añadido que la tecnología iSIM puede proveer a sectores de monitorización y vigilancia, tanto industrial como de defensa y seguridad.

Se presentan en este artículo la tecnología iSIM, el proyecto de vuelo hasta el lanzamiento y los resultados de la campaña de validación en órbita, junto con las lecciones aprendidas y el porfolio de aplicaciones en varios sectores de interés.

Palabras clave: geo-inteligencia, miniaturización, observación de la Tierra, resolución submétrica, vigilancia.

1. Introducción

SATLANTIS MICROSATS S.L. [1] es una pyme tecnológica con sede en el País Vasco que diseña, fabrica y comercializa cargas útiles de altas prestaciones para pequeños satélites y *CubeSats* de observación terrestre.

Su tecnología central es iSIM (*integrated Standard Imager for Microsatellites*), una familia de varios modelos de telescopios binoculares de alta resolución óptica resultado de la combinación de diferentes tecnologías: óptica, mecánica, electrónica y algoritmos de procesamiento. La tecnología desarrollada por SATLANTIS es disruptiva dentro del mercado de cámaras de observación de la Tierra dado que ofrece prestaciones superiores a las de instrumentos convencionales en un formato más compacto y ligero. Las cámaras iSIM son particularmente adecuadas para monitorizar estructuras lineales como costas, oleoductos, fronteras y ofrecen una resolución submétrica en las bandas VNIR (450-900nm) y mejor que 3m en SWIR (900-1450nm).

La cámara iSIM-170 ha sido validada en órbita en una misión IOD (*In Orbit Demonstration*) a la ISS, que fue lanzada desde el Centro Espacial de Tanegashima en Japón, con la misión HTV-9, el pasado 20 de mayo 2020 (hora española). La cámara fue posteriormente integrada en la plataforma externa (i-SEEP) del módulo japonés Kibo de la Estación Espacial Internacional (ISS). Esta ha sido la primera vez que una carga útil extranjera se instalara en la plataforma japonesa i-SEEP. JAXA [2], la agencia espacial japonesa, ha coordinado el proyecto y se ha encargado de ponerla en órbita. Para ello, la cámara iSIM-170 fue sometida a una estricta campaña de ensayos, revisiones y evaluaciones por parte de las agencias JAXA y NASA terminada en diciembre 2019 con la aprobación final para la Estación Espacial.

Además de esta misión, SATLANTIS está trabajando en dos misiones adicionales, una a través del Departamento de Defensa estadounidense y otro con la Agencia Espacial Europea (ESA), cuyos lanzamientos al espacio están previstos para 2021 y 2022. Además, la empresa está involucrada en varios proyectos de investigación y desarrollo, incluido un estudio de viabilidad de una carga útil para la Agencia de Defensa Europea (EDA).

2. Materiales y métodos

2.1 La tecnología iSIM

Los aspectos innovadores de la tecnología que desarrolla SATLANTIS han sido resultado de la transferencia exitosa de tecnología y técnicas astrofísicas, al sector de observación de la Tierra. Los aspectos innovadores pueden resumirse en los siguientes puntos.

Óptica: sistema binocular que consiste en dos canales idénticos, cuyo diseño está basado en la configuración óptica Maksutov-Cassegrain. Esta configuración se

usa tradicionalmente en la astrofísica para construir telescopios compactos de alta resolución. SATLANTIS ha adaptado esta configuración para conseguir un diseño novedoso para satélites de observación de la Tierra pequeños, es decir, por debajo de los 100 kg. Esto permite obtener imágenes al límite de difracción en un rango muy amplio de longitudes de onda (de 450 a 900 nanómetros). Además, el diseño binocular proporciona mayor flexibilidad que un diseño monocular, pues permite doblar el campo de visión, el número de bandas espectrales, o usar distintos instrumentos en cada canal.

Mecánica: una estructura de aleación de alta precisión, ligera y robusta con barras de fibra de carbono, que proporciona un soporte idóneo al sistema óptico, y que facilita la integración y montaje de los elementos ópticos.

Detectores: unos detectores matriciales CMOS montados al final de cada canal óptico que permiten observar en múltiples bandas espectrales sin pérdida de resolución espacial y permiten tomas de imágenes y lecturas de alta velocidad en comparación con los detectores tradicionales.

Electrónica: un ordenador abordo reconfigurable y de alta precisión basado en procesadores de última generación para llevar a cabo el procesamiento de imágenes en tiempo real.

Super-resolución: algoritmos de procesamiento de imagen que permiten una mejora de la calidad óptica nativa de las imágenes de hasta un factor 2 en resolución espacial y reducen la cantidad de datos a descargar a Tierra en hasta un factor 8.



Figura 1. Modelo 3D de la optomecánica de iSIM-170 (izquierda) y arquitectura de la electrónica de iSIM-170 para el IOD (derecha).

La cámara iSIM170-IOD, validada este año en la ISS, tiene las siguientes especificaciones técnicas. Hay que destacar que estos valores se basan en el caso específico del IOD, en particular la altitud de la ISS (400-500 km), y en algunas limitaciones que han impuesto un cambio de diseño con un resultante aumento de masa y volumen.

Especificaciones de iSIM170-IOD	
GSD (m)	<1
Swath (km)	6
Rango espectral (nm)	450 – 900
Bandas espectrales	1 (PAN)
Masa (kg)	16.6
Volumen (mm ³)	595.0x476.8x345.5

Tabla 1. Especificaciones principales de iSIM170-IOD.

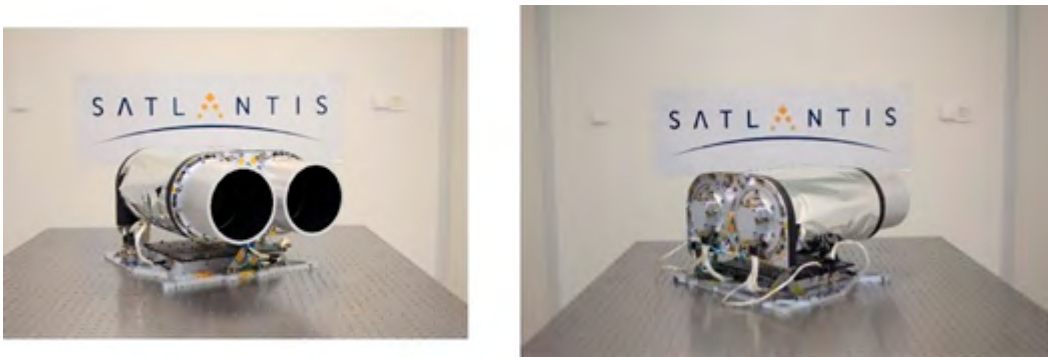


Figura 2. Modelo de vuelo de iSIM170-IOD.

2.2 La demostración en órbita (IOD)

La preparación concreta para la misión de demostración en órbita (IOD) comenzó en enero de 2019 con la firma de un contrato de lanzamiento con la empresa japonesa Space BD [3], contratista de la agencia espacial japonesa JAXA para la utilización del módulo japonés Kibo de la Estación Espacial Internacional (ISS).

SATLANTIS superó todas las fases de revisión oficiales de NASA y JAXA, entregando el modelo de vuelo de iSIM-170 a Japón para el lanzamiento a tiempo, en un proceso que duró solamente un año. El éxito fue fruto de una estrecha colaboración entre SATLANTIS y Space BD, para cumplir los requisitos de funcionamiento, lanzamiento y operaciones de la carga útil.

El lanzamiento tuvo lugar el pasado día 20 de mayo de 2020 desde el Centro Espacial de Tanegashima en Japón (hora española), utilizando el lanzador H-IIB. Cinco días más tarde, el 25 de mayo de 2020, a bordo del vector cargo HTV-9, iSIM-170 llegó a la ISS. El comandante de la Estación, el astronauta de NASA Chris Cassidy, fue el encargado de integrar la cámara en la plataforma i-SEEP del módulo japonés Kibo el día 9 de junio

de 2020, finalizando su instalación por medio de un brazo robótico el día 11 de junio de 2020 en las instalaciones externas de la ISS.

Desde ese momento, la cámara ha estado operando satisfactoriamente desde la ISS. Las operaciones, con duración prevista de tres meses, han estado condicionadas por algunas limitaciones impuestas por la Estación, como, por ejemplo:

- el apuntado fijo de la ISS a Tierra, sin posibilidad de mover la cámara;
- capacidad máxima de descarga de datos a 149 GB;
- número limitado de envío de comandos y descargas de datos (21);
- utilización exclusiva de la banda de comunicación Ku de la ISS;
- capacidad limitada de ancho de banda a 50 Mbps.

A eso se suma que cualquier operación tiene que ser previamente autorizada por NASA con cadencia semanal, y el sistema de comunicación de la ISS tiene que ser compartido con otros experimentos y *payloads* a bordo de la Estación, con la consecuencia que el plan de operaciones ha tenido que ser adaptado cada semana, con mínima antelación, y haciendo frente a posibles interferencias y cambios debidos a otras misiones o a exigencias de las agencias espaciales que trabajan en la Estación.

El plan de operaciones, que se basa en el plan de calibración y validación de la cámara, se ha diseñado desde el principio teniendo en cuenta las dos limitaciones mayores:

el apuntado fijo de la ISS, que afecta el área de observación y reduce la selección de objetivos. Por eso, SATLANTIS ha efectuado simulaciones de la propagación de órbita de la ISS, mapeando el área de cobertura prevista y eligiendo targets para iSIM;

la falta de contacto directo con la carga útil ya que todas las operaciones se han realizado a través de JAXA, lo cual ha obligado al planteamiento de las operaciones con gran antelación, incluidas las instrucciones y los comandos de adquisición de imágenes, basados en la previsión de cobertura.

Frente a estas limitaciones, SATLANTIS ha desarrollado herramientas de planificación y gestión de las operaciones, en concreto:

- *IOD Realtime Telemetry Website* para monitorear las operaciones;
- *Operations GUI* para generar los comandos;
- *IOD Image Acquisition Website* para visualizar y gestionar las imágenes;
- *LEXH Simulator* para ensayos en tierra.



Figura 3. IOD Realtime Telemetry Website desarrollado por SATLANTIS.

3. Resultados y discusión

3.1 Resultados

El resultado de la misión ha sido un éxito. Tras tres meses de operaciones a bordo de la plataforma externa de la ISS, la carga útil iSIM-170 ha alcanzado la adquisición de más que 60.000 imágenes, descargadas en 21 ventanas útiles, que han permitido realizar enteramente el plan de calibración y validación de la cámara y demostrar sus capacidades, en particular la super-resolución, es decir los algoritmos de procesado de imagen que per-

miten una mejora de la calidad óptica nativa de las imágenes de hasta un factor 2, como muestran las siguientes imágenes de un aeropuerto militar en San Diego (EE.UU.).

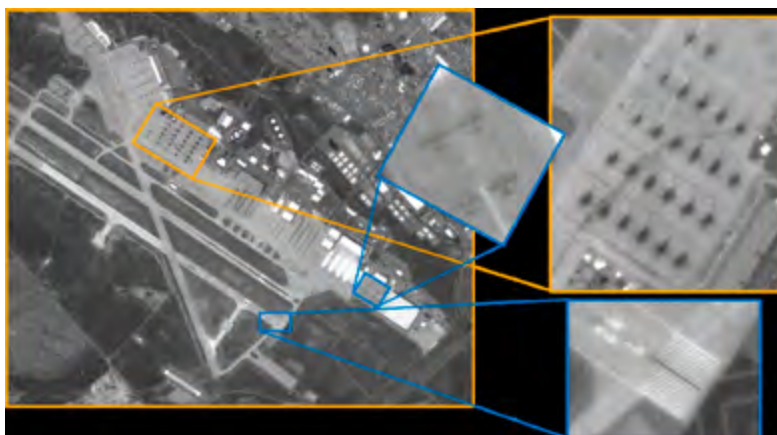


Figura 4. Aeropuerto en San Diego, antes de la Super-Resolución. (Crédito: SATLANTIS.)

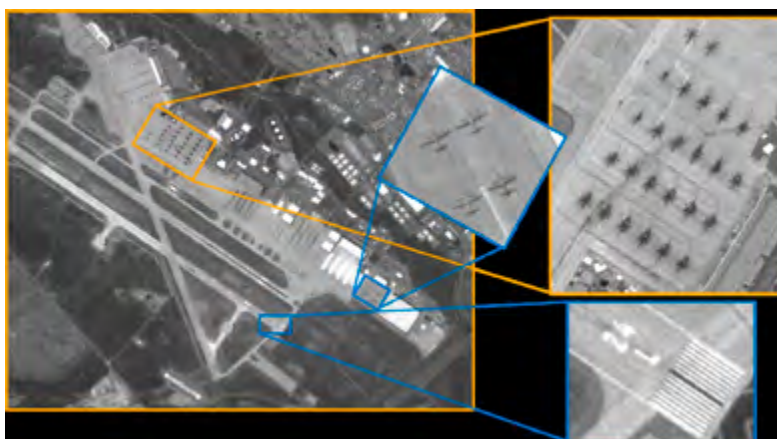


Figura 5. Aeropuerto en San Diego, después de la Super-Resolución. (Crédito: SATLANTIS.)

La mejora de resolución se nota en particular en objetos como aviones, helicópteros y franjas del aeropuerto, tal que, tras la super-resolución, se pueden claramente distinguir, respectivamente, los perfiles de las alas, las hélices y la distancia entre franjas. Además, considerando el tamaño de las hélices de los helicópteros, estimado en 0,9 m, se puede deducir la resolución final de iSIM como submétrica.

Otro ejemplo de la aplicación de super-resolución se puede observar en la siguiente comparativa, que analiza la zona del embarcadero en San Francisco (EE. UU.).

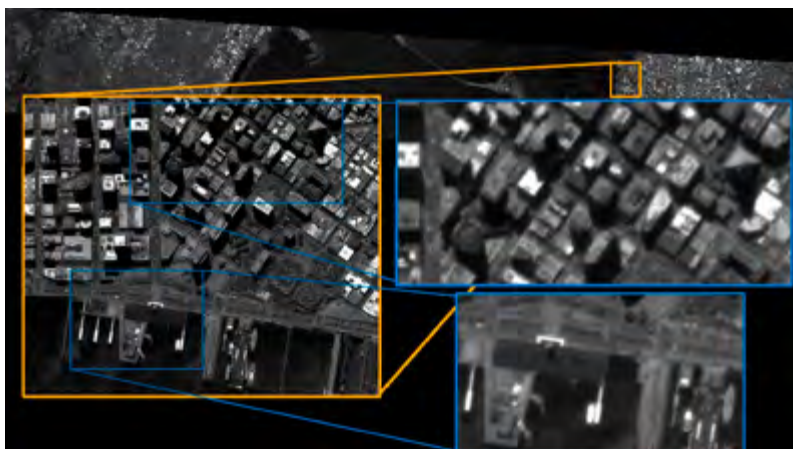


Figura 6. El Embarcadero en San Francisco, antes de la Super-Resolución. (Crédito: SATLANTIS.)

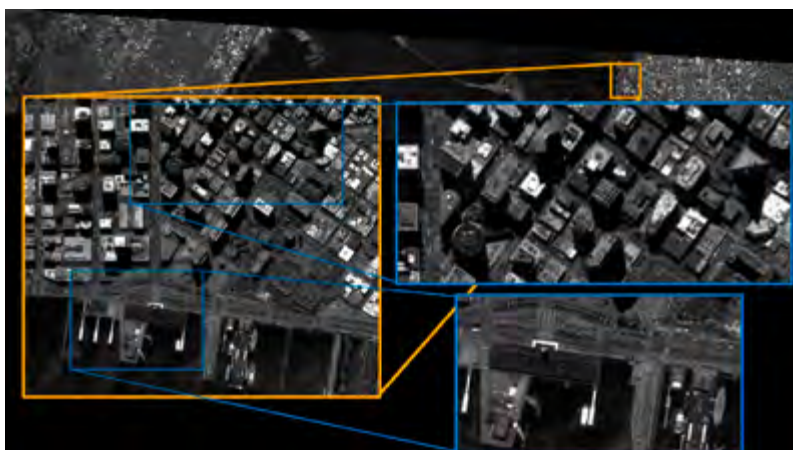


Figura 7. El Embarcadero en San Francisco, después de la Super-Resolución. (Crédito: SATLANTIS.)

Empezando por la franja matricial que se ve en segundo plano en las figuras, la comparativa toma como ejemplo pequeñas áreas y evidencia la mejora tras super-resolución. La segunda figura permite apreciar detalles del área urbano como edificios, calles, hasta los barcos y setos.

3.2 Aplicaciones y misiones futuras

Las cámaras iSIM son particularmente adecuadas para monitorizar estructuras lineales o puntuales con muy alta resolución, ofreciendo un gran potencial para sectores de defensa y seguridad. De hecho, SATLANTIS está planteando integrar sus cámaras en una constelación de pequeños satélites, lo que permitiría ofrecer un tiempo de revisita

muy alto, incluso hasta poder ofrecer informaciones en tiempo real. Esto, junto con su muy alta resolución, sería un valor añadido para aplicaciones C4ISR a través de plataformas aeroespaciales.

Además, con la misión CASPR, SATLANTIS está a punto de validar en órbita la agilidad, es decir, la capacidad de seguir con precisión estructuras lineales como costas, oleoductos, fronteras, con el fin de ofrecer servicios de vigilancia sobre áreas estratégicas, tanto para la industria como para agencias gubernamentales.

El Departamento de Defensa estadounidense, reconociendo el potencial de la tecnología iSIM, ha seleccionado una de las cámaras de SATLANTIS como *payload* primario de la misión STP-H7, que volará a la Estación Espacial Internacional en 2021.

En ámbito europeo, precisamente por el programa de desarrollo tecnológico EDIDP, SATLANTIS ha sido adjudicataria de un contrato para el estudio y el diseño de una carga útil que ofrezca una muy alta resolución espacial, junto a la agilidad, a una capacidad de descarga de datos de muy alta prestación y a un filtro activo para la alerta en tiempo real de detección y reconocimiento de objetivos específicos, con aplicaciones en ámbito marítimo.

4. Conclusiones

En conclusión, se puede afirmar que la demostración en órbita de iSIM-170 ha conseguido varios éxitos: ha sido la primera misión de observación de la Tierra desde la plataforma externa de la ISS con capacidades submétricas, ha sido la primera carga útil no japonesa en ser instalada en dicha plataforma y además ha representado un banco de pruebas para las posibles aplicaciones de iSIM en sectores de defensa y seguridad. Las misiones que SATLANTIS tiene planteadas a corto y medio plazo confirmarán las capacidades de la tecnología iSIM en ámbito de geointeligencia, monitorización y reconocimiento.

Agradecimientos

This project has received funding from the European Union's Horizon 2020 research and innovation programme under grant agreement No 768278.

Referencias

1. <https://satlantis.com/>.
2. <https://global.jaxa.jp/>.
3. <https://space-bd.com/en/business/>.

Identificación de agentes de guerra biológica mediante inmunobiosensado de exoproteínas dependientes del sistema de secreción bacteriano de tipo 2

Dabbagh-Escalante, Nushin Alba^{1,*}; Gil-García, Matilde¹; Peraile-Muñoz, Inés²; Cabria-Ramos, Juan Carlos¹ y Lorenzo-Lozano, Paloma¹

¹ Área de Defensa Biológica. Departamento de Sistemas de Defensa NBQ. Instituto Nacional de Técnica Aeroespacial «Esteban Terradas» (INTA).

² Ingeniería de Sistemas para la Defensa de España (ISDEFE). iperaile@isdefe.es

* Autor principal; dabbaghena@inta.es

Resumen: en los últimos años, la amenaza del uso de agentes de guerra biológica ha surgido como uno de los principales problemas de seguridad. Como consecuencia, la mayoría de los países se han visto obligados a aumentar recursos en investigación, diseño y desarrollo de tecnologías para detección e identificación de estos agentes.

La detección precoz en un incidente NBQ es crucial para la evaluación de riesgos, la toma de decisiones y la mitigación de la amenaza. Por ello, el diseño de dispositivos de identificación *in situ*, específicos y sensibles, fáciles de usar, de bajo coste y miniaturizables, se ha convertido en un objetivo prioritario.

Una de las tecnologías más relevantes actualmente para el desarrollo de estos dispositivos son los inmunobiosensores, basados en técnicas inmunológicas, ya que permiten una detección sensible y altamente específica.

En la presente investigación se plantea como novedad la utilización de mecanismos bacterianos de comunicación intercelular encaminados al desarrollo de un sistema de detección e identificación de agentes de guerra biológica mediante inmunobiosensado de exoproteínas dependientes de sistemas de secreción bacterianos, lo que permitirá un diagnóstico rápido y específico.

Palabras clave: agentes de guerra biológica, detección, inmunobiosensores, *Quorum Sensing*, sistema de secreción bacteriano, *Vibrio cholerae*.

1. Introducción

Los agentes biológicos constituyen un riesgo creciente en la actualidad ya que, debido al desarrollo tecnológico, son relativamente sencillos de producir. Algunos de estos agentes, por sus características de infectividad, toxicidad, patogenicidad y transmisibilidad, así como su facilidad para ser diseminados, son susceptibles de ser utilizados como agentes de guerra biológica contra la población, el ganado o las cosechas, provocando efectos devastadores en la salud y en la economía [1].

Por ello, la mayoría de los países invierten en la investigación y el desarrollo de técnicas de detección e identificación de dichos agentes, siendo un objetivo prioritario el desarrollo de dispositivos que permitan una identificación temprana *in situ* lo más específica y sensible posible [2]. Dentro de estos dispositivos, los inmunobiosensores, basados en técnicas inmunológicas, resultan ser los candidatos de elección, puesto que la unión antígeno-anticuerpo es sensible, rápida y altamente específica [3].

Las bacterias patógenas codifican proteínas asociadas a patogénesis (factores de virulencia) cuya secreción extracelular se realiza a través de diferentes sistemas de secreción (TTSS–*Sec Protein Secretion System*) regulados por *Quorum Sensing*. La autoinducción o *Quorum Sensing* es un mecanismo de regulación de la expresión genética en respuesta a la densidad de población celular. Las células involucradas producen y excretan sustancias, llamadas autoinductores, que sirven de señal química para inducir la expresión genética colectiva (figura 1) [4]. Las bacterias Gram positivas y Gram negativas usan estos mecanismos de comunicación para regular una gran variedad de actividades fisiológicas. Estos procesos incluyen producción de antibióticos, motilidad, esporulación, formación de biopelículas y virulencia. Existen dos grupos principales de autoinductores característicos de bacterias Gram positivas y Gram negativas, que son oligopéptidos cíclicos y derivados de homoserín lactona, respectivamente [5].

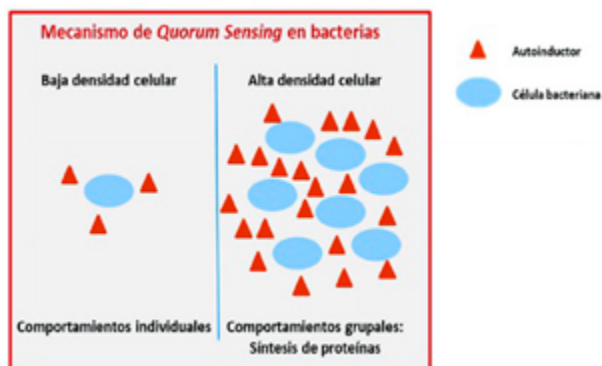


Figura 1. Mecanismo de Quorum Sensing en bacterias.

En este trabajo se propone la detección de factores de virulencia para el desarrollo de un sistema de inmunobiosensado que permita identificar de manera específica, rápida y sensible agentes de guerra biológica.

2. Materiales y métodos

2.1 Elección del agente biológico y analito a detectar

El agente biológico de elección debe ser una bacteria patógena que cumpla las condiciones para ser usado como agente de guerra biológica. Debe secretar al medio extracelular factores de virulencia, lo que posibilitará su unión a un biorreceptor específico inmovilizado en el inmunobiosensor y, por ende, su identificación.

El analito a detectar debe ser específico de especie y secretado al medio extracelular (no inyectado al interior de la célula hospedadora) a través de un sistema de secreción muy conservado que esté regulado por moléculas señal o autoinductores a través de un mecanismo de *Quorum Sensing*.

2.2 Mecanismo de comunicación celular

Para el estudio de los mecanismos de regulación bacteriana se llevará a cabo la caracterización de la curva de crecimiento del agente biológico a detectar. Se elegirán las técnicas cromatográficas más adecuadas para estudiar los niveles de autoinductores, se tomarán alícuotas del cultivo celular en distintos puntos de su fase de crecimiento exponencial y se analizará la concentración de autoinductores mediante técnicas cromatográficas. De este modo, será posible establecer las concentraciones de autoinductores y el tiempo mínimo para inducir la expresión y secreción del analito a detectar.

2.3 Optimización del método inmunológico para la detección del analito seleccionado

Se definirán las pautas de inmunización de los ratones y el análisis de la respuesta inmune para la obtención de hibridomas. Se realizará la purificación de anticuerpos específicos y la unión antígeno-anticuerpo se detectará mediante marcadores colorimétricos y/o luminiscentes (RPE y FITC).

2.4 Desarrollo de los procesos de biofuncionalización y adaptación al sistema de sensado

Se utilizará el sistema BSA/ anti BSA para optimizar la unión del elemento de reconocimiento a la superficie del transductor. Se ensayarán tres modelos:

1. Absorción pasiva. Incubación del anticuerpo sobre la superficie planar en *buffer* fosfato salino (PBS 1x overnight).

2. Unión covalente a través de glutaraldehído. Incubación *overnight* del anticuerpo marcado con FITC sobre la superficie planar previamente tratada con glutaraldehído al 5 % durante 15 minutos. Tras un bloqueo con PBS 1x-caseína, se incuba 1 hora el antígeno (BSA marcada con ficoeritrina).
3. Unión orientada a través de proteína mediadora A/G. Incubación *overnight* de una solución de proteína A/G a diferentes concentraciones (10 μ L, 50 μ g/mL y 100 μ g/mL). Tras un bloqueo con PBS-caseína, se incuba 1 hora el anticuerpo marcado con FITC y, después, se incuba el antígeno (BSA marcada con RPE) durante otra hora.

3. Resultados y discusión

3.1 Elección del agente biológico y analito a detectar

Las bacterias patógenas utilizan diversos métodos para invadir a sus huéspedes, siendo la secreción de proteínas que actúan como factores de virulencia una de las principales estrategias para ello. La secreción extracelular se realiza a través de diferentes sistemas de secreción que presentan distintos mecanismos de acción (figura 2) [6].

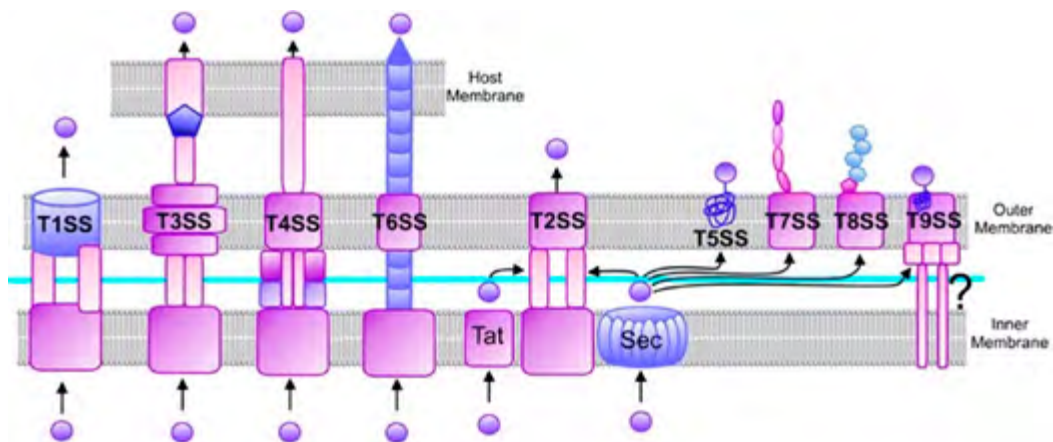


Figura 2. Sistemas de Secreción de las bacterias Gram negativas; adaptada de [7].

Estos se clasifican en tres grandes grupos, aquellos que liberan proteínas al medio extracelular (T1SS, T2SS y T5SS), los que inyectan proteínas al interior de las células hospedadoras eucariotas (T3SS, T4SS y T6SS) o al de bacterias (T6SS) y los que están implicados en la biosíntesis de apéndices en la superficie de la bacteria (T7SS-específico de bacterias Gram positivas, T8SS y T9SS) [7].

Se utilizará *Vibrio cholerae*, agente etiológico del cólera, como agente patógeno a detectar y la exoproteína hemaglutinina proteasa (HA/P), factor de virulencia específico de dicho agente [8], como analito empleado para su detección.

Vibrio cholerae, potencial arma de guerra biológica, está clasificado en la categoría B por el Centro para el Control de Enfermedades de Atlanta (CDC), es fácil de diseminar y provoca tasas de morbilidad y de mortalidad moderadas [9]. *V. cholerae* regula la producción de sus factores de virulencia, entre los que se encuentra la HA/P, mediante un sistema de *Quorum Sensing* mediado por dos autoinductores, CAI-1 (S-3-hidroxitridecan-4-ona) y AI-2 (furanosil borato diéster) a través de un mecanismo de *Quorum Sensing* [5].

La HA/P de *Vibrio cholerae* es una exoproteína específica de especie que se secreta al medio extracelular (no se inyecta al interior de la célula hospedadora, lo que permite su unión a un biorreceptor específico y, por ende, su detección) a través de un sistema de secreción bacteriano de tipo 2 (T2SS). T2SS es un sistema de secreción muy conservado regulado por moléculas señal o autoinductores que desencadenan la expresión de los genes que codifican para las exoproteínas, en la misma célula que los liberó y sobre el resto de la población (mecanismos bacterianos de comunicación intercelular o *Quorum Sensing*) [10].

3.2 Estudio del mecanismo de regulación por *Quorum Sensing* en la producción de factores de virulencia en *Vibrio cholerae*

In vivo en ausencia de autoinductores, *V. cholerae* asume el programa de expresión génica de baja densidad celular (LCD), en el que las bacterias se agregan en pequeños grupos de biopelículas en el epitelio intestinal y se sintetiza la toxina colérica. Cuando aumenta la densidad celular (HCD), los autoinductores se sintetizan, se inhibe la expresión de los genes de virulencia y se expresa HA/P induciendo la disociación de las células epiteliales del huésped. La diarrea severa causada por la toxina colérica moviliza las bacterias colonizando nuevas áreas del tracto intestinal o son expulsadas junto con las heces pudiendo infectar a otros huéspedes [8].

Por analogía a lo que ocurre con otras especies de género *Vibrio* [5], la concentración de autoinductor depende de la densidad poblacional y su acumulación desencadena la producción de HA/P. El primer paso para el estudio de estos mecanismos de comunicación intercelular en función de la densidad celular es la caracterización de la curva de crecimiento de *V. cholerae* (figura 3).

Un ejemplo de procesos regulados a través de estos mecanismos de comunicación intercelular, es la producción de bioluminiscencia en *Vibrio harveyi*. Cuando la densidad celular es elevada la concentración de autoinductores es máxima y se produce inducción

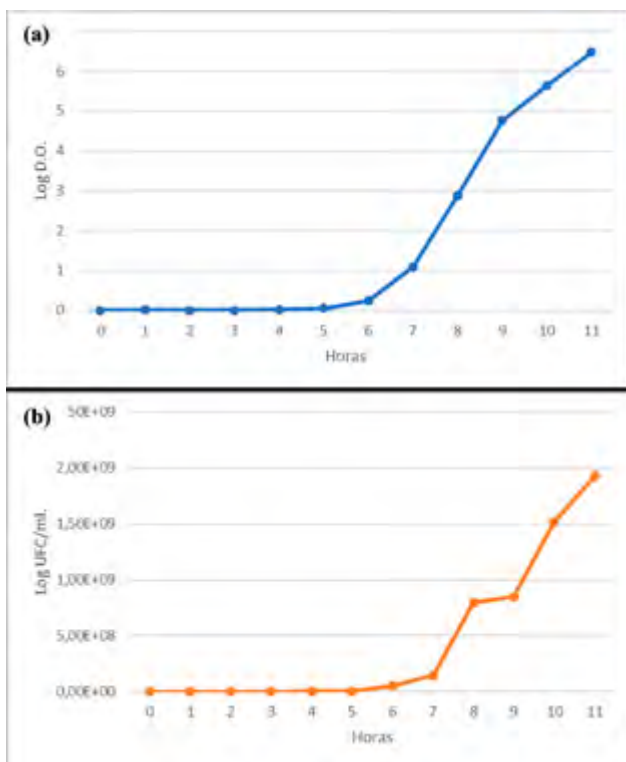


Figura 3. Curva de crecimiento de *V. cholerae* con medidas de densidad óptica (D.O.) (a) y de unidades formadoras de colonias por ml. (UFC/ml.) (b).

de luciferasa, enzima que cataliza el proceso de producción de bioluminiscencia. Para demostrar que la concentración de autoinductor dependía de la densidad poblacional y que su acumulación desencadenaba la producción de bioluminiscencia, Nealson [11] cultivó *V. harveyi* hasta su fase de crecimiento exponencial, donde la concentración de autoinductor debía ser máxima. Tras extraer y filtrar la fracción del medio con autoinductor, se añadió a un nuevo cultivo recién inoculado, comprobando que, tras su incorporación, se inducía la producción de bioluminiscencia (figura 4) [5].

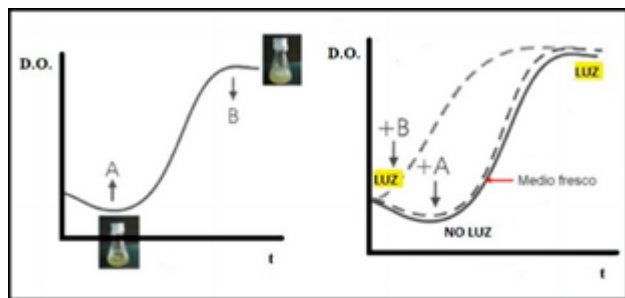


Figura 4. Esquema del experimento en el que se demostró que la bioluminiscencia de *V. harveyi* estaba causada por la acumulación de un autoinductor; adaptada de [5].

Los niveles de autoinductores se determinarán mediante cromatografía líquida de alta resolución acoplada a espectrómetro de masas (HPLC-MSD) a partir de alícuotas del cultivo celular tomadas en distintos puntos de su fase de crecimiento exponencial. De este modo, será posible establecer las concentraciones de autoinductores y el tiempo mínimo para inducir la expresión y secreción del analito a detectar, la proteína HA/P, al medio extracelular.

3.3 Optimización del método inmunológico para la detección del analito seleccionado

Se tomarán diferentes alícuotas de la fase de crecimiento exponencial de *V. cholerae* y se realizarán inmunoensayos para determinar la concentración de la exoproteína HA/P utilizando anticuerpos anti HA/P *in house*. La unión antígeno-anticuerpo se detecta mediante marcadores fluorescentes que permiten cuantificar la concentración de HA/P presente en el cultivo. El método se puso a punto utilizando el sistema BSA/anti BSA (figura 5).

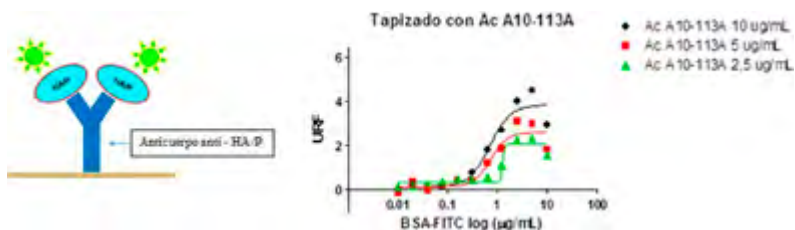


Figura 5. Inmunoensayo modelo para validar el sistema de detección con fluorescencia [12].

3.4 Desarrollo de los procesos de biofuncionalización y adaptación al sistema de sensado

Para la validación de los procesos de inmovilización se hará uso de métodos inmunológicos convencionales, utilizando anticuerpos y antígenos conjugados con marcadores fluorescentes (RPE y FITC) (figura 6).

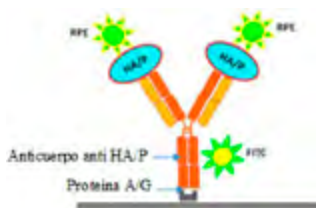


Figura 6. Validación del modelo de inmunocaptura.

La eficiencia del sistema de inmunocaptura mediante unión covalente con glutaraldehído, es estadísticamente inferior a la adsorción pasiva. La inmovilización del anticuerpo a través de proteínas mediadoras A/G, incrementa la eficiencia de unión de la

BSA en el sistema siendo estadísticamente superior a todos los demás modelos de biofuncionalización (figura 7) [12].

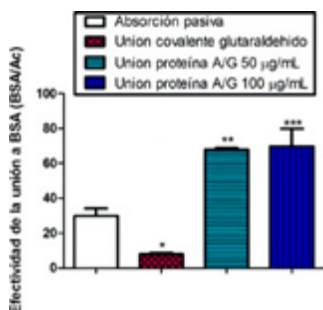


Figura 7. Efectividad de la unión a la BSA en los distintos modelos de biofuncionalización.* $p < 0,05$, ** $p < 0,01$, *** $p < 0,001$. Diferencias respecto al grupo "Absorción pasiva". (One Way ANOVA seguido del test de Comparación Múltiple Newman Keuls) [12].

3.5 Adaptación del sistema de detección a otros agentes de guerra biológica

Una vez optimizado el proceso de inmunocaptura, podrá incorporarse este método de detección en dispositivos de alerta temprana desarrollados para la identificación *in situ* de agentes de guerra biológica. Dado que los autoinductores son genéricos para cada grupo de bacterias, únicamente sería necesario seleccionar los anticuerpos específicos contra el factor de virulencia que se quiera detectar en cada caso.

4. Conclusiones

La exoproteína HA/P de *Vibrio cholerae* es específica de especie y secretada al medio extracelular utilizando un sistema de secreción bacteriano altamente conservado, regulado por un mecanismo de *Quorum Sensing*.

Combinando el uso de inmunobiosensores con estos mecanismos de comunicación celular es posible desarrollar un sistema de detección de agentes patógenos con alta sensibilidad, rapidez y especificidad.

La tecnología propuesta en este trabajo es adaptable a otros agentes patógenos, ya que permite incorporar diferentes biorreceptores específicos.

Este trabajo aborda la detección e identificación de agentes de guerra biológica desde un punto de vista novedoso que permitiría disponer de técnicas de diagnóstico con sensibilidades y especificidades muy elevadas que podrían ser incorporadas en diferentes dispositivos para un gran número de agentes patógenos.

Agradecimientos

Los trabajos recogidos en esta comunicación están incluidos en el proyecto coordinado PHOLOW PID2019-106965RB-C22 financiado con fondos MICINN.

Contrato predoctoral Resolución de la Dirección General del Instituto Nacional de Técnica Aeroespacial «Esteban Terradas», de 30 de enero de 2019.

Referencias

1. D. Thavaselvam, R. Vijayaraghavan, «Biological warfare agents», *J Pharm Bioallied Sci*, vol. 2, no. 3, pp. 179-188, 2010.
2. *Detección e identificación de agentes de guerra biológica. Estado del arte y tendencia futura*. Madrid: Ministerio de Defensa, Dirección General de Relaciones Institucionales, 2010.
3. S. Sharma, H. Byrne, R. J. O'Kennedy, «Antibodies and antibody-derived analytical biosensors», *Essays Biochem*, vol. 60, no. 1, pp. 9-18, 2016.
4. T. R. Kievit, B. H. Iglewsky, «Bacterial Quorum Sensing in Pathogenic Relationships», *Infect Immun*, vol. 68, no. 9, pp. 4839-4849, 2000.
5. D. Marquina, A. Santos, «Sistemas de *quorum sensing* en bacterias», *Reduca (Biología)*, Serie Microbiología, vol. 3, no. 5, pp. 39-55, 2010.
6. E. R. Green, J. Meccas, «Bacterial Secretion Systems – An overview», *Microbiol Spectr*, vol. 4, no. 1, 2016. doi: 10.1128/microbiolspec.VMBF-0012-2015.
7. K. M. Bocian-Ostrzycka, M. J. Grzeszczuk, A. M. Banaś, E. K. Jagusztyn-Krynicka, «Bacterial thiol oxidoreductases – from basic research to new antibacterial strategies», *Appl Microbiol Biotechnol*, vol. 101, no. 10, pp. 3977-3989, 2017. doi: 10.1007/s00253-017-8291-8.
8. M. B. Miller, K. Skorupski, D. H. Lenz, R. K. Taylor, B. L. Bassler, «Parallel Quorum Sensing Systems Converge to Regulate Virulence in *Vibrio cholerae*», *Cell*, vol. 110, pp. 303-314, 2002.
9. <https://www.cdc.gov/cholera/index.html>.
10. K. Wooldridge, *Bacterial Secreted Proteins: Secretory Mechanisms and Role in Pathogenesis*. Nottingham: Caister Academic Press, 2009.
11. K.H. Nealson, «Early observations defining quorum-dependent gene expression», en *Cell-cell signaling in Bacteria*. Washington, D.C.: American Society for Microbiology, 1999.
12. I. Peraile, M. Gil, C. E. Guamán, L. González, J. C. Cabria, P. Lorenzo, «Optimización del proceso de inmovilización de anticuerpos en inmunobiosensores», *Rev Sanid Milit*, vol. 74, no. 3, pp. 151-156, 2018.

Influencia de la emisividad superficial en la medida de la temperatura de transición vítrea en el análisis dinamomecánico de propulsantes sólidos

Jiménez Pérez, Javier¹; López Sánchez, Raúl^{2,*}; Gómez López, Miguel Ángel²;
Rodríguez Pérez, Jesús¹

¹ Grupo de investigación Durabilidad e Integridad Mecánica de Materiales Estructurales (DIMME), Universidad Rey Juan Carlos (URJC). Calle Tulipán, s/n, 28933, Móstoles, Madrid, Spain

² Departamento de Optoelectrónica y Misilística, campus «La Marañosa», Instituto Nacional de Técnica Aeroespacial (INTA). Ctra M301 km 10,5, edif 8, 28330, San Martín de la Vega (Comunidad de Madrid), España

* Autor principal; lopezsr@inta.es

Resumen: el análisis dinamomecánico, DMA, es una de las técnicas de caracterización más utilizadas para determinar la transición vítrea, T_g , en materiales poliméricos. En los propulsantes sólidos de material compuesto, una determinación precisa de la T_g es de especial relevancia puesto que marca un límite entre dos regímenes de rigidez muy diferentes. Se considera que, a parte de la necesaria calibración del equipo, también es necesario conocer y entender las interferencias que se pueden producir en la cuantificación de la T_g . En ediciones anteriores de este congreso, se presentaron metodologías para realizar una calibración apropiada, así como estudios experimentales que demostraban desviaciones sistemáticas superiores a los 15°C por efecto combinado, tanto del tipo de ensayo, como por propiedades del material.

En este trabajo se presenta un modelo teórico, basado en las leyes de calentamiento de Newton y de radiación de Stefan-Boltzman, que explica estas desviaciones, validándose con datos experimentales. El objetivo final es el de minimizar algunas desviaciones experimentales, a partir de un mayor conocimiento del funcionamiento del DMA así como de la influencia de las propiedades del material analizado.

Palabras clave: DMA, propulsante, *composite*, T_g , emisividad.

1. Introducción

La determinación precisa de las propiedades viscoelásticas en propulsores sólidos de tipo composite, o compositas, tiene una especial relevancia para las pruebas de vigilancia sobre la munición, particularmente en cohetes de gran tamaño y envejecidos. En estos materiales, la principal propiedad viscoelástica es la temperatura de transición vítrea, T_g , habitualmente determinada mediante analizadores dinamomecánicos, DMA, acorde al STANAG 4540 [1]. En los polímeros y materiales compuestos de matriz polimérica, la T_g es la temperatura a la cual se observa un cambio en las propiedades de la fase amorfa, sin que exista influencia de los rellenos, ni de la adhesión entre matriz y relleno [2].

1.1 Importancia de la temperatura de transición vítrea en propulsores compositas

Las compositas más habitualmente empleadas en cohetes suelen ser poliuretanos, que pueden obtenerse por reacción de curado entre polibutadienos (i.e. HTPB) e isocianatos. Además, el propulsante suele cargarse con oxidantes sólidos en más de un 60 %. La reacción típica de curado puede observarse en la figura 1, siendo habitual que esta reacción se realice de forma incompleta, aumentando la viscoelasticidad del propulsante y exacerbando la evolución de algunas propiedades con el envejecimiento, entre las que destaca la T_g .

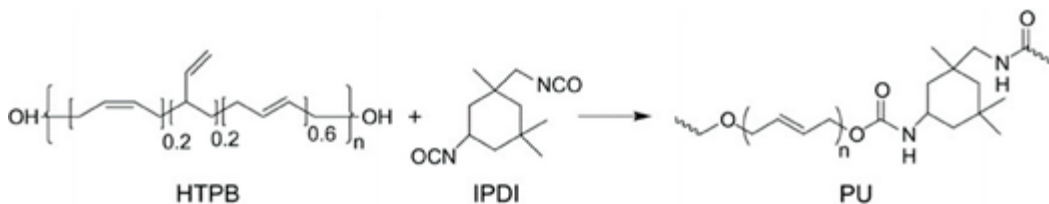


Figura 1. Curado de un polibutadieno con isocianatos, para proporcionar poliuretanos.

Este post-curado progresivo y natural, que se induce intencionadamente durante la fabricación, es uno de los procesos implicados en el envejecimiento natural de los propulsores durante su almacenamiento. Es habitual encontrar estudios que aplican envejecimientos acelerados (habitualmente térmicos) para relacionar el envejecimiento del propulsante con un cambio en la T_g , habitualmente hacia valores más altos [2].

Es conocido que gran parte de la respuesta viscoelástica de un propulsante sólido de tipo poliuretano se relaciona con el doble enlace de los polibutadienos (ver figura 1), el cual suele perderse con el envejecimiento y, por lo tanto, se produce una modificación del comportamiento mecánico. Un cohete, con un propulsante envejecido, disparado a una temperatura ambiental inferior a la T_g sufrirá una especial rigidización que aumentaría la probabilidad de fallo catastrófico del grano.

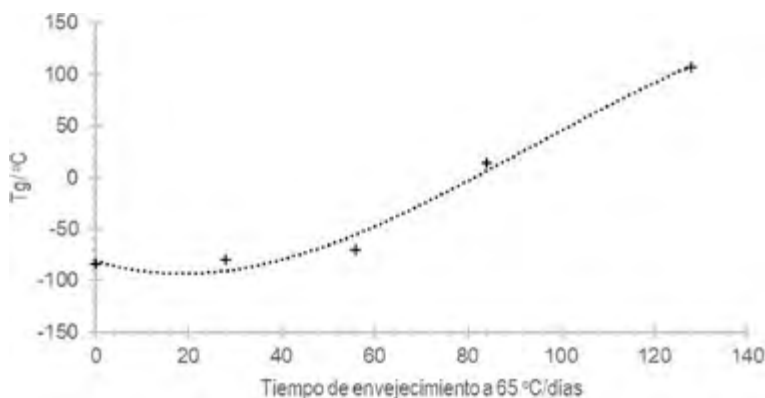


Figura 2. Evolución de la Tg con el envejecimiento, acelerado, en un propulsante de HTPB [2].

1.2 Análisis dinamomecánico y errores sistemáticos debidos a la calibración del equipo

La mayor parte de los equipos dinamomecánicos suelen presentar severas deficiencias en la periodicidad de sus calibraciones, a pesar de que informan continuamente de valores cuantitativos relacionados con la temperatura. De manera habitual, e incorrecta, se suele considerar suficiente la calibración del termopar en el horno, aludiendo a la dificultad para obtener patrones trazables [3]. La realización de medidas, sin una correcta calibración del DMA, puede llegar a introducir desviaciones sistemáticas superiores a 25°C.

1.3 Desviaciones debidas al método de ensayo y a propiedades del material

Un desconocimiento sobre la influencia en la determinación de la T_g , tanto de las condiciones de ensayo como de las propiedades del material, suele tener como consecuencia que el valor informado sea el que proporciona el equipo, sin corrección alguna. En ediciones anteriores del DESEI+D [4] ya se presentaron pruebas empíricas que justificaban la necesidad de calibrar correctamente estos equipos y que, además, aconsejaban corregir la respuesta del DMA.

El presente documento ahonda en las importantes desviaciones sobre los resultados experimentales causada por dos variables, habitualmente subestimadas: la rampa de temperatura aplicada en ensayos dinámicos y la emisividad superficial de la probeta analizada. Estos factores contribuyen a las desviaciones de una manera mucho menos relevante que un fallo en la calibración, pero su contribución no debería considerarse despreciable. Las citadas correcciones, aunque aún no se encuentren consideradas en las normativas militares [1], ya existen normas civiles [5] que recomiendan correcciones

como la debida a la velocidad de la rampa, pero obviando interferencias debidas a propiedades del material ensayado.

2. Materiales, modelos y métodos

Este trabajo se basa en datos experimentales anteriormente estudiados desde una perspectiva fenomenológica [4]. Ahora, el análisis se realiza con herramientas de simulación que, desde una perspectiva teórica, permita justificar los resultados obtenidos.

2.1 Estudio experimental: equipamiento y muestras analizadas

El equipo utilizado ha sido un DMA Q800 de TA Instruments, trabajando en modo doble empotramiento, frecuencia de trabajo de 1 Hz, amplitud de desplazamiento de 15 μm , intervalo de ensayo de -150 a 200 $^{\circ}\text{C}$, rampas de calentamiento de hasta 5 $^{\circ}\text{C}/\text{min}$, y probetas de dimensiones 60 x 3 x 10,2 mm³. Para este trabajo, la temperatura de transición vítrea se calcula cuantificando el máximo del módulo de pérdidas, E'' , por ser la estimación más definida y reproducible.

Las muestras utilizadas para este estudio han sido de polimetilmetacrilato, PMMA, con diferentes aditivos de color en su composición y estados superficiales modificados [4], así como probetas específicas para la calibración del DMA [3], con objeto de minimizar la introducción involuntaria de desviaciones sistemáticas.

2.2 Simulación: modelo y software

La ecuación 1 es la base del modelo teórico sobre el comportamiento térmico de la probeta en un horno de DMA con convección forzada. Los principales mecanismos de transmisión de calor, en ensayos dinámicos, son la convección y la radiación. Por lo tanto, el modelo debe combinar las leyes de calentamiento de Newton y de radiación de Stefan-Boltzman. La ecuación obtenida es de tipo diferencial, ordinaria y sin resolución analítica.

$$\frac{dT}{dt} = -k[\alpha\gamma(T - T_{\text{horno}}) - \varepsilon\sigma\lambda(\tau \cdot T_{\text{resistencias}}^4 - T^4)]; \text{ siendo } k = \frac{S}{\rho V C_e} \quad (1)$$

donde T es la temperatura de la probeta, α el coeficiente de transmisión térmica, T_{horno} la temperatura del horno, ε la emisividad de la probeta, σ el coeficiente de Stefan Boltzmann, $T_{\text{resistencias}}$ la temperatura superficial de las resistencias del horno, k es una constante de la probeta donde S es su superficie, ρ su densidad, V el volumen, C_e el calor específico del material, γ , λ , τ son factores de corrección que se explicarán a continuación.

El modelo propuesto combina dos leyes relacionadas con la transferencia de energía: el primer término está basado en la Ley del Calentamiento de Newton, la cual está principalmente influida por la rampa de temperatura del ensayo dinámico. Por otra parte, el segundo término se focaliza en el calor absorbido por radiación, siguiendo la Ley

de Stefan-Boltzman y, en este caso, la propiedad del material más relevante es la emisividad superficial. Además, algunas de las propiedades del material han sido tomadas de la bibliografía y estas aproximaciones al valor verdadero han requerido la introducción de compensaciones, en forma de factores de corrección al factor de corrección de convección, γ , factor de radiación λ y a la temperatura de la resistencia τ .

Mediante algoritmos implementados en Matlab® se han realizado simulaciones que proporcionan resultados similares a las observaciones experimentales. La instrucción principal, para la resolución del modelo diferencial se presenta en la ecuación 2.

$$[t,y]=ode45(@(t,y)odefcn(t,y,a,b,...),tspan,y0,options) \quad (2)$$

$$\text{function } dydt=odefcn(t,y,a,b,...) \quad (3)$$

donde t e y son los argumentos de entrada, *ode45* es la función utilizada para resolver la EDO, *odefcn* es una función recogida en la ecuación 3, *tspan* es el intervalo de tiempo en el que se desea resolver la ecuación y por último $y0$ el valor inicial de temperatura.

3. Resultados y discusión

Para la validación del modelo se comparan los resultados experimentales con los obtenidos mediante el modelo matemático. Las simulaciones han necesitado de los parámetros recogidos en la tabla 1.

$\alpha \left[\frac{W}{m^2 \cdot K} \right]$	$S[m^2]$	$\rho \left[\frac{kg}{m^3} \right]$	$V[m^3]$	$C_e \left[\frac{J}{K \cdot kg} \right]$	$\sigma \left[\frac{W}{m^2 \cdot K^4} \right]$
6.3E1	1.2E-3	1.2E3	1.8E-6	1.5E3	5.7E8

Tabla 1. Propiedades del PMMA obtenidas de la bibliografía.

Las simulaciones realizadas, para el valor de la T_g se enmarcan en 3 tres grupos principales: la determinación del tiempo límite para diferentes velocidades de ensayo, la influencia de la velocidad de calentamiento en la T_g medida y la desviación producida por la emisividad.

3.1 Tiempo límite

Durante el ensayo, la temperatura de la probeta se va distanciando del régimen teórico impuesto hasta que, llegado un momento, la distancia entre temperatura de probeta y horno se mantiene constante hasta el final del ensayo, ver figura 3. El instante en que esta diferencia de temperaturas se mantiene constante, dentro de una tolerancia,

determina la duración mínima del ensayo para obtener medidas reproducibles, denominándose *tiempo límite*.

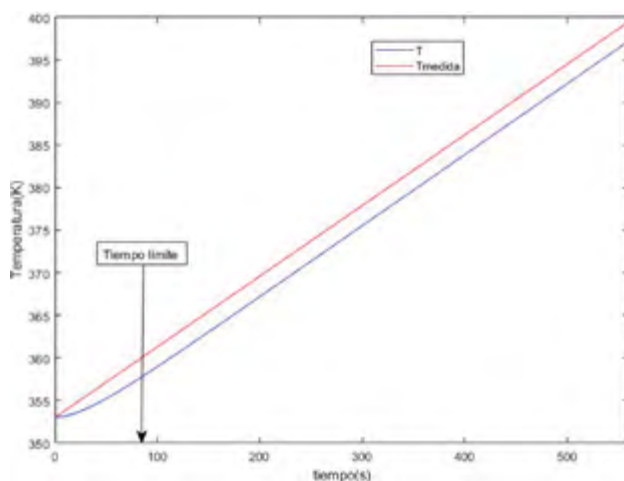


Figura 3. Variación de la temperatura del horno (roja) y de la probeta (azul).

3.2 Velocidad de calentamiento

La velocidad de calentamiento en los ensayos dinámicos, influye en la medida recogida, provocando que el valor registrado se desvíe más del valor verdadero, cuanto mayor sea la velocidad de calentamiento [5]. Así mismo, como se muestra en la figura 4, suele considerarse que este valor verdadero de T_g corresponde con la menor velocidad posible del ensayo, es decir, con la ordenada en el origen de una representación de T_g medida frente a velocidad de ensayo.

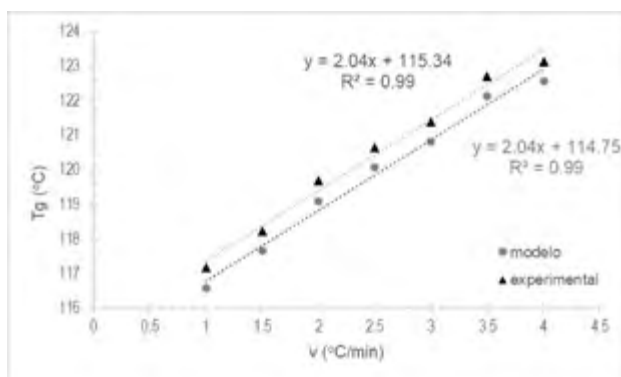


Figura 4. Influencia de la velocidad de calentamiento en el valor de la T_g .

La linealidad observada viene garantizada por dos aspectos recogidos en el modelo: que el ensayo se ha realizado en condiciones de reproducibilidad ($t_{\text{ensayo}} > t_{\text{límite}}$) y por el comportamiento de Ley de Calentamiento de Newton. La ordenada en el origen del ajuste (115,28°C) corresponde al valor estimativo de la del material.

3.3 Estado superficial

Experimentalmente se observa que la emisividad superficial de la probeta es capaz de introducir desviaciones sistemáticas en la medida de la T_g . Este mismo comportamiento se ha reproducido con simulaciones sobre el modelo propuesto. El estudio experimental consistió en el análisis de siete probetas blancas a las que se les modificó su estado superficial. Las medidas de emisividad se realizaron con un fotodiodo diseñado *ad hoc* [4]. Los resultados de la tabla 2 resumen los resultados más importantes.

	Probeta Blanca	Sputtering de Platino	Lacado Plateado	Lacado Morado	Lacado Negro Brillo	Lacado Amarillo	Lacado Negro Mate
ϵ	0.1	0.17	0.33	0.42	0.49	0.55	0.55

Tabla 2. Tratamiento superficial sobre probetas de PMMA y su emisividad.

El estudio experimental mostró diferencias de T_g de hasta 3°C entre tratamientos superficiales, ver figura 5, que pueden llegar a ser de más de 4°C para emisividades cercanas a 1, precisamente las de los propulsores sólidos de tipo *composite*.

La aplicación del modelo a estos datos experimentales coincide con la tendencia observada experimentalmente, es decir, se observa una disminución de la T_g con el aumento de la emisividad, tal y como predice la Ley de Stefan Boltzman. Nuevamente, la extrapolación a un valor de emisividad nula correspondería con la mejor estimación del valor de la T_g del material, que en este caso sería de 115,3°C.

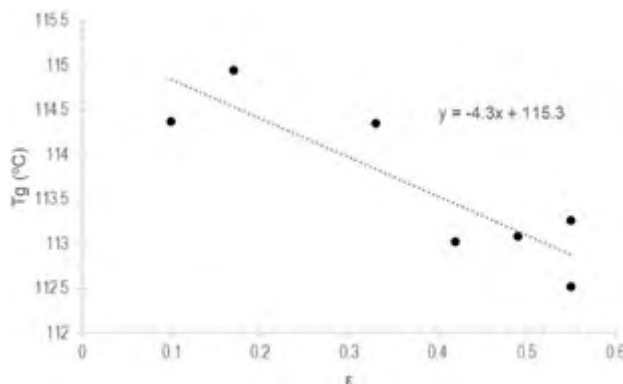


Figura 5. Efecto de la emisividad en la medida de la T_g del material.

Tras la comparación entre simulaciones y datos experimentales, se obtienen los siguientes factores de ajuste del modelo γ , λ y τ , recogidos en la tabla 3.

γ	λ	τ
$60 \cdot v$	$60 \cdot v$	1.01

Tabla 3. Factores de ajuste.

Donde v es la velocidad del ensayo en $^{\circ}\text{C} \cdot \text{min}^{-1}$.

4. Conclusiones

En el presente documento se ha profundizado en algunas de las causas más importantes que introducen, involuntariamente, errores sistemáticos en la determinación de la T_g mediante DMA. Estos estudios son de especial relevancia para la seguridad de los propulsores sólidos de tipo *composite* durante las pruebas de vigilancia sobre la munición.

Aunque la causa principal de los mayores errores sistemáticos suele estar relacionada con calibraciones inapropiadas del equipo, no se pueden considerar despreciables las desviaciones debidas al desconocimiento de la metodología empleada y de la influencia de las propiedades del material. En este documento se analizan dos de estas propiedades, la primera relacionada con el tipo de ensayo, la velocidad de calentamiento del horno, mientras que la segunda es una propiedad del material bajo estudio: la emisividad superficial.

Un análisis teórico de los mecanismos de transmisión de calor que se producen dentro del horno del DMA, unido a estudios experimentales exhaustivos, ha permitido postular un modelo matemático capaz de justificar el comportamiento térmico de la probeta en un ensayo dinámico. Para ello, se ha considerado que los dos fenómenos principales de la transmisión de calor son convección y radiación. Adicionalmente a la postulación del modelo, se ha realizado una validación del mismo con datos experimentales en tres etapas: determinación del tiempo límite, influencia de la velocidad de calentamiento y de la emisividad superficial.

De los efectos estudiados, la velocidad de calentamiento del horno es la que presenta una contribución más importante en la desviación del valor de la T_g pudiendo llegar a ser, según nuestros estudios, hasta de 10°C para velocidades de $5^{\circ}\text{C}/\text{min}$. Por su parte, con una desviación mucho más modesta, la emisividad superficial presenta desviaciones de hasta 5°C para emisividades de entre 0 y 1, siendo la máxima emisividad la que mayor desviación produce. La desviación mínima se asocia a un comportamiento ideal, mientras que la desviación máxima corresponde a propulsores de tipo *composite*,

principalmente cuando llevan elevadas cargas de aluminio micronizado (i.e motores de aceleración).

En conclusión, la desviación producida por la metodología empleada y por propiedades del material debería ser corregida, si realmente buscamos una determinación precisa de las propiedades del material. No se considera que un error sistemático de 15°C pueda ser considerado despreciable en la determinación de la T_g cuando nos referimos a propulsores, sobre todo cuando el disparo pueda realizarse a temperaturas ambientales del entorno de la temperatura de transición vítrea.

Agradecimientos

Este trabajo ha sido posible gracias al programa de prácticas en empresa, en el marco del convenio de la URJC con el INTA. Agradecer al grupo DIMME de la URJC por su tutela científica y al departamento de Optoelectrónica y Misilística del INTA por el acceso a las instalaciones y recursos.

Referencias

1. STANAG 4540, Explosives, procedures for dynamic mechanical analysis (DMA) and determination of glass transition temperature, NSA NATO, 2002.
2. J. L. de la Fuente, «An analysis of the thermal aging behaviour in high-performance energetic composites through the glass transition temperature», Polymer Degradation and Stability, vol. 94, pp. 664-669, 2009.
3. R. López Sánchez, Á. Rico García, F. J. del Olmo Navarro y J. Rodríguez Pérez, «Probeta de calibración en temperatura de equipos analizadores dinamomecánicos». España Patente ES 2 561 952 A8, 01 03 2016.
4. I. Vaz Fernández, R. López Sánchez y J. Rams Ramos, «Influencia de la emisividad superficial en la determinación de la T_g mediante DMA», de DESEI+D, Valladolid, 2018.
5. «ISO 6721-11:2012. Plastics-Determination of dynamic mechanical properties. Part 11: Glass transition temperature» 2012.

Formación física y eficacia operativa en unidades paracaidistas

Serrano Muñoz, Alberto^{1*}; Tejada Vázquez, Álvaro Nicolás²; y Gómez Cabello, Alba María³

¹ Brigada Paracaidista (BRIPAC). Carr. M-108, km. 4,4, 28860, Paracuellos de Jarama. museralb@hotmail.com

² Brigada Paracaidista (BRIPAC). Carr. M-108, km. 4,4, 28860, Paracuellos de Jarama. atejada37@hotmail.com

³ Centro Universitario de la Defensa (CUD). Carr. de Huesca s/n, 50090, Zaragoza. ago-meiz@unizar.es

* Autor principal; museralb@hotmail.com

Resumen: la eficacia operativa en las unidades paracaidistas se relaciona íntimamente con el salto ya que, tras este, el personal debe estar preparado para el combate.

Para intentar llevar a cabo una mejora en dicha eficacia operativa, se identificaron las principales respuestas psicofisiológicas del cuerpo humano en la realización de un salto. Asimismo, se determinaron qué factores, modificables y no, afectan a dichas respuestas. Para la obtención de estos datos se llevó a cabo un estudio en el que se realizaron mediciones tanto antes como durante el salto.

Todas estas mediciones, tras un análisis estadístico, permitieron extraer una serie de hallazgos. En primer lugar, se ha visto que durante el salto se produce un aumento tanto en las pulsaciones como en el estrés y la ansiedad, además de un empeoramiento en las funciones ejecutivas. En segundo lugar, se han observado una serie de factores que afectan a las respuestas psicofisiológicas del salto: la vida sedentaria, número de saltos, tiempo entre saltos y número de años en las Fuerzas Armadas.

Finalmente, se ha propuesto una serie de acciones que pueden modificar los factores que influyen en las respuestas psicofisiológicas del salto. Estas acciones son: promover un estilo de vida saludable, pudiendo introducir pausas activas en actividades sedentarias; reducir el tiempo entre saltos, aumentando de esta manera también el número de saltos que realizan; y, promover un entrenamiento tipo HIIT (*High Intensity Interval Training*), ya que este entrenamiento produce una serie de adaptaciones corporales que podrían ayudar a la realización del salto.

Palabras clave: eficacia operativa, estilo de vida saludable, HIIT, respuestas psicofisiológicas, salto paracaidista.

1. Introducción

1.1 Eficacia operativa en unidades paracaidistas

La eficacia operativa se puede definir como la capacidad que tienen las unidades para entrar en combate a través de la utilización de sus medios y llevar a cabo la misión encomendada. Es por esto que es importante definir qué significa eficacia operativa en las unidades paracaidistas.

Cuando se habla de eficacia operativa en las unidades paracaidistas se debe tener en cuenta que en el cumplimiento de la misión se introduce un factor diferente a otras unidades, el salto paracaidista. Es por esto que la eficacia operativa de la BRIPAC está plenamente relacionada con el salto.

La BRIPAC necesita que, tras la realización de un salto, lo cual es considerado como una situación muy demandante y estresante, sus soldados estén en plenas condiciones para el combate y para el cumplimiento de la misión.

1.2 Formación física en el Ejército de Tierra (ET)

La formación física (FF) siempre ha sido un aspecto de gran importancia para los ejércitos, de hecho, etimológicamente hablando, la palabra ejército y ejercicio provienen de la misma palabra en latín, «exercitus».

La FF es de tal relevancia que se refleja dentro de las Reales Ordenanzas (RROO) de las Fuerzas Armadas (FAS). Las RROO forman un código deontológico que comprende los principios éticos y las reglas de comportamiento del militar español. En particular, el artículo 40, relacionado con el cuidado de la salud dice así: «Considerará la educación física y las prácticas deportivas como elementos básicos en el mantenimiento de las condiciones psicofísicas necesarias para el ejercicio profesional y que, además, favorecen la solidaridad y la integración» [1].

Del mismo modo, es tal la importancia de la FF en el ET que, para cumplir el art. 40 de las RROO y garantizar que los militares tengan una buena condición física (CF), existe el Test General de la Condición Física (TGCF) [2].

Finalmente, la implementación de la FF en las unidades se ve reflejada en una hora diaria por la mañana. Del mismo modo, tanto las FAS, como el ET y las unidades fomentan la realización de competiciones deportivas que ayudan a los militares a interesarse en su forma física y que les permite llevar a cabo más horas de entrenamiento diarias que, por tanto, contribuirán a la mejora de su salud, así como de su rendimiento físico y mental en el trabajo. Todo esto, en última instancia, podrá repercutir positivamente en la eficacia operativa de aquellas misiones o tareas que les sean encomendadas.

1.3 Equipo y salto paracaidista

Antes de hablar del salto paracaidista, es necesario exponer de qué se compone el equipo que se lleva durante el salto: casco (alrededor de 2 kg), chaleco anti-fragmentos con portacargadores y cargadores de HK G-36 (7-12 kg), mochila de combate o mochila «Altus» con el fusil HK G-36 C acoplado (15-33 kg), paracaídas y paracaídas de reserva (3-5 kg).

Como observamos, el personal paracaidista porta una gran cantidad de material, lo que suma un gran peso al salto. Esto supone una mayor velocidad de caída, por lo que los paracaidistas deben estar completamente preparados y entrenados para no cometer fallos que puedan desembocar en una lesión o accidente mayor.

Una vez expuesto el material, se debe explicar cómo transcurre un salto paracaidista y cuáles son los hitos previos y posteriores al mismo. Una vez los paracaidistas llegan a la base aérea, se deben equipar, con los paracaídas y la mochila, y deben esperar hasta que les permitan embarcar en el avión (unas dos horas de espera normalmente). Posteriormente al embarque, los paracaidistas son trasladados a la zona de caída donde, si las condiciones meteorológicas son buenas, saltarán.

El salto, debido a que normalmente es realizado a baja altura (400-500 metros), es de corta duración, unos 40 segundos. Durante este tiempo el paracaidista debe realizar una serie de pasos: ver que las cuerdas de su paracaídas no están enredadas, encontrar a su binomio, observar si va a colisionar o si tiene debajo a algún otro paracaidista, soltar el bulto y prepararse para tomar tierra.

Como se puede observar, los paracaidistas tienen que llevar a cabo numerosas acciones de alta complejidad técnica de una forma correcta en un periodo muy corto de tiempo y manteniendo la calma. Por todo esto, es de gran importancia conocer e identificar qué respuestas psicofisiológicas se dan durante el salto y qué factores están relacionados con estas respuestas para así poder tenerlos en cuenta y poder mejorarlos.

1.4 Objetivos

Los objetivos de este estudio fueron: (1) identificar las principales respuestas psicofisiológicas que afectan durante el salto paracaidista, (2) conocer qué factores afectan a dichas respuestas y (3) determinar una propuesta de forma de vida y entrenamiento que permita mejorar las respuestas al salto.

2. Materiales y métodos

En este proyecto, el cual forma parte del trabajo de fin de grado realizado en el marco de las prácticas externas en la BRIPAC, participaron 37 paracaidistas, a los cuales se les realizó una entrevista individual, mediciones en estado basal y en situación de

estrés, y mediciones de los posibles factores modificables. A partir de estas mediciones se creó una base de datos para su posterior análisis estadístico.

2.1 Entrevista individual con los paracaidistas

En esta entrevista se les entregó un cuestionario a partir del cual se obtuvieron los datos personales que podrían ser factores no modificables que afectarían al salto: edad, altura (en metros), empleo, años en las FAS y años en la BRIPAC.

Asimismo, de este cuestionario se obtuvieron algunos de los diferentes factores modificables que podrían afectar a las respuestas psicofisiológicas que se observarían posteriormente en el salto paracaidista: peso, número de saltos, número de saltos en los últimos tres meses, tiempo (en meses) transcurrido entre los dos últimos saltos, número de horas sentado al día, horas de entrenamiento semanal y última marca del TGCF.

2.2 Mediciones en estado basal y en situación de estrés

En este apartado se muestran las mediciones que se realizaron en estado basal y en situación de estrés para así poder llevar a cabo posteriormente un análisis comparativo de los datos. Las variables seleccionadas, a partir de entrevistas con expertos, fueron:

- Frecuencia cardíaca. Como indicadora de la respuesta fisiológica del salto. La medición de las pulsaciones por minuto (ppm) se llevó a cabo utilizando un reloj Polar M400® y un sensor de frecuencia cardíaca y se tomó en dos momentos: nada más despertarse, sin haberse levantado de la cama; y durante el salto.
- Estrés psicológico y valoración de las funciones ejecutivas. Como indicadoras de la respuesta psicológica del salto. Para la medición del estrés se utilizó el cuestionario validado STAI (*State-Trait Anxiety Inventory*) [3] y para la valoración de las funciones ejecutivas se utilizó el test validado TESEN (test de los senderos) [4]. Ambos test se realizaron en estado basal y en estado de estrés, en el área de embarque.

2.3 Medidas de los posibles factores modificables

En este apartado se llevó a cabo la medición de una serie de factores modificables que podrían afectar a las respuestas psicofisiológicas. Los posibles factores seleccionados fueron los siguientes:

- Índice de masa corporal (IMC). Método utilizado para estimar la grasa corporal de una persona. Se calcula como el peso partido de la talla al cuadrado (kg/m^2).
- Bioimpedancia eléctrica. Dado que el IMC es un cálculo de primer nivel y tiene numerosas limitaciones, se procedió a utilizar la bioimpedancia para calcular la grasa corporal de los participantes.

- O₂ en sangre. Esta medición permite conocer si alguno de los participantes posee alguna enfermedad respiratoria relacionada con la saturación de O₂ que pudiera afectar en el salto.
- Dinamometría manual. Test que permite medir la fuerza muscular isométrica máxima de las extremidades superiores, la cual se relaciona con diversos aspectos de salud y rendimiento.
- TGCF. Test que muestra la CF de un militar a través de cuatro pruebas (flexo-extensiones de brazos en suelo, extensiones de tronco, 6.000 m lisos y un circuito de agilidad-velocidad), asignando un máximo de 100 puntos a cada prueba.

2.4 Análisis estadístico

Una vez obtenidos todos los datos y refundados en un Excel, se realizó un análisis estadístico utilizando el programa Statistical Package for the Social Sciences (SPSS) en su versión 22.0.

Se ha fijado un nivel de significación $P < 0,05$ y las variables cuantitativas se muestran como la media $\pm$ la desviación estándar. Para comprobar la normalidad de las variables, se llevó a cabo la prueba de Kolmogorov-Smirnov, asumiendo finalmente el teorema del límite central (al tener una muestra mayor de 30). Para ver las diferencias en las variables que fueron registradas antes y durante el salto paracaidista se realizó una prueba t de muestras relacionadas. Asimismo, la relación entre variables se estudió mediante correlaciones parciales de Pearson, ajustando por la covariable edad.

3. Resultados y discusión

3.1 Participantes

En el estudio participaron de manera voluntaria 37 paracaidistas (4 cuadros de mando y 33 caballeros legionarios paracaidistas), todos varones con una media de edad de $25,9 \pm 3,3$ años.

Cabe destacar que la muestra está formada por gente joven, con pocos años de experiencia en las FAS (una media de 3,3) y, debido al gran tiempo que suele transcurrir entre saltos (4,2 meses de media) la mayoría no llevan realizados muchos saltos (17 de media). Asimismo, tienen una grasa corporal considerada sana (13,9 de media) y altas marcas en el TGCF (una puntuación media de 351,4).

3.2 Resultados de los efectos psicológicos en el salto

En la tabla 1 se muestran como varían las respuestas psicofisiológicas antes y durante el salto:

	M ± DE	M ± DE	P
	RELAJADO (basal)	SALTO (estrés)	
Frecuencia cardiaca (ppm)	55,2 ± 6,6	169,8 ± 17,4	< 0,001
STAI A-E	8,8 ± 6,6	15,7 ± 8,1	< 0,001
TESEN tiempo 2 (min)	1,2 ± 0,2	1,4 ± 0,2	< 0,001
TESEN fallos 2	0,1 ± 0,4	0,2 ± 0,4	0,375
TESEN tiempo 3 (min)	1,3 ± 0,3	1,5 ± 0,3	< 0,001
TESEN fallos 3	0,4 ± 0,9	0,3 ± 0,7	0,573

Tabla 1. Análisis de los factores psicofisiológicos en estado basal y estado de estrés (salto paracaidista).

Como se puede observar en la tabla, durante el salto, las pulsaciones de los paracaidistas aumentan hasta situarse en una media de 170 ppm. Según el libro *Sobre el Combate* [5], que estudia las respuestas psicofisiológicas del cuerpo humano ante situaciones de estrés altas, a partir de 175 ppm, el cuerpo entra en lo denominado como la «fase negra». Al entrar en esta fase el cuerpo sufre un deterioro del proceso cognitivo, perdidas de visión periférica, de percepción de profundidad y de visión cercana, lo que afecta negativamente durante el salto.

Asimismo, observando el test TESEN observamos cómo los paracaidistas sufren un aumento de su tiempo de reacción y una mayor lentitud en sus funciones ejecutivas. Finalmente, observamos cómo existe un alto incremento de la ansiedad-estado reflejado en el test STAI. Aunque es cierto que niveles ligeros de estrés pueden ser beneficiosos para el rendimiento físico y mental, unos valores como los que se pueden observar en el salto pueden ser contraproducentes (ver figura 1) [6].



Figura 1. Rendimiento físico-psicológico.
Fuente: Entrenamientos en ambientes extremos 2.

3.3 Resultados de los factores que afectan a las respuestas psicofisiológicas

De las correlaciones parciales de Pearson se han obtenido varias tablas que son de interés para el estudio de este proyecto y permiten cumplir con el objetivo número 2.

	Horas sentado al día	
	r	P
STAI A-E salto	0,585	0,035
Ppm medias del salto	0,647	0,017

Tabla 2. Relación entre las horas de estar sentado al día y la A-E y ppm medias del salto.

En esta tabla podemos observar cómo las horas que pasa sentado un paracaidista están correlacionadas positivamente con la ansiedad y con las ppm medias que el paracaidista experimenta durante el salto. Esto puede provocar malas consecuencias durante el salto, al existir una mayor probabilidad de entrar en la «fase negra» o de experimentar un nivel de estrés contraproducente.

	Años en las FAS	
	r	P
STAI A-E salto	-0,466	0,029
TESEN relajado fallos 2	-0,444	0,039
TESEN salto fallos 2	-0,473	0,026

Tabla 3. Relación entre el número de años en las FAS y la ansiedad y funciones ejecutivas.

La tabla superior muestra cómo existe una correlación inversa entre el número de años en las FAS y la A-E en el salto y los fallos en el TESEN (tanto en estado relajado como en el momento previo al salto). Por todo esto, podemos llegar a la conclusión de que a más años en las FAS, menos problemas debería tener un paracaidista en el salto ya que experimenta menos estrés y una menor cantidad de errores de sus funciones ejecutivas.

	Núm. saltos últimos 3 meses		Tiempo entre los 2 últimos saltos	
	r	P	r	P
Ppm medias del salto	-0,453	0,034	0,472	0,027

Tabla 4. Relación del número de saltos en los últimos tres meses y el tiempo entre los dos últimos saltos con las ppm medias en el salto.

Finalmente, tal y como se observa en la tabla 5, la experiencia en los saltos se relaciona con una menor frecuencia cardíaca, lo que podría evitar sufrir problemas de visión

al no superar las 175 ppm. Asimismo, también podemos ver en la tabla superior que, además de aumentar el número de saltos, es importante reducir al máximo el tiempo entre ellos.

3.4 *Propuestas para la mejora de la respuesta del personal al salto*

Una vez estudiados todos los resultados obtenidos de las diferentes mediciones y análisis realizados se proponen tres acciones que podrían mejorar las respuestas psicofisiológicas que afectan a los paracaidistas en los saltos.

Por un lado, ya que las horas de estar sentado al día tienen impacto en las ppm medias y la A-E durante el salto, se propone una modificación del estilo de vida del paracaidista, intentando que sea lo menos sedentario posible. Para ello, por ejemplo, para cuando el personal de las FAS necesite permanecer muchas horas sentado (limpieza de armamento, reuniones, etc.), se debería fomentar la inclusión de pausas activas.

Por otro lado, debido a que la experiencia en saltos paracaidistas y el tiempo entre estos está relacionado con una menor frecuencia cardiaca durante el salto, las unidades paracaidistas deberían tratar de aumentar el número de saltos que realiza su personal. Asimismo, sería recomendable reducir el tiempo entre saltos.

Finalmente, debido a que durante el salto el personal suele tener unas pulsaciones altas que afectan a la realización del mismo, se propone llevar a cabo entrenamientos tipo HIIT. El uso de este tipo de ejercicio, cada vez más extendido, es fruto de las adaptaciones corporales que surgen de este tipo de ejercicios (muy útiles en el salto paracaidista) como es un aumento de las ppm máximas, de la fuerza explosiva de las extremidades inferiores (lo que puede ayudar a evitar lesiones); así como un descenso del peso corporal [7].

4. Conclusiones

De este proyecto se han conseguido extraer varios hallazgos. En primer lugar, *cuáles son las respuestas psicofisiológicas que se observan durante el salto*. Gracias a las mediciones tomadas en estado basal y en el salto se ha podido descubrir de qué manera varía el cuerpo humano en la fase de un salto paracaidista. Especialmente se ha visto un aumento de las pulsaciones y del estrés y la ansiedad, haciendo que las funciones ejecutivas fallen y sean más lentas.

En segundo lugar, se han identificado *qué factores*, algunos modificables y otros no, *se relacionan con dichas respuestas psicofisiológicas*. Como factores modificables que afectan se han hallado tres: horas de estar sentado que pasa al día un paracaidista, saltos en los últimos tres meses y tiempo entre los dos últimos saltos. Por otro lado, otro factor influyente en las respuestas psicofisiológicas del salto, pero no modificable, es el número de años que el sujeto lleva formando parte de las FAS.

Finalmente, se han propuesto *una serie de acciones a llevar a cabo para contribuir a la mejora de la respuesta de los sujetos al salto y, por tanto, a la eficacia operativa del mismo*. Las acciones propuestas han sido: promover un estilo de vida saludable (introduciendo descansos activos en actividades sedentarias), realizar entrenamientos tipo HIIT y reducir el tiempo entre saltos, ayudando de esta manera a aumentar el número de saltos que realiza el personal.

En conclusión, la mejora de las respuestas psicofisiológicas de los saltadores ante una situación de estrés como es el salto, pasa no tanto por un aumento de las horas de formación física o un aumento de la carga en estas sesiones, sino que debe ser fruto de dos cosas: una vida saludable y no sedentaria; y un mayor número de saltos y un menor espaciamiento entre los mismos. Además, se propone la implantación de entrenamientos tipo HIIT como posible sistema de FF encaminado a la mejora de la eficacia operativa del salto.

Agradecimientos

Este proyecto pudo ser realizado gracias al apoyo recibido por parte del profesorado perteneciente al Centro Universitario de la Defensa y a la implicación y voluntariedad del personal perteneciente a la Bandera «Roger de Flor» I de Paracaidistas.

Referencias

1. Reales Ordenanzas de las Fuerzas Armadas. Real Decreto 96/2009. Publicado en el BOE número 33, del 7 de febrero de 2009.
2. Mando de Adiestramiento y Doctrina. Instrucción Técnica 03/15, actualización 2019, Test General de la Condición Física (TGCF).
3. R. D. Spielberger, R. L. Gorsuch, R. E. Lushene, *STAI cuestionario de ansiedad estado-rasgo*. Manual, 9.ª ed. Madrid: tea, 2015.
4. A. Portellano, R. Martínez Arias, *TESEN test de los senderos para evaluar las funciones ejecutivas*. Manual. Madrid: tea, 2014.
5. Teniente coronel D. Grossman, L. W. Christensen, *Sobre el combate*, 2.ª ed. Nueva York: Melusina, 2014, pp. 90-93.
6. L. M. López Mojares, *Entrenamiento para ambientes extremos 2*. Madrid: Ministerio de Defensa, 2015, p. 15.
7. J. G. Tornero Aguilera, V. J. Clemente Suárez, «Resisted and Endurance High Intensity Interval Training for Combat Preparedness», *Aerospace Medicine and Human Performance*, vol. 90, no. 1, pp. 32-35, Ene, 2019.

Modelización de cargas explosivas para el diseño de estructuras blindadas mediante simulación

Negro, Alberto^{1,*}; Carlón, M. Julia¹

¹ Fundación Cidaut. Parque Tecnológico de Boecillo, pl. Vicente Aleixandre Campos n.º2, 47151, Boecillo (Valladolid). info@cidaut.es

* Autor principal; albneg@cidaut.es

Resumen: la normativa STANAG 4569 describe los requerimientos exigibles a los vehículos blindados de cara a obtener un determinado nivel de protección de los ocupantes frente a diferentes tipos de amenazas del tipo tanto balístico como explosivo. El diseño de estructuras capaces de resistir cargas extremas es un reto tecnológico de gran magnitud, debido a la complejidad de los fenómenos físicos implicados en los eventos de impacto a alta velocidad. Es por ello que el diseño de protecciones y blindajes novedosos implica la utilización de materiales de última generación, cuya realización de prototipos tiene un alto coste asociado. Por otro lado, las instalaciones necesarias para realizar ensayos que permitan la validación de esas estructuras ante las cargas para las que han sido diseñadas resultan asimismo sumamente complejas y costosas.

En este contexto, las herramientas de simulación numérica pueden convertirse en un potente aliado en las tareas de diseño y optimización de estructuras ante cargas extremas. La clave de la calidad de los resultados de una simulación por ordenador es que sea capaz de reflejar de forma fidedigna tanto las condiciones de carga que actúan sobre una estructura como la respuesta mecánica de los materiales y los mecanismos de fractura que tienen lugar al ser sometidos a cargas tan sumamente severas. El presente artículo se centra en la modelización numérica del problema físico de condiciones de carga extremas del tipo explosivo, analizando y comparando diferentes métodos y algoritmos existentes en los paquetes de *software* comerciales disponibles en la industria.

Palabras clave: ALE, blindaje, explosión, materiales, PBM, simulación.

1. Introducción

1.1 Impactos a alta velocidad. Ámbitos de aplicación

Actualmente el diseño de estructuras para resistir cargas extremas es algo cada vez más habitual, no solo en el caso de aplicaciones militares o de defensa, sino también en el caso del sector aeronáutico o en la ingeniería civil. Además, las soluciones a implementar resultan técnicamente más complejas, ya que en muchas ocasiones es necesario emplear materiales de última generación bien sea por su ligereza o por sus prestaciones mecánicas.

Así, en el ámbito de la defensa y antiterrorismo se busca por ejemplo diseñar blindajes en materiales ligeros como los composites, pero que incluyan protección antiminas, balística e IED, diseñar protecciones y equipamientos militares con materiales como el Kevlar e incluso garantizar la integridad de las instalaciones e infraestructuras (edificios, puentes, plantas nucleares,...) ante ataques.

En el ámbito de la aeronáutica y las aplicaciones aeroespaciales, se busca conseguir estructuras (fuselaje, alas, hélices, parabrisas, trenes de aterrizaje...) con capacidad para soportar impactos a alta velocidad como por ejemplo impactos de pájaro (*bird strike*), de fragmentos de motor o hélices, de granizo, de trozos de hielo desprendidos del fuselaje o de objetos que se encuentren en la pista de aterrizaje, así como resistir choques y aterrizajes de emergencia, explosiones en cabina, etc.

Esta complejidad también se refleja en el ámbito civil, por ejemplo en el diseño de palas para aerogeneradores capaces de resistir impactos a alta velocidad, soldadura y conformado de materiales por explosión (*shock*), diseño de instalaciones especialmente sensibles frente a accidentes (centrales nucleares), etc.

1.2 Problemática asociada a las cargas explosivas

En el caso de estructuras sometidas a los efectos de una carga explosiva, las ondas resultantes de una explosión se caracterizan por presentar discontinuidades en las magnitudes de presión, densidad, velocidad y temperatura a lo largo de una onda de choque, a diferencia de las ondas acústicas [1]. Además las ondas de choque viajan a velocidades superiores a la velocidad del sonido en el medio en el que se transmiten y a medida que el frente de onda avanza e interacciona con aquellos elementos que encuentra a su paso, se producen fenómenos de atenuación, reflexión, etc. El comportamiento de materiales explosivos, la interacción entre explosivos y estructuras y su modelización mediante técnicas numéricas se describen en [2].

Hay que tener en cuenta que cuando se diseña una estructura o blindaje frente a cualquiera de los niveles de protección definidos en la norma STANAG 4569 habrá que considerar tanto efectos debidos a carga explosiva como los debidos a impactos balísti-

cos, además de los resultantes como combinación de ambos tipos de carga (como se da, por ejemplo, en el caso de los IED).

1.3 La simulación numérica como herramienta

El diseño y optimización de estructuras ante estas situaciones requiere de la realización de test de validación en instalaciones de ensayo altamente complejas en las que verificar el correcto comportamiento de dichas estructuras ante las cargas para las que han sido diseñadas. Debido a la sensibilidad de los ensayos que es necesario realizar (manejo de cargas explosivas, proyectiles...), en la mayoría de las ocasiones es necesario contar con licencias específicas para poder llevar a cabo estas pruebas, además de disponer tanto de un personal como de un equipamiento e instrumental altamente especializados, lo que reduce el número de laboratorios capaces de realizar estos ensayos. Además, la fabricación de prototipos siempre es una tarea costosa.

Es en este punto en el que la simulación numérica por ordenador se convierte en una potente herramienta a tener en cuenta a la hora del diseño de estructuras en situaciones de carga altamente dinámica. Esta herramienta ya se emplea de manera exhaustiva para el desarrollo de nuevos productos en la mayoría de sectores industriales, así como para todo tipo de investigaciones avanzadas dentro del marco académico y científico, gracias al desarrollo y evolución de los sistemas computacionales, tanto a nivel de la capacidad del *hardware* como al desarrollo de nuevos algoritmos y *softwares* especializados de muy diferentes características y por tanto adaptados para poder trabajar sobre diferentes aplicaciones prácticas.

En el mercado existen diferentes paquetes comerciales de simulación que permiten afrontar esta problemática. El trabajo presentado en el presente artículo ha sido desarrollado mediante el *software* LS-Dyna [3].

2. Materiales y métodos

En todo modelo de simulación numérica es necesario hacer una correcta discretización espacial del problema (mallado). Aunque una malla convencional (lagrangiana) es adecuada para la modelización de las estructuras implicadas en situaciones de deformación a alta velocidad, la complejidad de los fenómenos físicos ocurridos durante las situaciones de carga extrema hace que, para la modelización de las cargas explosivas, sean necesarias otras técnicas de simulación específicas, como los métodos de discretización eulerianos o eulerianos-lagrangianos (ALE), o los métodos «*meshless*» (partículas).

La combinación de varias técnicas de simulación en un mismo modelo requiere de complejas interacciones fluido-estructura. Además, debido a la violencia de las cargas, se producen condiciones de impacto, contacto y comportamiento de los materiales alta-

mente no lineales. Todo ello genera la necesidad de disponer de modelos de simulación extremadamente detallados que implican un alto coste computacional.

Por otro lado, en los modelos de simulación numérica la calidad de los resultados obtenidos siempre depende directamente de la correcta definición de los materiales que se implementan en los mismos, convirtiendo la caracterización de los materiales que se vayan a emplear en el punto clave del proceso de diseño mediante herramientas de simulación.

En situaciones de carga extrema, la definición correcta de los materiales en simulación incluye, por un lado, la definición de modelos constitutivos de material que permitan reflejar tanto el efecto de endurecimiento del material debido a la velocidad de deformación (*strain rate*) como el ablandamiento debido al calentamiento del material causado por la elevada velocidad a la que se deforma. Por otro lado, en estas situaciones de impacto, las presiones internas que se producen en los materiales exceden la tensión de fluencia en dos o más órdenes de magnitud. El comportamiento de los materiales en estas condiciones, incluso tratándose en muchos casos de metales y aceros de alta resistencia, es hidrodinámico, es decir, se comportan como si fuesen fluidos, por lo que es necesario incluir una ecuación de estado (EOS) para su correcta definición en los modelos numéricos. Además, debido a las condiciones críticas de la carga, es de especial importancia reflejar en los modelos la mecánica de fractura de los materiales a altas velocidades.

En el presente artículo se han empleado dos metodologías de simulación diferentes para evaluar la misma situación de carga (explosión de una mina enterrada de 8 kg de TNT bajo un vehículo blindado) con el fin de identificar sus ventajas y sus puntos débiles. Para el explosivo, por un lado se emplea una modelización ALE, mientras que en el otro caso se evalúa una modelización de partículas.

En ambos casos la modelización del vehículo es la misma, realizándose mediante elementos de tipo sólido (se han empleado unos 855.000 elementos), de forma que cada una de las partes que lo componen se ha discretizado empleando un mínimo de tres elementos sólidos a lo largo de su espesor para capturar correctamente la flexión de las piezas durante la carga. El material utilizado en la estructura es fundamentalmente acero 4340, el cual se ha modelizado mediante una ley de Johnson-Cook (MAT_015 en LS-Dyna), incluyendo parámetros que definen la rotura del material cuyos datos han sido obtenidos de [4] y una ecuación de estado de Gruneisen [5].

Las particularidades de los principales métodos de modelización de explosivos se detallan a continuación.

3. Modelo LOAD_BLAST_ENHANCED (LOAD_BLAST)

Ambos métodos de modelización de explosivos en LS-Dyna están basados en las bases del *software* CONWEP (*CONventional WEaPons*) desarrollado en 1988. Se trata de

modelos de explosión empíricos desarrollados para TNT. Si la detonación a estudiar es de otro explosivo diferente, es necesario calcular la carga de TNT equivalente [6].

El método `LOAD_BLAST_ENHANCED` es una mejora del original (`LOAD_BLAST`) que permite la presencia de múltiples cargas explosivas (distintos orígenes) en el mismo cálculo, permite la modelización de la reflexión de la explosión contra el suelo en el caso de explosiones aéreas (mientras que el modelo original solo permitía explosiones aéreas sin reflexiones o explosiones de cargas semiesféricas apoyadas sobre el suelo), además permite tener en cuenta la fase negativa de la presión típicamente observada en la presión de las ondas explosivas (ecuación de Friedlander) y la modelización de cargas explosivas moviéndose a altas velocidades [7].

Esta modelización podría usarse como primera aproximación a la evaluación de cargas explosivas en las primeras fases de diseño o como condición de contorno de un modelo más detallado de ALE en el que se pretenda estudiar la propagación de las ondas de choque. En este estudio no se ha empleado por su simplicidad.

4. Modelo ALE

En los códigos numéricos existen diferentes algoritmos para determinar la relación entre la deformación de un medio y la malla de elementos finitos, siendo las dos descripciones clásicas más utilizadas las denominadas lagrangiana y la euleriana. En la descripción lagrangiana los nodos de la malla se deforman junto con el material, lo cual puede dar lugar a problemas de inexactitud y convergencia del cálculo cuando existen grandes deformaciones que den lugar a una gran distorsión de la malla. Por el contrario, la descripción euleriana utiliza una malla fija en el espacio, lo que le permite gestionar grandes deformaciones del material, pero el seguimiento de la superficie del mismo se hace más complejo, y debiéndose considerar advección sobre los límites de los elementos. Es por ello que la descripción lagrangiana se utiliza para el análisis estructural, mientras que la euleriana tiene mayor aplicación para el estudio de dinámica de fluidos.

El método ALE (*Arbitrary Lagrangian Eulerian*) permite combinar las ventajas de ambas descripciones, ya que en este caso los nodos pueden moverse como en la descripción lagrangiana o bien permanecer fijos como en la euleriana, así como moverse sin seguir el material continuo para minimizar la advección sobre los contornos y obtener de esta manera una solución más precisa. En el caso considerado de la simulación del efecto de un explosivo sobre una determinada estructura o blindaje, la aproximación utilizada consiste en modelizar como lagrangianas las partes estructurales y como eulerianas el aire circundante y el explosivo, así como el suelo en que la mina puede estar potencialmente enterrada, debiendo existir asimismo un algoritmo que posibilite el acoplamiento entre ambas definiciones y que permita la transferencia de momentum y energía entre ambas modelizaciones.

Los datos para la modelización del aire y explosivo han sido obtenidos de [8].

5. Modelo de partículas

El modelo de partículas para explosivos (PBM de sus siglas en inglés *Particle Blast Method*) es una extensión del CPM (*Corpuscular Particle Method*), desarrollado para la simulación de despliegues de *airbags* basado en la teoría cinética molecular (KMT), que es el estudio de las moléculas de gas y su interacción a nivel macroscópico. A diferencia del CPM, el PBM es capaz de simular un gas que se rija por la ley real a presiones y temperaturas extremas [9].

Cada partícula representa varias moléculas, y se modelan como esferas rígidas para facilitar la definición de los contactos entre ellas y con la estructura. En el caso de gases monoatómicos, la colisión entre partículas es perfectamente elástica, mientras que para el caso de gases no monoatómicos, para cada partícula se establece un balance entre la energía cinética y la energía de spin-vibracional que define la compresibilidad del gas asumiendo el equilibrio térmico local [6].



Figura 1. Comparativa de resultados de simulación de la explosión de una mina enterrada mediante PBM con y sin aire del entorno.

La modelización tanto del aire del entorno como del explosivo (volumen y características) se define directamente en la propia carta (**PARTICLE_BLAST*). En el caso de la simulación de minas enterradas, la modelización del aire puede ser omitida sin afectar significativamente a los resultados [10], dando lugar a modelos más rápidos y fáciles de analizar posteriormente. Sin embargo en el caso de explosiones de cargas aéreas sí que tiene importancia su modelización.

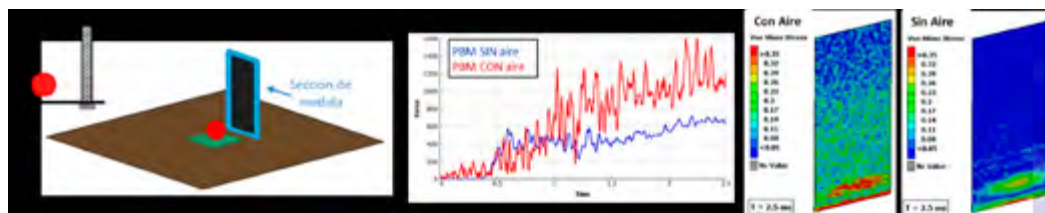


Figura 2. Influencia de la modelización del aire del entorno en la simulación de una explosión aérea contra una pared mediante el método de partículas (PBM).

En ambos casos, la simulación de minas enterradas permite incluir la modelización del suelo y su interacción con el explosivo, bien sea mediante la modelización del mismo con ALE o en el caso de PBM mediante los denominados *Discrete Elements*. En este estudio no se ha tenido en cuenta esta interacción con el suelo, sino que se ha modelado el caso de una mina confinada para tratar de eliminar la influencia que pueda tener la diferente modelización del suelo.

6. Resultados y discusión

La definición del modelo de ALE es notablemente más compleja que la del modelo PBM, no solo porque requiera de más parámetros y controles, sino por la dificultad de ajustar correctamente la interacción fluido-estructura. Sin embargo, la creación de un modelo de partículas estable es bastante más sencilla.

Método	Parámetros	Tiempo de cálculo [segundos]
PBM	No aire, 100k partículas	5.640
ALE	Sí aire, 608.080 elementos ALE	9.793

Tabla 4. Resumen de parámetros y tiempos de cálculo de las simulaciones.

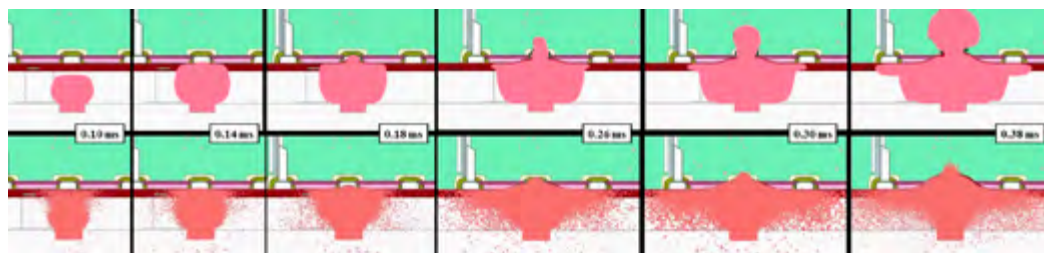


Figura 3. (a) Resultados del modelo de mina con ALE. (b) Resultados del modelo con PBM.

En la figura 4 puede verse que la cinemática del explosivo en ambos casos es algo diferente, por lo que la deformación que se obtiene del vehículo varía en ambos modelos. El modelo ALE presenta un comportamiento de la estructura del vehículo más frágil que en el caso del PBM, lo que conduce a una rotura más temprana y también a una deformación de la panza del vehículo más localizada.

Cada una de las modelizaciones tiene sus ventajas y sus inconvenientes. Mientras que el modelo ALE es más complejo de definir y de controlar, principalmente en lo que se refiere a la interacción entre los fluidos (explosivo, aire, suelo) y la estructura del vehículo, el modelo de partículas es un modelo bastante estable y sencillo de definir. Por otro lado, el tiempo de cálculo en el caso del modelo de partículas sería hasta un 40 %

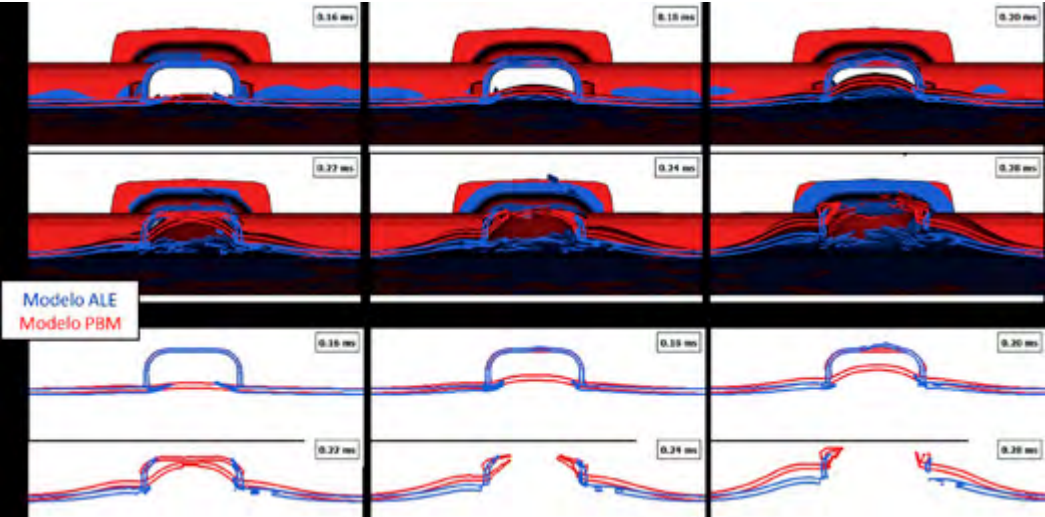


Figura 4. Comparativa de la sección del vehículo en el centro de la carga explosiva para ambos modelos.

menor que en el caso del modelo ALE. Sin embargo el método de partículas, aunque es bastante prometedor, es un método de reciente creación [9], por lo que está aún sujeto a posibles mejoras y avances en la modelización. Actualmente no ofrece resultados sobre el comportamiento del explosivo, mientras que el método ALE sí.

Método	Ventajas	Desventajas
PBM	Facilidad a la hora de definir el modelo	-
ALE	-	La definición del modelo es compleja
PBM	-	Método aún en evolución
ALE	Método ampliamente desarrollado y evaluado	-
PBM	Interacción con el entorno fácil de definir	-
ALE	-	Complejidad a la hora de definir la interacción con el entorno
PBM	-	Sin resultados visuales para el comportamiento del explosivo
ALE	Disponibilidad de resultados visuales para el comportamiento del explosivo	-

Tabla 5. Comparativa de ventajas y desventajas de ambos métodos.

Según lo observado en el presente estudio, los modelos de simulación basados en el método de partículas pueden resultar adecuados para las primeras fases de diseño de estructuras ante impacto, dada la sencillez de definición y manejo que presentan, mientras que los modelos de ALE, que aportan una mayor información del comportamiento de los explosivos, pueden ser interesantes para las fases de diseño finales. Una futura línea de continuación del presente estudio podría incluir la validación de ambas técnicas respecto a situaciones reales, dado que para este estudio no se disponen de datos experimentales.

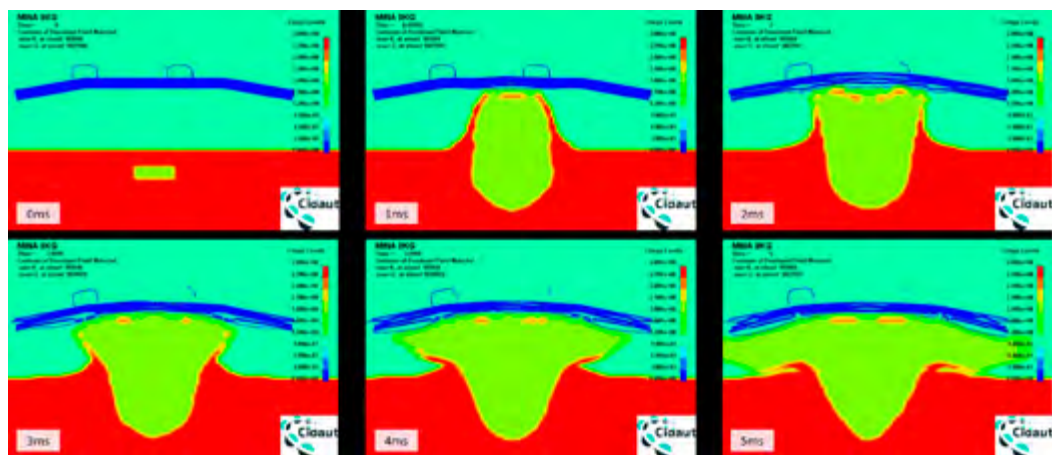


Figura 5. Ejemplo de los resultados obtenidos para las diferentes partes modeladas en ALE en un modelo de explosión, (en este caso explosivo, aire, suelo).

7. Conclusiones

La simulación numérica resulta ser una herramienta muy potente a la hora de diseñar estructuras capaces de soportar cargas altamente dinámicas como las originadas durante la explosión de minas enterradas. La correcta definición del modelo numérico dependerá en gran medida de la habilidad del usuario y de los datos de partida disponibles: datos validados sobre el comportamiento de materiales, definición adecuada de los explosivos y de las interacciones de estos con la estructura en el caso de modelos ALE.

Los métodos actualmente implementados en el *software* LS-Dyna ofrecen opciones adaptables a la complejidad del modelo que se quiera definir y a la información que el usuario desee obtener del cálculo. El método PBM resulta más sencillo y rápido para el usuario, pero debido a su reciente implementación en el *software* presenta todavía algunas limitaciones, sobre todo a la hora de analizar los resultados. Por su parte el método ALE está ya ampliamente desarrollado y consolidado, pero resulta más costoso tanto a nivel de modelización y puesta a punto como a nivel computacional.

Referencias

1. L. Schwer, «Air Blast – Engineering Models», Blast and Penetration Modeling, 2007.
2. J. A. Zukas, *High Velocity impact dynamics*. New York: John Wiley & Sons, Inc., 1990.
3. «Ls-Dyna Keyword Users Manual», Livermore Software Technology Corporation (LSTC), 2016.
4. «Penetration and Perforation», Blast-Penetration Class, LS-Dyna Training Class, 2007.
5. B. Adams, «Simulation of Ballistic impacts on armored civil vehicles», Final Thesis, Dep. Mechanical Engineering, Eindhoven University of Technology.
6. D. Hilding, «Methods for modelling Air blast on structures in LS-Dyna», in 14th German LS-DYNA Forum, October 2016.
7. «Air Blast – Engineering Models», Blast-Penetration Class, LS-Dyna Training Class, 2007.
8. H. B. Rebelo, C. Cismasiu, «A Comparison between Three Air Blast Simulation Techniques in LS-DYNA», in 11th European LS-DYNA Conference, 2017.
9. H. Teng, J. Wang, «Particle Blast Method (PBM) for the Simulation of Blast Loading», in 13th International LS-DYNA Conference, 2014.
10. H. Teng, «Keyword Setup for Particle Blast Method», Livermore Software Technology Corp. Guidelines updated on Jun 24th, 2016.

Análisis de sentimiento en las redes sociales producido por las Fuerzas Armadas españolas

Ortega de los Ríos, Alejandro¹, Fernández Gavilanes, Milagros^{2,*}; Suárez García, Andrés²

¹ Alnte. Juan de Borbón (F-102) Arsenal de Ferrol. Rúa Argüelles, 4, 15401, Ferrol, La Coruña. aortde4@mde.es

² Centro Universitario de la Defensa (CUD). Plaza de España, 2. 36920, Marín, Pontevedra. {mfgavilanes;andres.suarez}@cud.uvigo.es

* Fernández Gavilanes, Milagros; mfgavilanes@cud.uvigo.es

Resumen: la percepción sobre la seguridad, la defensa y las Fuerzas Armadas ha sido un tema no exento de discusión a lo largo de estos años en nuestra sociedad. Esto se debe en gran medida a la historia de España, que ha provocado en la sociedad cierta percepción. Ahora bien, en la actualidad existe la necesidad de superar esas viejas concepciones que hacen necesaria una política de expansión de la cultura de defensa.

El entorno actual dominado por el uso de redes sociales hace que esta labor de difusión del trabajo realizado por los ejércitos deba involucrar activamente a la sociedad civil. Este arduo trabajo de comunicación social, en constante cambio y evolución, ha propiciado que sus lectores escriban sobre sus experiencias personales y expresen sus opiniones. En consecuencia, el análisis de estos comentarios podría, por ejemplo, permitir predecir el nivel de aceptación de determinadas actividades llevadas a cabo por las Fuerzas Armadas [1].

En este trabajo se propone un estudio utilizando diferentes algoritmos de clasificación para predecir el sentimiento en *tweets*, aprovechando una variedad de técnicas de procesamiento del lenguaje natural y características derivadas. Partiendo de las predicciones creadas a partir de estos modelos, se propone la creación de dos metaclassificadores que mejoren los resultados obtenidos de forma individual. Las pruebas realizadas sobre un conjunto de datos real de referencia confirman el desempeño competitivo y la solidez del sistema.

Palabras clave: análisis de sentimiento, *machine learning*, procesamiento del lenguaje natural, metaclassificador.

1. Introducción

Este trabajo es un sucinto resumen del trabajo fin de grado (TFG) del alumno egresado don Alejandro Ortega de los Ríos titulado «Análisis de sentimiento en las redes sociales producido por las Fuerzas Armadas españolas» [2]. Debido a los límites de extensión impuestos por el formato del congreso, se recomienda a la persona interesada en profundizar en el mismo que lo descargue del repositorio institucional Centro Universitario de la Defensa de la Escuela Naval Militar.

En los últimos años, el crecimiento explosivo de las redes sociales ha permitido a la sociedad, y más concretamente a los individuos y organizaciones, escribir sobre sus experiencias personales y expresar opiniones. En este contexto, determinar la opinión de los ciudadanos acerca de las Fuerzas Armadas podría permitir definir mejores estrategias para la comunicación de la cultura de defensa, aunque siendo conscientes que la clasificación automática del sentimiento de estos comentarios es una tarea en general ardua.

Teniendo estos aspectos en cuenta, el campo del análisis de sentimientos (AS) ha recibido una gran atención en los últimos años [3, 4], particularmente debido al auge de las redes sociales, que ha permitido tanto a personas como a organizaciones escribir mediante lenguaje compacto y coloquial sobre sus experiencias y opiniones. Esta nueva forma de comunicación es a su vez una fuente de información extremadamente valiosa. Por ejemplo, Twitter, una de las redes sociales más populares, ha crecido de 5.000 nuevos *tweets* por día en 2007 a 500 millones en 2019 ¹, pasando por menos de 30 millones de usuarios activos mensualmente en 2007 frente a los 330 millones en 2019.

Sin embargo, extraer y analizar la opinión en esos grandes volúmenes de datos resulta tremendamente difícil y costoso, razón por la cual la predicción automática es tan atractiva [5]. AS representa un desafío interdisciplinario que aprovecha una gran variedad de técnicas de procesamiento del lenguaje natural (PLN) y de machine learning (ML) para determinar el sentimiento expresado en los textos y decidir si estos son positivos, negativos o neutros.

2. Estado del arte

PLN es el campo que estudia el lenguaje natural (español, inglés, francés, etc.) desde un punto de vista computacional. A día de hoy existen numerosos ejemplos de *software* basados en PLN como Siri o Cortana que son capaces de imitar el lenguaje humano e interactuar con los usuarios. Como norma general se puede afirmar que ciertas tareas han sido resueltas gracias a la implementación de modelos de PLN. Entre otras podemos

¹ Extraído de «Twitter by the numbers 2020». Disponible en <https://www.omnicoreagency.com/twitter-statistics/>. [Accedido el 25/09/2020].

destacar detectores de *spam* y etiquetado de palabras en una frase (sustantivo, adjetivo, adverbio,...). Por otro lado, existen otras tareas que aún no han sido completamente resueltas pero que están siguiendo un buen progreso. Ejemplos de estas son el AS y la traducción de textos.

ML es el campo de estudio que dota a las máquinas de la capacidad de aprender de la experiencia sin haber sido explícitamente programadas para ello. ML se encuentra íntimamente relacionado con la inteligencia artificial (IA). Existe una gran cantidad de problemas que pueden ser resueltos mediante métodos de programación tradicionales. Se programa un determinado algoritmo que cubra todos los escenarios posibles para realizar la tarea en cuestión. Sin embargo, se pueden encontrar problemas que no solo sean de gran complejidad, sino también que no exista un algoritmo conocido que permita resolverlo o que este solo valga para casos muy específicos. El análisis de sentimiento resulta ser uno de ellos.

3. Sistema desarrollado

El principal objetivo de esta investigación es predecir si un mensaje procedente de una red social de las Fuerzas Armadas expresa un sentimiento positivo, negativo o neutro sin necesidad de la supervisión por parte de un operador. El enfoque que proponemos se fundamenta en una idea sencilla e intuitiva: crear un modelo partiendo de otros modelos. Para ello se utilizaron técnicas de PLN para capturar peculiaridades lingüísticas que luego se utilizaron para mejorar el rendimiento de detección de sentimientos en cada uno de los modelos propuestos. Este enfoque consta de varias etapas que se describirán a continuación.

3.1 Preprocesamiento

Se plantean varios desafíos en PLN según el lenguaje que se use en redes sociales. Uno de los principales problemas es el uso de ciertas características ortográficas y tipográficas en las palabras, como por ejemplo la duplicación de letras y/o palabras; o el uso de términos irreconocibles en un diccionario, como el argot. Por lo tanto, el primer paso es realizar un preprocesado de estos datos, con el objetivo de restaurar, en la medida de lo posible, el lenguaje eliminando expresiones atípicas para minimizar el ruido en etapas posteriores.

Por tanto, se procedió a la eliminación de las URL, saltos de línea, nombres de usuarios, menciones, números y palabras de parada (se trata de palabras que aparecen con gran frecuencia en el texto y que carecen de semántica, tales como «de» o «y»). La eliminación de estas palabras de parada se justifica debido a sus altas frecuencias en los *tweets*, provocando que los algoritmos centren demasiado su atención en ellos creyendo que aportan gran significado al texto cuando no es así.

También se procedió a remplazar las mayúsculas por minúsculas y a las palabras con tilde por su equivalente sin ella. Haciéndolo de este modo, se ha buscado simplificar el proceso, evitando la incorporación de un corrector ortográfico. Por otro lado, se ha considerado necesario sustituir aquellas expresiones ampliamente utilizadas por los usuarios de redes sociales que no se encuentran recogidas en la RAE, tales como el argot, por su equivalente textual. Ejemplos de ello son: «LOL», «WTF», «k wapo» y otros muchos más.

Finalmente, se obtuvo la raíz de cada palabra, con el objetivo de normalizar sus distintas derivaciones léxicas, agrupando familias léxicas en su raíz. Así, el algoritmo considerará las palabras «hacer», «haciendo» e «hice» como una misma unidad: «hacer».

3.2 Extracción de características

Una vez realizado el preprocesado, es necesario extraer las características para entrenar a un modelo. Se utiliza la métrica *Term frequency - inverse document frequency* (*Tf-idf*). Se trata de una medida que expresa el peso de una palabra en el documento. La frecuencia del término (*tf*) es básicamente la salida de un modelo de bolsa de palabras, haciendo referencia al número de veces que aparece en ese documento; mientras que el término (*idf*) consiste en un valor que mide la frecuencia con la que aparece la palabra en el conjunto de documentos, expresado como el logaritmo de la inversa de su frecuencia de aparición.

3.3 Aprendizaje con supervisión

El objetivo de este trabajo es la obtención de un metaclassificador robusto a partir del diseño de tres modelos de clasificación diferentes, tales como *Gaussian Naive Bayes* [6] (se trata de un algoritmo de aprendizaje basado en el teorema de Bayes, en el que se asume que cada una de las variables aleatorias son linealmente independientes), *Support Vector Machine* (SVM) [7] (consiste en un algoritmo de clasificación que representa elementos de un flujo de datos como vectores pertenecientes a un espacio n -dimensional, donde se calcula un hiperplano de dimensión $n-1$ que divide los distintos vectores a un lado u otro del hiperplano, actuando como frontera, de forma que se obtenga una serie de conjuntos disjuntos de vectores, asociando cada uno de ellos a una clase), y *K-nearest neighbors* (KNN) [8, 9] (es un algoritmo de clasificación por comparación con los k vecinos más cercanos, de forma que un elemento será etiquetado con una etiqueta en función de la cantidad de vecinos que pertenezcan a esa clase). Estos clasificadores simples que analizan los textos de entrada generan cada uno de ellos un modelo simple.

Tras diseñar estos tres modelos de clasificación diferentes, se diseña una primera aproximación que busque combinar entre sí las diferentes matrices de predicción. El modo de combinación de estos para obtenerla será considerando la media ponderada de las matrices de los modelos anteriores:

$$= \frac{1}{\alpha_0 + \alpha_1 + \alpha_2} (\alpha_0 X_0 + \alpha_1 X_1 + \alpha_2 X_2)$$

Ecuación 1. Mecanismo de combinación para la obtención de la primera aproximación.

siendo $\alpha_i \in [0,1]$ el valor F-1 del modelo i , y X_i su matriz de predicción. De este modo se da más peso a modelos más robustos en detrimento de los menos fiables. El resultado obtenido en la ecuación 1 es una matriz con coeficientes reales en el intervalo $[-1,1]$.

A continuación, se diseña una segunda aproximación que consiste en utilizar las salidas proporcionadas por los modelos simples como entrada del nuevo sistema, es decir, crear un modelo partiendo de otros modelos. Con ese objetivo, se utiliza un algoritmo de clasificación multiclase SVM, que se nutre de los resultados de los algoritmos anteriores para mejorar la exactitud del nuevo modelo combinando los clasificadores simples. Este algoritmo multiclase sigue una estrategia *One versus Rest*.

4. Evaluación y resultados experimentales

4.1 Conjunto de datos usados como referencia

Para probar el funcionamiento y el rendimiento del sistema desarrollado, se ha procedido a la extracción de un conjunto de datos por medio de las herramientas para desarrolladores que dispone Twitter. Más concretamente, se trata de la librería de *Python* denominada *Tweepy* que permite acceder a la API de Twitter, realizar búsqueda de *tweets* y finalmente su descarga. Es importante señalar que, para evitar la saturación de sus servidores, de acuerdo con [10], Twitter limita a una búsqueda máxima de 1.000 *tweets* cada 15 minutos, considerando hasta 7 días atrás en el tiempo; lo cual es muy restrictivo para cuentas que no publiquen diariamente gran cantidad de *tweets*.

Por tanto, el objetivo de esta fase ha sido extraer 3.000 *tweets* (evitando *retweets*) escritos en español durante diciembre de 2019 y febrero de 2020, tomando las cuentas *@Armada_esp*, *@EjercitoTierra* y *@EjercitoAire*. Se ilustra en la tabla 1 algún ejemplo recuperado.

<i>Tweets</i>	Fuente
Se sabe horarios de visita? E S	Armada
ojala me llevaran de excursion a algun sitio del ejercito E A E A	Armada
mira loca con mis hijosnotemetas	Armada

Tabla 1. Ejemplo de *tweets* recuperados para la cuenta *@Armada_esp*.

Posteriormente, a cada uno de estos *tweets* se le ha asignado manualmente una etiqueta asociada a su sentimiento: *POS* para el caso positivo, *NEU* para el caso neutro y *NEG* para el negativo, quedando su distribución tal y como se observa en la Figura 1. De este modo, este conjunto de datos junto con su etiqueta es lo que se denominará el conjunto de referencia.

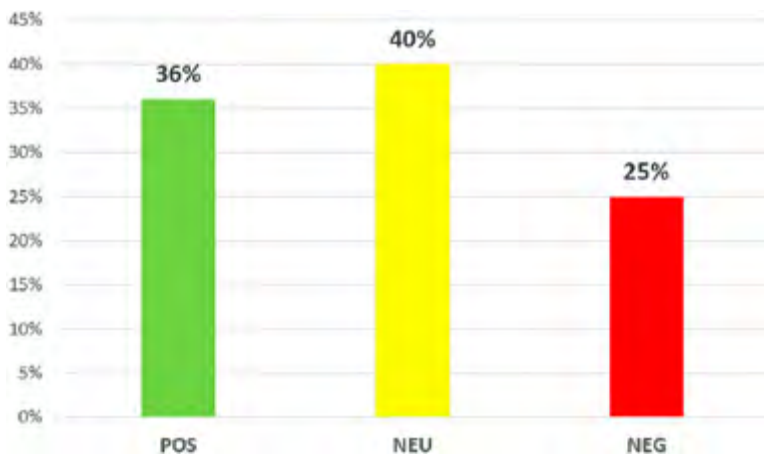


Figura 1: Distribución de la muestra atendiendo a su clase.

En la Figura 1 se observa que existe un cierto sesgo en los *tweets* negativos, aunque este conjunto de datos resulta no ser significativamente desbalanceado.

4.2 Métricas utilizadas

Antes de entrar a valorar los resultados obtenidos, es necesario indicar las métricas empleadas para su estimación:

- Precisión. Razón entre el número de elementos clasificados correctamente dentro de una clase y el total de elementos clasificados dentro de la clase.
- Exhaustividad (*Recall*). Relación entre los documentos clasificados correctamente dentro de una clase y la suma de todos los elementos de la clase.
- Valor F-1 (*F1-Score*). Media armónica de los anteriores. Suele emplearse como referencia para comparar el rendimiento entre modelos.

Como en este caso el problema cuenta con tres posibles etiquetas, estas métricas se estimarán por cada una de estas individualmente y se combinarán para obtener un valor global. A la vista de que el conjunto de datos no se encuentra totalmente balanceado, se seguirá una ponderación en base al peso de cada clase (*weighted-averaging*) [11].

4.3 Resultados para el conjunto de datos

Además de comparar la precisión de cada modelo por separado, se muestran los resultados obtenidos tras realizar validación cruzada (*cross-validation*). Esta técnica es empleada para evitar un sobreajuste² (*overfitting*) o un subajuste³ (*underfitting*). A pesar de que el conjunto de datos no muestra signos alarmantes de desbalanceo, se ha optado por seguir esta técnica para probar que los resultados obtenidos con el modelo no se deben a un sobre entrenamiento con el conjunto.

Para ello, se dividió la muestra de datos en diez particiones, utilizando para el aprendizaje nueve de estas como conjunto de entrenamiento para luego aplicar las métricas sobre la sobrante. En cada paso, se ha ido variando el conjunto de entrenamiento y el de prueba. Así, en la Figura 2 (a) se pueden apreciar los resultados obtenidos del valor F-1 de cada modelo utilizados de forma independiente, observando una escasa variabilidad en cada una de las particiones.

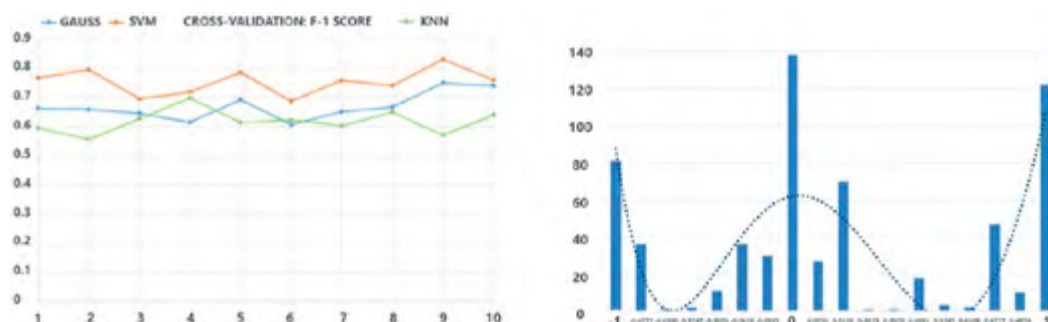


Figura 2: (a) Resultados del Valor F-1 tras realizar validación cruzada por cada partición. (b) Distribución de los valores predichos por la primera aproximación.

En la Figura 2 (b) se puede observar el histograma que describe la distribución de los datos a la salida de la primera aproximación. Debido a que esta distribución dista de ser normal (los tres clasificadores tienen tendencia al consenso y por lo tanto los valores están concentrados en torno al -1, 0 y 1), se focaliza en la parte central de la muestra (la que engloba los puntos de corte con el eje horizontal más cercanos al valor 0), que sí parece comportarse como una distribución normal, con el objetivo de definir el rango de valores para los cuales un tweet será considerado como neutro. Por medio de la función de distribución acumulada, ese intervalo queda definido como $[-0,32, 0,36]$.

- 2 El sobreajuste consiste en sobreentrenar a un algoritmo de aprendizaje con datos para los que se conoce el resultado. De este modo el algoritmo se vuelve muy preciso para ese conjunto de datos, pero es incapaz de generalizar el aprendizaje, convirtiéndolo en un modelo poco fiable.
- 3 El subajuste se produce cuando no existe suficiente información en el conjunto de entrenamiento, o cuando este sea poco representativo, haciendo que el modelo ignore parámetros relevantes para su clasificación.

Teniendo esto en cuenta, la tabla 2 muestra los valores de precisión, exhaustividad y valor F-1 de cada modelo de forma independiente en función de su polaridad, así como los valores obtenidos por ambos metaclasificadores. A la vista de estos resultados, se observa con claridad la superioridad del segundo metaclasificador con respecto al resto de modelos. Además, también se deduce una fiabilidad menor de *GaussNB* y *KNN* respecto al resto de modelos. Este resultado es razonablemente lógico, dado que el primero parte de la presunción de que los elementos del conjunto de datos presentan una distribución normal y que cada una de las variables es independiente respecto al resto; y en el segundo, por tratarse de un algoritmo de *lazy learning*.

Algoritmo	POSITIVO			NEUTRO			NEGATIVO		
	P	R	F-1	P	R	F-1	P	R	F-1
GaussNB	0,69	0,72	0,71	0,66	0,53	0,59	0,64	0,74	0,69
SVM	0,78	0,68	0,73	0,71	0,8	0,75	0,83	0,79	0,81
KNN	0,91	0,49	0,63	0,54	0,93	0,68	0,91	0,57	0,7
Metaclasificador inicial	0,8	0,66	0,72	0,67	0,81	0,73	0,85	0,78	0,81
Metaclasificador final	0,83	0,88	0,85	0,76	0,76	0,76	0,84	0,75	0,79

Tabla 2. Comparación de los resultados obtenidos de los modelos con supervisión.

La Figura 3 muestra mediante un histograma el valor F-1 ponderado (aplicando *weighted-averaging*) de cada uno de los modelos, en el que se observa una superioridad de un 5 % del segundo metaclasificador creado, mientras que el primero está dentro de los rangos del mejor clasificador simple, es decir, el SVM.

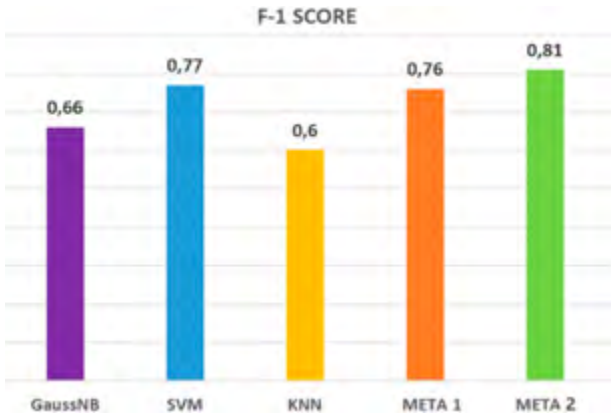


Figura 3: Diagrama ilustrativo del valor F-1 medio.

5. Conclusiones

En los últimos años, las redes sociales se han convertido en una nueva forma de comunicación y expresión para todo tipo de usuarios. Esto hace que organizaciones, como en el caso de las Fuerzas Armadas, deban dedicar recursos para el análisis y monitorización de esta información en Internet. En nuestro caso, estos análisis permitirían ayudar a desarrollar estrategias para fomentar la cultura de defensa.

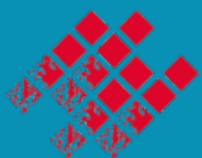
En vista a los resultados obtenidos, se puede concluir que se han generado una serie de modelos de aprendizaje con supervisión lo suficientemente precisos como para considerarlos ciertamente fiables, obteniendo en el mejor caso un valor F-1 de más del 80 %. Además, se ha visto cómo la utilización del segundo metaclassificador mejoraba los resultados de los clasificadores simples para cada una de las etiquetas de forma individual.

Referencias

1. B. J. Jansen, M. Zhang, K. Sobel y A. Chowdury, «Twitter power: Tweets as electronic word of mouth», *Journal of the American Society for Information Science and Technology*, vol. 60, n.º 11, pp. 2169-2188, 2009.
2. A. Ortega de los Ríos, «Calderón (repositorio institucional del CUD)», abril 2020. [En línea]. Available: <http://calderon.cud.uvigo.es/handle/123456789/376>. [Último acceso: septiembre 2020].
3. E. Cambria, S. Poria, A. Gelbukh y M. Thelwall, «Sentiment Analysis Is a Big Suitcase», *IEEE Intelligent Systems*, vol. 32, n.º 6, pp. 74-80, 2017.
4. B. Liu, *Sentiment analysis and opinion mining*, Morgan & Claypool, 2012.
5. E. Bothos, D. Apostolou y G. Mentzas, «Using social media to predict future», *IEEE Intelligent Systems*, vol. 25, n.º 6, pp. 50-58, 2010.
6. «Scikit-Learn», 2017. [En línea]. Available: https://scikitlearn.org/stable/modules/naive_bayes.html. [Último acceso: enero 2020].
7. «Scikit-Learn», 2017. [En línea]. Available: <https://scikitlearn.org/stable/modules/svm.html>. [Último acceso: enero 2020].
8. «Scikit-Learn», 2017. [En línea]. Available: <https://scikitlearn.org/stable/modules/neighbors.html>. [Último acceso: enero 2020].
9. A. P. a. C. Williams, «brilliant», [En línea]. Available: <https://brilliant.org/wiki/k-nearest-neighbors/>. [Último acceso: febrero 2020].
10. «Tweepy», [En línea]. Available: <http://docs.tweepy.org/en/latest/api.html>. [Último acceso: 26 septiembre 2020].

11. J. C. Sobrino Sande, «Análisis de sentimiento en Twitter», Máster Universitario en Ingeniería Informática - Inteligencia Artificial, Barcelona, 2018.
12. Scikit-Learn, [En línea]. Available: <https://scikit-learn.org/stable/modules/multi-class.html>.

DESEi+d 2020



Isdefe



GOBIERNO
DE ESPAÑA

MINISTERIO
DE DEFENSA

SUBSECRETARÍA DE DEFENSA
SECRETARÍA GENERAL TÉCNICA

SUBDIRECCIÓN GENERAL
DE PUBLICACIONES
Y PATRIMONIO CULTURAL