



**INSTITUTO POTOSINO DE INVESTIGACIÓN
CIENTÍFICA Y TECNOLÓGICA, A.C.**

POSGRADO EN GEOCIENCIAS APLICADAS

**Detección de derrames de petróleo mediante
imágenes de satélites ópticos y activos e
inteligencia artificial**

Tesis que presenta

Rubicel Trujillo Acatitla

Para obtener el grado de

Doctor en Geociencias Aplicadas

Codirectores de la Tesis:

Dr. José Tuxpan Vargas

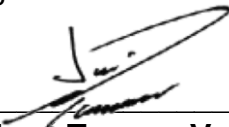
Dr. Cesaré Moisés Ovando Vázquez

San Luis Potosí, S.L.P. Septiembre 2023



Constancia de aprobación de la tesis

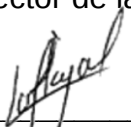
La tesis “***Detección de derrames de petróleo mediante imágenes de satélites ópticos y activos e inteligencia artificial***” presentada para obtener el Grado de Doctor en Geociencias Aplicadas fue elaborada por **Rubicel Trujillo Acatitla** y aprobada el **04 de Septiembre de 2023** por los suscritos, designados por el Colegio de Profesores de la División de Geociencias Aplicadas del Instituto Potosino de Investigación Científica y Tecnológica, A.C.



Dr. José Tuxpan Vargas
Codirector de la tesis



Dr. Cesaré Moisés Ovando Vázquez
Codirector de la tesis



Dr. José Noel Carvajal Pérez
Miembro del Comité Tutorial



Dr. Mariano J. J. Rivera Meraz
Miembro del Comité Tutorial



Dr. Francisco Javier Ocampo Torres
Miembro del Comité Tutorial



Créditos Institucionales

Esta tesis fue elaborada en el Laboratorio de Geomática y Modelación Numérica de la División Geociencias Aplicadas, y el Laboratorio de Bioinformática e Inteligencia Artificial (AI-Bio Lab) del Centro Nacional de Supercómputo del Instituto Potosino de Investigación Científica y Tecnológica, A.C., bajo la codirección de los doctores José Tuxpan Vargas y Cesaré Moisés Ovando Vázquez.

Durante la realización del trabajo el autor recibió una beca académica del Consejo Nacional de Humanidades, Ciencias y Tecnologías (Conahcyt) CVU: 699765, y del Instituto Potosino de Investigación Científica y Tecnológica, A. C.

Acta de examen

Dedicatorias

A mis padres, por su inquebrantable apoyo y amor incondicional a lo largo de esta larga travesía académica. Su sacrificio y confianza en mí han sido, y serán la fuente de mi motivación y éxito.

A Erandi mi amada esposa, quien siempre esta cuando más lo necesito, por ser mi luz, y no soltar mi mano a lo largo de todo este camino a pesar de todo.

A mi hermano, Jesús Adrián, que me motiva a siempre dar lo mejor y ser un ejemplo.

A mis suegros, Gilberto Monterrubio y Ana María Martínez, por su gran y valioso apoyo, así como sus consejos y enseñanzas.

A mis perros Toby, Maya, y Tiny, por estar siempre durante este camino.

Agradecimientos

A mis padres por todo su apoyo, amor, y comprensión en todos y cada uno de los pasos de esta larga trayectoria. Su esfuerzo, lucha, dedicación y palabras siempre fueron una guía y motivo de seguir adelante a pesar de todo, y siempre dar más. A mi hermano, para quien siempre busco ser un ejemplo y hermano excepcional. Jamás terminare de darles las gracias.

A mi esposa Erandi, por ser siempre mi apoyo y motivo de seguir. Todos tus consejos, paciencia, así como tu compañía fueron fundamentales para lograr todo, siempre fuiste la luz que iluminó mi camino, siempre me motivas a dar lo mejor. Gracias.

A mis suegros, Gilberto Monterrubio y Ana María Martínez, por todo su apoyo, consejos, y sabiduría. Gracias.

Al IPICYT por brindarme las herramientas necesarias para el desarrollo de esta tesis.

Al Consejo Nacional de Humanidades, Ciencias y Tecnologías (Conahcyt) por otorgarme la beca para el desarrollo de esta tesis.

Al Dr. José Tuxpan Vargas y al Dr. Cesaré Ovando Vázquez por su orientación experta, su paciencia infinita y su dedicación incansable. Su sabiduría y guía han sido fundamentales en mi camino hacia este logro. Gracias.

A mi comité, Dr. Noel Carbajal, Dr. Mariano Rivera, y al Dr. Francisco Ocampo, por sus grandes consejos para mejorar la elaboración de esta tesis. Su gran disposición fue fundamental para concluir satisfactoriamente esta etapa. Gracias.

A mis amigos, compañeros y colegas, Fermín Villalpando, Rodrigo Dávila, Adriana Martínez, Brenda Mendoza, Stephanie Martínez y Daniel Saldaña, quienes siempre estuvieron presentes para todo. Gracias.

Contenido

Constancia de aprobación de la tesis	II
Créditos Institucionales.....	III
Acta de examen	IV
Dedicatorias.....	V
Agradecimientos	VI
Contenido	VII
Lista de tablas	IX
Lista de figuras	XI
Resumen	XIX
Abstract	XX
Introducción	1
Sensores ópticos para el monitoreo de derrames.....	3
Radar de apertura sintética (SAR) para el monitoreo de derrames	5
IA en la observación de la Tierra.....	8
Mecanismo acoplado para el monitoreo de derrames mediante sensores ópticos y activos.....	13
Hipótesis	15
Objetivos.....	16
Método	17
Sensores pasivos.....	17
Datos.....	18
Análisis exploratorios y estadísticos	23
Modelos de Machine Learning.....	25
Postproceso	30
Sensores activos.....	30
Datos y preproceso	32
Conjunto de datos de Sentinel-1 SAR.....	35
Modelos de aprendizaje profundo (<i>Deep learning models</i>)	37
Características del hardware	51
Equipo utilizado para ML	51
Equipo utilizado para DL.....	51
Resultados	52
Sensores pasivos.....	52

Diferencias espectrales entre superficies	52
Correlación y clustering de datos de materiales	54
Modelos lineales de datos de concentración	58
Modelos de ML.....	60
Sensores activos.....	76
CNN para clasificación de derrames de petróleo	76
Prueba del modelo CNN para clasificación de derrames de petróleo	80
Modelos de segmentación de derrames de petróleo.....	83
Estimación de área.....	95
Discusión.....	97
Sensores pasivos.....	97
Diferenciación de materiales.....	97
Precisión en la clasificación de superficies	100
Potencial de detección de aceite en imágenes satelitales ópticas.....	107
Datos espectrales e imágenes ópticas para la estimación de la concentración	109
Sensores activos.....	111
Importancia del contexto y forma en la discriminación de aceite.....	111
Comparación de modelos de segmentación	116
Modelo U-net con pérdida focal como alternativa para segmentación de aceite	118
Necesidad de un esquema clasificación y segmentación de derrames en dos etapas	122
Consideraciones	127
Sensores pasivos.....	127
Sensores activos.....	129
Esquema de monitoreo conjunto.....	130
Conclusiones.....	132
Sensores pasivos.....	132
Sensores activos.....	133
Perspectivas del proyecto	135
Referencias	138
Anexos.....	158

Lista de tablas

Tabla 1. Ejemplo de base de datos utilizada para el modelo de clasificación de imágenes. Las bandas se utilizaron como características (X), y cada material con su respectiva etiqueta se utilizó como la salida deseada (Y).....	19
Tabla 2. Detalles de las imágenes utilizadas. La información incluye el satélite, la fecha (formato ISO), el Sistema de Referencia Mundial-2 (WRS-2 Path/WRS-2 Row) y el identificador de escena Landsat.....	21
Tabla 3. Valores e hiperparámetros utilizados para el entrenamiento de los modelos de clasificación.....	61
Tabla 4. Valores e hiperparámetros utilizados para los modelos de clasificación con características importantes en el proceso de entrenamiento.	67
Tabla 5. Métricas y puntuaciones obtenidas en la evaluación del rendimiento del modelo de regresión Random Forest y del modelo lineal estadístico. Las primeras columnas corresponden a las métricas de Random Forest, y la última a un modelo lineal.....	72
Tabla 6. Exactitud y Pérdida de los cinco modelos con los scores más altos. Se observa que el modelo de dos capas ocultas con 20 neuronas y 6 convoluciones con 32 filtros fue el que obtuvo los mejores valores (en negritas)	79
Tabla 7. Diferentes modelos U-net entrenados y validados. Se implementaron distintas combinaciones de parámetros y arquitectura para encontrar el modelo con las mejores aproximaciones. Se implemento desde el modelo U-net original, posteriormente se le cambio la función de pérdida por la de Pérdida Focal, la siguiente modificación fue la función Conv2DTranspose por Upsampling2D, y las últimas modificaciones fue el número de filtros y la profundidad (convoluciones). Se observa que el modelo con el IoU promedio más bajo fue la U-net original (0.50 y 0.40 para entrenamiento y validación), y el que obtuvo las mejores puntuaciones fue el modelo con Pérdida Focal, Upsampling2D y una mayor cantidad	

de filtros y profundidad (última fila), IoU promedio de 0.96 para entrenamiento y 0.90 para validación. 88

Tabla 8. Matriz de confusión con los resultados de clasificación del modelo KNN para las imágenes con aceite y sin aceite. El modelo obtiene una precisión del 100% para ambos casos. 102

Tabla 9. Trabajos realizados para la segmentación de derrames de petróleo con imágenes SAR. Se observa el modelo propuesto, así como algunas métricas que utilizaron. 120

Lista de figuras

Figura 1. Esquema general del método. En el paso 1, se encuentra la obtención de las imágenes, corrección a TOA (Top of Atmosphere), apilamiento, enmascaramiento nubes y reestructuración. En el paso 2, la imagen ingresa al modelo de clasificación, previamente entrenado y validado con datos públicos de respuesta espectral, la salida del modelo se reestructuró de nuevo para obtener la imagen clasificada. Posteriormente, sólo si había petróleo, se continuo con el paso 3, que es el modelo de estimación de la concentración. Finalmente, la salida se reestructuró para obtener la imagen del resultado final con la presencia de petróleo y estimación de su concentración. 17

Figura 2. Reestructuración de conjunto de bandas para crear una matriz, donde cada columna corresponde a las bandas y cada fila son los valores de los pixeles. 23

Figura 3. El esquema general del método abarca varias etapas. En primer lugar, se obtienen las imágenes Sentinel-1 SAR y se aplican técnicas de preprocesamiento, como la calibración, eliminación de ruido y correcciones, para obtener imágenes Sigma0 en decibels (db). A continuación, se divide cada imagen en recortes de tamaño 2048x2048 para las dos polarizaciones (VV, VH). Después, se crea una máscara para cada recorte, lo que permite generar un conjunto de datos completo. Luego, se procede a entrenar y validar modelos de aprendizaje profundo, redes neuronales convolucionales (CNN), perceptrón multicapa (MLP) y la red de segmentación U-Net. Por último, se obtiene un esquema de dos pasos para la detección y segmentación de derrames de petróleo. Este esquema utiliza los modelos entrenados para identificar y delimitar los derrames de petróleo en las imágenes procesadas, proporcionando así una herramienta eficiente para detectar y segmentar los derrames de petróleo en imágenes Sentinel-1 SAR. 32

Figura 4. Ubicación de las 685 imágenes SAR Sentinel-1 GRD descargadas con base en los reportes de la NOAA y de la EMSA. 33

Figura 5. Método aplicado para el preproceso de las imágenes, desde la descarga hasta su almacenamiento como Sigma0 en db. El preproceso de las imágenes se realizó siguiendo el método descrito en (Filipponi, 2019). Se aplicó el filtro Refined Lee (5x5) para reducir el ruido de speckle (ruido multiplicativo). Además, se utilizó el modelo digital de elevación SRTM 1Sec HGT para corregir las distorsiones topográficas. Las imágenes resultantes se almacenaron en formato Sigma0 en decibels (db). 34

Figura 6. Subimágenes de 2048x2048 para ambas polarizaciones. La primera fila corresponde a un derrame de petróleo con su ground truth, la segunda fila es la imagen libre de aceite con su ground truth, y la tercera fila es look-alike con su ground truth. Se observa que el ground truth para el aceite tiene el valor de 1 y el fondo (background) de 0, y el ground

truth para los otros dos casos es únicamente 0. En el caso de las imágenes libres de aceite y look-alikes, ambas se consideraron como No aceite. 36

Figura 7. Esquema general de una CNN, donde se observa la entrada, seguida de las capas de convolución (Conv2D) y de agrupación (Pooling). Posteriormente las salidas de la última capa de convolución son aplanadas y conectadas a una red densa, donde cada capa puede tener diferente número de neuronas, y por último la neurona de salida con función de activación sigmoide la cual dará valores de probabilidad entre 0 y 1. 40

Figura 8. Esquema general de un modelo de MLP. Para las entradas se utilizaron los recortes de las imágenes Sentinel-1, en conjunto con su ground truth. Estas imágenes se aplanaron para obtener un vector para cada polarización con su respectiva etiqueta, por lo que la entrada al modelo son dos polarizaciones (VV, VH). La siguiente parte del modelo son las capas ocultas, las cuales cuentan con activación tanh, y por último la capa de salida con una neurona con activación sigmoide. La salida de este modelo será un vector de probabilidad de pertenencia a la clase (Aceite, No aceite) el cual será convertido nuevamente en matriz para tener la posición del píxel ya clasificado dentro de la imagen. 44

Figura 9. A) Las distribuciones de reflectancia de los cinco materiales estudiados mostraron diferencias estadísticas significativas en el análisis ANOVA ($p = 9.15e-140$). Según el análisis post hoc de Tukey, se encontraron similitudes entre el aceite y el suelo ($p = 6.05e-2$), así como entre el suelo y la vegetación ($p = 8.58e-01$). B) Se muestra la respuesta espectral de cada material en cada banda: aceite (negro), plástico (amarillo), suelo (marrón), vegetación (verde) y agua (azul). 53

Figura 10. La matriz de correlación muestra las correlaciones entre todas las observaciones de los materiales, con 35 observaciones para cada clase. Se utiliza un gradiente de colores para representar los valores de correlación: rojo para correlaciones positivas (1.0), azul para correlaciones negativas (-1.0) y amarillo para datos no correlacionados (~ 0). Las clases de cada muestra se indican en los lados de las columnas y las filas, con negro para aceite, amarillo para plástico, marrón para suelo, verde para vegetación y azul para agua. El dendrograma en las filas y columnas muestra la agrupación jerárquica basada en las similitudes o diferencias entre los datos. 55

Figura 11. Matriz de correlación entre las bandas utilizadas. Se utiliza un gradiente de colores para representar los valores de correlación, donde los valores superiores a 0.85 se muestran en azul oscuro, los valores entre 0.75 y 0.80 se muestran en azul claro, y los valores inferiores a 0.70 se muestran casi blancos. Las bandas infrarrojas se representan en tonos de gris, mientras que las bandas RGB se representan en rojo, verde y azul. El dendrograma presente en las filas y columnas muestra la agrupación jerárquica de las bandas en función de sus similitudes o diferencias. 56

Figura 12. Análisis PCA y T-SNE de los datos analizados para las cinco clases trabajadas. A) PCA, sólo se utilizaron dos componentes, ya que juntos contienen el 90% de la variabilidad de los datos. B) T-SNE, sólo se utilizaron dos dimensiones. En ambos gráficos, los puntos negros y las elipses representan el petróleo, el amarillo el plástico, el marrón el suelo, el verde la vegetación y el azul el agua..... 58

Figura 13. Los modelos lineales para los datos de concentración de petróleo. El eje X es la reflectancia y el eje Y es la concentración. A) Modelo general, construido con la reflectancia media de las seis bandas utilizadas, la pendiente de este modelo fue estadísticamente significativa (valor $p = 3.23e-11$), con valor de ajuste $R^2 = 0.67$, y tendencia negativa (-270). B) El modelo de concentración lineal para cada banda utilizada. Se observa una tendencia negativa para todas las bandas, que es significativa excepto para la banda SWIR2 (valor $p = 0.54$). Los valores de ajuste R^2 , así como el valor p y la ecuación lineal, se muestran en los gráficos para cada banda..... 59

Figura 14. Métricas calculadas para el conjunto de datos de entrenamiento y validación, para los cinco modelos utilizados. La línea azul representa las puntuaciones de los datos de entrenamiento y la línea naranja los datos de validación. Los valores más altos corresponden a Random Forest y los más bajos a SVM..... 62

Figura 15. Resultados de la clasificación para la imagen del derrame de DeepWater Horizon del 09 de mayo de 2010. En la primera fila se muestra un mapa de localización de la imagen Landsat 5TM y su composición en RGB, con la ubicación de la plataforma DeepWater Horizon marcada con una gota negra y las letras DWH. En las filas siguientes se presentan los resultados de los modelos de clasificación: Random Forest, SVM RBF, KNN y Decision Tree. Las clases se muestran en diferentes colores: nubes en rojo, agua en azul, petróleo en negro, plástico en amarillo, vegetación en verde y suelo en marrón..... 64

Figura 16. Imágenes de prueba de Landsat 8 OLI. La primera imagen muestra la ubicación de las dos imágenes seleccionadas para las pruebas. La imagen superior derecha (2014/01/12) corresponde a una zona cercana a la costa de Luisiana, similar a las imágenes de la Figura 15. La imagen inferior derecha (2017/12/08) corresponde a una zona cercana a la península de Yucatán, en México. En ambas imágenes se enmascararon las nubes (rojo) y la tierra, por lo que esta última se presenta tal y como se ve en la imagen RGB..... 65

Figura 17. Importancia de las variables según el método de disminución media de impurezas de Gini. Se observa que las bandas con mayor importancia son las bandas infrarrojas NIR, SWIR2 y SWIR1, y en cuarto lugar se encuentra la banda azul. 66

Figura 18. Métricas calculadas para el conjunto de datos de entrenamiento y validación, para los cinco modelos utilizados con las cuatro variables seleccionadas. La línea azul representa las puntuaciones de los datos de entrenamiento, y la línea naranja las de los datos de

validación. Los valores más altos corresponden a Random Forest, y los más bajos a SVM.
 68

Figura 19. Resultados de los modelos de clasificación entrenados sólo con las partes del espectro correspondientes a las bandas Azul, NIR, SWIR1 y SWIR2. La imagen seleccionada para esta prueba es de Landsat 5TM del 9 de mayo de 2010, capturada durante el derrame de DeepWater Horizon (la plataforma se localiza con una gota negra y las letras DWH). En las filas segunda y tercera se presentan los resultados de los modelos de clasificación Random Forest, SVM RBF, KNN y Decision Tree. Las nubes se representan en rojo, el agua en azul, el aceite en negro, el plástico en amarillo, la vegetación en verde y el suelo en marrón. 70

Figura 20. Imágenes de prueba Landsat 8 OLI para el modelo KNN entrenado sólo con las bandas Azul, NIR, SWIR1 y SWIR2. La primera imagen muestra la ubicación de las dos imágenes seleccionadas para la prueba. La imagen superior derecha (2014/01/12) corresponde a una zona cercana a la costa de Luisiana, una ubicación como las imágenes de la Figura 19. La imagen inferior derecha (2017/12/08) corresponde a una zona cercana a la península de Yucatán, en México. En ambas imágenes se enmascararon las nubes (rojo) y la tierra, por lo que esta última se presenta tal y como se ve en la imagen RGB..... 71

Figura 21. Prueba del modelo de estimación de la concentración utilizando una imagen sintética creada con los datos originales. A) Se muestra la banda azul de la imagen sintética, dentro de cada píxel está el valor de concentración original. Los datos se organizaron en bandas como una imagen de satélite. B) Resultado de la predicción del modelo, dentro de cada píxel está el valor de concentración estimado (%). 73

Figura 22. Resultados de la estimación de la concentración utilizando el regresor Random Forest. En la primera columna está la imagen RGB con la zonificación de petróleo identificada por el modelo de clasificación KNN para las variables importantes. La segunda columna es la imagen RGB con la concentración estimada para la zona identificada como petróleo. En ambos casos, las nubes están enmascaradas e indicadas en rojo. 74

Figura 23. A) Transectos trazados para observar los perfiles de concentración de la imagen de mayo 09 del 2010. Las líneas azules representan los perfiles, los cuadrados verdes el inicio del perfil y las elipses amarillas el final. B) Perfil de estimación de la concentración para el transecto 1. C) Perfil de estimación de la concentración para el transecto 2. 75

Figura 24. Resultados de los modelos CNN. La primera columna especifica el número de capas ocultas y neuronas en cada capa. En la segunda columna se encuentran los valores de Pérdida, y en la tercera los de Exactitud. En rojo los datos de validación, y en azul los de entrenamiento. El modelo con los mejores scores (Pérdida y Exactitud) fue el de dos capas ocultas con 20 neuronas y seis capas convolucionales (6Conv) con 32 filtros, exactitud de

entrenamiento 0.9999, exactitud de validación 0.9824, pérdida de entrenamiento 0.0018, pérdida de validación 0.0573..... 77

Figura 25. Arquitectura del modelo CNN con los mejores scores. La entrada al modelo son imágenes Sentinel-1 de tamaño 512x512 píxeles con dos polarizaciones (VV, VH), por lo que la dimensión de entrada es de 512x512,2. El modelo cuenta con seis capas convolucionales (Conv2D) de tamaño de kernel de 3x3 con 32 filtros, seguida de MaxPooling2D de 2x2, la función de activación utilizada para todas las capas convolucionales fue ReLU. Para la red densa, después del flatten, se utilizaron dos capas ocultas con 20 neuronas cada una y con activación ReLU. La capa de salida solo contaba con una neurona con activación sigmoide. 80

Figura 26. A) Matriz de confusión de las pruebas realizadas, se observa que la exactitud del modelo para la clase aceite fue de 0.98, y para la clase No Aceite de 0.96. B) Curva ROC de las pruebas realizadas, el valor de Area Under Curve (AUC) fue de 0.99. 81

Figura 27. Resultado de las pruebas realizadas con el modelo CNN. Se observa cada una de las imágenes Sentinel-1 utilizadas, las cuales tienen dimensión 512x512,2, ya que cuentan con ambas polarizaciones (VV, VH), se muestra únicamente la polarización VV ya que es en la que a simple vista se observa la presencia de derrames, así como los look-alikes. En la parte superior de cada imagen se encuentra su Etiqueta, la cual tiene valores de 0.00 para No Aceite, y 1.00 para Aceite, además se encuentra el valor de predicción en probabilidad. Observamos que las imágenes correspondientes a Aceite tienen valor de probabilidad superior a 0.9, y las de No aceite de entre $2.2e-28$ y 0.0049. Cabe señalar que en dentro de la clase No aceite se encuentran imágenes de look-alikes. 82

Figura 28. Exactitud and Pérdida para entrenamiento y validación de los modelos de MLP. Se observa que se comienza con un modelo de regresión logística, el cual tiene exactitud en entrenamiento de 0.9331 y exactitud en validación de 0.933. Posteriormente los modelos con una capa oculta, donde el modelo con 500 neuronas obtuvo una exactitud en entrenamiento de 0.9454 y exactitud en validación de 0.944. Después se tienen los modelos con una capa oculta con diferente cantidad de neuronas, donde el modelo con el exactitud más alta fue el de 20 neuronas con exactitud en entrenamiento de 0.9513 y exactitud en validación de 0.9512. Y por último los modelos con tres capas ocultas, donde el modelo con 20 neuronas en cada capa obtuvo la exactitud en entrenamiento y validación más alta, de 0.9514 y 0.9513 respectivamente. De la misma manera, dichos modelos obtuvieron los valores más bajos en Pérdidas. 84

Figura 29. Resultados de algunas de las pruebas realizadas con los diferentes modelos. La primera columna corresponde a la imagen Sentinel-1, la segunda columna es el ground truth, posteriormente se encuentran las predicciones realizadas por los distintos modelos, cada una de las predicciones cuenta con su exactitud e IoU. Debido a que la salida del modelo da como

resultado valores de probabilidad, se utilizó un valor mayor o igual a 0.99 para realizar la asignación de clase (binarización). Se observa que los modelos con el mejor desempeño son el logístico y el de tres capas ocultas, sin embargo, en la imagen con look-alike el modelo segmenta erróneamente, dando como resultados los valores más bajos en las métricas..... 86

Figura 30. Puntajes de entrenamiento y validación de las métricas de las distintas arquitecturas basadas en el modelo U-net. Los modelos que se utilizaron fueron el modelo Attention U-net (Attt_Unet2D), la arquitectura U-net para segmentación de imágenes de un concurso de kaggle (Unet2D Kaggle), el modelo U-net con las modificaciones realizadas (Unet2D), y el modelo U-net++ (Unet2D Plus). Donde el modelo con los scores más altos fue Att_Unet2D, y el modelo con las puntuaciones más bajas fue Unet2D Plus..... 91

Figura 31. Resultado de las pruebas de los modelos Unet con imágenes con presencia de aceite. En la primera columna se encuentran las imágenes Sentinel-1, la segunda columna son las máscaras o ground truth, las columnas siguientes corresponden a los resultados para cada modelo. En la parte inferior de cada imagen se encuentra su valor de IoU. 92

Figura 32. Resultados de las pruebas realizadas para imágenes sin aceite (primeras tres filas) y con look-alikes (últimas tres filas), ambos catalogados como no aceite (negativos). Se observa que los resultados de todos los modelos obtuvieron puntajes de IoU = 1.0, ya que en ninguno se segmentó la superficie como aceite..... 93

Figura 33. Resultado de las pruebas de los modelos con distintos sets de datos. Se utilizaron imágenes únicamente de aceite, un conjunto de imágenes con y sin aceite, y otro conjunto con imágenes con aceite, sin aceite y con presencia de look-alikes. Se observa que para los primeros dos conjuntos el modelo Unet2D obtuvo el valor más alto de IoU (0.90), y para el set donde se encuentran los look-alikes obtuvo un valor de 0.85..... 94

Figura 34. Resultados de la estimación de aérea posterior a la segmentación por el modelo Unet2D. 96

Figura 35. A) Diferencia de métricas para los modelos Random Forest, SVM RBF, KNN y Decision Tree. El modelo con las menores diferencias fue KNN, con diferencias de 0.1 en todas las métricas calculadas. B) Análisis de validación cruzada para las métricas utilizadas. Este análisis se realizó para los conjuntos de entrenamiento y validación. Observamos que el modelo KNN presenta una menor variación en los scores obtenidos mediante este método. 101

Figura 36. Resultados del modelo KNN de clasificación de variables importantes para imágenes con petróleo. La primera imagen es la localización de las imágenes, que corresponden a la zona de la plataforma DeepWater Horizon. A continuación, la columna de la izquierda muestra las imágenes RGB obtenidas, y la columna de la derecha muestra los

resultados de clasificación del modelo. El color azul corresponde al agua, negro al petróleo, marrón al suelo, verde a la vegetación y amarillo al plástico. Las nubes aparecen enmascaradas en rojo en todas las imágenes. 105

Figura 37. Resultado del modelo KNN de variables importantes. La primera imagen corresponde a la localización de todas las imágenes sin aceite. A continuación, se muestran los resultados de la clasificación. La parte de agua sólo se asignó a la clase de agua, y no se observó petróleo. La Tierra se enmascaró y se observa en color verdadero, y las nubes se enmascaran en rojo en todas las imágenes. 106

Figura 38. A) Imagen RGB del 25 de mayo de 2010, los recuadros corresponden a las zonas donde hay emulsiones (Lu et al., 2020). B) Imágenes RGB del acercamiento a las emulsiones WO y OW. C) Imágenes en falso color (SWIR1, NIR, R) de las emulsiones. D) Resultado de la clasificación con KNN para bandas importantes, ambas emulsiones se clasificaron parcialmente como aceite. 108

Figura 39. Gráficos de los valores de reflectancia de los perfiles de la Figura 23 para cada banda utilizada. Sólo se utiliza la imagen de mayo de 09. El cuadrado verde representa el inicio, y la elipse amarilla representa el final del perfil en la imagen A) Corresponde al transecto 1 en la Figura 23A. B) Corresponde al transecto 2 de la Figura 23A. 110

Figura 40. Extracción de características en algunas convoluciones del modelo CNN para la imagen Sentinel-1 SAR con presencia de aceite. Se observa que en la primera convolución se extraen características sencillas como el borde e intensidad; y en las otras convoluciones son más complejas. Se observa que, para el aceite, el modelo logra extraer y aprender un patrón regular y lineal correspondiente a la mancha. 112

Figura 41. Extracción de características en algunas convoluciones del modelo CNN para la imagen Sentinel-1 SAR sin aceite. En general se observa que el modelo aprende aspectos y características sin un patrón aparente. 113

Figura 42. Extracción de características en cada una de las convoluciones del modelo CNN para la imagen Sentinel-1 de look-alike. Para este tipo de superficie se observa que el modelo extrae y aprende características que no presentan patrones muy definidos, es decir son muy irregulares. 115

Figura 43. Resultado de la segmentación realizada por los modelos MLP para una imagen de look-alike. Se observa que los modelos de una y dos capas ocultas no segmentan nada, pero el modelo logístico y de tres capas ocultas segmentan la superficie como aceite, generando un falso positivo. 116

Figura 44. Comparación de las estrategias de segmentación posterior a la identificación de aceite dentro de las imágenes. En la primera columna se encuentra la imagen utilizada para

prueba, en la segunda columna se encuentra la máscara o ground truth, en la tercera columna se encuentra la predicción realizada mediante el modelo CNN, en el cual para los ejemplos utilizados se detectó la presencia de aceite. La cuarta columna corresponde a la predicción realizada por el modelo MLP de tres capas ocultas, donde se observa que los valores de IoU se encuentran entre 0.36 y 0.61. La última columna es la predicción realizada por el modelo Unet2D, dando valores de IoU de 0.81, 0.82, y 0.96 para el tercer caso. 117

Figura 45. Esquema de la arquitectura del modelo U-net con las mejores aproximaciones para la segmentación de aceite. 119

Figura 46. Ejemplos de imágenes con presencia de look-alikes. Se observa la imagen con su correspondiente mascara, así como la predicción realizada por el modelo Unet2D, la cual muestra errores en la segmentación..... 122

Figura 47. Resultados del modelo CNN. Las imágenes utilizadas fueron las mismas que en la Figura 46. Se observa que la etiqueta verdadera tiene valor de 0, ya que son imágenes que pertenecen a la clase No aceite, y el resultado de la predicción hecha por el modelo CNN ubica a las imágenes dentro de dicha clase al dar valores de probabilidad de 0 o cercanos a 0. 123

Figura 48. Esquema general de clasificación, segmentación y estimación de área para imágenes Sentinel-1 SAR. El esquema abarca desde la obtención, el preproceso y división de la imagen, posteriormente se encuentra la detección por el modelo CNN, y en dado caso de encontrarse la presencia de aceite pasara al modelo de segmentación. Por último, se realiza la estimación del área y localización de la mancha de aceite. 124

Figura 49. Mapa realizado a partir del esquema de clasificación por CNN y segmentación por el modelo U-net. Asimismo, se observa en la leyenda las coordenadas del derrame de petróleo y la extensión estimada en el recuadro de la derecha. 126

Resumen

Detección de derrames de petróleo mediante imágenes satelitales ópticos y activos e inteligencia artificial

Los derrames de petróleo marinos representan un problema a nivel mundial, causando daños al ecosistema y pérdidas económicas. Por tanto, es crucial contar con un mecanismo que mejore la toma de decisiones ante estos desastres. Actualmente, la disponibilidad de datos de sensores remotos, y el incremento en las capacidades computacionales que permiten implementar técnicas de inteligencia artificial (IA), pueden ayudar al monitoreo terrestre y prevención de desastres. En este trabajo se utilizan estas herramientas para generar un mecanismo capaz de clasificar y segmentar los derrames de petróleo, así como extraer su área. Se realizaron dos esquemas, uno para sensores pasivos, y otro para activos. Para los sensores pasivos se desarrolló un método para clasificar derrames de petróleo y estimar su concentración mediante la respuesta espectral de las superficies, y técnicas de aprendizaje automático (ML) en imágenes Landsat. Para la parte de activos proponemos un mecanismo en dos etapas usando aprendizaje profundo en imágenes Sentinel-1 para detectar y segmentar derrames de petróleo. Generamos un conjunto de datos con imágenes de diversas regiones del mundo y entrenamos redes neuronales convolucionales (CNN) para clasificación. Para la segmentación, experimentamos con modelos como logístico, perceptrón multicapa (MLP) y U-net, usando métricas y funciones de pérdida como *Pérdida focal (Focal Loss)* e *Intesección sobre la Unión (Intersection over Union (IoU))*. Los resultados para los sensores pasivos en la tarea de clasificación, el modelo *K-Nearest Neighbor (KNN)* obtuvo las mejores aproximaciones en la detección de petróleo utilizando las bandas Azul (0,453-0,520 μm), NIR (0,790-0,891 μm), SWIR1 (1,557-1,717 μm) y SWIR2 (1,960-2,162 μm) para imágenes del derrame de la plataforma DeepWater Horizon en el Golfo de México en 2010. En el modelo de concentración, el error absoluto medio (MAE) fue de 1,41 y 3,34, para entrenamiento y validación. Al probar el modelo de estimación de concentración en imágenes en las que se detectó petróleo, la estimación obtenida se situó entre el 40 y el 60%. Para los sensores activos el modelo CNN alcanzó una exactitud (Accuracy) del 0.99 en entrenamiento, 0.98 en validación y 0.96 en pruebas. Con el modelo U-net con más capas convolucionales y filtros, junto con Focal Loss e IoU, obtuvimos una exactitud del 0.99 y un IoU de 0.96. Al combinar el enfoque de dos etapas (CNN-Unet), en las pruebas obtuvimos una exactitud del 0.95 y un IoU del 0.90. Con estas perspectivas del trabajo se demuestra el potencial de las técnicas de ML y datos de respuesta espectral para detectar y estimar la concentración de vertidos de petróleo con imágenes ópticas. Mientras que para la parte de SAR se demuestra que nuestro enfoque con aprendizaje profundo en imágenes Sentinel-1 SAR es una prometedora alternativa para monitorear derrames de petróleo y apoyar la toma de decisiones. Y en conjunto nuestra metodología puede contribuir a un mecanismo efectivo de detección y seguimiento de estos desastres.

Palabras clave: Derrame de petróleo, Concentración de petróleo, Inteligencia artificial (IA), Detección y segmentación, Landsat, Sentinel-1

Abstract

Detection of oil spills using optical and active satellite imagery and artificial intelligence.

Marine oil spills are a global problem, causing ecosystem damage and economic losses. Therefore, it is crucial to have a mechanism to improve decision making in the face of these disasters. Currently, the availability of remote sensing data and the increase in computational power, which allows the implementation of artificial intelligence (AI) techniques, can help in land monitoring and disaster prevention. In this work, these tools are used to generate a mechanism capable of classifying and segmenting oil spills, as well as extracting their area. Two schemes have been developed, one for passive sensors and the other for active sensors. For the passive sensors, a method has been developed to classify oil spills and estimate their concentration using the spectral response of surfaces and machine learning (ML) techniques on Landsat images. For the active part, we propose a two-step mechanism using deep learning on Sentinel-1 imagery to detect and segment oil spills. We generate a dataset of images from different regions of the world and train convolutional neural networks (CNNs) for classification. For segmentation, we experimented with models such as logistic, multilayer perceptron (MLP), and U-net, using loss metrics and functions such as focal loss and *Intersection over Union* (IoU). Results for passive sensors in the classification task, the *K-Nearest Neighbor* (KNN) model gave the best approximations for oil detection using the blue (0.453-0.520 μm), NIR (0.790-0.891 μm), SWIR1 (1.557-1.717 μm), and SWIR2 (1.960-2.162 μm) bands for images from the 2010 DeepWater Horizon oil spill in the Gulf of Mexico. For the concentration model, the mean absolute error (MAE) was 1.41 and 3.34 for training and validation, respectively. When evaluating the concentration estimation model on images where oil was detected, the estimate obtained was between 40 and 60%. For active sensors, the CNN model achieved an accuracy of 0.99 in training, 0.98 in validation, and 0.96 in testing. Using the U-net model with more convolutional layers and filters, together with Focal Loss and IoU, we obtained an accuracy of 0.99 and an IoU of 0.96. When combining the two-stage approach (CNN-U-net), in tests we obtained an accuracy of 0.95 and an IoU of 0.90. With these perspectives of the work, the potential of ML techniques and spectral response data to detect and estimate the concentration of oil spills with optical imaging is demonstrated. And for the SAR part, it is shown that our deep learning approach on Sentinel-1 SAR images is a promising alternative for oil spill monitoring and decision support. Our methodology can contribute to an effective mechanism for detecting and monitoring natural disasters.

Keywords: Oil spill, Oil concentration, Artificial Intelligence (AI), Detection and segmentation, Landsat, Sentinel-1

Introducción

En toda la historia de la industria petrolera han ocurrido derrames petroleros de gran magnitud, ocasionando daños y pérdidas a sectores ambientales, sociales, económicos e incluso políticos. Estos derrames han ocurrido en diferentes etapas de la cadena de procesos de la industria (exploración, producción, refinamiento, almacenamiento y/o distribución), ya sea en medios marinos o terrestres, siendo los primeros los de mayor impacto (Burgherr, 2007; Ivshina *et al.*, 2015; Daly *et al.*, 2016; Li *et al.*, 2016; Yang *et al.*, 2021).

Los derrames de petróleo en el medio marino son un problema crónico en todo el mundo, se tiene registro de la ocurrencia de derrames de gran magnitud como el Amoco Cádiz en 1978, el Ixtoc-I en 1979, Exxon Valdez en 1989, los ocurridos durante la Guerra del Golfo en Kuwait en 1991, el Erica en Francia en 1999, el del Mar Egeo en 1992, el Prestige en España y Francia en 1999, y más recientemente en el año 2010 en la plataforma DeepWater Horizon en el Golfo de México, en Siria en 2021, y en las costas de Perú en 2022, solo por mencionar algunos (Hooper, 1982; de la Huz, Lastra and López, 2011; White *et al.*, 2012; Kourdi, Salem and Qiblawi, 2021; Mega, 2022).

Se estima que aproximadamente dos millones de toneladas de petróleo ingresan al medio marino cada año, debido a accidentes o fallas en plataformas petroleras, oleoductos, refinerías, barcos petroleros, así como pequeñas embarcaciones (Brekke and Solberg, 2005; Kim, Moon and Kim, 2010; Briggs and Briggs, 2018; Huz, Lastra and López, 2019). Además, se suman las descargas ilegales efectuadas

por diversas embarcaciones que encuentran ventajas económicas en el vertido de sus desechos (Gauthier *et al.*, 2007; Solberg, 2012). El riesgo esperado de enfrentar multas es menor al costo del cumplimiento, y la probabilidad de ser detectados en estas acciones es prácticamente nula (Solberg, 2012).

Las consecuencias generadas por todos estos derrames abarcan desde aspectos sociales, económicos, e incluso políticos (Burgherr, 2007; Ivshina *et al.*, 2015; Daly *et al.*, 2016). Pero los impactos y consecuencias que generan sobre los ecosistemas costeros y marinos son especialmente destructivos (Solberg, 2012), ya que pueden requerirse décadas para poder restaurar el ecosistema dañado (White *et al.*, 2012). Por ello, surge la necesidad de un mecanismo de detección rápida y precisa que ayude a una toma de decisión y una respuesta de contención más eficiente y eficaz, que, en medida de lo posible, evite la dispersión del contaminante. Esto para salvaguardar el ecosistema marino y prevenir todos los daños consecuentes (Solberg, 2012; Wang *et al.*, 2023).

Durante los derrames se implementaron mecanismos para detectar y monitorear las trayectorias del aceite, como observaciones humanas desde el aire y uso de sensores aerotransportados (Leifer *et al.*, 2012; Oscar Garcia-Pineda *et al.*, 2013; Garcia-Pineda *et al.*, 2017). Sin embargo, debido a la gran extensión que pueden abarcar los derrames, los enfoques de vigilancia comenzaron a ser costosos, aunado a los costos de remediación y mitigación de daños (Leifer *et al.*, 2012; Topouzelis *et al.*, 2015). Por ello, se comenzó una búsqueda de un método para mapear la ubicación y realizar el seguimiento de los derrames de petróleo (Clark *et al.*, 2010; Leifer *et al.*, 2012).

En este sentido, se ha reconocido la importancia de la tecnología de teledetección como herramienta de apoyo a la respuesta ante diversas catástrofes, como los derrames de petróleo (Kim, Moon and Kim, 2010). El uso de estas tecnologías puede ayudar a identificar derrames menores antes de que causen daños generalizados, y en el caso de accidentes más grandes, son de excelente ayuda para obtener una visión general de la extensión del derrame y poder seguir su desarrollo (Solberg, 2012; Fingas and Brown E., 2015). Otra ventaja es el hecho de que, estas tecnologías están disponibles en todo el mundo, y son proporcionados por diversas agencias, instituciones y gobiernos en caso de presentarse un derrame (Bessis, 2003; Leifer *et al.*, 2012).

Sensores ópticos para el monitoreo de derrames

Se han llevado a cabo múltiples investigaciones con sensores remotos pasivos, en las cuales se han realizado caracterizaciones espectrales aprovechando las diferencias en los rasgos espectrales de las manchas de petróleo y del agua (Hu *et al.*, 2009; Bulgarelli and Djavidnia, 2012; Leifer *et al.*, 2012; Sun, Hu and Tunnell, 2015). Además, la rugosidad que presenta el aceite en la superficie del agua redistribuye el reflejo de la luz solar, provocando que las manchas parezcan más brillantes u oscuras que el agua circundante (Hörig *et al.*, 2001; Winkelmann, 2005; Adamo *et al.*, 2009; Clark *et al.*, 2010; Leifer *et al.*, 2012; Oscar Garcia-Pineda *et al.*, 2013; Sun, Hu and Tunnell, 2015; Fingas and Brown, 2018; Lu *et al.*, 2019).

Siguiendo este principio se han hecho trabajos realizando la caracterización espectral de las manchas de aceite y su detección con sensores hiperespectrales

(Winkelmann, 2005; Clark *et al.*, 2010; Leifer *et al.*, 2012; Kokaly *et al.*, 2017; Hu *et al.*, 2018; Lu *et al.*, 2019, 2020), y sensores ópticos multiespectrales como MODIS, MERIS, y Landsat (Bonn Agreement Accord de Bonn, 2007; Hu *et al.*, 2009, 2018; Srivastava and Singh, 2010; Taravat and Del Frate, 2012; Leifer *et al.*, 2012; Polychronis and Vassilia, 2013; Sun, Hu and Tunnell, 2015; Xing *et al.*, 2015; Liu *et al.*, 2016; Li *et al.*, 2017; Arslan, 2018; Luciani 1 and LANEVE 1, 2018; Kolokoussis and Karathanassi, 2018; Lu *et al.*, 2019, 2020; Ozigis, Kaduk and Jarvis, 2019; Al-Ruzouq *et al.*, 2020; Balogun *et al.*, 2020; De Kerf *et al.*, 2020; Mohammadi *et al.*, 2021; Chowdhury, Evans and Shipman, 2021).

Sin embargo, al realizar la detección de derrames, se debe considerar que existen factores que modifican las propiedades físicas y químicas del petróleo desde el primer momento de su liberación (Daling *et al.*, 1990, 2014; Daling and Strøm, 1999; Lu *et al.*, 2019; Mohammadiun *et al.*, 2021). Este proceso se conoce como meteorización y provoca que, al encontrarse en agua, el aceite se pueda emulsionar y/o cambiar su concentración, causando que la radiación reflejada se combine con la proveniente del agua, y que por ende se observen cambios en su respuesta espectral (Daling and Strøm, 1999; Bonn Agreement Accord de Bonn, 2007; Daling *et al.*, 2014; Ivshina *et al.*, 2015; Fingas and Brown, 2017).

Considerando esto, se han realizado estudios de laboratorio y campo para conocer el proceso de degradación, formación de emulsiones, cambios en la concentración de aceite, y su detección mediante imágenes de satélites con sensores pasivos (Daling and Strøm, 1999; Clark *et al.*, 2010; Leifer *et al.*, 2012; Daling *et al.*, 2014; Ivshina *et al.*, 2015; Sun, Hu and Tunnell, 2015; Svejksky *et al.*, 2016; Hu *et al.*,

2018; Lu *et al.*, 2019, 2020). Y en algunos de ellos se ha demostrado la existencia de una relación entre la reflectancia y la concentración de aceite, así como en las diferencias espectrales en las emulsiones, que posibilitan su detección, y monitoreo en el tiempo (Daling and Strøm, 1999; Clark *et al.*, 2010; Oscar Garcia-Pineda *et al.*, 2013; Lu *et al.*, 2019, 2020).

Debido a estas variaciones en la respuesta espectral del aceite, existe una necesidad en la que la gestión de derrames marinos cubra tanto la detección como el seguimiento (Ivshina *et al.*, 2015; Svejksky *et al.*, 2016; Lu *et al.*, 2019, 2020; Mohammadiun *et al.*, 2021). Donde, el primer paso es recopilar datos sobre petróleo, como características espectrales y sus variaciones, así como datos de imágenes ópticas, y posteriormente desarrollar técnicas computacionales confiables basadas en dichos datos (Leifer *et al.*, 2012; Ivshina *et al.*, 2015; Mohammadiun *et al.*, 2021).

Radar de apertura sintética (SAR) para el monitoreo de derrames

El radar de apertura sintética (SAR), al ser un sensor activo, ha resultado particularmente útil en el monitoreo de la superficie terrestre, brindando información en una amplia gama de condiciones. A diferencia de los sensores pasivos, el SAR no se ve afectado por condiciones atmosféricas, ni por la ausencia de radiación solar. En consecuencia, el SAR es capaz de proporcionar imágenes tanto de día como de noche, independientemente de las condiciones climáticas (Kim, Moon and Kim, 2010; Solberg, 2012).

Una característica importante que considerar es que los SAR se clasifican en bandas dependiendo de sus longitudes de onda: banda X (2,4 a 3,75 cm), banda C (3,75 a 7,5 cm), banda S (7,5 a 15 cm), banda L (15 a 30 cm) o banda P (30 a 100 cm). Cada banda tiene una capacidad de penetración distinta, y se presentan de manera independiente en los sistemas SAR satelitales o aerotransportados. Cada tipo de SAR puede tener una, dos o cuatro polarizaciones (VV, VH, HV, HH), dependiendo de cómo envíe y reciba el pulso de microondas (Kim, Moon and Kim, 2010; Solberg, 2012). Por tanto, la detectabilidad de un derrame de petróleo en la superficie del mar dependerá de la frecuencia y la polarización del sensor SAR (Gade *et al.*, 1998).

En general, debido a las ondas cortas y ondas capilares, el mecanismo de dispersión de la superficie del mar está dominado por el fenómeno de dispersión de Bragg, lo que conduce a que el pulso proveniente del SAR se disperse en todas las direcciones, ocasionando que el agua limpia se observe rugosa en la imagen SAR (Gao, Li and Liu, 2022).

Cuando el petróleo se derrama sobre la superficie del mar, se forma una fina capa que dificulta la generación de ondas de gravedad cortas y ondas capilares, lo que a su vez obstaculiza el movimiento normal del agua (Solberg, 2012; Wang *et al.*, 2022). Este efecto conduce a una disminución significativa en el fenómeno de dispersión de Bragg, resultando en un mecanismo que no es de Bragg (Gao, Li and Liu, 2022). Como consecuencia, la superficie cubierta de petróleo se presenta más suave y se observe como áreas o parches oscuros en las imágenes SAR (Fiscella

et al., 2000; Solberg, 2012; Li *et al.*, 2018; Alpers, Liu and Wu, 2019; Wang *et al.*, 2022).

Este fenómeno, de la disminución en la dispersión de Bragg, es más notorio en longitudes de onda más corta, ya que dicha reducción es más intensa en comparación con longitudes de onda más larga (Dasari, Anjaneyulu and Nadimikeri, 2022). En ese sentido, se ha observado que para la detección de derrames de petróleo las bandas X y C, al tener las longitudes de onda más corta, son las más efectivas (Solberg, 2012; Dasari, Anjaneyulu and Nadimikeri, 2022). Ya que, en ellas, los valores de retrodispersión provenientes del petróleo son, notoriamente, menores, en comparación con su medio circundante (Solberg, 2012).

Algunos enfoques para el monitoreo de petróleo, basados en la retrodispersión, han utilizado imágenes de los satélites ERS, ENVISAT ASAR y RADARSAT. En estos enfoques, se describe un método que implica la identificación de puntos oscuros, seguida de transformaciones matemáticas para describir el contraste y la homogeneidad de dichos puntos. El objetivo de este método es clasificar las manchas como aceite a través de umbralización (Fiscella *et al.*, 2000; Solberg, Brekke and Husoy, 2007; Kim, Moon and Kim, 2010; Solberg, 2012; Xu *et al.*, 2017; Chaturvedi, Banerjee and Lele, 2020; Abou Samra and Ali, 2022; Babagolimatikolaei, 2022; Jones, 2023).

Sin embargo, existen fenómenos naturales, como áreas con poco o viento nulo, floraciones algales, o presencia de celdas microconvectivas, que suavizan la superficie del mar y generan manchas oscuras con valores de retrodispersión

similares a los del aceite (Solberg, 2012; Alpers, Liu and Wu, 2019). Este hecho plantea un desafío en el desarrollo de un sistema de detección y seguimiento de derrames de petróleo, ya que las películas naturales, conocidas como similares (look-alikes en inglés), pueden ser clasificadas erróneamente como aceite. Estos "look-alikes" pueden generar falsos positivos en un sistema de monitoreo o detección, lo que afecta la precisión y confiabilidad del sistema (Solberg, Brekke and Husoy, 2007; Solberg, 2012).

Debido a esto los algoritmos de detección de derrames de petróleo incorporaron, además de la detección de áreas con valores bajos de retrodispersión, una búsqueda y extracción de características y/o transformaciones que proporcionen información sobre el contraste, forma, bordes y contexto de las zonas oscuras (Fiscella *et al.*, 2000; Brekke and Solberg, 2005; Xu *et al.*, 2017; Babagolimatikolaei, 2022). Esto con la finalidad de lograr una clasificación precisa de los manchones como aceite o *look-alikes* (Solberg, Brekke and Husoy, 2007; Kim, Moon and Kim, 2010).

IA en la observación de la Tierra

Los modelos basados en computación e inteligencia artificial (IA) han ganado popularidad en la actualidad, ya que ofrecen un alto grado de certeza y generan información de calidad que respalda la toma de decisiones. Estos modelos han demostrado ser eficaces en diversos procesos, como descripción, predicción, clasificación, segmentación, entre otros (Simon, 1996; Mitchell, 1997; Singha, Bellerby and Trieschmann, 2013; Lary *et al.*, 2018; Mohammadiun *et al.*, 2021).

En este sentido se ha iniciado el uso de herramientas de IA, específicamente de aprendizaje automático (*Machine Learning*, ML), que emplea ingeniería de características o variables convencionales para la discriminación, clasificación o predicción (Topouzelis and Psyllos, 2012; Wan and Cheng, 2013; Guo and Zhang, 2014; Mera *et al.*, 2017; Cantorna *et al.*, 2019; Dasari, Anjaneyulu and Nadimikeri, 2022; Ma *et al.*, 2023; Wang *et al.*, 2023).

En el área de observación de la Tierra se han aplicado distintas técnicas de ML a diferentes aspectos como monitoreo de agua, suelo y vegetación (Lee *et al.*, 2016; Azzari and Lobell, 2017; Lary *et al.*, 2018; Teluguntla *et al.*, 2018; Wolanin *et al.*, 2019; Balogun *et al.*, 2020; Bar, Parida and Pandey, 2020; Cao *et al.*, 2020; Ghorbanian *et al.*, 2020; Zhou *et al.*, 2021; Lin *et al.*, 2021; Singh *et al.*, 2021). Asimismo, ya se han desarrollado algunos modelos que realizan la detección de derrames de petróleo en imágenes SAR (Topouzelis *et al.*, 2007; Topouzelis and Psyllos, 2012; O Garcia-Pineda *et al.*, 2013; Singha, Bellerby and Trieschmann, 2013; Svejksky *et al.*, 2016; Zou *et al.*, 2016; Sudha and Vijendran, 2021), y también en imágenes de sensores pasivos (Kubat, Holte and Matwin, 1998; Cococcioni *et al.*, 2012; Taravat and Del Frate, 2012; Xing *et al.*, 2015; Liu *et al.*, 2016; Li *et al.*, 2018; Ozigis, Kaduk and Jarvis, 2019; Al-Ruzouq *et al.*, 2020; De Kerf *et al.*, 2020; Lu *et al.*, 2020; Garain, Mitra and Das, 2021; Mohammadi *et al.*, 2021; Trujillo-Acatitla, Tuxpan-Vargas and Ovando-Vázquez, 2022), además unos cuantos tratan de desarrollar técnicas que permitan discernir las diferentes formas de petróleo, ya sea en emulsión, diferentes espesores y/o concentraciones, así como volúmenes en imágenes satelitales de sensores ópticos (Clark *et al.*, 2010;

Svejksky *et al.*, 2016; Lu *et al.*, 2019, 2020; Trujillo-Acatitla, Tuxpan-Vargas and Ovando-Vázquez, 2022).

Sin embargo, las técnicas de ML requieren gran intervención humana, y para realizar las tareas de monitoreo se busca que el mecanismo tenga poca o nula intervención, por lo que las técnicas de aprendizaje profundo (*Deep Learning*, DL) pueden ser una buena alternativa. En comparación con las técnicas de ML, las técnicas de DL son métodos de representación-aprendizaje con múltiples niveles de representación obtenidos mediante la composición de módulos simples, pero no lineales, donde cada uno transforma la representación de un nivel superior (comenzando con la entrada sin procesar) en una representación cada vez más compleja o abstracta (LeCun, Bengio and Hinton, 2015; Vasudevan *et al.*, 2021; Martín Abadi *et al.*, 2022). El aspecto clave del DL es que aprenden representaciones (características) complejas y/o abstractas de los datos que no están diseñadas por ingenieros humanos expertos (LeCun, Bengio and Hinton, 2015; Vasudevan *et al.*, 2021).

En el ámbito de los derrames de petróleo, los de DL podrían aprender aspectos esenciales en la discriminación del aceite de los *look-alikes*. Estos aspectos incluyen la forma de los derrames, que tienden a tener una forma regular, especialmente los derrames provenientes de fuentes en movimiento. Además, los modelos podrían aprender a identificar el tipo y cantidad de aceite, el tiempo transcurrido entre el derrame y la toma de la imagen, así como la homogeneidad tanto de la mancha oscura como de su entorno. Por ejemplo, es más probable que una mancha aislada en un entorno homogéneo sea petróleo que la misma mancha en un entorno

heterogéneo (Brekke and Solberg, 2005; Solberg, Brekke and Husoy, 2007; Solberg, 2012).

Desde este punto de vista, se han realizado estudios para la detección de petróleo en imágenes SAR utilizando algoritmos de DL como perceptrón multicapa (MLP) (Del Frate *et al.*, 2000; Singha, Bellerby and Trieschmann, 2012, 2013; Taravat and Del Frate, 2012; Guo and Zhang, 2014; Ma *et al.*, 2014; Mera *et al.*, 2017; Dasari, Anjaneyulu and Nadimikeri, 2022) hasta algoritmos más complejos como Redes Neuronales Convolucionales (CNN) (Nieto-Hidalgo *et al.*, 2018; Zeng and Wang, 2020; Shaban *et al.*, 2021; Feinauer *et al.*, 2022), Res-Net (Temitope Yekeen, Balogun and Wan Yusof, 2020; Wang *et al.*, 2023), Mask R-CNN (Temitope Yekeen, Balogun and Wan Yusof, 2020), UNET (Nieto-Hidalgo *et al.*, 2018; Krestenitis *et al.*, 2019; Shaban *et al.*, 2021; Basit *et al.*, 2022; Zhai *et al.*, 2022; Zhu *et al.*, 2022; Dehghani-Dehcheshmeh, Akhoondzadeh and Homayouni, 2023), AlexNet (Wang *et al.*, 2021), Faster R-CNN (Huang *et al.*, 2022).

Sin embargo, algunos de estos estudios se han centrado en derrames específicos y en el uso de ciertos tipos de SAR, como ERS-1, ERS-2 (Del Frate *et al.*, 2000; Singha, Bellerby and Trieschmann, 2012; Taravat and Del Frate, 2012; Topouzelis and Psyllos, 2012), ES-2 SAR, ENVISAT ASAR (Singha, Bellerby and Trieschmann, 2012, 2013; Ma *et al.*, 2014; Hasimoto-Beltran *et al.*, 2023), RADARSAT-2 (Huang *et al.*, 2022), e incluso sensores aerotransportados como SLAR (Nieto-Hidalgo *et al.*, 2018), los cuales tienen disponibilidad limitada, y por tanto disminuye la posibilidad de generalización de los modelos para el monitoreo en distintas partes del mundo. Solo algunos utilizan imágenes con una disponibilidad más amplia, como

Sentinel-1 (Krestenitis *et al.*, 2019; Huang *et al.*, 2022; Zhai *et al.*, 2022; Zhu *et al.*, 2022; Dehghani-Dehcheshmeh, Akhoondzadeh and Homayouni, 2023).

Otro aspecto a destacar es que estos estudios emplean la estrategia de *Transfer-Learning* (TL), la cual utiliza modelos previamente entrenados para cumplir una tarea específica, y que pueden ser utilizados para resolver otra tarea similar; esta estrategia omite ciertas etapas por lo que el tiempo de entrenamiento se reduce (Chollet and others, 2015; Ng *et al.*, 2015; Zhuang *et al.*, 2019). Sin embargo, para poder utilizar esta estrategia, es un requisito indispensable que los datos con los que se quiere trabajar tengan las mismas dimensiones que los datos con los que fue entrenado el modelo (Chollet and others, 2015; Ng *et al.*, 2015; Zhuang *et al.*, 2019), como es el caso de las imágenes, ya que los modelos suelen utilizar como entrada tres canales, rojo, verde, y azul (Zhuang *et al.*, 2019).

Esto genera un inconveniente, ya que las imágenes SAR no siempre tienen las dimensiones necesarias, en ocasiones solo cuentan con una o dos polarizaciones (canales) como el caso de Sentinel-1 (VV, VH), por lo que se tendría que realizar alguna combinación lineal o no lineal, para crear uno o dos canales sintéticos, generando la necesidad de intervención humana en el preproceso. En este sentido se debe tener claro que la detección de derrames de petróleo, y el monitoreo constante de los mares y océanos debe ser lo más, eficiente y eficaz posible, tratando que la intervención humana sea mínima (Brekke and Solberg, 2005; Solberg, 2012).

A pesar de los esfuerzos que se han realizado, aún se reporta la problemática que generan los *look-alikes* al clasificarse como aceite (Temitope Yekeen, Balogun and Wan Yusof, 2020; Zhang *et al.*, 2020; Ma *et al.*, 2023), por lo que aún se requiere un mecanismo capaz de eliminar, o minimizar este aspecto (Solberg, 2012). En este sentido, algunos estudios se han enfocado en analizar la posibilidad un mecanismo de dos etapas para la detección y segmentación de derrames de petróleo utilizando mecanismos de DL (Nieto-Hidalgo *et al.*, 2018; Shaban *et al.*, 2021). Sin embargo, al utilizar datos SAR de regiones específicas y TL mediante la búsqueda de un canal sintético, se dificulta su uso en otras regiones (Ng *et al.*, 2015; Nieto-Hidalgo *et al.*, 2018; Zhuang *et al.*, 2019).

Mecanismo acoplado para el monitoreo de derrames mediante sensores ópticos y activos

Aprovechando las ventajas de los sensores ópticos y activos, además del potencial que proporcionan las herramientas de IA, se podría realizar un mecanismo acoplado en el cual se incrementa la resolución temporal de un sistema de monitoreo de derrames de petróleo. Esto con la finalidad de desarrollar un método de detección remota para mapear la ubicación y extensión de un derrame de petróleo, su concentración y con ello realizar el seguimiento de la mancha (Clark *et al.*, 2010; De Padova *et al.*, 2017; Fingas and Brown, 2018). Además, que genere confiabilidad al contar con mecanismos físicos entendidos e información significativamente mejorada, en relación a otros enfoques utilizados (Burgherr, 2007; Leifer *et al.*, 2012; Ivshina *et al.*, 2015; Ye *et al.*, 2019).

Para los sensores ópticos se buscó integrar datos de respuesta espectral del aceite proveniente del derrame del año 2010 en el Golfo de México, así como datos de otros materiales, con los que se generó un modelo mediante técnicas de ML para detectar la presencia de aceite y diferenciarlo de otros materiales presentes en imágenes ópticas. Además, utilizando datos de respuesta espectral de petróleo con distintas concentraciones, se generó un modelo capaz de realizar estimaciones en la concentración de los derrames de petróleo.

Para la parte de SAR se desarrolló e implementó un mecanismo en dos etapas para la detección y segmentación de derrames de petróleo mediante métodos de DL y el uso de imágenes Sentinel-1 SAR de distintas partes del mundo, con la finalidad de abarcar la mayor variación de los tipos y cantidad de petróleo en los derrames. Asimismo, esta estrategia de DL utilizó únicamente las dos polarizaciones presentes en el sensor (VV, VH), lo cual minimiza la intervención humana en la etapa del preproceso. Por tanto, los modelos utilizados fueron modificados para cumplir con las tareas de detección y segmentación. Además, con esta estrategia se busca minimizar los falsos positivos al eliminar, en el proceso de detección, todas aquellas escenas con presencia de *look-alikes*. Todo ello con la finalidad de crear un marco completo de detección y segmentación de derrames de petróleo, con la capacidad de utilizarlo donde se cuente con imágenes de libre acceso Sentinel-1 SAR.

Finalmente, se presenta una herramienta potencialmente útil para el monitoreo, detección y segmentación de derrames de petróleo en tiempo casi-real, que pueda ser de utilidad para los planes de respuesta a emergencia en futuros derrames. Además, se incrementa el conocimiento y estado del arte existente en este tema.

Hipótesis

Las herramientas de inteligencia artificial son capaces de realizar detección, segmentación y monitoreo, así como estimación de concentración de derrames de petróleo en imágenes provenientes de satélites con sensores pasivos y activos de manera precisa, a partir de las propiedades físicas, de forma y contexto del petróleo.

Objetivos

-Objetivo general-

Desarrollar un sistema de detección y monitoreo de derrames de hidrocarburos aplicando sensores remotos y técnicas basadas en inteligencia artificial.

-Objetivos específicos-

- Detectar derrames de hidrocarburos mediante sensores remotos pasivos y activos
- Establecer y validar mecanismos algorítmicos de inteligencia artificial para el monitoreo de derrames de petróleo
- Identificar, clasificar, extraer y estimar área del fluido derramado
- Generar un mecanismo de monitoreo de derrames de petróleo

Método

Sensores pasivos

El esquema general del método se describe en la Figura 1, la cual muestra las principales etapas del trabajo: a) adquisición y preprocesamiento de imágenes ópticas, b) procesamiento para la detección y cuantificación del petróleo utilizando técnicas de aprendizaje automático (ML) y datos espectrales, y c) resultados del modelo.

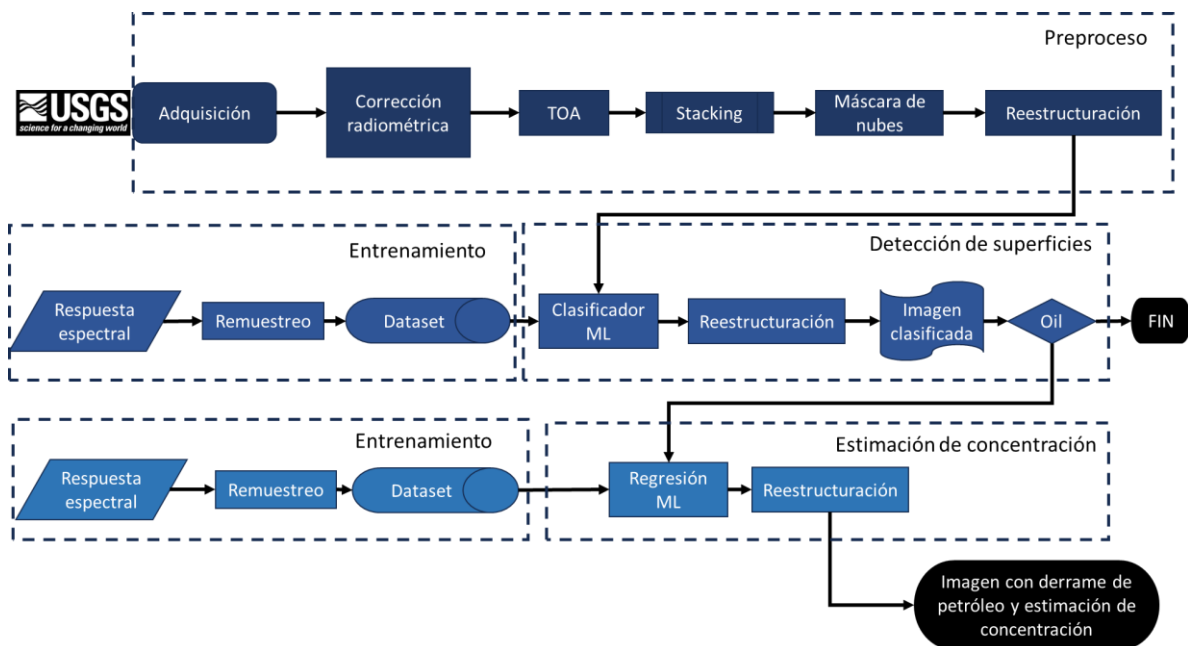


Figura 1. Esquema general del método. En el paso 1, se encuentra la obtención de las imágenes, corrección a TOA (Top of Atmosphere), apilamiento, enmascaramiento nubes y reestructuración. En el paso 2, la imagen ingresa al modelo de clasificación, previamente entrenado y validado con datos públicos de respuesta espectral, la salida del modelo se reestructuró de nuevo para obtener la imagen clasificada. Posteriormente, sólo si había petróleo, se continuo con el paso 3, que es el modelo de estimación de la concentración. Finalmente, la salida se reestructuró para obtener la imagen del resultado final con la presencia de petróleo y estimación de su concentración.

Datos

Datos espectrales para la segmentación de superficies

Se obtuvieron un total de 175 datos de respuesta espectral de cinco materiales/superficies que se utilizaron para el entrenamiento y la validación de los modelos de ML. Estos materiales fueron plástico, agua, suelo, vegetación y aceite, cada uno con 35 observaciones. Los datos se obtuvieron de Spectral library versión 7 (Kokaly *et al.*, 2017), ECOSTRESS Spectral library versión 1.0 (Meerdink *et al.*, 2019), Aster spectral library versión 2.0 (Baldrige *et al.*, 2009), y la Spectral Signature Library generada y proporcionada por la Comisión Nacional de Actividades Espaciales (CONAE, no date).

Los datos de las librerías se basan en información obtenida mediante diversas técnicas, como espectroscopía de laboratorio, espectroscopía de campo y datos de satélite. El propósito principal de estas mediciones es la cartografía de minerales, suelos, líquidos, compuestos orgánicos, compuestos artificiales y vegetación. Estos conjuntos de datos abarcan un amplio rango espectral, que incluye el ultravioleta, el infrarrojo cercano e infrarrojo de onda corta, con una cobertura desde los 200 nm hasta los 2500 nm. Esta amplia gama de mediciones permite capturar información detallada sobre las características y propiedades de los diferentes materiales y superficies analizadas en el estudio.

Los datos de respuesta espectral de las bibliotecas originales mostraron variaciones continuas en diferentes rangos espectrales. Con el fin de estandarizar los datos y adaptarlos a los rangos espectrales de interés en este estudio, se realizó un

remuestreo considerando los siguientes rangos espectrales de seis bandas: Azul:0.453-0.520, Verde:0.528-0.592, Rojo:0.635-0.682, NIR:0.790-0.891, SWIR1:1.557-1.717, SWIR2:1.960-2.162 μm . Estos rangos espectrales fueron seleccionados con base en las capacidades espectrales de los satélites Landsat 4-5 TM, Landsat 7 ETM+ y Landsat 8 OLI, utilizados en este estudio.

Los datos espectrales de interés se organizaron y almacenaron en una base de datos, donde cada columna es la representación lexicográfica de cada banda (*características X*), mientras que las filas corresponden a las observaciones de los materiales seleccionados junto con sus etiquetas respectivas (*Y*) (Tabla 1). Esta estructura de base de datos permitió utilizar los datos en los modelos de clasificación de ML, donde se utilizaron las características espectrales (*X*) como variables predictoras para clasificar y predecir la etiqueta de los materiales (*Y*). Al tener los datos estructurados de esta manera, fue posible entrenar y evaluar los modelos de clasificación.

Tabla 1. Ejemplo de base de datos utilizada para el modelo de clasificación de imágenes. Las bandas se utilizaron como características (X), y cada material con su respectiva etiqueta se utilizó como la salida deseada (Y)

Azul (X1)	Verde (X2)	Rojo (X3)	NIR (X4)	SWIR1 (X5)	SWIR2 (X6)	Nombre	Etiqueta
325	...	...	...	...	...	Aceite	1
258	...	...	...	...	...	Plástico	2
165	...	...	...	...	...	Suelo	3
485	...	...	...	...	...	Vegetación	4
260	...	...	...	...	...	Agua	5

Datos espectrales para la estimación de concentración

La estimación de concentración de petróleo en las imágenes utilizó como referencia un total de 40 datos de respuesta espectral del petróleo, los cuales abarcan diferentes concentraciones y espesores, tal como se detalla en la USGS Spectral Library Version 7 (Kokaly *et al.*, 2017). Para este estudio, se utilizaron concentraciones específicas de petróleo que incluyeron valores de 1, 23, 40, 60, 75 y 92%, el porcentaje restante corresponde a agua.

Imágenes satelitales

Se utilizaron imágenes del sensor TM (Thematic mapper) del satélite Landsat 5, ya que su cobertura temporal nos permitió disponer de imágenes relevantes del derrame ocurrido en la plataforma Deepwater Horizon en 2010. También se utilizaron imágenes del sensor OLI (Operational Land Imager) de Landsat8 de zonas libres de petróleo con el objetivo de corroborar la funcionalidad del modelo de aprendizaje automático (ML).

Estas imágenes de satélite están disponibles de forma gratuita para su descarga en la plataforma USGS Earth Explorer. En la Tabla 2 se proporcionan las fechas y los detalles de identificación de las imágenes utilizadas.

Tabla 2. Detalles de las imágenes utilizadas. La información incluye el satélite, la fecha (formato ISO), el Sistema de Referencia Mundial-2 (WRS-2 Path/WRS-2 Row) y el identificador de escena Landsat.

	Satélite	Fecha (YYYY-MM-DD)	WRS-2 Path/Row	LANDSAT_SCENE_ID
Imágenes con aceite. Derrame del 2010	Landsat 5TM	2010-05-09	021/040	LT50210402010129GNC01
		2010-05-25	021/040	LT50210402010145EDC00
		2010-06-26	021/040	LT50210402010177GNC01
		2010-06-10	021/040	LT50210402010161EDC00
		2010-07-12	021/040	LT50210402010193EDC00
	Landsat 7ETM+	2010-05-01	021/040	LE70210402010121EDC00
		2010-05-17	021/040	LE70210402010137EDC00
		2010-07-04	021/040	LE70210402010185EDC00
Imágenes sin aceite	Landsat 8OLI	2014-01-12	021/040	LC80210402014012LGN01
		2017-12-08	019/045	LC80190452017342LGN00
		2018-01-07	021/046	LC80210462018007LGN00
		2018-01-07	021/047	LC80210472018007LGN00
		2019-12-17	040/040	LC80400402019351LGN01
		2021-01-05	015/040	LC80150402021005LGN01
		2021-01-31	021/045	LC80210452021031LGN00
		2021-02-23	038/038	LC80380382021054LGN00
		2021-11-26	018/045	LC80180452021330LGN00
		2022-01-24	015/037	LC80150372022024LGN00

Las imágenes obtenidas fueron preprocesadas por el método descrito por Chander *et al.*, 2009 para obtener valores de reflectancia TOA (Top of Atmosphere) sin unidades. Además, con el objetivo de enmascarar las áreas cubiertas por nubes y minimizar los errores causados en los modelos de ML, se aplicó un método de detección de nubes utilizando la implementación CloudMasking de QGIS 3.10.6-A Coruña (QGIS Development Team, 2018). Esta implementación se basa en los métodos propuestos por Zhu and Woodcock (2012); Zhu, Wang and Woodcock (2015), que permiten identificar y enmascarar las áreas afectadas por la presencia de nubes en las imágenes.

Las bandas calibradas se apilaron (*stacking*), lo que generó una pila de imágenes con dimensiones de m píxeles $\times$ n píxeles $\times$ 6 bandas (por ejemplo, 3 píxeles $\times$ 3 píxeles $\times$ 6 bandas). Esta pila de bandas se reestructuró para obtener una representación lexicográfica del apilamiento de bandas, en una matriz de dimensión de m' filas $\times$ 6 columnas (por ejemplo, 9 filas $\times$ 6 columnas, donde $m' = m \times n$) (Figura 2). En esta matriz, cada columna corresponde a una banda o característica específica, y cada fila representa el valor espectral del píxel en relación con cada banda. Esta representación lexicográfica de los datos se utilizó como entrada para los modelos de ML.

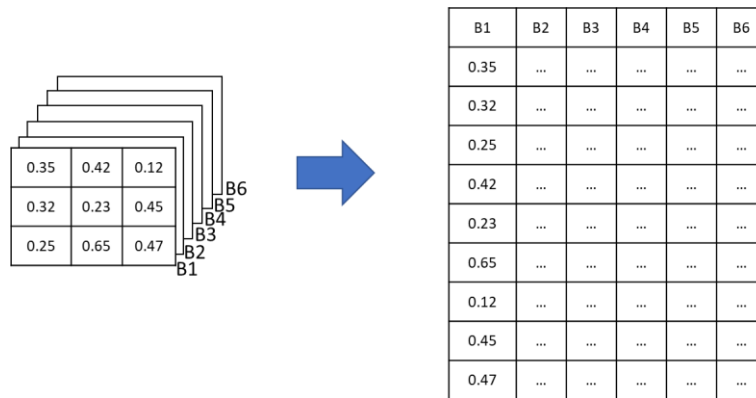


Figura 2. Reestructuración de conjunto de bandas para crear una matriz, donde cada columna corresponde a las bandas y cada fila son los valores de los píxeles.

Las imágenes se manipularon como matrices numpy para facilitar su manejo; por ello, se utilizaron las librerías NumPy 1.20.2 (Harris *et al.*, 2020), Pandas 1.2.3 (McKinney, 2010) y Matplotlib 3.4.1 (Hunter, 2007). Al tratarse de imágenes con aspectos geoespaciales, también se utilizó la librería GDAL 3.1.4 (GDAL/OGR contributors, 2021), todo ello en el entorno Anaconda 4.10.0 (Distribution, 2020), y Python 3.7.10 (Van Rossum and Drake, 2009).

Análisis exploratorios y estadísticos

Se llevaron a cabo análisis exploratorios y pruebas estadísticas para observar las diferencias y similitudes en la respuesta espectral entre las superficies estudiadas. Para determinar la diferencia de medias entre los materiales, se utilizó un análisis de varianza (ANOVA) basado en el método de Fisher (Fisher, 1992). Este análisis fue realizado utilizando el paquete "stats" en su versión 4.0.4 (R Core Team, 2021).

En caso de que el resultado del ANOVA fuera estadísticamente significativo, se procedió a realizar un análisis post hoc de Tukey (Tukey, 1977) para determinar los

grupos que presentaban diferencias significativas entre ellos. Para llevar a cabo este análisis post hoc, se utilizó el paquete "agricolae" en su versión 1.3.3 (de Mendiburu, 2020).

Adicionalmente, se realizó un análisis de correlación utilizando el método de Pearson (Freedman, Pisani and Purves, 2007) para determinar la relación entre los datos correspondientes a cada material. Este análisis fue realizado utilizando el paquete "stats" en su versión 4.0.4 (R Core Team, 2021). Los resultados de las correlaciones fueron representados visualmente en un mapa de calor (heatmap), mientras que el dendrograma fue dibujado en las filas y columnas utilizando la función "heatmap.2" de la biblioteca "gplots" en su versión 3.1.1 (Warnes *et al.*, 2020).

Para analizar la agrupación de las clases con métodos no supervisados, se llevó a cabo un análisis de componentes principales (principal component analysis, PCA) (Pearson, 1901). Además, se utilizó el análisis *t-Distributed Stochastic Neighbor Embedding* (T-SNE) propuesto por van der Maaten and Hinton, (2012). Estos análisis fueron realizados utilizando el paquete "stats" en su versión 4.0.4 (R Core Team, 2021) y el paquete Rtsne en su versión 0.15 (Krijthe, 2015). El objetivo de estos análisis fue comprender el comportamiento de los datos, tanto mediante transformaciones lineales (PCA) como no lineales (T-SNE).

Para evaluar la distancia de separación entre los grupos o clústeres resultantes, se utilizó el análisis silhouette. Este análisis, realizado utilizando el paquete "cluster" en su versión 2.1.0 (Maechler *et al.*, 2015), permite medir la coherencia de los

patrones de agrupación. Los coeficientes de silhouette cercanos a 1 indican que las muestras están bien separadas de los conglomerados vecinos; un valor de 0 indica que la muestra se encuentra en o cerca del límite de decisión entre dos conglomerados vecinos; y valores negativos indican una asignación incorrecta de las muestras a un conglomerado (Rousseeuw, 1987; Pedregosa *et al.*, 2011).

Para analizar la relación entre la concentración y la respuesta espectral del aceite se utilizó un modelo de regresión lineal (Everitt, 1992). Este análisis se realizó tanto de manera general, multivariada, como para cada banda por separado, univariado. El paquete "stats" en su versión 4.0.4 (R Core Team, 2021) se empleó para llevar a cabo este análisis. El objetivo fue determinar si existe una relación lineal entre la concentración de aceite y los valores de respuesta espectral.

Modelos de Machine Learning

Los modelos de ML se entrenan para establecer una relación entre las entradas (datos de reflectancia) y salidas (superficies). Estos modelos se ajustan directamente a los datos, minimizando el error de predicción y evitando el sobreajuste (Verrelst *et al.*, 2012). El objetivo es encontrar un modelo que pueda generalizar y clasificar de manera precisa las respuestas espectrales en las clases correspondientes.

En contraste con los modelos paramétricos o estadísticos que dependen de una función de entrada-salida con un conjunto fijo de parámetros, el aprendizaje automático se basa en modelos no paramétricos, no lineales y flexibles (Verrelst *et al.*, 2012; Wolanin *et al.*, 2019). Esto significa que el aprendizaje automático tiene la

capacidad de buscar relaciones entre las entradas y las salidas sin estar limitado por una estructura paramétrica predefinida (Verrelst *et al.*, 2012). Esta flexibilidad y capacidad de adaptación son ventajas clave del enfoque de aprendizaje automático en comparación con los modelos paramétricos o estadísticos más tradicionales.

Detección de superficies con seis bandas

La primera etapa del procesamiento se enfocó en la detección y clasificación de cinco superficies: aceite, plástico, vegetación, suelo y agua. Para lograr esto, se utilizó un enfoque de clasificación supervisada y multiclase, ya que había más de dos clases a distinguir (Pedregosa *et al.*, 2011).

En este enfoque, se aplicó la estrategia de clasificación uno contra el resto (*one-vs-rest*), la cual implica ajustar un clasificador para cada clase, donde cada clasificador se entrena para distinguir la clase de interés de todas las demás clases. Dicho enfoque es ampliamente utilizado en tareas de clasificación multiclase debido a su eficiencia computacional (Pedregosa *et al.*, 2011).

En este estudio, se implementaron y evaluaron varios algoritmos de aprendizaje supervisado: *Support Vector Machine* (SVM) con *kernel* lineal (SVM-lineal) y *kernel* radial (SVM RBF), *K-Nearest Neighbors* (KNN), *Decision Tree* (DT) y *Random Forest* (RF), utilizando la librería Scikit-learn en su versión 0.24.1 (Pedregosa *et al.*, 2011).

Detección de superficies con selección de variables

En la segunda parte del procesamiento, se llevó a cabo la selección de las variables más relevantes utilizando el método Gini del modelo Random Forest, con el cual se identifican las variables con mayor contribución a la precisión del modelo (Pedregosa *et al.*, 2011).

Una vez seleccionadas las variables más importantes, se utilizaron los enfoques de aprendizaje supervisado y uno contra el resto descritos previamente. Estos modelos fueron implementados utilizando la librería Scikit-learn en su versión 0.24.1, que proporciona una amplia gama de herramientas y funcionalidades para el aprendizaje automático (Pedregosa *et al.*, 2011). El uso de estos enfoques y modelos permitió evaluar la precisión de la clasificación y determinar cuál de ellos ofrecía el mejor rendimiento después de la selección de variables más relevantes.

Estimación de concentración

En la tercera etapa del procesamiento de las imágenes ópticas, se llevó a cabo la estimación de la concentración de petróleo a partir de los datos de respuesta espectral, una vez detectada la presencia de aceite en la imagen. Dado que esta tarea implicaba la predicción de un valor continuo, se trató como un problema de regresión. Para ello, se utilizó un modelo de regresión basado en el enfoque de *Random Forest* implementado en la librería Scikit-learn en su versión 0.24.1 (Pedregosa *et al.*, 2011).

Es importante destacar que esta parte del trabajo solo se lleva a cabo si se detecta la presencia de petróleo. En ese caso, se estima la concentración de petróleo en la misma zona donde se ha detectado la mancha.

Entrenamiento y búsqueda de hiperparámetros

En este estudio, los conjuntos de datos se dividieron en un conjunto de entrenamiento y un conjunto de validación. El conjunto de entrenamiento representó el 67% de los datos, mientras que el conjunto de validación comprendió el 33% restante.

Dado que no se conocen los mejores valores para los hiperparámetros de los modelos, se realizó una búsqueda exhaustiva de todas las combinaciones posibles de valores de parámetros predefinidos mediante el método de búsqueda en gradilla (GridSearchCV) proporcionado por el paquete "model_selection" de la librería Scikit-learn en su versión 0.24.1 (Pedregosa *et al.*, 2011). Esta técnica permitió probar y evaluar múltiples combinaciones de hiperparámetros para los modelos de clasificación y concentración.

Evaluación de precisión de los modelos

La evaluación de la precisión es una parte fundamental para evaluar el rendimiento de los modelos de aprendizaje automático en clasificación (Bar, Parida and Pandey, 2020). En este estudio, se utilizaron conjuntos de entrenamiento y validación para evaluar cada modelo en términos de precisión.

Se calcularon las métricas Exactitud (*Accuracy*), *Exactitud Balanceada* (*Balanced Accuracy*), *Kappa de Cohen*, *Precisión y Recuperación* (*Precision-Recall*), y F1, utilizando el paquete metrics proporcionado por la librería Scikit-learn en su versión 0.24.1 (Pedregosa *et al.*, 2011). Estas métricas proporcionan una medida cuantitativa del rendimiento del modelo en términos de su capacidad para clasificar correctamente las muestras.

Para evaluar el rendimiento del modelo de regresión lineal y el modelo de aprendizaje automático para la estimación de la concentración, se utilizaron Error Absoluto Medio (MAE), *Error Cuadrático Medio* (MSE), *Raíz del Error Cuadrático Medio* (RMSE), R^2 y *Error Máximo* (ME).

Estas métricas se calcularon para todo el conjunto de datos, lo que proporciona una medida general del rendimiento del modelo. Para el modelo de aprendizaje automático, también se calcularon estas métricas por separado para los conjuntos de entrenamiento y validación, lo que permitió evaluar su rendimiento en cada conjunto de datos.

Adicionalmente, se utilizó una imagen sintética (matriz) con seis bandas, que coincidían con los rangos de longitud de onda del sensor Landsat, para probar el modelo de aprendizaje automático en la estimación de la concentración. Estos datos fueron ordenados de mayor a menor concentración para asegurar que el gradiente de concentración fuera evidente en la imagen sintética.

Postproceso

En la tarea de clasificación, el resultado obtenido fue un vector que asignaba a cada observación la etiqueta de la clase con la mayor probabilidad de pertenencia. Esto se determinó según los valores de respuesta espectral presentes en la imagen.

En el caso de la presencia de la clase de aceite en la imagen, se obtuvo un resultado adicional para la tarea de regresión. En este caso, cada observación del resultado contenía el valor estimado de concentración correspondiente, es decir, se estimó la concentración de petróleo para aquellos píxeles que se clasificaron como aceite.

Una vez obtenidos los vectores resultantes de la clasificación y regresión, se llevó a cabo una reestructuración inversa para generar una matriz cuadrada. Esta reestructuración tuvo como objetivo crear una imagen clasificada donde se pudiera ubicar cada una de las cinco superficies utilizadas en el entrenamiento del modelo. Del mismo modo, se realizó para el resultado del modelo de regresión, lo que permitió obtener la distribución y el gradiente de concentración estimado en la mancha de aceite. Sin embargo, esta reestructuración solo se llevó a cabo si se detectaba la presencia de aceite en la imagen.

Sensores activos

En la Figura 3 se presenta el esquema general del método utilizado para los sensores activos en este estudio. El proceso comienza con la adquisición de las imágenes Sentinel-1 SAR, las cuales se someten a un preproceso para obtener

imágenes Sigma0 con valores en decibeles (db). A continuación, se dividen las imágenes en recortes de tamaño 2048x2048 para ambas polarizaciones (VV, VH).

Después de obtener los recortes, se realizó el enmascaramiento o *ground truth* (GT), lo que resulta en un conjunto completo de datos que incluye imágenes Sentinel-1 Dual-Pol Sigma0 en db junto con su correspondiente GT. A continuación, se llevó a cabo el entrenamiento, validación y prueba de varios modelos de aprendizaje profundo, incluyendo *Redes Neuronales Convolucionales* (CNN), *Perceptrón Multicapa* (MLP) y una red compuesta por un *codificador-decodificador* (*encoder-decoder*) conocida como U-Net.

Finalmente, se seleccionaron los modelos con los mejores resultados para generar el esquema de detección-segmentación de derrames de petróleo. Este esquema se basa en la combinación de los modelos seleccionados y se utiliza para identificar y delimitar los derrames de petróleo en las imágenes procesadas.

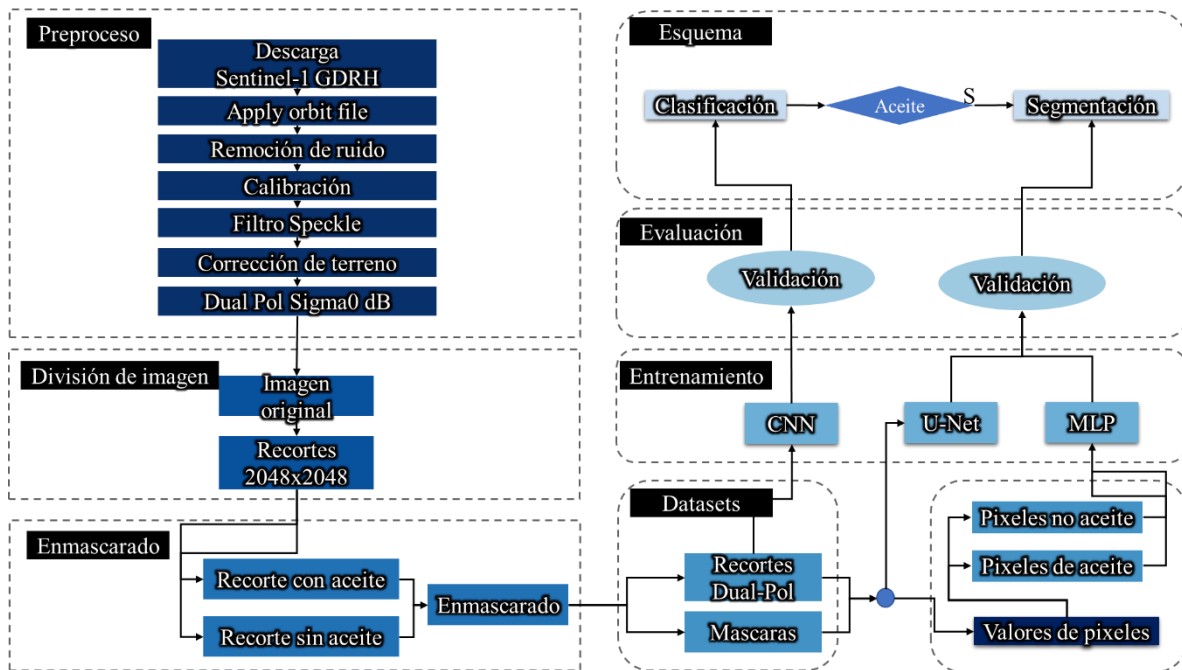


Figura 3. El esquema general del método abarca varias etapas. En primer lugar, se obtienen las imágenes Sentinel-1 SAR y se aplican técnicas de preprocesamiento, como la calibración, eliminación de ruido y correcciones, para obtener imágenes Sigma0 en decibeles (db). A continuación, se divide cada imagen en recortes de tamaño 2048x2048 para las dos polarizaciones (VV, VH). Después, se crea una máscara para cada recorte, lo que permite generar un conjunto de datos completo. Luego, se procede a entrenar y validar modelos de aprendizaje profundo, redes neuronales convolucionales (CNN), perceptrón multicapa (MLP) y la red de segmentación U-Net. Por último, se obtiene un esquema de dos pasos para la detección y segmentación de derrames de petróleo. Este esquema utiliza los modelos entrenados para identificar y delimitar los derrames de petróleo en las imágenes procesadas, proporcionando así una herramienta eficiente para detectar y segmentar los derrames de petróleo en imágenes Sentinel-1 SAR.

Datos y preproceso

Se obtuvieron un total de 685 imágenes Copernicus Sentinel-1 a través del portal ASF DAAC (<https://search.asf.alaska.edu/>), procesadas por ESA, y distribuidas en diferentes ubicaciones alrededor del mundo Figura 4. Para la selección y descarga de las imágenes, se utilizó la base de datos *Raw Incident Data*

(<https://incidentnews.noaa.gov/raw/index>) de la *National Oceanic and Atmospheric Administration* (NOAA), así como las detecciones realizadas por el sistema CleanSeaNet de la *European Maritime Safety Agency* (EMSA) (<https://www.emsa.europa.eu/>). Esto se hizo con el propósito de verificar la presencia de derrames de petróleo, ya que estas bases de datos cuentan con coordenadas XY de los derrames de petróleo reportados.



Figura 4. Ubicación de las 685 imágenes SAR Sentinel-1 GRD descargadas con base en los reportes de la NOAA y de la EMSA.

Cada una de las imágenes obtenidas fue sometida a un preproceso utilizando Sentinel Application Platform SNAP (ESA, 2018), siguiendo el procedimiento descrito en Filipponi (2019). En la etapa de filtrado de *speckle*, se utilizó el filtro Refined Lee (5x5) para reducir el ruido y mejorar la calidad de las imágenes. Este filtro es ampliamente utilizado en el procesamiento de imágenes SAR para mitigar el efecto de *speckle* (ruido multiplicativo) y mejorar la interpretación de las características presentes en las imágenes como los bordes (Filipponi, 2019).

Además, se utilizó el modelo digital de elevación (DEM) SRTM 1Sec HGT para corregir las distorsiones topográficas en las imágenes. El DEM proporciona información precisa sobre la elevación del terreno, lo que permite corregir las variaciones en la geometría de las imágenes SAR y obtener una representación más precisa de la superficie (Filipponi, 2019). Las imágenes se almacenaron en formato Sigma0 en valores de decibeles (db) con una resolución de 10 metros (Figura 5).

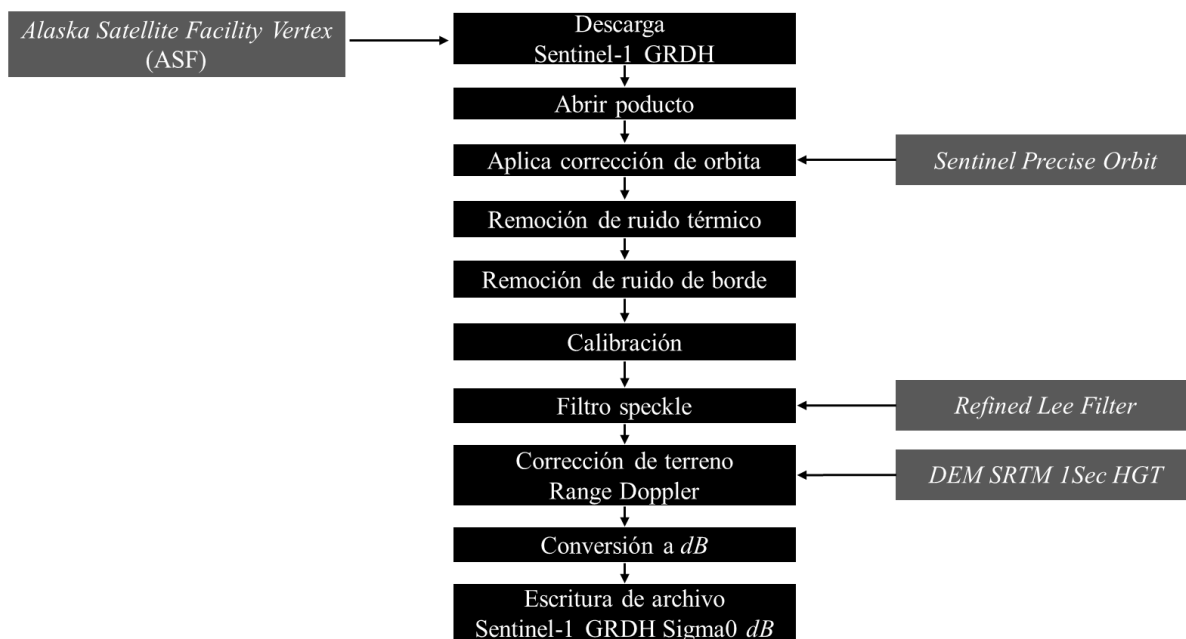


Figura 5. Método aplicado para el preproceso de las imágenes, desde la descarga hasta su almacenamiento como Sigma0 en db. El preproceso de las imágenes se realizó siguiendo el método descrito en (Filipponi, 2019). Se aplicó el filtro Refined Lee (5x5) para reducir el ruido de speckle (ruido multiplicativo). Además, se utilizó el modelo digital de elevación SRTM 1Sec HGT para corregir las distorsiones topográficas. Las imágenes resultantes se almacenaron en formato Sigma0 en decibeles (db).

Posteriormente, cada imagen se dividió en subimágenes de tamaño 2048x2048 para las dos polarizaciones (VV, VH) utilizando una implementación de GDAL versión 4.1.1 (GDAL/OGR contributors, 2021) en el lenguaje de programación Python versión 3.8 (Van Rossum and Drake, 2009). Estas subimágenes fueron seleccionadas de acuerdo a diferentes categorías: imágenes con presencia de aceite, basándose en las coordenadas de los reportes de la NOAA y la EMSA; imágenes sin aceite provenientes de tierra y zonas de mar abierto; e imágenes que presentaban zonas oscuras en mar abierto (conocidas como "look-alikes" en inglés), que no coincidían con las coordenadas de los reportes de derrames, las cuales podrían corresponder a diversas superficies y condiciones como viento bajo o nulo, celdas microconvectivas o surfactantes biológicos.

Conjunto de datos de Sentinel-1 SAR

Para las subimágenes de tamaño 2048x2048, se generaron máscaras utilizando la herramienta de etiquetado Labkit (Arzt *et al.*, 2022). Se enmascaró la parte correspondiente al derrame asignando la etiqueta de 1 (*foreground*), mientras que el resto de la imagen se etiquetó como 0 (*background*). Esto permitió crear máscaras para las imágenes con presencia de aceite, y crear el conjunto de máscaras verdaderas o *Ground Truth* (GT).

En cuanto a las imágenes sin aceite y *look-alikes*, como no se realizó una segmentación multiclase en este estudio, las máscaras se etiquetaron únicamente con el valor de 0 (*background*), indicando la ausencia de derrames. Ambos casos se consideraron como "No aceite" en la clasificación (Figura 6).

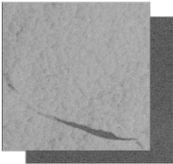
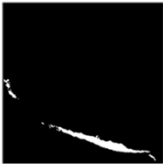
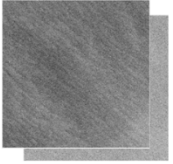
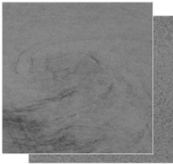
	Imagen	Ground truth
Aceite		 0 1 No aceite Aceite
No aceite		 0 1 No aceite Aceite
Look-alike (No aceite)		 0 1 No aceite Aceite

Figura 6. Subimágenes de 2048x2048 para ambas polarizaciones. La primera fila corresponde a un derrame de petróleo con su ground truth, la segunda fila es la imagen libre de aceite con su ground truth, y la tercera fila es look-alike con su ground truth. Se observa que el ground truth para el aceite tiene el valor de 1 y el fondo (background) de 0, y el ground truth para los otros dos casos es únicamente 0. En el caso de las imágenes libres de aceite y look-alikes, ambas se consideraron como No aceite.

El conjunto de datos completo utilizado para el entrenamiento y validación cuenta con 2,400 imágenes en total, divididas en 1,200 imágenes con presencia de petróleo (Aceite) y 1,200 imágenes libres de petróleo (No aceite), estas últimas están compuestas por imágenes libres de aceite y *look-alikes*. Cada imagen en el conjunto de datos está acompañada de su correspondiente máscara o *ground truth*.

Además, se obtuvieron 450 imágenes adicionales, compuestas por 150 imágenes para la clase Aceite, 150 imágenes de No Aceite y 150 imágenes *look-alikes*, para ser utilizadas como conjunto de prueba de los modelos. Estas imágenes tienen las mismas dimensiones que las utilizadas en el entrenamiento y validación. En total, el conjunto de datos completo cuenta con 2,850 imágenes, cada una con su respectivo *ground truth*, y todas tienen una resolución espacial de 10 metros.

Modelos de aprendizaje profundo (*Deep learning*)

Modelo de CNN para clasificación

Las redes neuronales convolucionales (CNN por sus siglas en inglés) son modelos de aprendizaje profundo que se utilizan para procesar datos en forma de matriz, como imágenes. Estos modelos están inspirados en la organización de la corteza visual animal y se diseñan para aprender características de bajo y alto nivel (Yamashita *et al.*, 2018).

Una CNN consta típicamente de tres tipos de capas o bloques: la capa de convolución, la capa de agrupación (*pooling*) y las capas totalmente conectadas (Figura 7). La capa de convolución realiza operaciones de convolución en la imagen de entrada para extraer características relevantes. La capa de agrupación reduce la dimensionalidad de las características obtenidas, preservando las características más importantes. Por último, las capas totalmente conectadas mapean las características extraídas a las clases de interés, generando una clasificación (O'Shea and Nash, 2015; Indolia *et al.*, 2018; Yamashita *et al.*, 2018).

La capa de convolución (Figura 7) utiliza kernels o filtros de diferentes tamaños, como 3x3, 5x5 o 7x7, para extraer características de los datos de entrada. Estos kernels son optimizables, lo que significa que sus pesos se ajustan durante el proceso de entrenamiento del modelo (Indolia *et al.*, 2018; Taye, 2023). Al aplicar los kernels a la imagen de entrada, se recorren vertical y horizontalmente las entradas para extraer características específicas. Esto da lugar a múltiples filtros y mapas de características, donde cada filtro representa una característica distinta de la imagen de entrada (Indolia *et al.*, 2018).

Después de aplicar los kernels en la capa de convolución, las salidas pasan a través de una función de activación no lineal. Esta función se utiliza para ajustar, saturar o limitar la salida generada por los kernels (Albawi, Mohammed and Al-Zawi, 2017; Yamashita *et al.*, 2018). La función de activación toma la decisión de activar o desactivar una neurona, creando así la salida correspondiente (Alzubaidi *et al.*, 2021).

La función de activación comúnmente utilizada en las capas de convolución es la *Rectified Linear Unit* (ReLU). Esta función, dada por la fórmula $\text{ReLU}(x) = \max(0, x)$, asigna un valor de cero a todas las entradas negativas y mantiene las entradas positivas sin modificar (Alzubaidi *et al.*, 2021). La elección de ReLU se debe a su simplicidad computacional y a su capacidad para introducir no linealidades en el modelo.

La capa de *Agrupación (Pooling)* es la segunda capa de las CNN (Figura 7), y tiene la función principal de realizar submuestreo de los mapas de características, manteniendo la información más relevante. También se le conoce como capa de reducción de dimensionalidad, ya que ayuda a reducir la complejidad de las capas posteriores y la cantidad de parámetros aprendibles (Albawi, Mohammed and Al-Zawi, 2017; Alzubaidi *et al.*, 2021).

El tipo más comúnmente utilizado de capa de agrupación es el Max-Pooling, el cual divide la imagen en subregiones, y se toma únicamente el valor máximo de cada subregión. El tamaño más utilizado para esta operación es de 2x2 (Albawi, Mohammed and Al-Zawi, 2017; Taye, 2023), lo que significa que reduce la dimensión de los mapas de características en el plano en un factor de 2, pero no altera el número de mapas de características (Yamashita *et al.*, 2018).

En la etapa final de las CNN, las características extraídas por las capas de convolución y agrupación se aplanan y se conectan a una o más capas completamente conectadas. Cada neurona de estas capas está conectada a todas las neuronas de la capa posterior. Luego de esta conexión, se aplica una función no lineal, como ReLU (Figura 7).

La capa final de la red neuronal debe tener el mismo número de neuronas que el número de clases en el problema de clasificación (Albawi, Mohammed and Al-Zawi, 2017; Yamashita *et al.*, 2018). La función de activación utilizada en esta capa final puede variar según la tarea de clasificación. Por ejemplo, en una tarea de

clasificación binaria, se suele utilizar la función sigmoide, que normaliza los valores a probabilidades entre 0 y 1 (Yamashita *et al.*, 2018).

Para optimizar los parámetros del modelo, se utiliza el método de optimización ADAM, que se basa en el algoritmo de descenso de gradiente estocástico. ADAM es ampliamente utilizado debido a su eficiencia y buen desempeño en la convergencia de los modelos de redes neuronales (Kingma and Ba, 2014).

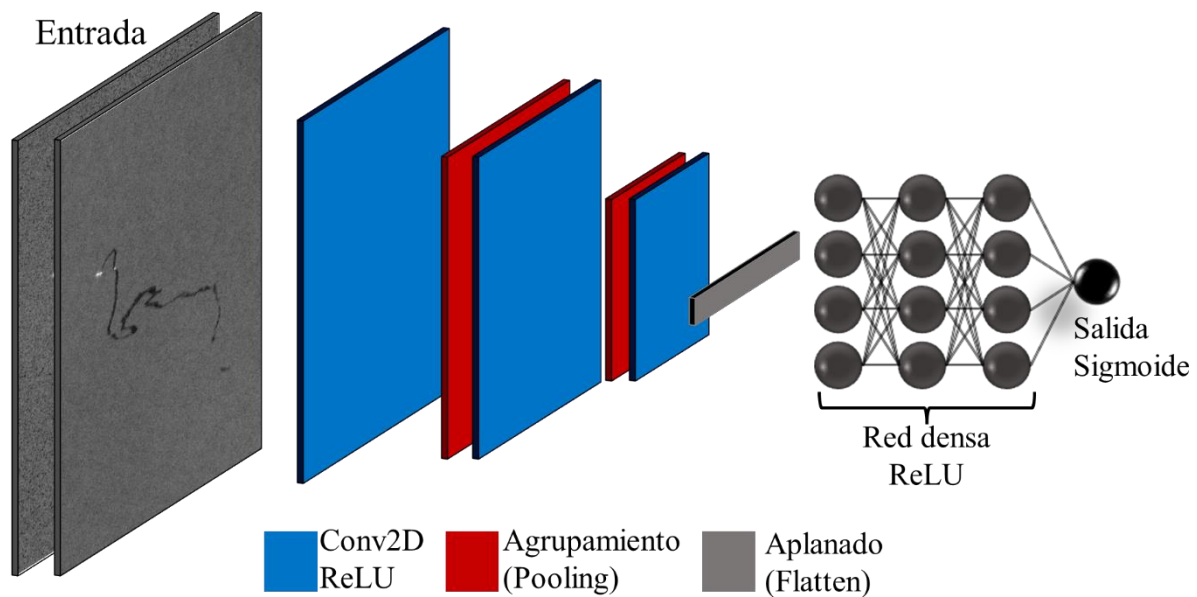


Figura 7. Esquema general de una CNN, donde se observa la entrada, seguida de las capas de convolución (Conv2D) y de agrupación (Pooling). Posteriormente las salidas de la última capa de convolución son aplanadas y conectadas a una red densa, donde cada capa puede tener diferente número de neuronas, y por último la neurona de salida con función de activación sigmoide la cual dará valores de probabilidad entre 0 y 1.

Configuración del modelo CNN

Se utilizó un modelo CNN para la clasificación de derrames de petróleo. Las imágenes de entrada correspondían a las subimágenes de Sentinel-1 Sigma0 db del conjunto de datos de entrenamiento, a las cuales se les aplicó una reducción de dimensión (*resize*) a un tamaño de 512x512 píxeles (40m de resolución) con la librería Numpy (Harris *et al.*, 2020), con la finalidad de reducir los costos computacionales, ya que el utilizar las imágenes de 2048x2048 saturaban la memoria RAM del equipo de cómputo. El conjunto de imágenes se dividió en set de entrenamiento (66%) y validación (34%) con el módulo Train Test Split de Scikit-learn (Pedregosa *et al.*, 2011).

Se realizaron varios experimentos con diferente número de capas convolucionales y de *agrupación* (*pooling*) que variaron de una hasta seis. También se exploraron diferentes números de mapas de características (filtros), incluyendo 16, 32 y 64. Se utilizó un kernel de tamaño 3x3 en las capas convolucionales y de 2x2 en las capas de agrupamiento.

Para las capas densas, se incluyeron una o dos capas ocultas con un número de neuronas que variaba entre 10 y 20 en cada capa. La función de activación utilizada en las capas convolucionales y densas fue ReLU, mientras que la capa de salida empleó la función de activación sigmoide.

Durante el entrenamiento de los modelos, se utilizó la métrica de *Exactitud* (*Accuracy*) para evaluar el rendimiento de los modelos en la clasificación de las imágenes. Esta métrica se calcula como la proporción de predicciones correctas en

comparación con el total de muestras (Goutte and Gaussier, 2005; Elkan, 2012). Además, se empleó la función de pérdida *Entropía Cruzada Binaria* (*Binary CrossEntropy*), que penaliza los errores de clasificación y busca minimizar la discrepancia entre las etiquetas reales y las predicciones del modelo en cada época de entrenamiento. Esto permite que los modelos se ajusten de manera adecuada y optimicen sus parámetros para mejorar la precisión de las predicciones (Ruby and Yendapalli, 2020). El entrenamiento de los modelos se realizó con 600 épocas.

Para el entrenamiento y validación de los modelos de CNN, se utilizó la biblioteca TensorFlow en su versión 2.10 (Martín Abadi *et al.*, 2022), con el lenguaje de programación Python en su versión 3.8 (Van Rossum and Drake, 2009). De todas las combinaciones que se realizaron, en total se entrenaron y validaron 90 modelos de CNN para la clasificación de imágenes con y sin derrame de petróleo.

Perceptrón multicapa para segmentación semántica

La segmentación semántica es una tarea en la que se asigna una etiqueta a cada píxel de una imagen, con el objetivo de crear un mapa detallado de las distintas clases u objetos presentes en la imagen (Ronneberger, Fischer and Brox, 2015). Para la tarea de segmentación se puede optar por el uso de algoritmos de ML o DL, basándose únicamente en los valores de los píxeles, o elegir enfoques más complejos que incorporan información contextual y espacial (Mo *et al.*, 2022). La amplia variedad de algoritmos disponibles ha abierto un abanico de posibilidades para aplicar la segmentación semántica en diversas áreas, como visión por computadora, robótica, medicina e incluso en aspectos de monitoreo y desastres.

Las redes neuronales artificiales (ANNs) o perceptrones multicapa (MLP) son modelos que se encuentran en las categorías de ML y DL debido a que pueden tener configuraciones simples y complejas. Estos modelos utilizan nodos o neuronas interconectadas en capas para realizar cálculos y aprender de los datos. Utilizando el algoritmo de retropropagación (backpropagation), los MLP pueden ajustar los pesos de las conexiones entre las neuronas y mejorar su rendimiento mediante la optimización iterativa. Esto les permite resolver una variedad de problemas, incluyendo la clasificación a nivel de píxeles, lo que resulta en una segmentación semántica de la imagen (Mehlig, 2019; Mo *et al.*, 2022).

La arquitectura de un MLP consta de una capa de entrada que recibe los datos de entrada, una o varias capas ocultas que procesan y transforman la información, y una capa de salida que proporciona el resultado final. La cantidad de capas ocultas y el número de neuronas en cada capa pueden variar dependiendo de la complejidad del problema. En el caso de un problema de clasificación binaria, la capa de salida consta de un solo nodo que produce una salida binaria (0 o 1) (Figura 8). Esta arquitectura puede ser tanto simple, con solo una capa de entrada y una de salida, como compleja, con múltiples capas ocultas y múltiples neuronas en cada capa (Mehlig, 2019).

Para poder utilizar un MLP para el análisis de imágenes, es necesario convertir la matriz de la imagen de entrada en un vector o tensor de orden uno, esta operación se conoce como “aplanado” (flatten) (Mehlig, 2019). Cada capa oculta en el MLP está interconectada mediante pesos, y cada capa oculta utiliza una función de activación, como ReLU o Tangente hiperbólica (Tanh). La capa de salida también

tiene una función de activación, que en el caso de un problema de clasificación binaria suele ser la función sigmoide, que produce valores entre 0 y 1 (Figura 8) (Mehlig, 2019).

Durante el entrenamiento de los modelos MLP, se utiliza el mecanismo de retropropagación, que se basa en el cálculo del gradiente descendente para ajustar los pesos y sesgos del modelo, permitiendo que el modelo aprenda y mejore sus predicciones (Kingma and Ba, 2014). También se utilizan algoritmos de optimización, como ADAM, para actualizar los parámetros de forma eficiente; se definen funciones de pérdida que miden la discrepancia entre las predicciones del modelo y las etiquetas reales; y se establecen métricas para evaluar el rendimiento del modelo (Rumelhart, Hinton and Williams, 1986).

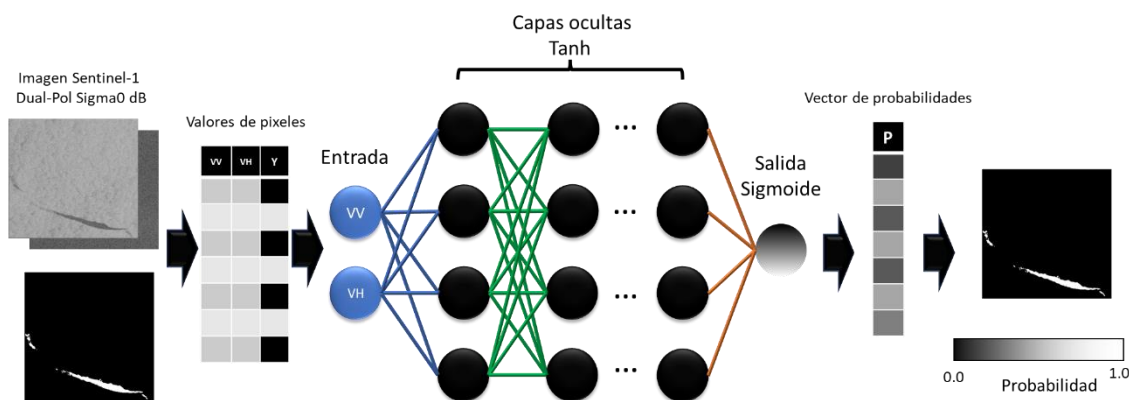


Figura 8. Esquema general de un modelo de MLP. Para las entradas se utilizaron los recortes de las imágenes Sentinel-1, en conjunto con su ground truth. Estas imágenes se aplanaron para obtener un vector para cada polarización con su respectiva etiqueta, por lo que la entrada al modelo son dos polarizaciones (VV, VH). La siguiente parte del modelo son las capas ocultas, las cuales cuentan con activación tanh, y por último la capa de salida con una neurona con activación sigmoide. La salida de este modelo será un vector de probabilidad de pertenencia a la clase (Aceite, No aceite) el cual será convertido nuevamente en matriz para tener la posición del píxel ya clasificado dentro de la imagen.

Configuración del modelo MLP

Para la implementación del modelo de MLP en la tarea de segmentación, se utilizó un conjunto de 50 imágenes Sentinel-1 en las polarizaciones VV y VH. Cada imagen se aplanó para obtener dos vectores correspondientes a cada polarización, y se generó un vector adicional para la máscara que contiene las etiquetas de los píxeles. Debido a un desbalance en la cantidad de píxeles entre las clases de aceite y no aceite, se realizó un submuestreo aleatorio de la clase no aceite basado en el número de ejemplos de clase de aceite para equilibrar los datos.

En total, se generó un conjunto de 61,598,228 de datos para cada polarización, divididos entre la clase de aceite (30,799,114 datos) y la clase no aceite (30,799,114 datos). Posteriormente, el conjunto de datos se dividió en un conjunto de entrenamiento (66%) y un conjunto de validación (34%) con el módulo Train Test Split de Scikit-learn (Pedregosa *et al.*, 2011).

Se implementaron arquitecturas de MLP con diferentes configuraciones de capas ocultas y número de nodos. Se probaron modelos con una, dos y tres capas ocultas, y se varió el número de nodos en cada capa entre 5, 10, 20, 50, 100, 150, 300 y 500. Se utilizó la función de activación tanh para las capas ocultas y la sigmoide para la capa de salida en todos los modelos. La función de pérdida utilizada fue Entropía Cruzada para problemas binarios (*Binary CrossEntropy*), y el entrenamiento de cada modelo se realizó con 1000 épocas.

Se realizaron pruebas utilizando imágenes del conjunto de prueba asignado. Estas imágenes fueron aplanadas para tener la misma dimensión que los datos de entrenamiento. La salida del modelo, que es un vector de probabilidad, se transformó en una matriz para que tuviera la misma dimensión que la imagen original (Figura 8).

Para el entrenamiento, validación y prueba de los modelos MLP, se utilizó la biblioteca TensorFlow en su versión 2.10 (Martín Abadi *et al.*, 2022), con el lenguaje de programación Python en su versión 3.8 (Van Rossum and Drake, 2009).

Modelo Dual-Pol U-net para segmentación semántica

En la tarea de segmentación, es importante considerar la información contextual y espacial de las imágenes, ya que la clasificación basada únicamente en los valores de los píxeles puede no ser suficiente. Para abordar esta necesidad, los modelos más adecuados son las redes neuronales convolucionales (CNN), que tienen la capacidad de capturar características locales y aprender patrones espaciales en las imágenes. Sin embargo, estos modelos clasifican toda la imagen, no nos da la localización precisa de lo que nos interesa (Yamashita *et al.*, 2018).

En este sentido, el modelo U-net es una arquitectura de red neuronal convolucional ampliamente utilizada en tareas de segmentación semántica de imágenes. Fue propuesto por Ronneberger, Fischer y Brox en 2015, y desde entonces ha demostrado su eficacia en la segmentación precisa de objetos en diversas áreas.

La arquitectura U-net se caracteriza por su capacidad para capturar información contextual y espacial de las imágenes, y está compuesta por dos partes principales: la parte de contracción y la parte de expansión. En la parte de contracción, se utilizan capas de convolución seguidas de funciones de activación y operaciones de agrupación (pooling) para extraer características de la imagen. Cada capa convolucional va incrementando progresivamente el número de características aprendidas, lo que permite capturar detalles y patrones relevantes en la imagen (Ronneberger, Fischer and Brox, 2015).

En la parte de expansión, se utilizan técnicas de *upsampling* y concatenación para fusionar información contextual con detalles espaciales. Esto permite una reconstrucción precisa de la imagen segmentada. A medida que el modelo se adentra en la parte de expansión, se van reduciendo gradualmente las características aprendidas en la parte de contracción, lo que ayuda a preservar la resolución y los detalles de la imagen (Ronneberger, Fischer and Brox, 2015).

La última capa convolucional utiliza una convolución 1x1 para asignar a cada mapa de características la etiqueta correspondiente a la clase de interés. Esto resulta en un mapa de segmentación por píxeles donde cada píxel se asigna a una clase específica. Originalmente, la U-net tiene 23 capas convolucionales en total (Ronneberger, Fischer and Brox, 2015).

En general, el modelo U-net se ha convertido en una elección popular para tareas de segmentación debido a su capacidad para capturar características de alto nivel y realizar una segmentación precisa. Su arquitectura flexible y adaptable permite su

aplicación en una variedad de dominios, incluyendo la detección y segmentación de derrames de petróleo en imágenes satelitales.

Configuración del modelo U-net para segmentación

Se realizaron experimentos y pruebas exhaustivas para encontrar la mejor configuración del modelo U-net en la tarea de segmentación de derrames de petróleo. Se evaluaron diferentes métricas de desempeño y funciones de pérdida, ajustando parámetros clave como el número de capas y filtros. El objetivo fue maximizar la precisión y calidad de la segmentación obtenida en las imágenes satelitales.

Inicialmente se utilizó el modelo U-net original propuesto por Ronneberger, Fischer y Brox en 2015, como base para el desarrollo del modelo de segmentación de derrames de petróleo. Se mantuvieron elementos clave del modelo original, como el número de capas convolucionales, tamaño del kernel (3x3), cantidad de mapas de características (64, 128, 256, 512, 1024), y operación de pooling (2x2); en la parte de expansión se utilizó la operación Conv2DTranspose de Keras (Chollet and others, 2015). Se emplearon la función de activación ReLU, la función de pérdida *Entropía Cruzada Binaria* y la métrica de evaluación *Exactitud*. Únicamente se realizó la modificación en las dimensiones de entrada del modelo para aceptar imágenes de dos canales correspondientes a las dos polarizaciones (VV, VH).

Se utilizó el conjunto de datos de imágenes asignado como entrenamiento y validación. Estas imágenes se sometieron a un proceso de reducción de dimensión (*resize*), reduciendo su tamaño de 2048x2048 píxeles (tamaño de píxel de 10m) a

512x512 píxeles (tamaño de píxel 40m) para disminuir la carga computacional. Luego, se aplicó una normalización a las imágenes utilizando el módulo Normalize de Keras (Chollet and others, 2015). Con las imágenes preparadas, se realizó la adaptación de la capa de entrada del modelo U-net para que fuera compatible con las imágenes de tamaño 512x512 y dos canales correspondientes a las polarizaciones (VV, VH).

Luego, se procedió a modificar la función de pérdida utilizada para el entrenamiento del modelo. Se optó por utilizar la función llamada Pérdida Focal (*Focal Loss*) (Lin et al., 2017). Esta función de pérdida penaliza los ejemplos fáciles de clasificar, y se enfoca en los ejemplos difíciles, además aborda el desequilibrio de clases entre el *foreground* (derrame de petróleo) y el *background* (sin derrame) durante el entrenamiento para la tarea de segmentación. El objetivo es enfocar el entrenamiento y evitar que la gran cantidad de ejemplos negativos fáciles de clasificar (*background*) dominen el proceso de aprendizaje (Lin et al., 2017; Basit et al., 2022). Esta función de pérdida (*Pérdida Focal*) introduce un factor de modulación $(1 - p_t)^{\gamma}$ en la función de entropía cruzada, donde p_t es la probabilidad de predicción del modelo y γ es un parámetro de enfoque que puede ajustarse entre 0 y 5 (Lin et al., 2017).

Además, en lugar de utilizar la métrica *Exactitud*, se utilizó *Intersección sobre la Unión* (IoU), que cuantifica el grado de superposición entre el ground truth y la predicción del modelo. El IoU toma valores entre 0 y 1, donde 1 indica una coincidencia perfecta entre el *ground truth* y la predicción (Niwattanakul et al., 2013; Rahman and Wang, 2016).

Después, se realizaron modificaciones adicionales en el modelo U-net. Se reemplazó el uso del módulo Conv2DTranspose por el módulo Upsampling2D de Keras (Chollet and others, 2015) en la parte de expansión. Además, se modificó el número de filtros utilizados en cada capa, utilizando 16, 32, 64, 128 y 256. También se entrenó un modelo con un orden específico de mapas de características: 16, 32, 64, 128, 256, 512 y 1024. En todos estos casos, las imágenes de entrada tenían un tamaño de 512x512 y dos canales (VV, VH). El entrenamiento se realizó durante 1000 épocas utilizando la función de *Pérdida Focal* y para el monitoreo del desempeño se utilizó la métrica IoU.

Una vez obtenida la arquitectura U-net con los mejores resultados, se procedió a entrenar y evaluar otros modelos basados en ella. Se exploraron dos variantes: Attention U-net (Att Unet) (Oktay *et al.*, 2018) y Unet++ (Zhou *et al.*, 2018). El modelo Attention U-net agrega módulos de atención en la concatenación de la parte de expansión, manteniendo la misma arquitectura que la U-net original (Oktay *et al.*, 2018). Por otro lado, el modelo Unet++ es una versión anidada del U-net, que busca mejorar el rendimiento segmentando la imagen a múltiples escalas y combinando las características de diferentes niveles de la red (Zhou *et al.*, 2018). Estos modelos fueron entrenados durante 1000 épocas, y evaluados para determinar cuál de ellos ofrece una mejor aproximación en la tarea de segmentación de derrames de petróleo.

Las implementaciones para el entrenamiento y validación de los modelos U-net se realizaron utilizando TensorFlow en su versión 2.10 (Martín Abadi *et al.*, 2022) en el lenguaje de programación Python en su versión 3.8 (Van Rossum and Drake, 2009).

Estas herramientas proporcionan una base sólida y eficiente para el desarrollo y entrenamiento de modelos de redes neuronales convolucionales como el U-net.

Características del hardware

Para el entrenamiento y prueba de los modelos se utilizaros dos equipos de cómputo con distintas características, uno para los modelos de ML de los sensores pasivos, y otro para los modelos de DL de los sensores activos.

Equipo utilizado para ML

Para el entrenamiento y validación de los modelos de ML para la detección de derrames de petróleo en imágenes de sensores pasivos se utilizó una computadora portátil con un procesador Intel(R) Core (TM) i7-8750H 2.20GHz de Octava Generación con 6 núcleos y 12 procesadores lógicos. La memoria RAM equipada en este equipo fue de 32GB, y SSD de 1TB. Además, contaba con una GPU NVIDIA GForce GTX 1050Ti (4GB).

Equipo utilizado para DL

Para el entrenamiento y prueba de los modelos de DL para la clasificación y segmentación de derrames de petróleo en imágenes de sensores activos, se utilizó una estación de trabajo (workstation) con procesador AMD Ryzen Threadripper 2990WX 4200MHz con 32 núcleos y 64 procesadores lógicos. La memoria RAM equipada fue de 128GB, SSD 1TB, HDD 4TB. Además, contaba con una GPU Quadro RTX 5000 (16GB). Cabe destacar que el entrenamiento, y prueba de los modelos se realizó aprovechando el compute en GPU del equipo.

Resultados

Sensores pasivos

Diferencias espectrales entre superficies

En la Figura 9A se muestran las distribuciones de la respuesta espectral de los cinco materiales analizados. Se puede observar que el agua presenta los valores de reflectancia más bajos, mientras que el plástico exhibe una amplia variación con valores cercanos a 0 y 1. El aceite, la vegetación y el suelo muestran una dispersión similar en sus distribuciones. Sin embargo, el análisis ANOVA reveló diferencias estadísticamente significativas ($p = 9.15e-140$) entre todos los materiales. El análisis post hoc de Tukey indicó que el aceite y el suelo son similares ($p = 6.05e-2$), al igual que el suelo y la vegetación ($p = 8.58e-01$), mientras que los demás materiales son diferentes entre sí.

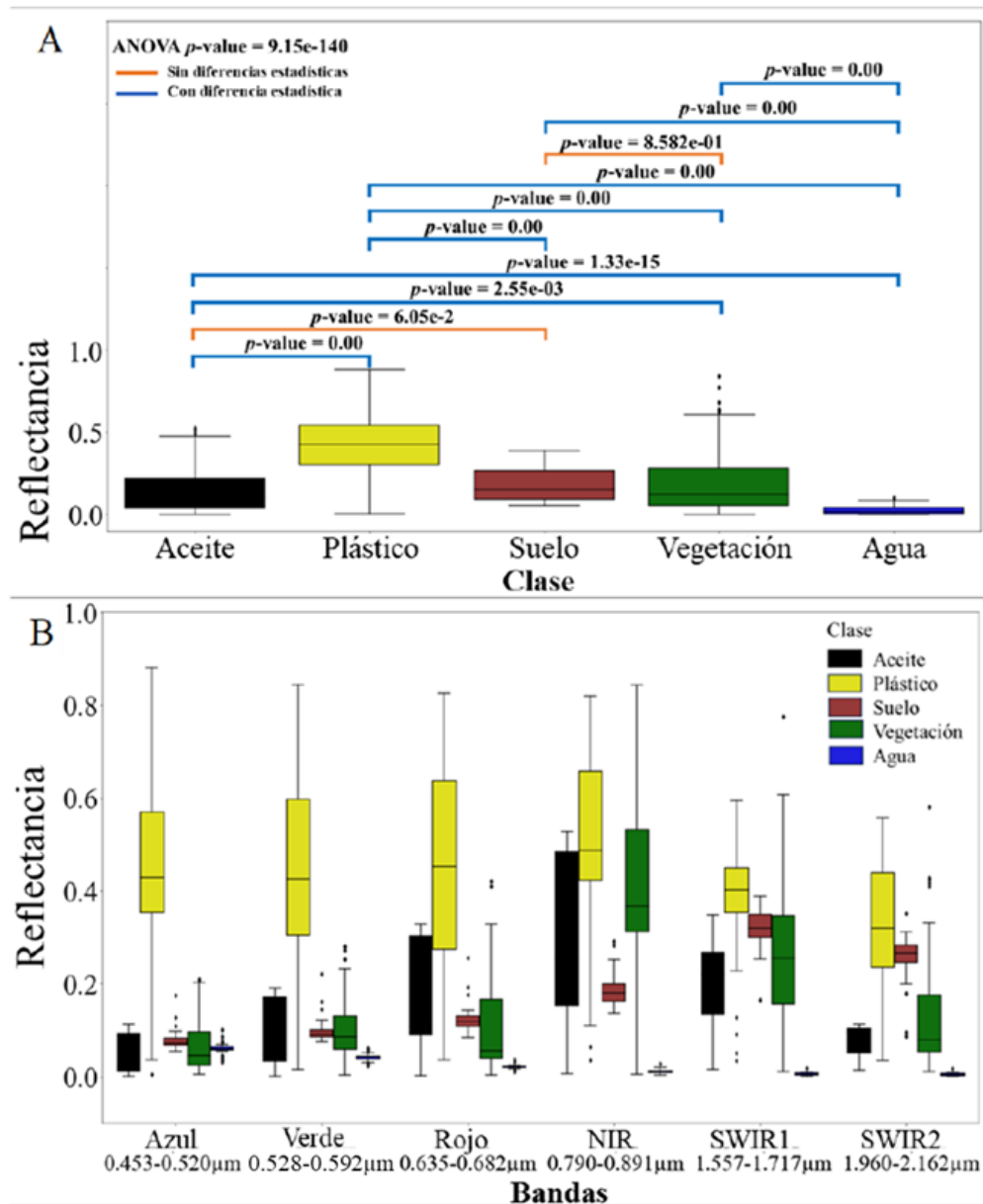


Figura 9. A) Las distribuciones de reflectancia de los cinco materiales estudiados mostraron diferencias estadísticas significativas en el análisis ANOVA ($p = 9.15\text{e-}140$). Según el análisis post hoc de Tukey, se encontraron similitudes entre el aceite y el suelo ($p = 6.05\text{e-}2$), así como entre el suelo y la vegetación ($p = 8.58\text{e-}01$). B) Se muestra la respuesta espectral de cada material en cada banda: aceite (negro), plástico (amarillo), suelo (marrón), vegetación (verde) y agua (azul).

Los análisis por banda (Figura 9B) revelaron que el plástico tiene la mayor dispersión en todas las bandas (0.0046 - 0.8884), mientras que el agua presenta los valores más bajos (0.0009 - 0.1443), especialmente en las bandas infrarrojas (NIR, SWIR1 y SWIR2). En las bandas azul, verde y roja, los materiales exhiben una dispersión similar, pero en las bandas infrarrojas se observa una mayor separación. En la banda NIR, se notaron diferencias entre suelo, vegetación y agua. El aceite, el plástico y la vegetación mostraron una dispersión similar en la banda NIR. Sin embargo, en las bandas SWIR se observaron diferencias más notables entre el aceite, el plástico y la vegetación. El suelo y el agua también exhibieron este comportamiento en las bandas SWIR1 y SWIR2.

Correlación y clustering de datos de materiales

Se realizó un análisis de correlación utilizando el coeficiente de Pearson y se visualizó en un mapa de calor (Figura 10). Se identificaron 6,747 observaciones con correlación positiva ($r > 0$), 8,450 con correlación negativa ($r < 0$) y 28 datos sin correlación ($r = 0$). Se observó que la mayoría de las correlaciones positivas correspondían a datos del mismo material, como los datos de agua que mostraron una alta correlación positiva ($r > 0.96$) y estaban agrupados. Se observó un comportamiento similar en algunos datos de suelo y vegetación ($r > 0.74$ y $r > 0.70$, respectivamente), formando grupos de datos correlacionados dentro de la misma categoría.

Se observaron correlaciones positivas entre algunos datos de diferentes clases, como agua con plástico, petróleo, vegetación y suelo ($r > 0.70$). También se

encontraron correlaciones positivas entre los datos de petróleo, plástico, suelo y vegetación ($r > 0.70$). Algunos datos de petróleo y plástico mostraron correlaciones positivas con otras clases ($r > 0.75$), formando grupos tanto dentro de la misma clase como con otras clases. Por otro lado, se encontraron correlaciones negativas entre los datos de vegetación y agua, que presentaron valores cercanos a -1.0 . También se observaron correlaciones negativas entre algunos datos de petróleo con agua ($r < -0.80$), plástico ($r < -0.70$) y suelo ($r < -0.71$), así como entre los datos de plástico con vegetación ($r < -0.78$) y suelo ($r < -0.72$).

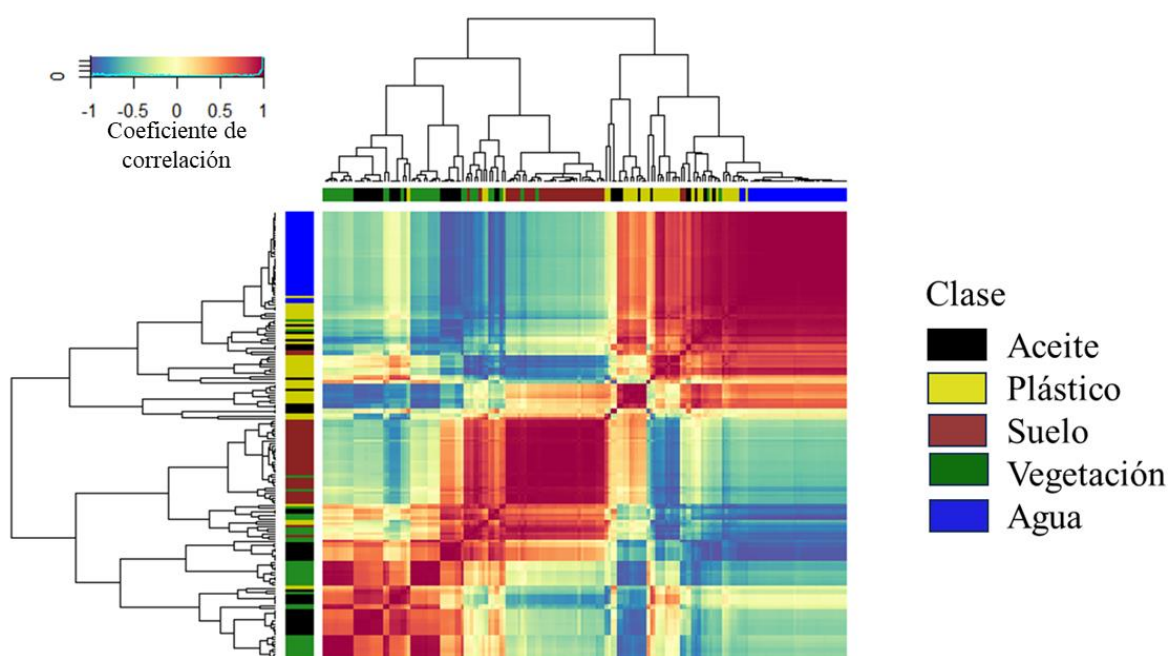


Figura 10. La matriz de correlación muestra las correlaciones entre todas las observaciones de los materiales, con 35 observaciones para cada clase. Se utiliza un gradiente de colores para representar los valores de correlación: rojo para correlaciones positivas (1.0), azul para correlaciones negativas (-1.0) y amarillo para datos no correlacionados (~0). Las clases de cada muestra se indican en los lados de las columnas y las filas, con negro para aceite, amarillo para plástico, marrón para suelo, verde para vegetación y azul para agua. El dendrograma en las filas y columnas muestra la agrupación jerárquica basada en las similitudes o diferencias entre los datos.

En el análisis de correlación de las bandas (Figura 11), se encontró que las bandas SWIR (*Short-Wave Infrared*) mostraron una alta correlación entre sí ($r = 0.90$), mientras que las bandas visibles formaron otro grupo correlacionado ($r > 0.88$). Sin embargo, se observó una correlación baja entre los dos grupos de bandas ($r < 0.68$). La banda NIR (*Near Infrared*) mostró una baja correlación tanto con las bandas visibles ($r < 0.75$) como con las bandas SWIR ($r < 0.77$).

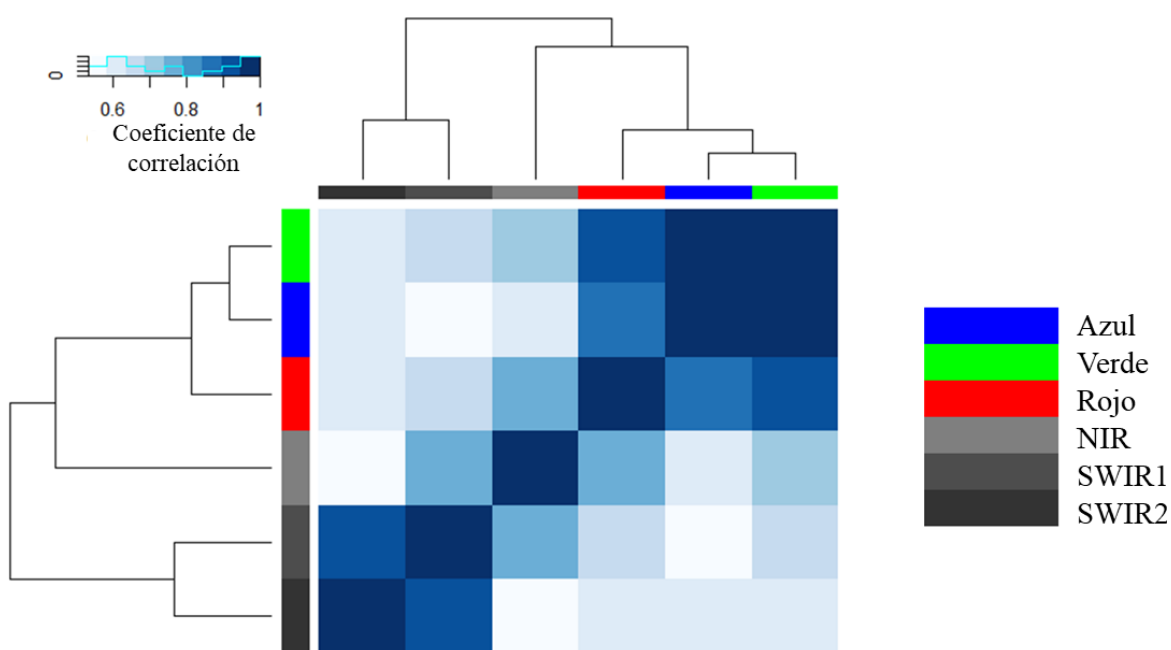


Figura 11. Matriz de correlación entre las bandas utilizadas. Se utiliza un gradiente de colores para representar los valores de correlación, donde los valores superiores a 0.85 se muestran en azul oscuro, los valores entre 0.75 y 0.80 se muestran en azul claro, y los valores inferiores a 0.70 se muestran casi blancos. Las bandas infrarrojas se representan en tonos de gris, mientras que las bandas RGB se representan en rojo, verde y azul. El dendrograma presente en las filas y columnas muestra la agrupación jerárquica de las bandas en función de sus similitudes o diferencias.

En el análisis PCA (Figura 12A), al utilizar dos componentes principales que explican el 90% de la variabilidad, no se observa una separación clara entre las clases, ya que la mayoría de los datos se encuentran cercanos en el espacio generado por los dos primeros componentes, a pesar de pertenecer a diferentes categorías. Sin embargo, los datos correspondientes al agua forman un grupo más compacto y se separan en cierta medida de las otras clases. Por otro lado, los datos asociados al plástico muestran una mayor variabilidad y dispersión en el espacio de PCA, superponiéndose con datos de otras clases. En cuanto a los datos de petróleo y vegetación, muestran una dirección similar en el espacio de PCA, mientras que los datos de suelo se encuentran próximos a los datos de plástico, vegetación y petróleo.

En el análisis T-SNE (Figura 12B), se observó una mayor agrupación y separabilidad de los datos, en comparación con el PCA. La puntuación media del Silhouette fue más alta en el T-SNE (0.38) que en el PCA (0.16), lo que indica una mejor estructuración de los grupos en el espacio. Sin embargo, se observaron algunas superposiciones de datos de petróleo con datos de agua y vegetación. Los datos de plástico formaron un grupo más cohesivo, aunque con algunas excepciones cerca de los datos de agua, vegetación, suelo y petróleo. Los datos de agua mostraron la mayor agrupación y no se encontraron datos de esta clase cerca de otras clases. Los datos de suelo y vegetación también estuvieron más cercanos entre sí.

En general, el análisis T-SNE demostró una mejor capacidad para diferenciar los materiales en comparación con el PCA, lo que sugiere que los enfoques no lineales pueden generar mejores resultados en la clasificación de los materiales.

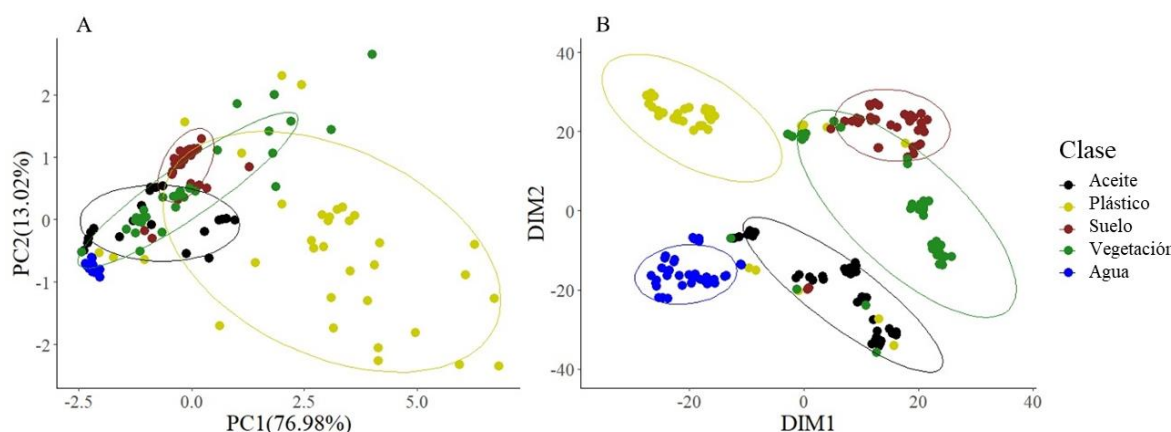


Figura 12. Análisis PCA y T-SNE de los datos analizados para las cinco clases trabajadas. A) PCA, sólo se utilizaron dos componentes, ya que juntos contienen el 90% de la variabilidad de los datos. B) T-SNE, sólo se utilizaron dos dimensiones. En ambos gráficos, los puntos negros y las elipses representan el petróleo, el amarillo el plástico, el marrón el suelo, el verde la vegetación y el azul el agua.

Modelos lineales de datos de concentración

Se utilizó un modelo lineal multivariado general para estimar la concentración de petróleo en las imágenes de satélite (Figura 13A). El modelo mostró una relación negativa entre la concentración de petróleo y la reflectancia, con una pendiente de -270. Esto indica que a medida que aumenta la reflectancia, la concentración de petróleo disminuye, y viceversa.

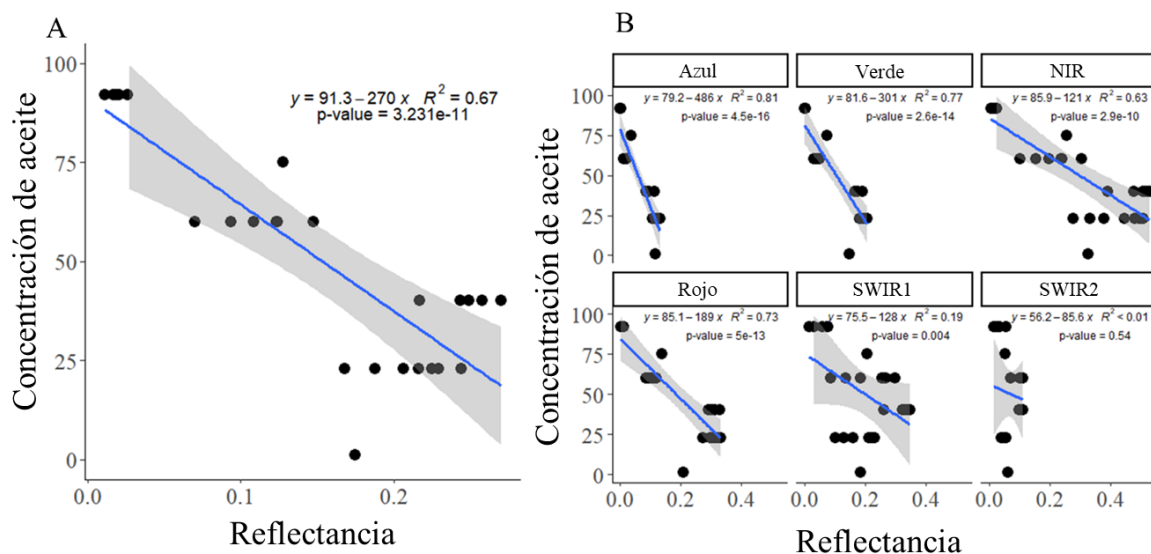


Figura 13. Los modelos lineales para los datos de concentración de petróleo. El eje X es la reflectancia y el eje Y es la concentración. A) Modelo general, construido con la reflectancia media de las seis bandas utilizadas, la pendiente de este modelo fue estadísticamente significativa (valor $p = 3.23e-11$), con valor de ajuste $R^2 = 0.67$, y tendencia negativa (-270). B) El modelo de concentración lineal para cada banda utilizada. Se observa una tendencia negativa para todas las bandas, que es significativa excepto para la banda SWIR2 (valor $p = 0.54$). Los valores de ajuste R^2 , así como el valor p y la ecuación lineal, se muestran en los gráficos para cada banda.

La tendencia negativa del modelo lineal general (Figura 13A) sigue presente en los modelos lineales univariados de concentración y reflectancia para cada banda estudiada (Figura 13B). Para la banda SWIR2, el modelo lineal no fue significativo (pendiente = 85.6, $p = 0.54$), pero para las demás bandas sí lo fue (valor p de la banda azul = $4.5e-16$, valor p de la banda verde = $2.6e-14$, valor p de la banda roja = $5e-13$, valor p de la banda NIR = $2.9e-10$, valor p de la banda SWIR1 = 0.004). Los valores de pendiente y R^2 fueron mayores para la banda azul (pendiente = 486, $R^2 = 0.81$), seguida de la verde (pendiente = 301, $R^2 = 0.77$), luego la roja (pendiente = 189, $R^2 = 0.73$), NIR (pendiente = 121, $R^2 = 0.63$) y SWIR1 (pendiente = 128, $R^2 = 0.19$).

Los modelos lineales utilizados mostraron una relación negativa entre la concentración de petróleo y la reflectancia. Sin embargo, estos modelos solo explicaron alrededor del 67% de la variación de los datos, lo que indica que existen otros factores que influyen en esta relación. Por lo tanto, se requiere explorar diferentes enfoques para mejorar la predicción de la concentración de petróleo basada en la reflectancia.

Modelos de ML

Los modelos ML utilizados se dividieron en tres secciones: detección de superficies utilizando las seis bandas de las imágenes Landsat, detección de superficies con modelos generados utilizando variables importantes según el método de impurezas de Gini, y estimación de la concentración de petróleo. En cada sección, se ajustaron hiperparámetros, se evaluaron las puntuaciones de precisión y se realizaron pruebas con los modelos seleccionados.

Modelo de detección de superficies con seis bandas

Hiperparámetros y scores

Los valores óptimos de los hiperparámetros que se obtuvieron mediante una búsqueda en gradilla se presentan en la Tabla 3. Para evaluar la precisión de los modelos de aprendizaje automático, se utilizaron seis métricas: *Exactitud*, *Exactitud Balanceada*, *Kappa*, *Precisión*, *Recuperación* y *F1*. Los resultados de estas métricas se muestran en la Figura 14.

Tabla 3. Valores e hiperparámetros utilizados para el entrenamiento de los modelos de clasificación

Modelo	Hiperparámetros y sus valores
<i>K-Nearest Neighbors</i> (KNN)	Leaf_size: 1 N_neighbors: 5 P: 2 Metric: Minkowski
<i>Support Vector Machines</i> (SVM_lineal)	C: 8 Degree: 1 Kernel: linear
<i>Support Vector Machines RBF</i> (SVM_rbf)	C: 17 Degree: 1 Gamma: 1
<i>Decision tree</i> (Dec_tree)	Max_Depth: 2 Min_simples_Split: 6
<i>Random Forest</i>	Max_Depth: 4 Min_simples_Split: 11 N_estimators: 18

Random Forest fue el modelo con mejor rendimiento en la predicción de las distintas clases, con los valores más altos en todas las métricas tanto para el conjunto de entrenamiento (0.97, 0.97, 0.96, 0.97, 0.97, 0.97) como de validación (0.90, 0.90, 0.87, 0.90, 0.90, 0.90), seguido de SVM RBF (0.93, 0.93, 0.91, 0.93, 0.93 y 0.93 para el conjunto de entrenamiento, y 0.88, 0.89, 0.85, 0.88, 0.88 y 0.88 para el conjunto de validación), y KNN (0.91, 0.92, 0.89, 0.91, 0.91 y 0.91 para el conjunto de entrenamiento, y 0.88, 0.89, 0.87, 0.88, 0.88 y 0.88 para el conjunto de validación). El modelo con los valores más bajos fue SVM (0.83, 0.85, 0.78, 0.83, 0.83 y 0.83) para el conjunto de entrenamiento y (0.85, 0.86, 0.81, 0.85, 0.85 y 0.85) para el de validación. *Decision Tree*, a pesar de tener valores altos en el conjunto de entrenamiento (0.96, 0.96, 0.95, 0.96, 0.96 y 0.96) obtuvo valores más bajos en el conjunto de datos de validación (0.85, 0.86, 0.81, 0.85, 0.85 y 0.85). En este sentido, este último modelo muestra una gran diferencia en los resultados de las métricas entre los conjuntos de entrenamiento y validación, en comparación con el

modelo KNN, donde las diferencias en las métricas entre ambos conjuntos son más cercanas a 0.

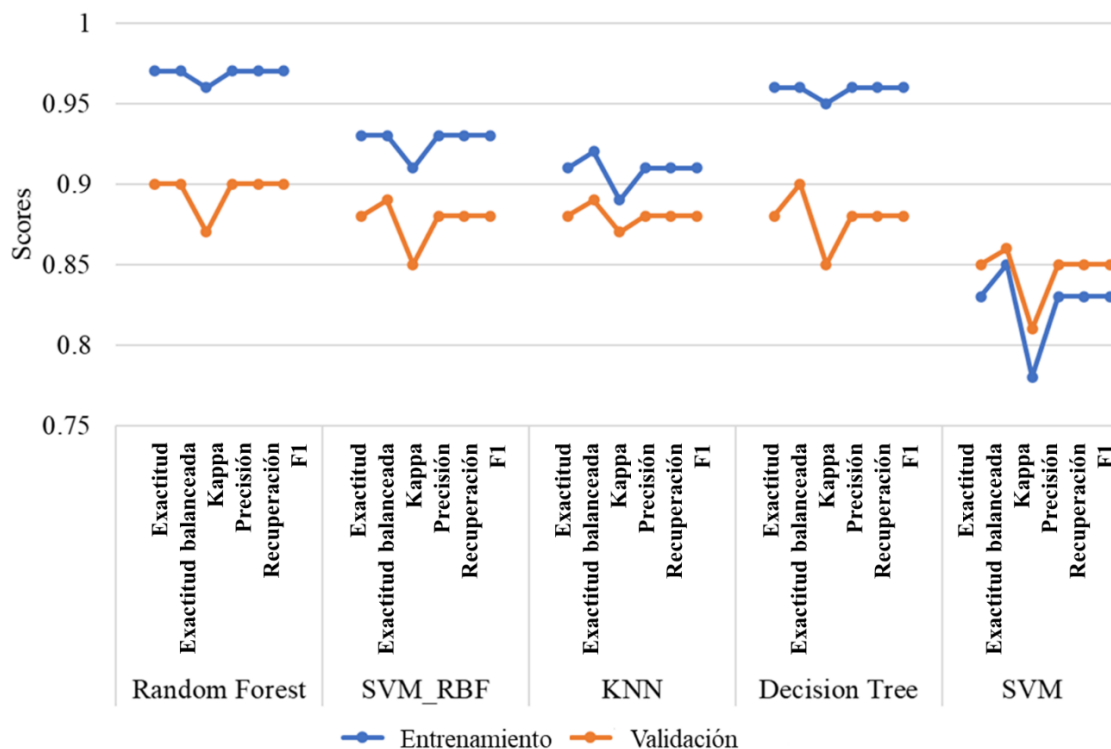


Figura 14. Métricas calculadas para el conjunto de datos de entrenamiento y validación, para los cinco modelos utilizados. La línea azul representa las puntuaciones de los datos de entrenamiento y la línea naranja los datos de validación. Los valores más altos corresponden a Random Forest y los más bajos a SVM.

Segmentación de superficies

Los modelos *Random Forest*, *SVM RBF*, *KNN* y *Decision Tree* se probaron con la imagen del derrame de petróleo del 09 de mayo del 2010, ocurrido en la plataforma *DeepWater Horizon* en el Golfo de México (Figura 15). La imagen RGB muestra una clara distinción entre las áreas con petróleo y las áreas sin petróleo, donde las zonas afectadas por el derrame se visualizaban como tonos grises en contraste con el agua sin petróleo. Además, se observan áreas de tierra y vegetación en el delta del

Misisipi, mientras que la mayor parte de la escena estaba compuesta por agua. El plástico fue el único material cuya presencia no se conocía previamente en la escena.

El resultado de la clasificación del modelo *Random Forest* (Figura 15) asignó la clase plástico (amarillo) y suelo (marrón) a la zona correspondiente a la marea negra, resultado similar en el modelo SVM RBF (Figura 15). El resultado del modelo KNN (Figura 15) clasificó la zona de aceite como aceite (negro) y algunas zonas como suelo (marrón). En el caso del modelo *Decision Tree* (Figura 15), una parte de la mancha de petróleo se clasificó como suelo (marrón).

El área de la imagen que corresponde a agua, los modelos SVM RBF y KNN la clasifican dentro de esta misma clase (azul), el modelo *Random Forest* clasifica una parte como vegetación (verde) y otra como aceite (negro), y el modelo *Decision Tree* asigna una gran parte de la escena a la clase aceite (negro) y otra porción a la clase agua (azul). En los resultados de los cuatro modelos se observa que la zona costera perteneciente al delta del Misisipi se clasifica como vegetación (verde), suelo (marrón) y algunas zonas como aceite (negro), excepto en el modelo SVM RBF, que clasifica la zona mayoritariamente como vegetación (verde).

Considerando que el objetivo principal del modelo es identificar derrames de petróleo se observó que, aunque el modelo *Random Forest* obtuvo las mejores puntuaciones, el modelo que presentó mejores aproximaciones para detectar petróleo en la imagen fue el KNN (Figura 15). Esto se debe a que la mancha clasificada como aceite (negro) corresponde a la zona de contraste positivo, lo que

refleja su efectividad en la detección de derrames. Sin embargo, se identificó una limitación en el modelo KNN, ya que algunas zonas de la mancha que presentan un mayor brillo fueron clasificadas como suelo (marrón), esto debido a que el suelo desnudo puede presentar alto brillo en comparación con otras superficies.

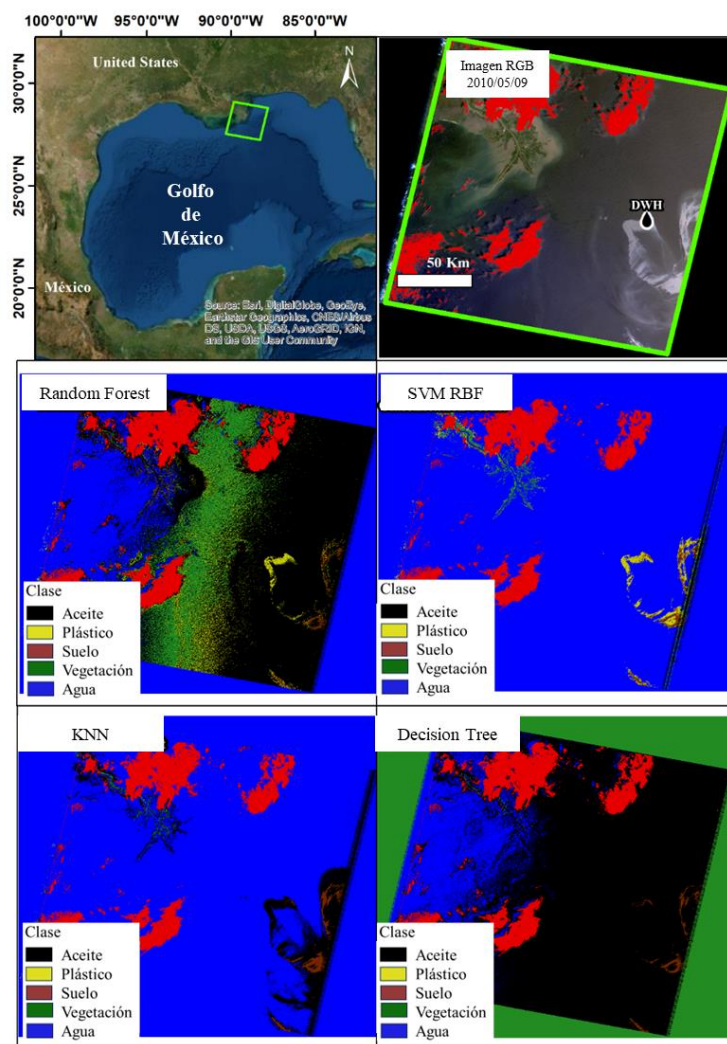


Figura 15. Resultados de la clasificación para la imagen del derrame de DeepWater Horizon del 09 de mayo de 2010. En la primera fila se muestra un mapa de localización de la imagen Landsat 5TM y su composición en RGB, con la ubicación de la plataforma DeepWater Horizon marcada con una gota negra y las letras DWH. En las filas siguientes se presentan los resultados de los modelos de clasificación: Random Forest, SVM RBF, KNN y Decision Tree. Las clases se muestran en diferentes colores: nubes en rojo, agua en azul, petróleo en negro, plástico en amarillo, vegetación en verde y suelo en marrón.

El modelo KNN, que mostró buen rendimiento en la detección de petróleo, fue evaluado con imágenes Landsat 8 OLI libres de petróleo (Figura 16). En la primera imagen, capturada el 12 de enero de 2014 en una zona costera de Luisiana, el modelo no detectó petróleo y clasificó correctamente la parte del Golfo de México como agua (color azul). En la segunda imagen, tomada el 08 de diciembre de 2017 en la Península de Yucatán, se obtuvo un resultado similar, donde el modelo clasificó adecuadamente la parte del mar como agua (color azul).

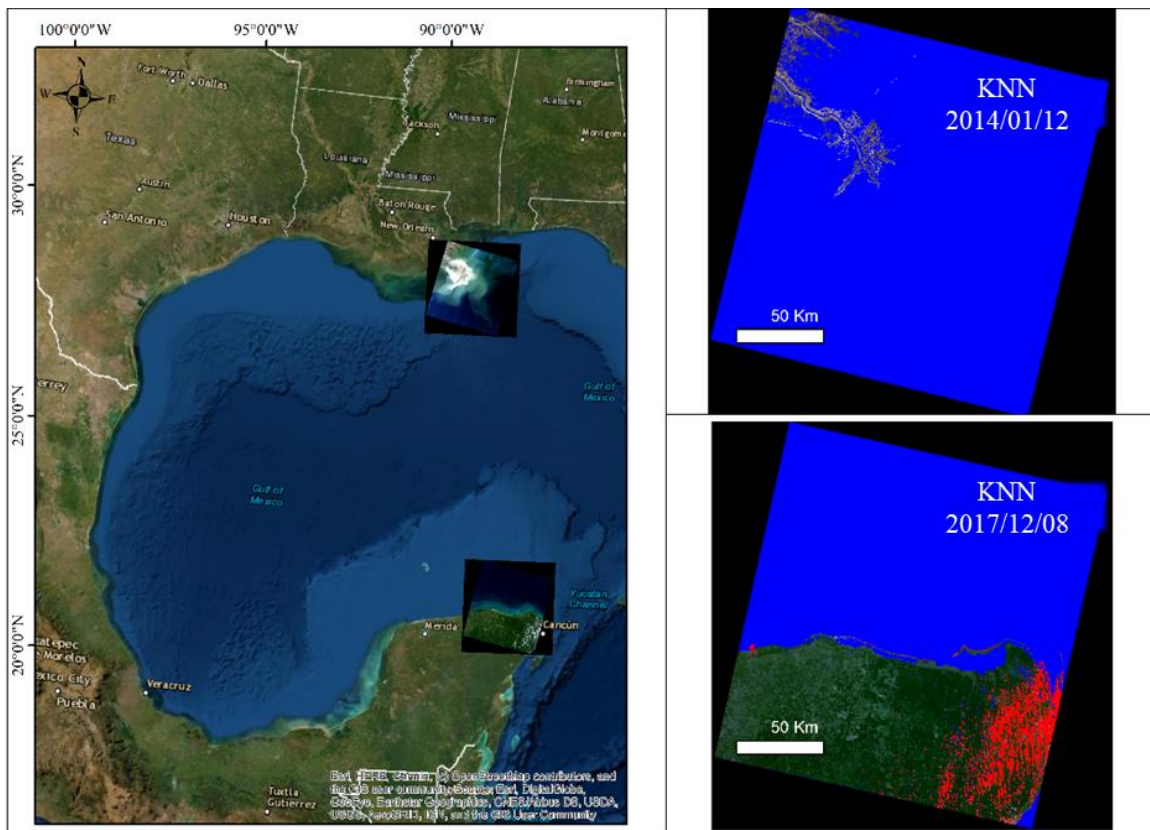


Figura 16. Imágenes de prueba de Landsat 8 OLI. La primera imagen muestra la ubicación de las dos imágenes seleccionadas para las pruebas. La imagen superior derecha (2014/01/12) corresponde a una zona cercana a la costa de Luisiana, similar a las imágenes de la Figura 15. La imagen inferior derecha (2017/12/08) corresponde a una zona cercana a la península de Yucatán, en México. En ambas imágenes se enmascararon las nubes (rojo) y la tierra, por lo que esta última se presenta tal y como se ve en la imagen RGB.

Los resultados de las pruebas realizadas con imágenes que contienen y no contienen petróleo respaldan la precisión del modelo. Estas pruebas ayudan a reducir la posibilidad de falsos positivos, ya que el modelo solo detecta correctamente el petróleo en las imágenes con derrame de 2010, y en las imágenes sin petróleo, el modelo no clasifica erróneamente otras superficies como petróleo.

Modelo de detección de superficies con bandas de importancia

Selección de bandas y precisión

En la selección de las variables importantes mediante el método de disminución de impurezas de Gini, se observó que el orden de las variables (Figura 17) fue SWIR2 (0.22), NIR (0.22), SWIR1 (0.19), Azul (0.17), Verde (0.12) y Rojo (0.06). Por lo que las bandas Azul, NIR, SWIR1 y SWIR2 son las más importantes en la separación de los datos y la diferenciación de las cinco clases.

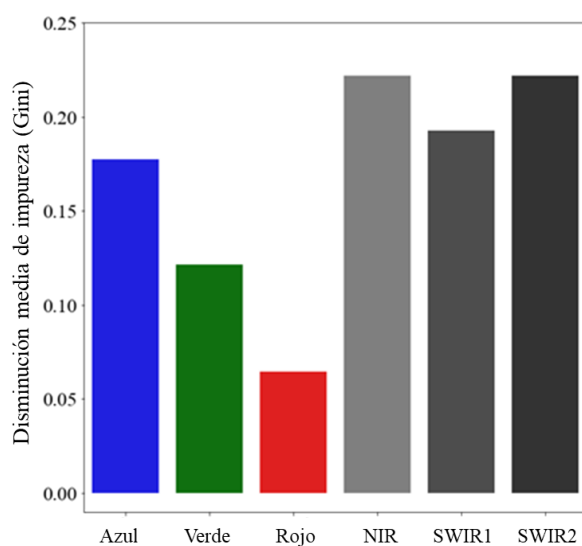


Figura 17. Importancia de las variables según el método de disminución media de impurezas de Gini. Se observa que las bandas con mayor importancia son las bandas infrarrojas NIR, SWIR2 y SWIR1, y en cuarto lugar se encuentra la banda azul.

Estas bandas (Azul, NIR, SWIR1 y SWIR2) se seleccionaron para crear los modelos de clasificación, por lo que se realizó una búsqueda en gradilla para optimizar los hiperparámetros. Los valores óptimos obtenidos para cada modelo se muestran en la Tabla 4.

Tabla 4. Valores e hiperparámetros utilizados para los modelos de clasificación con características importantes en el proceso de entrenamiento.

Modelo	Hiperparámetros y sus valores
<i>K-Nearest Neighbors</i> (KNN)	Leaf_size: 1 N_neighbors: 7 P: 1 Metric: Minkowski
<i>Support Vector Machines</i> (SVM_lineal)	C: 30 Degree: 1 Kernel: linear
<i>Support Vector Machines RBF</i> (SVM_rbf)	C: 30 Degree: 1 Gamma: 1 Kernel: rbf
<i>Decision tree</i> (Dec_tree)	Max_Depth: 7 Min_simples_Split: 14
<i>Random Forest</i>	Max_Depth: 4 Min simples Split: 3 N_estimators: 6

Para evaluar el rendimiento de los modelos creados con las variables importantes, se utilizaron las métricas *Exactitud*, *Exactitud balanceada*, *Kappa*, *Precisión*, *Recuperación* y *F1*. *Random Forest* fue el modelo que obtuvo las puntuaciones más altas para entrenamiento (0.98 en todas las métricas) y validación (0.94, 0.95, 0.94, 0.95, 0.95 y 0.95), seguido de SVM RBF (0.91, 0.91, 0.89, 0.91, 0.91 y 0.91 para el conjunto de datos de entrenamiento; y 0.87, 0.87, 0.83, 0.86, 0.86 y 0.87 para el conjunto de datos de validación) y KNN (0.89, 0.9, 0.86, 0.89, 0.89 y 0.89 para el entrenamiento; y 0.88, 0.89, 0.85, 0.88, 0.88 y 0.88 para el conjunto de datos de

validación). El modelo *Decision Tree*, a pesar de tener valores de entrenamiento altos (0.98 en todas las métricas), las puntuaciones de validación fueron bajas, y la diferencia entre las puntuaciones de ambos conjuntos es grande (entre 0.08 y 0.13) en comparación con los otros modelos, por ejemplo, la diferencia entre las puntuaciones obtenidas en las métricas de los conjuntos de entrenamiento y validación del modelo KNN es de 0.01. Por último, está el modelo SVM (0.77, 0.8, 0.72, 0.77, 0.77 y 0.77 para el entrenamiento; y 0.78, 0.81, 0.73, 0.78, 0.78 y 0.78 para la validación), que tiene un comportamiento similar al modelo con todas las variables (Figura 18).

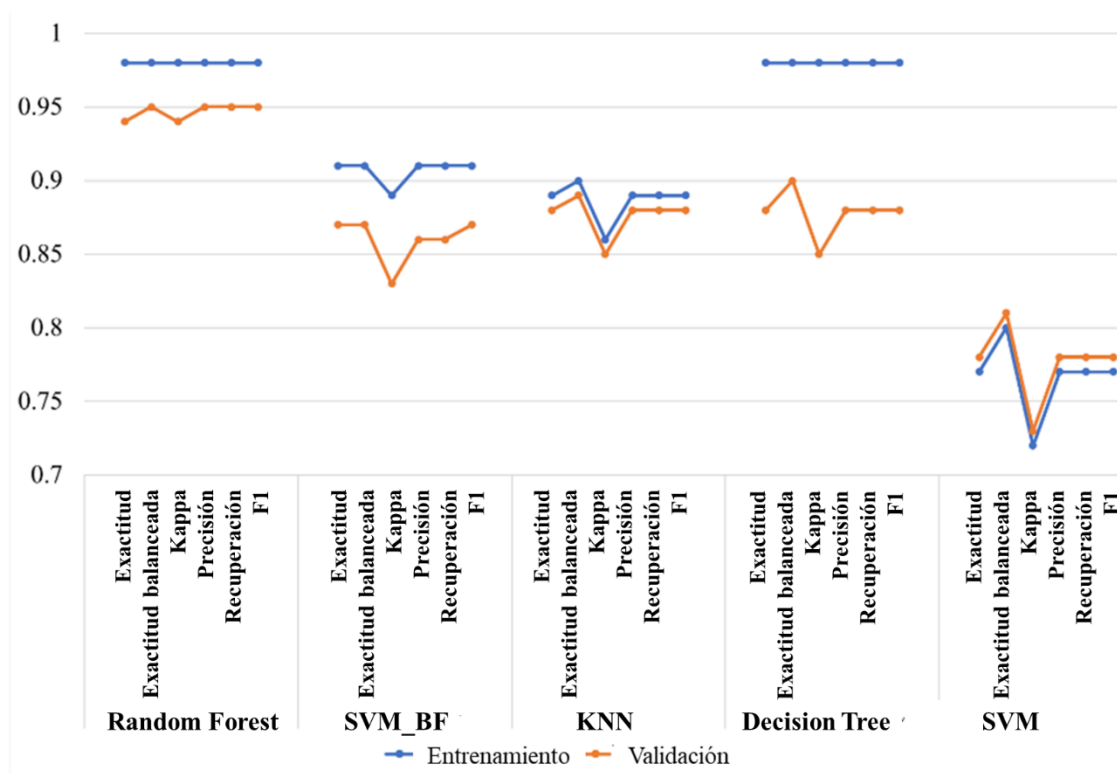


Figura 18. Métricas calculadas para el conjunto de datos de entrenamiento y validación, para los cinco modelos utilizados con las cuatro variables seleccionadas. La línea azul representa las puntuaciones de los datos de entrenamiento, y la línea naranja las de los datos de validación. Los valores más altos corresponden a Random Forest, y los más bajos a SVM.

Segmentación de superficies

Los cuatro modelos con mejor puntuación, Random Forest, SVM RBF, KNN y Decision Tree, se probaron en la escena del 09 de mayo de 2010. En los resultados del modelo Random Forest, se observó que la mancha de aceite se clasificó parcialmente dentro de la clase aceite (negro), otra parte se clasificó como plástico (amarillo) y tierra (marrón). Además, se observó que una gran parte de la zona perteneciente al Golfo de México se clasificó como plástico (amarillo), mientras que la zona costera se clasificó como aceite (negro) y plástico (amarillo). Por otro lado, el modelo SVM RBF clasificó la marea negra como plástico (amarillo), la zona del Golfo de México como agua (azul) y la zona costera como plástico (amarillo), aceite (negro) y suelo (marrón) (Figura 19).

En el caso del modelo KNN, observamos que la marea negra se clasificó como aceite (negro), la parte del Golfo de México se asignó a la clase agua (azul), y la zona costera se clasificó como vegetación (verde). Por último, el modelo Decision Tree clasificó la mancha como plástico (amarillo), gran parte del Golfo de México como vegetación (verde), así como la costa, sólo una parte del Golfo de México se clasificó como agua (azul) (Figura 19).

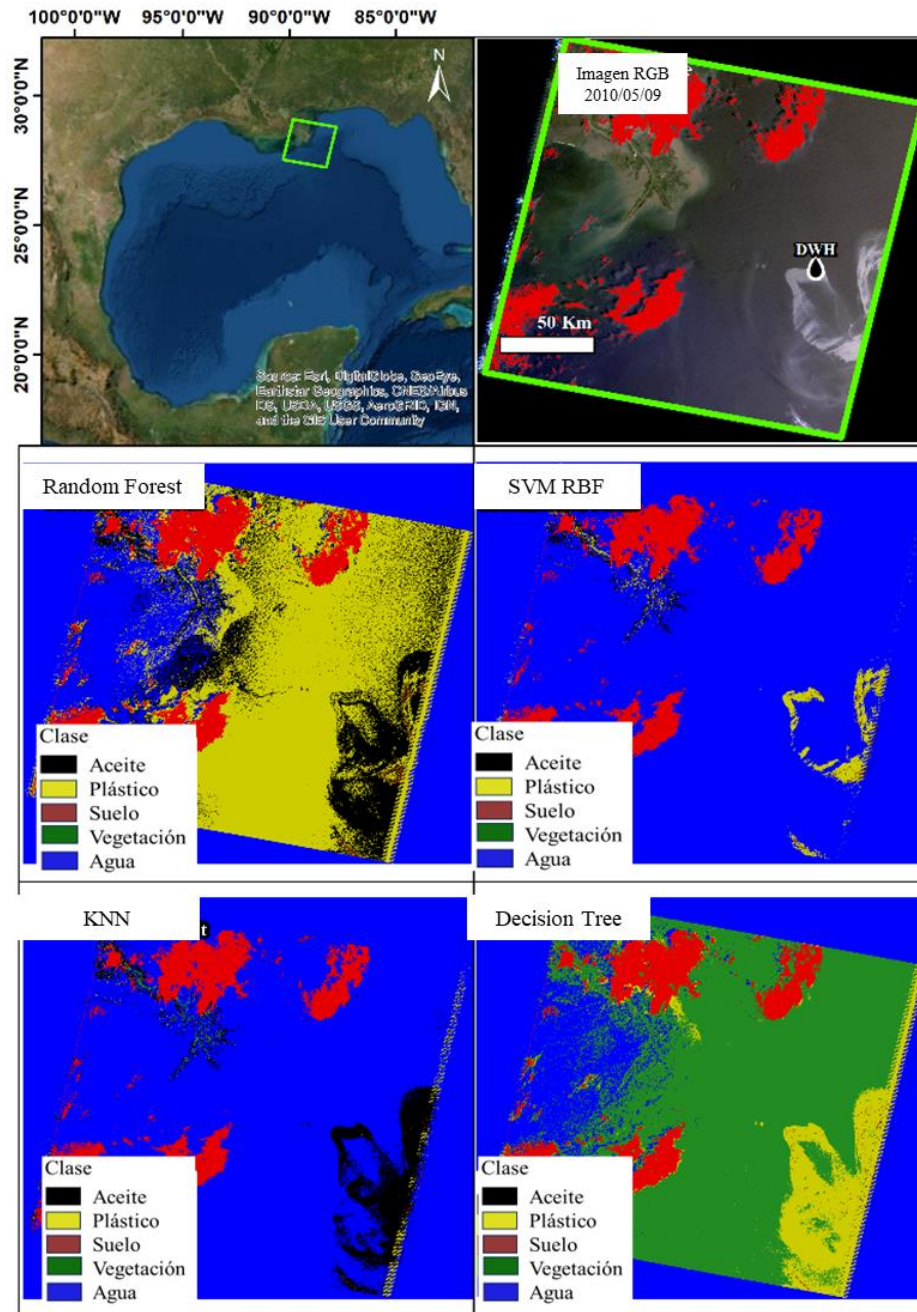


Figura 19. Resultados de los modelos de clasificación entrenados sólo con las partes del espectro correspondientes a las bandas Azul, NIR, SWIR1 y SWIR2. La imagen seleccionada para esta prueba es de Landsat 5TM del 9 de mayo de 2010, capturada durante el derrame de DeepWater Horizon (la plataforma se localiza con una gota negra y las letras DWH). En las filas segunda y tercera se presentan los resultados de los modelos de clasificación Random Forest, SVM RBF, KNN y Decision Tree. Las nubes se representan en rojo, el agua en azul, el aceite en negro, el plástico en amarillo, la vegetación en verde y el suelo en marrón.

En los resultados de estos modelos de clasificación utilizando sólo las variables importantes, se destaca que el modelo KNN logró las mejores aproximaciones en la identificación de la mancha de aceite en la imagen Landsat, a pesar de no haber obtenido las puntuaciones más altas entre los modelos evaluados. Posteriormente, se realizó una prueba utilizando este modelo en las mismas imágenes que en la prueba con las seis bandas. Los resultados fueron similares al modelo entrenado con las seis variables, es decir, no se detectó petróleo en el agua en ninguna imagen y la parte correspondiente al Golfo de México se clasificó únicamente como agua (azul) (Figura 20).

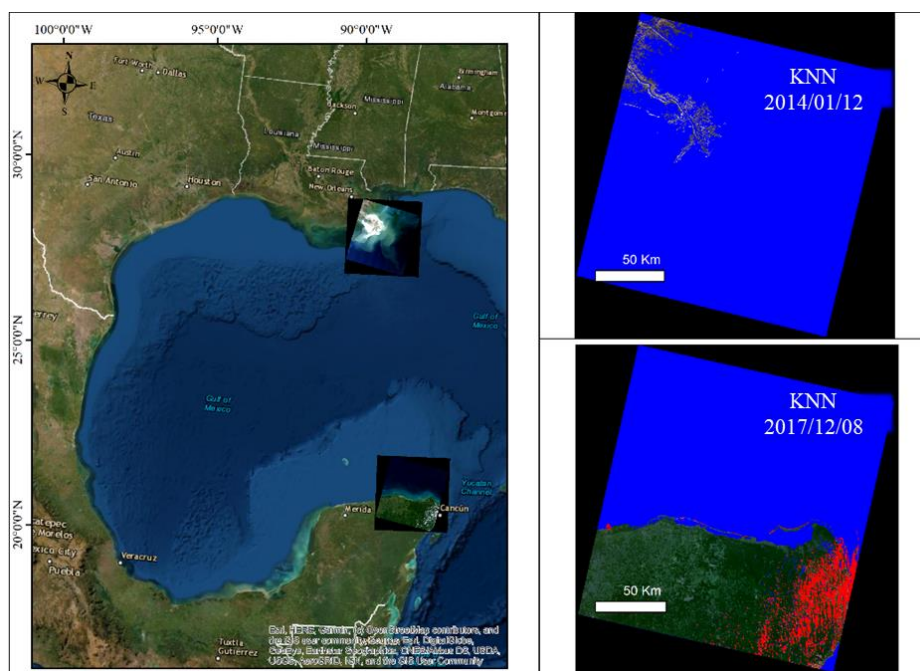


Figura 20. Imágenes de prueba Landsat 8 OLI para el modelo KNN entrenado sólo con las bandas Azul, NIR, SWIR1 y SWIR2. La primera imagen muestra la ubicación de las dos imágenes seleccionadas para la prueba. La imagen superior derecha (2014/01/12) corresponde a una zona cercana a la costa de Luisiana, una ubicación como las imágenes de la Figura 19. La imagen inferior derecha (2017/12/08) corresponde a una zona cercana a la península de Yucatán, en México. En ambas imágenes se enmascararon las nubes (rojo) y la tierra, por lo que esta última se presenta tal y como se ve en la imagen RGB.

Modelo de estimación de concentración

Hiperparámetros y scores

El regresor Random Forest fue entrenado utilizando los siguientes valores óptimos de hiperparámetros: $\text{max_depth} = 10$, $\text{max_features} = 2$, $\text{min_samples_split} = 3$ y $\text{n_estimators} = 50$, que fueron determinados mediante la búsqueda en gradilla. Se calcularon las puntuaciones para evaluar el rendimiento del modelo utilizando los conjuntos de entrenamiento, validación y el conjunto de datos completo. Estas puntuaciones se presentan en la Tabla 5.

De acuerdo con el valor R^2 de 0.90, el modelo genera buenas predicciones, en comparación con el modelo lineal general multivariado que tiene un valor R^2 de 0.67. Además, las métricas MAE, MSE, RMSE y ME del modelo fueron inferiores a las del modelo lineal, lo que indica un mejor rendimiento (Tabla 5).

Tabla 5. Métricas y puntuaciones obtenidas en la evaluación del rendimiento del modelo de regresión Random Forest y del modelo lineal estadístico. Las primeras columnas corresponden a las métricas de Random Forest, y la última a un modelo lineal.

	Entrenamiento Random Forest	Validación Random Forest	Todos los datos Random Forest	Todos los datos Modelo lineal
MAE	1.414	3.346	2.329	11.34
MSE	5.422	21.239	15.18	217.44
RMSE	2.329	4.609	3.896	14.74
R^2	0.99	0.97	0.977	0.67
ME	6.26	13.36	13.36	43.25

Estimación de concentración

El modelo *Random forest* se utilizó para la estimación de concentración debido a las puntuaciones obtenidas en comparación con el modelo lineal (Tabla 5). Se realizaron dos pruebas, la primera utilizando una imagen sintética creada con todos los datos. Se puede observar que los resultados de la predicción del modelo corresponden con la forma en que se ordenaron los datos de la imagen creada, de menor a mayor concentración (valor de concentración en porcentaje dentro de cada píxel) (Figura 21). En la fila de concentración de 1%, el modelo estimó una concentración de 6.3% en todos los casos, para la fila de 23% la estimación se situó entre 19% y 35%, pero la mayoría de los valores se situaron en el intervalo de 22% y 23%. En la fila de 40%, la estimación se situó entre 35% y 40%, para la fila de 60% la estimación se situó en el intervalo de 60% y 63%, con un dato de 52%. En la fila de 75%, la estimación fue idéntica a la de los datos originales, y en la última fila, de 92%, la estimación fue de 88%.

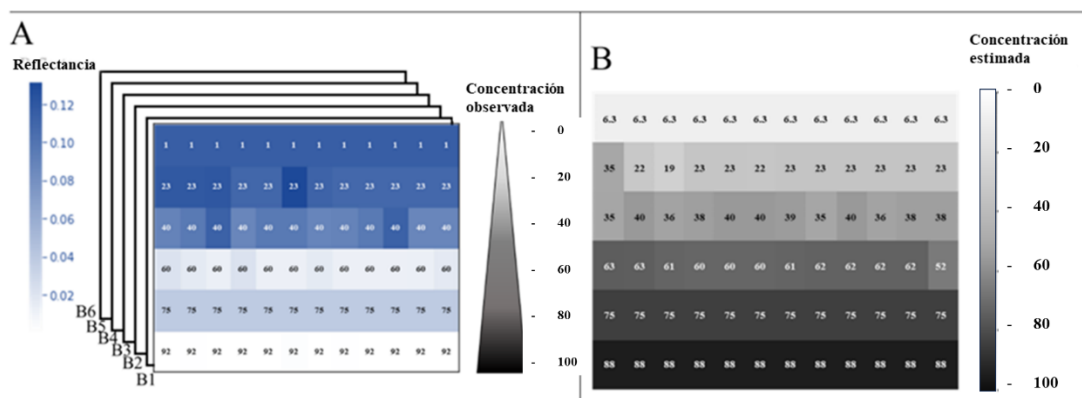


Figura 21. Prueba del modelo de estimación de la concentración utilizando una imagen sintética creada con los datos originales. A) Se muestra la banda azul de la imagen sintética, dentro de cada píxel está el valor de concentración original. Los datos se organizaron en bandas como una imagen de satélite. B) Resultado de la predicción del modelo, dentro de cada píxel está el valor de concentración estimado (%).

En la segunda prueba, se realizó la estimación de concentración de petróleo en las áreas clasificadas como petróleo por el modelo KNN para variables importantes en las imágenes del 9 de mayo y 12 de julio de 2010 (Figura 22). En las áreas identificadas como petróleo, la concentración estimada osciló entre 40% y 60% (Figura 22).

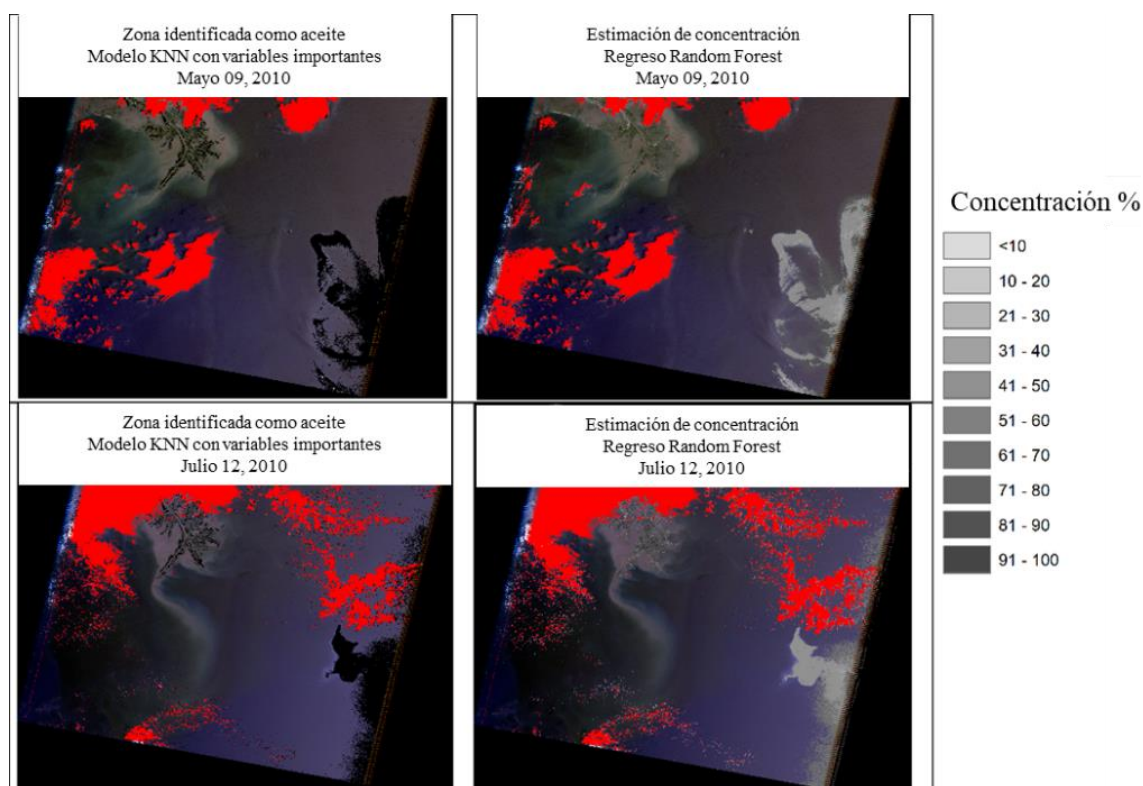


Figura 22. Resultados de la estimación de la concentración utilizando el regresor Random Forest. En la primera columna está la imagen RGB con la zonificación de petróleo identificada por el modelo de clasificación KNN para las variables importantes. La segunda columna es la imagen RGB con la concentración estimada para la zona identificada como petróleo. En ambos casos, las nubes están enmascaradas e indicadas en rojo.

Al realizar dos transectos en la zona donde se estimó la concentración de petróleo en la imagen del 9 de mayo, se observó una variación en la concentración (Figura 23A). En el transecto 1, se observó una concentración de 0% al inicio, ya que el modelo de clasificación no identificó petróleo en esa zona (Figura 23B).

Posteriormente, en la misma línea de perfil, el modelo detectó una concentración entre 40% y 45% (Figura 23B). Luego, la concentración volvió a ser 0% debido a que el modelo no detectó petróleo en esa parte de la imagen (Figura 23B). En el segundo transecto, la concentración estimada osciló entre 35% y 45% (Figura 23C). Aproximadamente en el píxel 700, la concentración aumentó y alcanzó valores cercanos 70% (Figura 23C). Finalmente, en el píxel 900, la concentración volvió a ser 0% porque no se detectó petróleo en ese punto específico (Figura 23C).

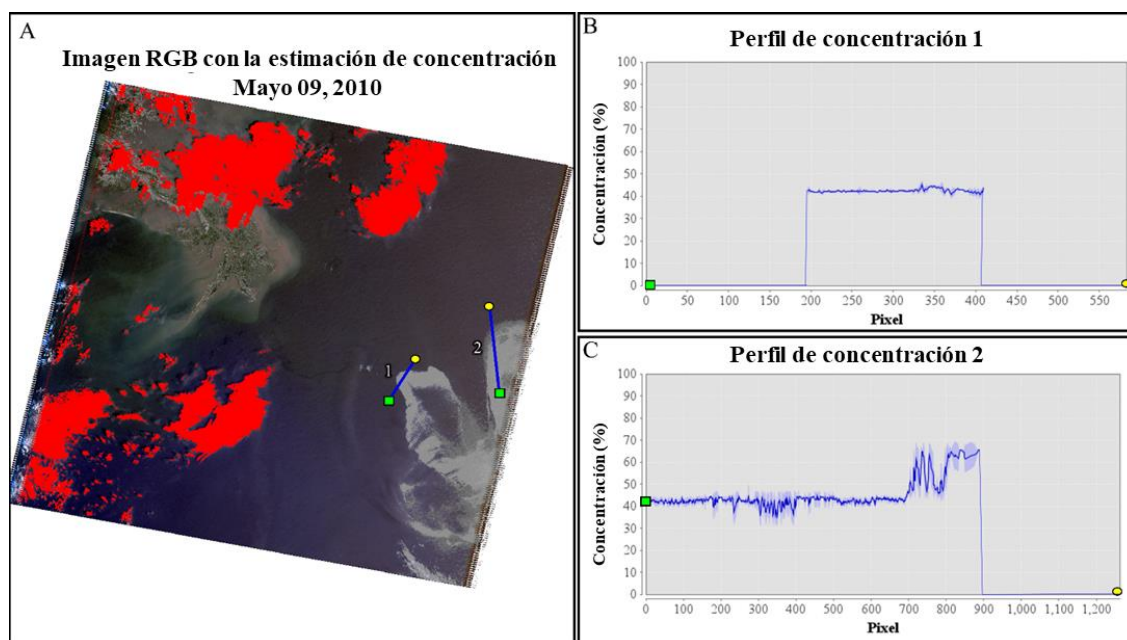


Figura 23. A) Transectos trazados para observar los perfiles de concentración de la imagen de mayo 09 del 2010. Las líneas azules representan los perfiles, los cuadrados verdes el inicio del perfil y las elipses amarillas el final. B) Perfil de estimación de la concentración para el transecto 1. C) Perfil de estimación de la concentración para el transecto 2.

Las concentraciones estimadas de petróleo se mantuvieron estables en el tiempo, oscilando entre 30% y 70%. No se encontraron concentraciones fuera de este rango en ninguna de las imágenes analizadas.

Sensores activos

CNN para clasificación de derrames de petróleo

Los resultados de *Exactitud* y *Pérdidas* de los modelos CNN para la clasificación de imágenes con y sin derrames de petróleo se observan en la Figura 24.

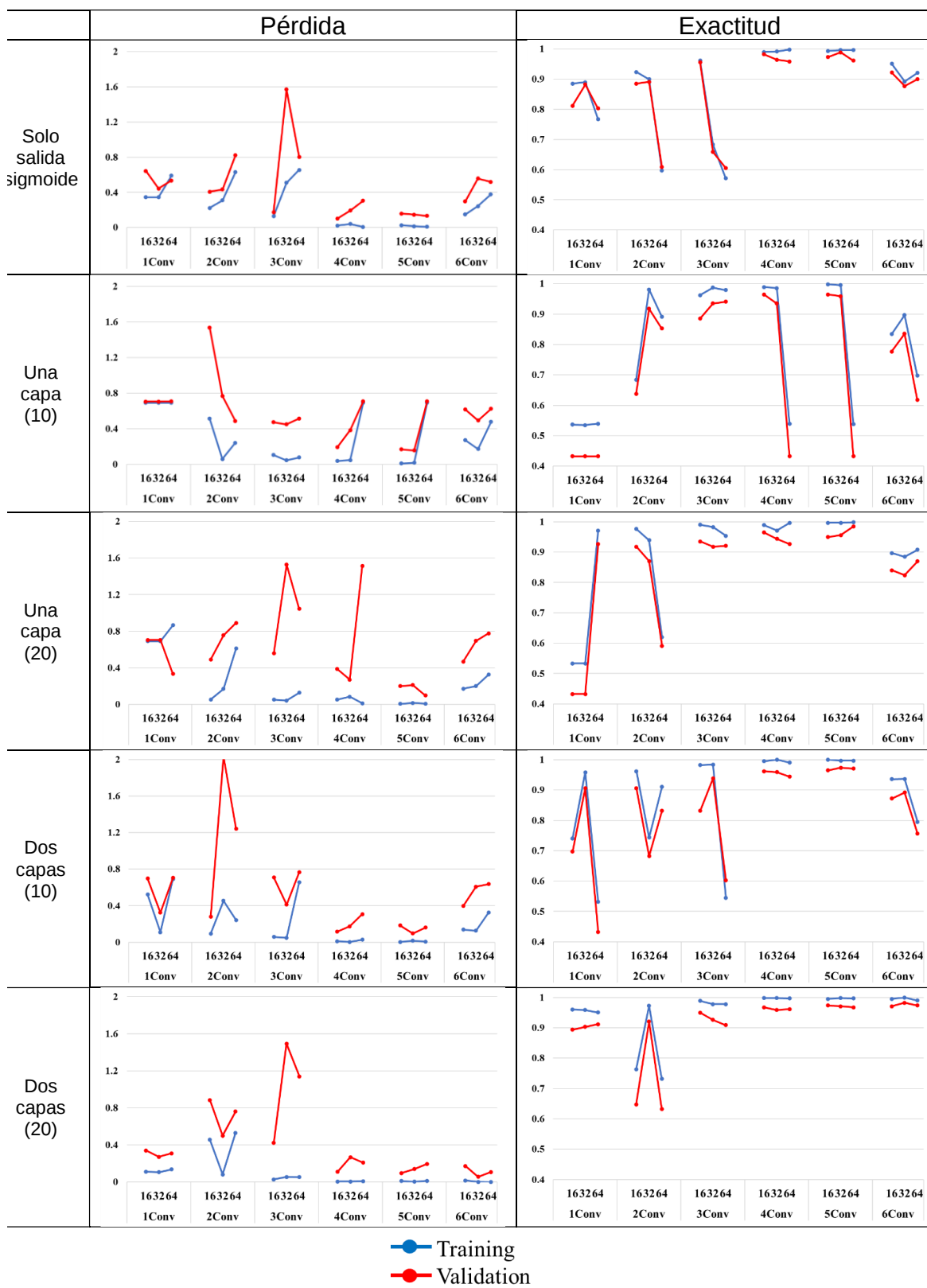


Figura 24. Resultados de los modelos CNN. La primera columna especifica el número de capas ocultas y neuronas en cada capa. En la segunda columna se encuentran los valores de Pérdida, y en la tercera los de

Exactitud. En rojo los datos de validación, y en azul los de entrenamiento. El modelo con los mejores scores (Pérdida y Exactitud) fue el de dos capas ocultas con 20 neuronas y seis capas convolucionales (6Conv) con 32 filtros, exactitud de entrenamiento 0.9999, exactitud de validación 0.9824, pérdida de entrenamiento 0.0018, pérdida de validación 0.0573.

Los modelos con los mejores scores fueron, el modelo sin capas ocultas con cinco capas convolucionales (5Conv) con 32 filtros, exactitud en entrenamiento 0.9968, exactitud en validación 0.9888, pérdida en entrenamiento 0.015, pérdida en validación 0.15; el modelo con una capa oculta con 10 neuronas y cinco convoluciones (5Conv) con 16 filtros, exactitud en entrenamiento 0.9984, exactitud en validación 0.9647, pérdida en entrenamiento 0.0088, pérdida en validación 0.1679. El modelo con una capa oculta con 20 neuronas y cinco convoluciones (5Conv) con 64 filtros exactitud en entrenamiento 0.9984, exactitud en entrenamiento 0.9853, pérdida en entrenamiento 0.0085, pérdida en validación 0.0979. El modelo con dos capas ocultas con 10 neuronas y cinco convoluciones (5Conv) con 32 filtros, exactitud en entrenamiento 0.9969, exactitud en validación 0.9735, pérdida en entrenamiento 0.0171, pérdida en validación 0.0964. Y el modelo con dos capas ocultas con 20 neuronas, y seis capas convolucionales (6Conv) con 32 filtros exactitud en entrenamiento 0.9999, exactitud en validación 0.9824, pérdida de entrenamiento 0.0018, pérdida en validación 0.0573 (Tabla 6).

Tabla 6. Exactitud y Pérdida de los cinco modelos con los scores más altos. Se observa que el modelo de dos capas ocultas con 20 neuronas y 6 convoluciones con 32 filtros fue el que obtuvo los mejores valores (en negritas)

Modelo	Exactitud		Pérdida	
	Entrenamiento	Validación	Entrenamiento	Validación
<i>Only output 5Conv 32 filters</i>	0.9968	0.9888	0.015	0.15
<i>One layer (10) 5Conv 16 filters</i>	0.9984	0.9647	0.0088	0.1679
<i>One layer (20) 5Conv 64 filters</i>	0.9984	0.9853	0.0085	0.0979
<i>Two layers (10) 5Conv 32 filters</i>	0.9969	0.9735	0.0171	0.0964
<i>Two layers (20) 6Conv 32 filters</i>	0.9999	0.9824	0.018	0.0573

Realizar las comparaciones de los modelos con distintas configuraciones, permitió conocer una vista general de su comportamiento, y con ello conocer el modelo con el mejor desempeño, el cual fue el modelo más complejo, con seis capas convolucionales (6Conv) con 32 filtros, y dos capas ocultas con 20 neuronas cada capa, y 600 épocas (Figura 25), el cual obtuvo la exactitud en entrenamiento más alto (0.9999), y pérdida de validación más bajo (0.0573), en comparación con los demás modelos (Tabla 6).

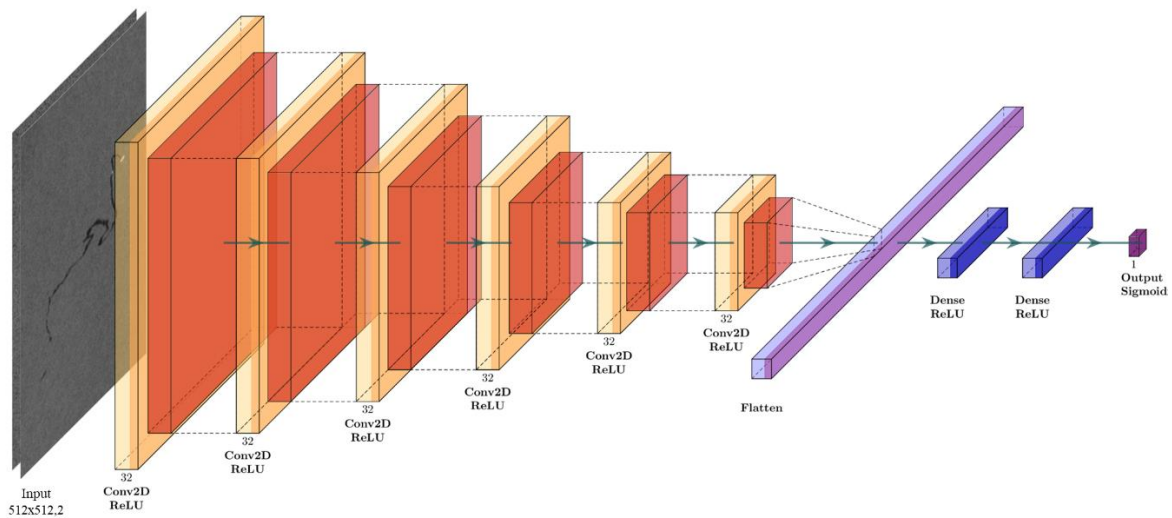


Figura 25. Arquitectura del modelo CNN con los mejores scores. La entrada al modelo son imágenes Sentinel-1 de tamaño 512x512 píxeles con dos polarizaciones (VV, VH), por lo que la dimensión de entrada es de 512x512,2. El modelo cuenta con seis capas convolucionales (Conv2D) de tamaño de kernel de 3x3 con 32 filtros, seguida de MaxPoling2D de 2x2, la función de activación utilizada para todas las capas convolucionales fue ReLU. Para la red densa, después del flatten, se utilizaron dos capas ocultas con 20 neuronas cada una y con activación ReLU. La capa de salida solo contaba con una neurona con activación sigmoide.

Prueba del modelo CNN para clasificación de derrames de petróleo

Para las pruebas se utilizó el modelo CNN con los mejores scores (Figura 25), se utilizaron 150 imágenes de la clase Aceite, 150 de No Aceite y 150 *look-alikes*. Las imágenes de No Aceite y *look-alikes* se utilizaron en conjunto como clase No Aceite. La Exactitud fue de 0.95 y Pérdida de 0.3487, en la Figura 26A se puede observar la matriz de confusión, la clase No aceite tuvo una exactitud de 0.96, y la clase Aceite de 0.98. En la curva ROC (Figura 26B) se obtuvo un AUC de 0.99, lo que indica una buena relación entre la sensibilidad y la especificidad del modelo.

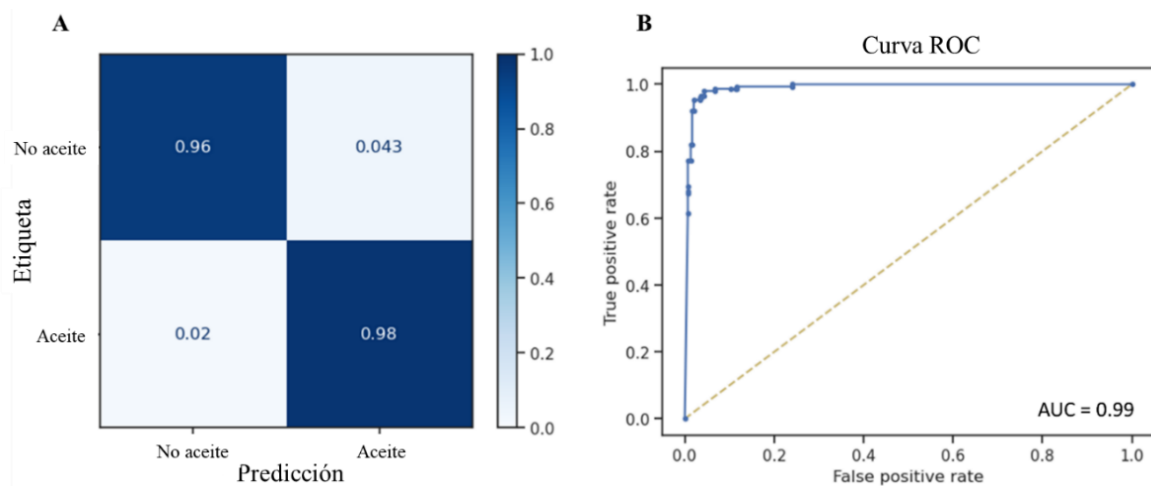


Figura 26. A) Matriz de confusión de las pruebas realizadas, se observa que la exactitud del modelo para la clase aceite fue de 0.98, y para la clase No Aceite de 0.96. B) Curva ROC de las pruebas realizadas, el valor de Área Bajo la curva (AUC) fue de 0.99.

Los resultados de algunas de las pruebas del modelo CNN se muestran en la Figura 27. Se observan imágenes de Aceite (Etiqueta verdadera = 1.00), y de No aceite (Etiqueta verdadera = 0.00), en esta última se encuentran también imágenes de *look-alikes* (Etiqueta verdadera = 0.00). Para la predicción de las imágenes de aceite los valores de probabilidad (*Predicción*) obtenidos se encuentran entre 0.90 y 1.00. Para las imágenes de No aceite, los valores de probabilidad estuvieron entre $2.2e-28$ y 0.0049.

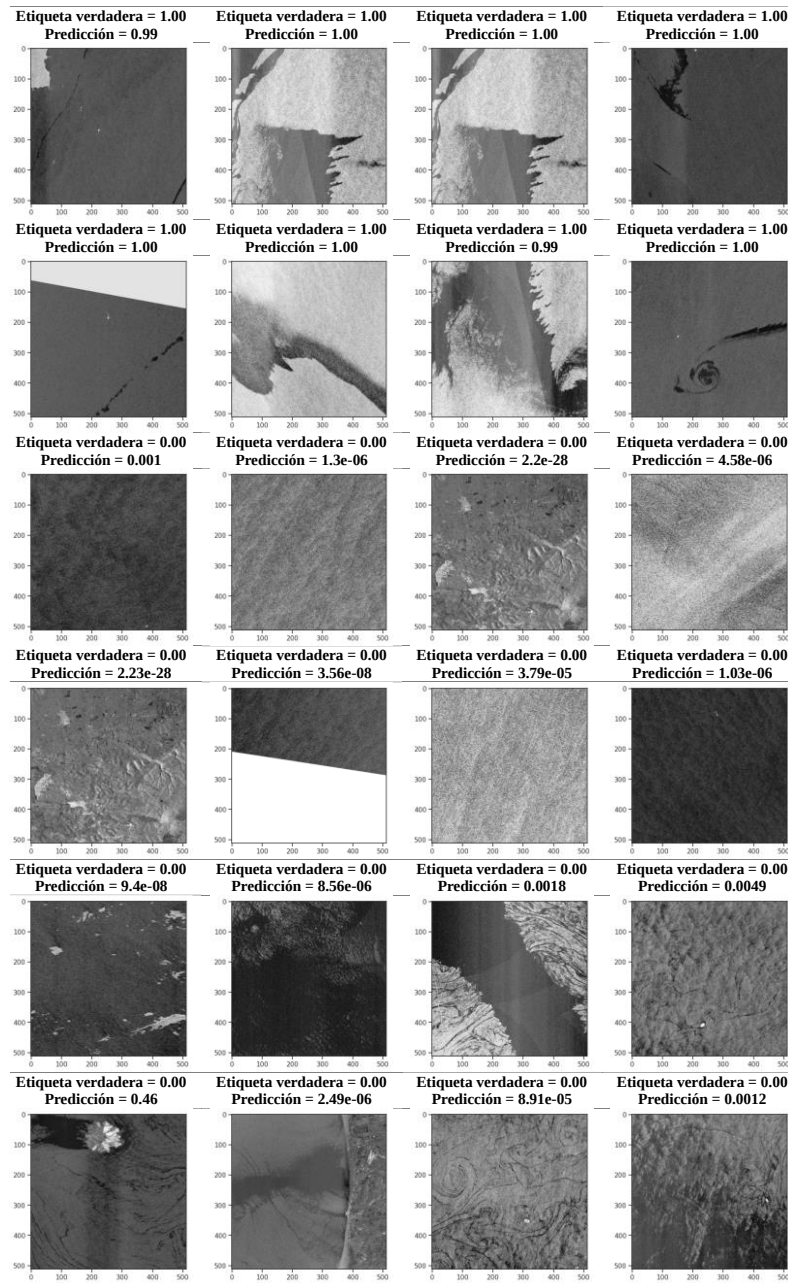


Figura 27. Resultado de las pruebas realizadas con el modelo CNN. Se observa cada una de las imágenes Sentinel-1 utilizadas, las cuales tienen dimensión 512x512, ya que cuentan con ambas polarizaciones (VV, VH), se muestra únicamente la polarización VV ya que es en la que a simple vista se observa la presencia de derrames, así como los look-alikes. En la parte superior de cada imagen se encuentra su Etiqueta, la cual tiene valores de 0.00 para No Aceite, y 1.00 para Aceite, además se encuentra el valor de predicción en probabilidad. Observamos que las imágenes correspondientes a Aceite tienen valor de probabilidad superior a 0.9, y las de No aceite de entre 2.2e-28 y 0.0049. Cabe señalar que en dentro de la clase No aceite se encuentran imágenes de look-alikes.

Estos resultados demuestran que el modelo CNN tiene buen desempeño en la clasificación de imágenes con presencia de derrames de petróleo, además a las imágenes sin aceite y *look-alikes* les asigna valores de probabilidad bajos, cercanos a cero, clasificándolas de manera correcta.

Modelos de segmentación de derrames de petróleo

Perceptrón multicapa

En la Figura 28 se observan los resultados de entrenamiento y validación para los modelos de MLP. Los modelos con los mejores valores, en función del número de capas ocultas y neuronas fueron, el modelo logístico (*Exactitud en entrenamiento y validación* de 0.9331 y 0.933, y *Pérdida en entrenamiento y validación* de 0.1959 y 0.1959), el modelo de una capa oculta con 500 neuronas (*Exactitud en entrenamiento y validación* de 0.9454y 0.944, y *Pérdida en entrenamiento y validación* de 0.1451 y 0.149), el modelo con dos capas ocultas y 20 neuronas en cada capa (*Exactitud en entrenamiento y validación* de 0.9513 y 0.9512, y *Pérdida en entrenamiento y validación* de 0.1331 y 0.1332), y el modelo con tres capas ocultas con 20 neuronas en cada capa (*Exactitud en entrenamiento y validación* de 0.9514 y 0.9513, y *Pérdida de entrenamiento y validación* de 0.1329 y 0.1329).

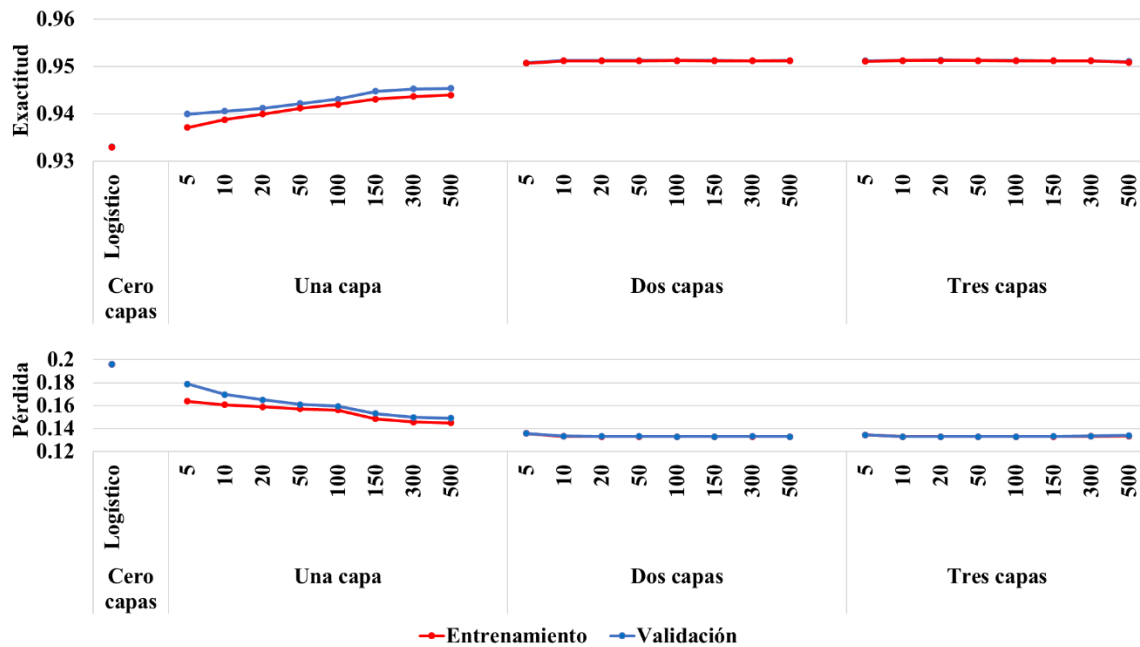


Figura 28. Exactitud and Pérdida para entrenamiento y validación de los modelos de MLP. Se observa que se comienza con un modelo de regresión logística, el cual tiene exactitud en entrenamiento de 0.9331 y exactitud en validación de 0.933. Posteriormente los modelos con una capa oculta, donde el modelo con 500 neuronas obtuvo una exactitud en entrenamiento de 0.9454 y exactitud en validación de 0.944. Después se tienen los modelos con una capa oculta con diferente cantidad de neuronas, donde el modelo con el exactitud más alta fue el de 20 neuronas con exactitud en entrenamiento de 0.9513 y exactitud en validación de 0.9512. Y por último los modelos con tres capas ocultas, donde el modelo con 20 neuronas en cada capa obtuvo la exactitud en entrenamiento y validación más alta, de 0.9514 y 0.9513 respectivamente. De la misma manera, dichos modelos obtuvieron los valores más bajos en Pérdidas.

La comparación de los modelos reveló que todos tuvieron un buen desempeño, con exactitud superior a 0.90. Sin embargo, el modelo de tres capas ocultas con 20 neuronas cada una, obtuvo las puntuaciones más altas. Es importante destacar que se observó una saturación en los valores de las métricas, lo que significa que no se observó un aumento en la precisión a medida que se incrementó la complejidad de los modelos. Esto sugiere que el modelo más complejo alcanzó un rendimiento óptimo y no se justificaría aumentar aún más su complejidad.

Prueba del modelo MLP para segmentación de derrames de petróleo

Los resultados de algunas de las pruebas realizadas con los cinco mejores modelos se muestran en la Figura 29. Debido a que la salida de los modelos MLP dan valores de probabilidad entre 0 y 1, se utilizó un umbral donde los pixeles que tuvieran valores mayores o iguales a 0.99 eran considerados como aceite, y los pixeles por debajo de este umbral se consideraron No aceite.

Para la primera imagen (primera fila de la Figura 29), derrame de aceite, el modelo logístico obtuvo una exactitud de 0.8588 y un IoU de 0.86, y el modelo de tres capas ocultas una exactitud de 0.7268 and IoU de 0.89 (Figura 29). Los modelos de una y dos capas ocultas obtuvieron una exactitud de 0.7493 y 0.7243 respectivamente, y un IoU de 0.50 cada uno. Para la segunda imagen, derrame de aceite, el modelo logístico obtuvo una exactitud de 0.6566 and IoU de 0.83, y el modelo de tres capas 0.6106 y 0.80 respectivamente (Figura 29).

Para la imagen sin presencia de aceite (tercera fila de la Figura 29) se observó que ningún modelo segmentó alguna parte como aceite, por lo tanto, los scores tuvieron valores de 1.0 (Figura 29). Y para la última imagen (cuarta fila de la Figura 29), look-alike, se observó que los modelos de una y dos capas no segmentaron nada como aceite (exactitud e IoU de 1.0), y los modelos de regresión logística y de tres capas ocultas segmentaron la zona como aceite (exactitud 0.40 y 0.39, IoU de 0.16 y 0.15) (Figura 29). Para esta última imagen, el comportamiento y los scores altos de los modelos de una y dos capas ocultas se atribuyó a que, en general, no segmentaban

nada, independientemente de la presencia o ausencia de aceite, por lo que dichos modelos fueron descartados.

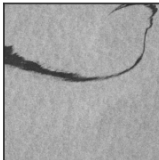
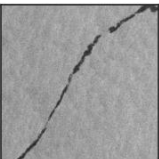
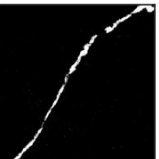
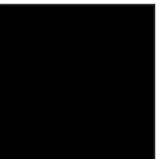
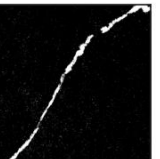
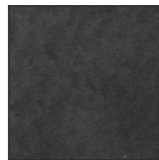
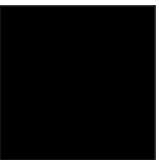
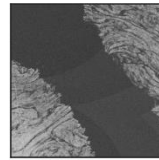
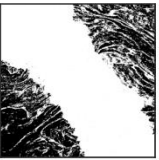
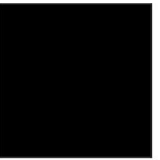
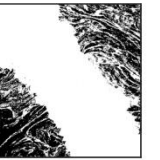
Imagen	Máscara (Ground truth)	Predicción logístico	Predicción Una capa	Predicción Dos capas	Predicción Tres capas
					
		Exac = 0.8588 IoU = 0.86	Exac = 0.7493 IoU = 0.50	Exac = 0.7243 IoU = 0.50	Exac = 0.7268 IoU = 0.89
					
		Exac = 0.6566 IoU = 0.83	Exac = 0.6027 IoU = 0.48	Exac = 0.6027 IoU = 0.48	Exac = 0.6106 IoU = 0.80
					
		Exac = 1.0 IoU = 1.0	Exac = 1.0 IoU = 1.0	Exac = 1.0 IoU = 1.0	Exac = 1.0 IoU = 1.0
					
		Exac = 0.40 IoU = 0.16	Exac = 1.0 IoU = 1.0	Exac = 1.0 IoU = 1.0	Exac = 0.39 IoU = 0.15

Figura 29. Resultados de algunas de las pruebas realizadas con los diferentes modelos. La primera columna corresponde a la imagen Sentinel-1, la segunda columna es el ground truth, posteriormente se encuentran las predicciones realizadas por los distintos modelos, cada una de las predicciones cuenta con su exactitud e IoU. Debido a que la salida del modelo da como resultado valores de probabilidad, se utilizó un valor mayor o igual a 0.99 para realizar la asignación de clase (binarización). Se observa que los modelos con el mejor desempeño son el logístico y el de tres capas ocultas, sin embargo, en la imagen con look-alike el modelo segmenta erróneamente, dando como resultados los valores más bajos en las métricas.

En general, se observa que los modelos de regresión logística y con tres capas ocultas tuvieron las mejores aproximaciones en la segmentación de aceite, y los modelos con una y dos capas ocultas no lograron segmentar ningún píxel como aceite.

Las comparaciones revelaron que, a pesar de obtener altos valores en las métricas, algunos modelos no lograron un buen desempeño en la segmentación del aceite, especialmente los de una y dos capas. El modelo logístico, siendo el más sencillo, tuvo buenas aproximaciones, al igual que el más complejo con tres capas ocultas. Además, fue notorio que la *exactitud* no reflejaba completamente el desempeño de los modelos, ya que, en ocasiones, aunque los valores fueran bajos (0.6 - 0.7), los scores de IoU fueron altos (> 0.8), indicando que IoU era una mejor métrica para evaluar la segmentación. También se destacó que en presencia de *imágenes look-alike*, los modelos cometían errores debido a la similitud de los valores de los píxeles con los del aceite, lo que resalta la necesidad de considerar el contexto y las formas para la segmentación de aceite.

Resultados de los experimentos del modelo U-net

Los resultados de los experimentos realizados para buscar la mejor combinación de parámetros del modelo U-net para la tarea de segmentación de derrames de petróleo se observan en la Tabla 7. La métrica para evaluar y monitorear los modelos fue Mean IoU, debido a que en los modelos de MLP se observó que la exactitud no representaba el comportamiento del modelo.

Tabla 7. Diferentes modelos U-net entrenados y validados. Se implementaron distintas combinaciones de parámetros y arquitectura para encontrar el modelo con las mejores aproximaciones. Se implemento desde el modelo U-net original, posteriormente se le cambio la función de pérdida por la de Pérdida Focal, la siguiente modificación fue la función Conv2DTranspose por Upsampling2D, y las últimas modificaciones fue el número de filtros y la profundidad (convoluciones). Se observa que el modelo con el IoU promedio más bajo fue la U-net original (0.50 y 0.40 para entrenamiento y validación), y el que obtuvo las mejores puntuaciones fue el modelo con Pérdida Focal, Upsampling2D y una mayor cantidad de filtros y profundidad (última fila), IoU promedio de 0.96 para entrenamiento y 0.90 para validación.

Modelo	Entrenamiento IoU promedio	Validación IoU promedio
U-net original	0.50	0.40
U-net Modificaciones: - Pérdida Focal	0.86	0.83
U-net Modificaciones: - Pérdida Focal - Upsampling2D	0.91	0.90
U-net Modificaciones: - Pérdida Focal - Upsampling2D - Filtros 16→32→64→128→256	0.92	0.90
U-net Modificaciones: - Pérdida Focal - Upsampling2D - Filtros 16→32→64→128→256→512→1024	0.96	0.90

En primera instancia se observa que el modelo U-net original, únicamente con la modificación de las entradas, obtuvo los scores más bajos de IoU promedio (0.50 y 0.40 para entrenamiento y validación). Posteriormente al modelo se le cambio la función de pérdida por Pérdida Focal, lo cual dio resultados favorables para el aprendizaje del modelo, dando un IoU promedio de 0.86 y 0.83 para entrenamiento y validación. En tercer lugar, además de agregar Pérdida Focal, se cambió la función

de Conv2DTranspose por Upsampling2D, esto por eficiencia computacional (reduce los tiempos de entrenamiento), lo cual también resultó favorable para la tarea de segmentación, dando IoU promedio de 0.91 para entrenamiento y 0.90 para validación. En tercer lugar, se modificaron el número de filtros en las capas convolucionales, utilizando 16→32→64→128→256, lo cual dio un IoU promedio de 0.92 y 0.90 para entrenamiento y validación respectivamente. Y, por último, se modificó el número de capas convolucionales, y el número de filtros, generando un modelo más profundo, utilizando el siguiente orden en los filtros 16→32→64→128→256→512→1024, lo cual dio valores de IoU promedio de 0.96 para entrenamiento y 0.90 para validación.

El modelo con la mayor cantidad de capas convolucionales, y número de filtros fue el que obtuvo el score más alto para entrenamiento (IoU promedio = 0.96), lo que indica que un modelo complejo y profundo, al extraer un mayor número de características, y cada vez más complejas, puede generar las mejores aproximaciones en la tarea de segmentación de derrames de petróleo. Además, con el resultado de estas experimentaciones, se logró conocer que la función de *Pérdida Focal*, al tener un parámetro de focalización, y el IoU, son una buena alternativa para la implementación de modelos de segmentación de derrames de petróleo, ya que, en la mayoría de las imágenes donde se presenta un derrame, lo que domina es el *fondo (background)* por lo que utilizar *Exactitud* y *Entropía Cruzada*, como métricas, podría ocasionar que el modelo no logre mejorar, ya que la métrica se satura por la clase dominante y genera scores altos a pesar de que el modelo no genere buenos resultados.

Modelo U-net para segmentación de derrames de petróleo

Después de obtener la combinación de parámetros que mostraba el mejor rendimiento en la arquitectura del modelo U-net para la tarea de segmentación, se realizaron pruebas con diferentes variantes de la U-net original, utilizando los mismos parámetros mencionados anteriormente. Los resultados de las métricas utilizadas para supervisar el proceso de entrenamiento y validación de los modelos se presentan en la Figura 30.

Los resultados muestran que todos los modelos lograron una *Exactitud* de 0.99 tanto en el entrenamiento como en la validación (Figura 30). Entre los modelos evaluados, el modelo Attention U-net (Att_Unet2D) destacó en el entrenamiento, con una Precisión de 0.96, Recuperación de 0.95, F1 de 0.95 e IoU de 0.96. En cuanto a la validación, los modelos Att_Unet2D, Unet2D Kaggle y Unet2D obtuvieron valores similares, con Precisión entre 0.9 y 0.91, Recuperación entre 0.89 y 0.90, F1 entre 0.88 y 0.89, y un IoU de 0.9. Por otro lado, el modelo U-net++ (Unet2D Plus) presentó los valores más bajos en la mayoría de las métricas, excepto en la *Exactitud*, donde alcanzó resultados similares a los demás modelos (Figura 30).

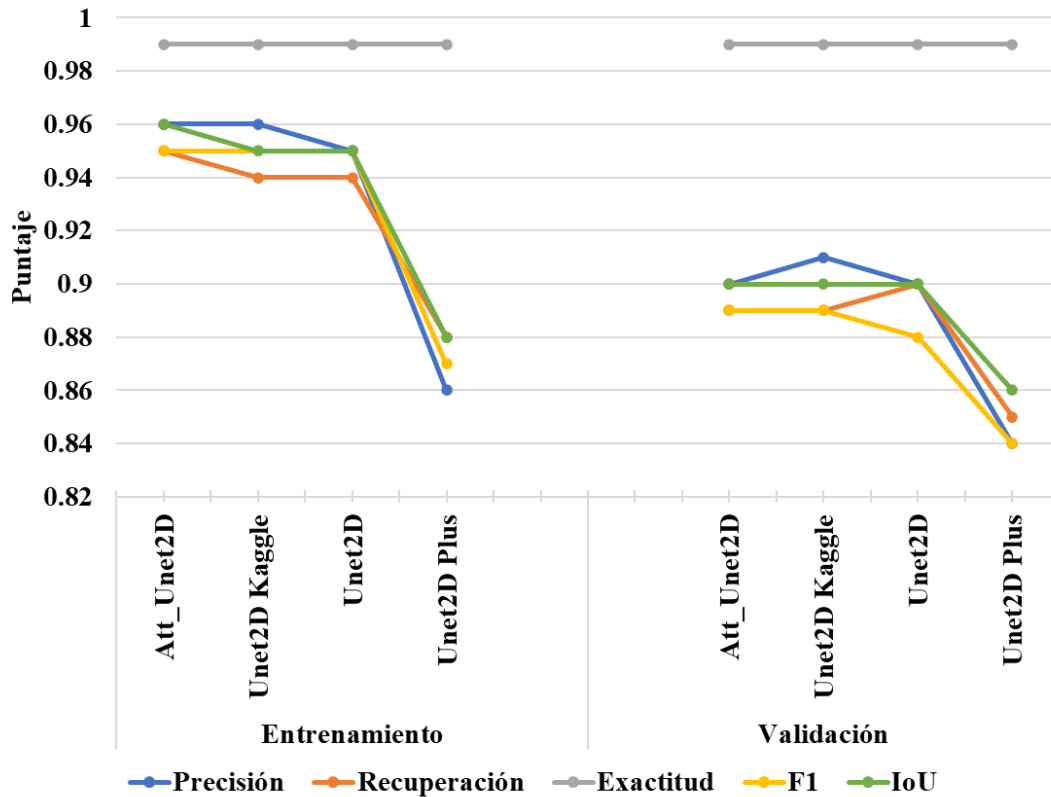


Figura 30. Puntajes de entrenamiento y validación de las métricas de las distintas arquitecturas basadas en el modelo U-net. Los modelos que se utilizaron fueron el modelo Attention U-net (Att_Unet2D), la arquitectura U-net para segmentación de imágenes de un concurso de kaggle (Unet2D Kaggle), el modelo U-net con las modificaciones realizadas (Unet2D), y el modelo U-net++ (Unet2D Plus). Donde el modelo con los scores más altos fue Att_Unet2D, y el modelo con las puntuaciones más bajas fue Unet2D Plus.

A partir de estas comparaciones, se puede observar que, excepto por el modelo Unet2D Plus, los demás modelos mostraron resultados similares en las métricas. Sin embargo, el modelo Att_Unet2D obtuvo las mejores métricas, seguido del modelo Unet2D Kaggle. Es importante tener en cuenta que, en ocasiones, altos puntajes en las métricas de entrenamiento o validación no siempre se traducen en un buen desempeño en pruebas reales. Es necesario evaluar el rendimiento de los modelos en situaciones de prueba más realistas para obtener una evaluación completa y precisa de su desempeño.

Prueba de los modelos U-net

Las pruebas de los modelos U-net se realizaron utilizando imágenes con aceite, sin aceite y con look-alikes. Los resultados de las pruebas con imágenes que contenían aceite se muestran en la Figura 31. El modelo Unet2D Plus obtuvo los scores más bajos, mientras que los demás modelos alcanzaron scores de IoU entre 0.90 y 0.99, a excepción del último ejemplo (sexta fila), que obtuvo valores de entre 0.83 y 0.86.

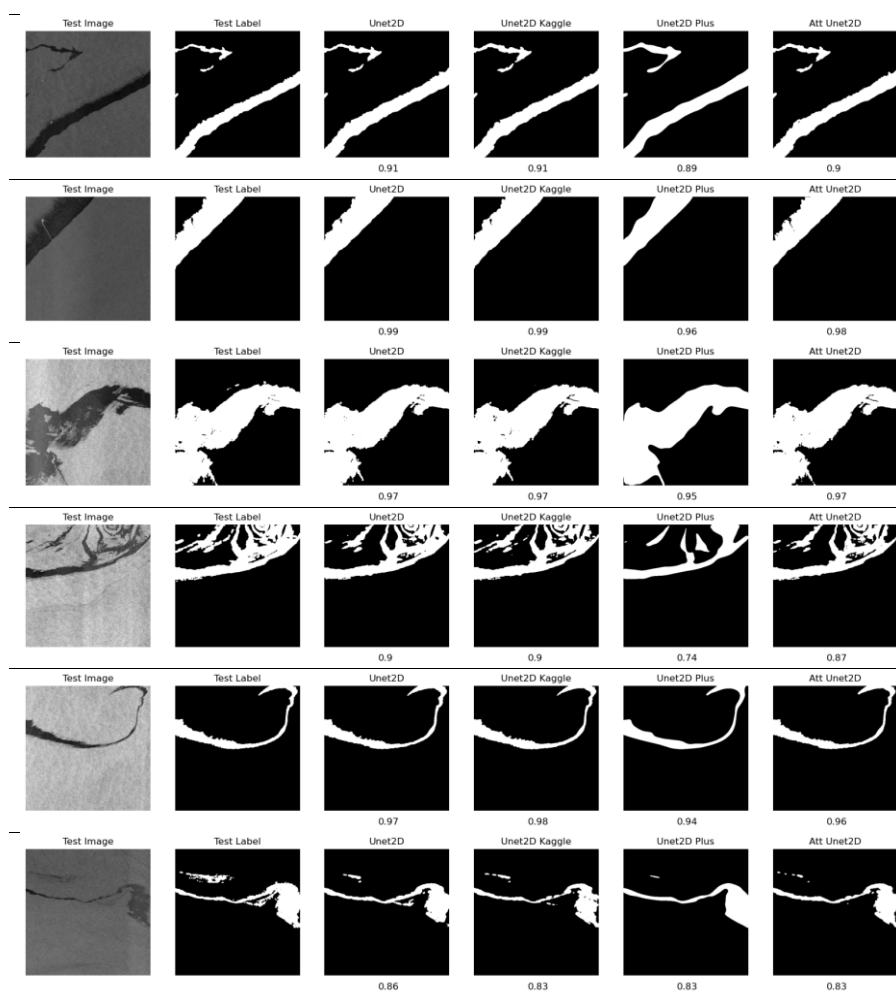


Figura 31. Resultado de las pruebas de los modelos Unet con imágenes con presencia de aceite. En la primera columna se encuentran las imágenes Sentinel-1, la segunda columna son las máscaras o ground truth, las columnas siguientes corresponden a los resultados para cada modelo. En la parte inferior de cada imagen se encuentra su valor de IoU.

Los resultados de las pruebas con imágenes sin aceite, y con look-alikes se presentan en la Figura 32. Se observa que todos los modelos obtuvieron un valor de $IoU = 1.0$, lo que indica una coincidencia perfecta con la máscara. En este caso, ninguno de los modelos detectó la presencia de aceite dentro de las escenas.

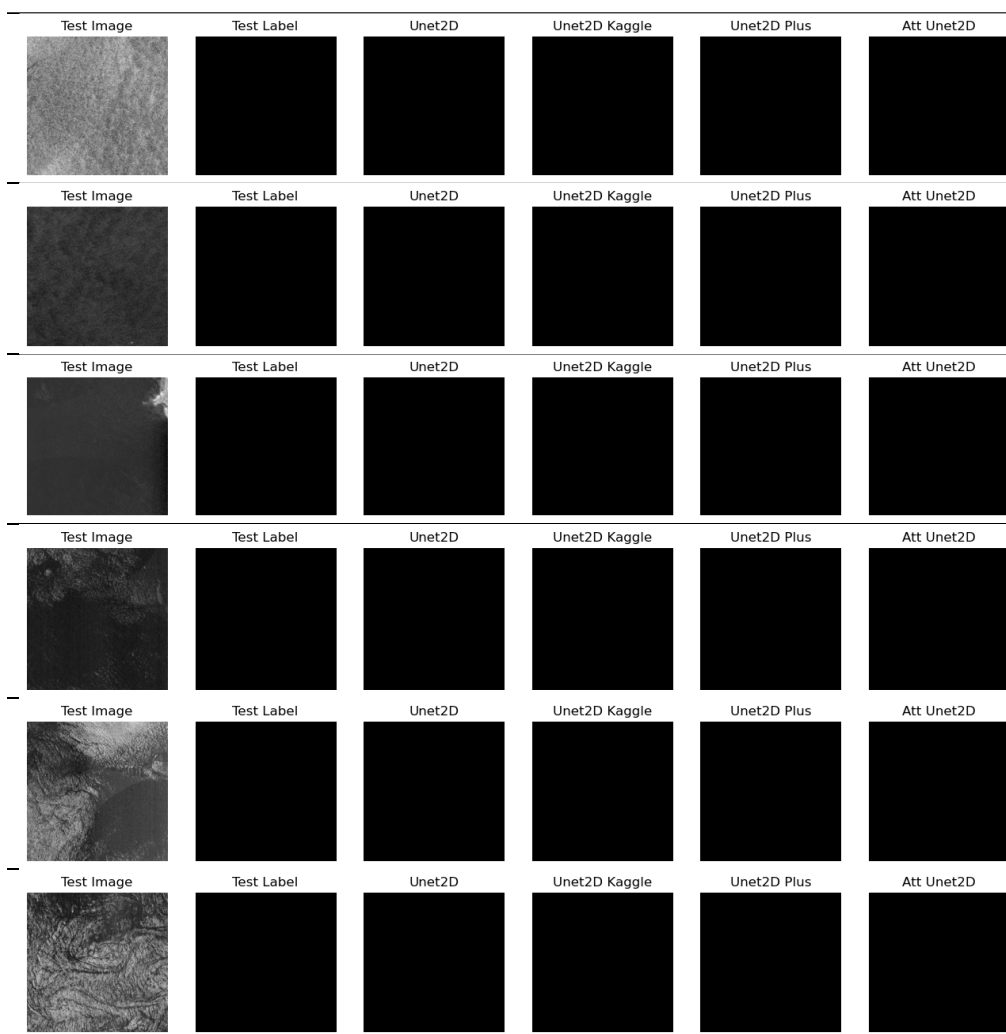


Figura 32. Resultados de las pruebas realizadas para imágenes sin aceite (primeras tres filas) y con look-alikes (últimas tres filas), ambos catalogados como no aceite (negativos). Se observa que los resultados de todos los modelos obtuvieron puntajes de $IoU = 1.0$, ya que en ninguno se segmentó la superficie como aceite.

El comportamiento promedio de todas las pruebas realizadas para los modelos, con los distintos conjuntos de datos se observa en la Figura 33. Se puede apreciar que para todos los casos el mejor score lo obtuvo el modelo Unet2D, con valores de IoU promedio de 0.90, 0.90, y 0.85 para cada conjunto de datos, y el modelo Unet2D Plus, presentó los valores de IoU promedio de 0.86 para los dos primeros conjuntos, y 0.87 con el tercer conjunto, siendo el modelo con los scores más bajos.

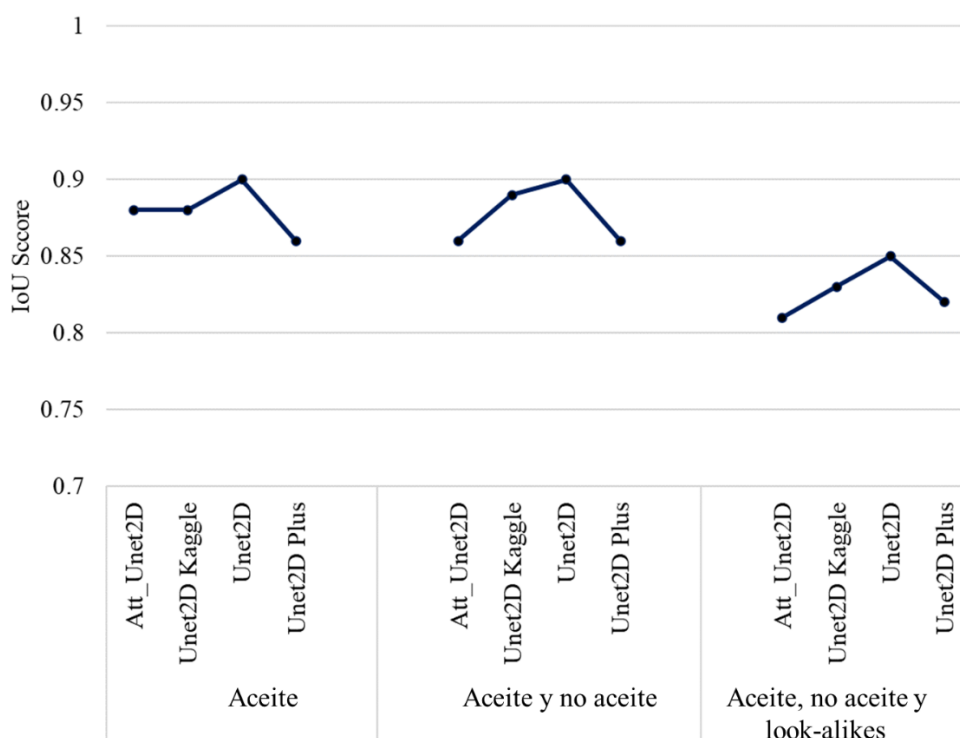


Figura 33. Resultado de las pruebas de los modelos con distintos sets de datos. Se utilizaron imágenes únicamente de aceite, un conjunto de imágenes con y sin aceite, y otro conjunto con imágenes con aceite, sin aceite y con presencia de look-alikes. Se observa que para los primeros dos conjuntos el modelo Unet2D obtuvo el valor más alto de IoU (0.90), y para el set donde se encuentran los look-alikes obtuvo un valor de 0.85.

Estimación de área

Para realizar la estimación de área de derrame se utilizó el resultado obtenido por el modelo de segmentación Unet2D, el cual generó las mejores aproximaciones para la segmentación de derrames de petróleo. Para esto se utilizaron seis imágenes de prueba, las cuales fueron clasificadas con presencia de aceite por el modelo CNN.

Para realizar la estimación de área, se consideró el área cubierta por un píxel de la imagen, la cual, por la resolución que se manejó para Sentinel-1, fue de 40mx40m, lo cual representa 1,600m². Teniendo esto, se realizó el conteo de los pixeles segmentados como áreas cubiertas por aceite, y luego se multiplicaron por el área correspondiente a cada píxel (Conteo*1600), lo que dio como resultado el área de la clase en metros cuadrados (m²). Finalmente se dividió este valor por 10,000 para obtener el área en hectáreas. Esta aproximación se logró realizar debido a que las imágenes cuentan con proyección geográfica, lo que posibilita la medición de áreas y distancias (Mathias Davidson, Michael Askne, 2012). Esto se puede resumir en la siguiente formula $((40*40)*(Conteo))/10,000$.

Las estimaciones realizadas se pueden observar en la Figura 34, donde los valores de estimación se encuentran entre 864 y 13,679.84 hectáreas (ha), lo cual corresponde al área de aceite dentro de cada imagen.

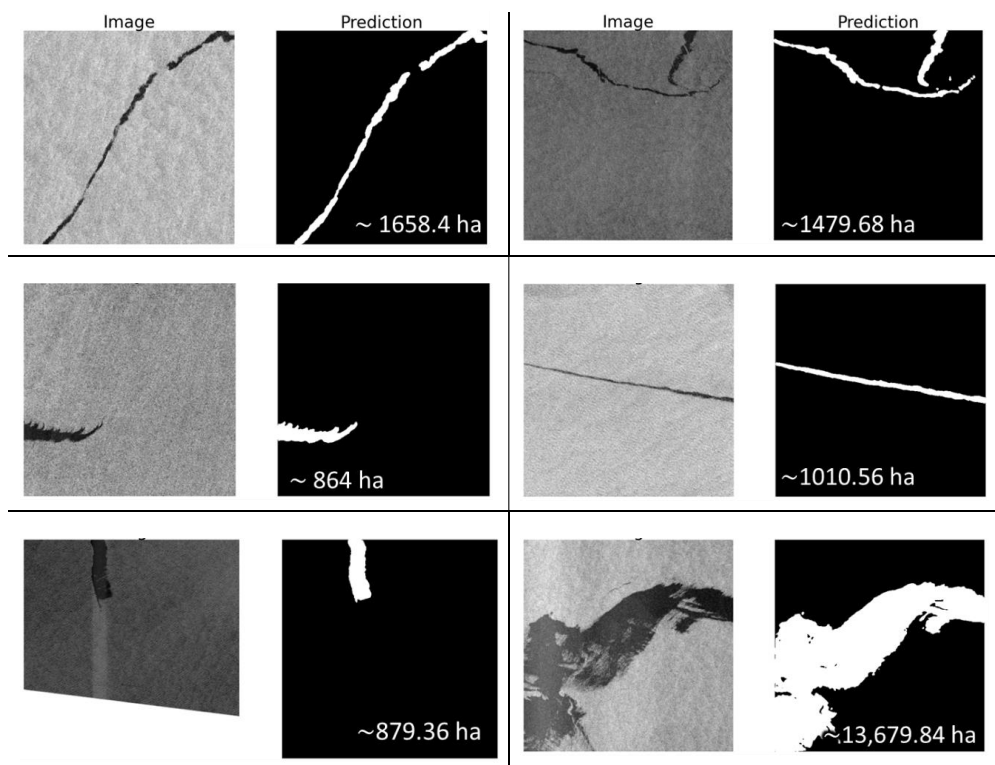


Figura 34. Resultados de la estimación de aérea posterior a la segmentación por el modelo Unet2D.

Discusión

Sensores pasivos

Diferenciación de materiales

Un análisis general de los datos espectrales de los cinco materiales demostró diferencias estadísticas en la mayoría de ellos, excepto en aceite y suelo ($p = 0,06$), y suelo y vegetación ($p = 0,85$) (Figura 9A). Sin embargo, este comportamiento por sí solo no demuestra completamente las diferencias entre los materiales porque las distribuciones, dispersión, rango y medias de cada clase son similares. Por este motivo, es necesario considerar el comportamiento espectral de los datos de cada clase en cada banda.

En el análisis por banda (Figura 9B), se puede observar un comportamiento distintivo para cada clase. La vegetación muestra un pico de reflectancia en la banda del infrarrojo cercano (NIR), mientras que la absorción del agua se evidencia en las bandas infrarrojas. Por otro lado, el suelo exhibe una reflectancia relativamente baja en todas las bandas, con aumentos en la parte infrarroja. En el caso del plástico, se observa una distribución amplia en todas las bandas. Esto se debe a que los datos no corresponden exclusivamente a un tipo de plástico, sino que incluyen diferentes variantes como plástico PET, claro, oscuro y con color. Cada uno de estos tipos de plástico presenta diferencias espectrales individuales, lo que resulta en una distribución amplia en todas las bandas cuando se consideran en conjunto.

Asimismo, se puede observar la respuesta espectral del aceite en cada banda (Figura 9B). En la parte visible (RGB) muestra una mayor reflectancia en comparación con el agua, pero no presenta tendencias específicas de absorción o reflectancia. Esto se debe a que las delgadas capas de aceite presentan una enorme variación en su respuesta espectral (Fingas and Brown, 2017). Sin embargo, se menciona que la reflectancia del aceite tiende a aumentar con la disminución de la longitud de onda, es decir, es mayor en la región azul-verde (Fingas and Brown, 2017). En este estudio, no se demostró completamente esta tendencia, posiblemente debido a que los datos analizados incluyen diferentes concentraciones y emulsiones que afectan su respuesta espectral. Aun así, es posible identificar la capacidad de absorción de este compuesto en la banda azul (Leifer *et al.*, 2012).

Estudios previos (Winkelmann, 2005; Leifer *et al.*, 2012; Lu *et al.*, 2019) han señalado que las bandas visibles no son las más adecuadas para diferenciar el aceite de otros materiales. En cambio, las bandas infrarrojas (NIR y SWIR) son más apropiadas debido a que las características espectrales del aceite cambian en la región NIR (700-1100 nm). Esto se debe al espesor del aceite, y a la relación agua-aceite ya que el contenido de agua en la mancha puede generar absorción y modificar los valores de reflectancia (Clark *et al.*, 2010). Además, los enlaces carbono-hidrógeno (C-H, C-H₂, y C-H₃), grupos hidroxilo (OH), dobles y triples enlaces de compuestos alifáticos y aromáticos, grupos carboxilo (C=O), éster (C-O-C) y grupos amina (N-H) tienen un comportamiento específico en esta región (720-

1730 nm, y 1750-1760 nm, y a 1190-1210 nm), lo cual es crucial para diferenciar el aceite de otros materiales. (Winkelmann, 2005; Leifer *et al.*, 2012; Lu *et al.*, 2019).

Así, podemos observar que el aceite, debido a sus propiedades físicas y químicas, muestra propiedades espectrales específicas en la parte visible e infrarroja, con picos de absorción y reflectancia, que en ocasiones puede generar coloraciones oscuras, grises, con brillo, opacas, e incluso con iridiscencia por la variación en la incidencia de la luz (Bonn Agreement Accord de Bonn, 2007; Lu *et al.*, 2019). Es por ello por lo que se deben incorporar todas las bandas para un análisis más completo, y lograr una mayor separabilidad.

Al seguir analizando la distribución en cada una de las bandas (Figura 9B), se observa que algunos datos presentan dispersión similar, ya que los cuartiles en ocasiones se sobreponen, lo que indica la presencia de observaciones similares a otros materiales. Esto es evidente al analizar de manera individual las observaciones de cada clase, ya que existe una correlación entre algunas observaciones de diferentes clases (Figura 10), y que estas se agrupan (Figura 12A), aun utilizando un método no lineal (Figura 12B). Este comportamiento podría deberse a que cada uno de los datos proviene de muestras distintas, ya que no se estableció un criterio específico en la selección de los datos de cada clase, simplemente fueron seleccionados por pertenecer a ellas. En el caso del aceite, se observa que en algunos casos presenta valores similares al agua, lo cual puede ser resultado del contenido variable de este material en las muestras, ya que los datos pertenecen a observaciones con diferentes valores de concentración.

Precisión en la clasificación de superficies

En la primera aproximación analizada, que utilizó los datos de las seis bandas, se encontró que el modelo KNN obtuvo las mejores aproximaciones en la detección de aceite, a pesar de no haber alcanzado las mejores puntuaciones en las métricas. Sin embargo, se observó confusión en el modelo al clasificar una parte de la mancha como suelo. Esto podría deberse a que al utilizar todas las bandas, el modelo no pudo realizar una separación adecuada en los datos, ya que la inclusión de datos del espectro visible puede generar confusión (Winkelmann, 2005; Lu *et al.*, 2019).

Después de analizar la importancia de las variables, se encontró que las bandas infrarrojas (NIR, SWIR1 y SWIR2), junto con la banda azul, tuvieron mayor influencia en la separación de los datos para cada clase. Esto coincide con estudios anteriores (Fingas and Brown E., 2015; Leifer *et al.*, 2012; Lu *et al.*, 2019; Winkelmann, 2005), que también mencionan que este rango espectral es el más adecuado para diferenciar el aceite de otros materiales.

El uso de sólo estas bandas en el entrenamiento y validación de los modelos mostró una ligera disminución en los scores, y aun así el modelo *Random Forest* obtuvo las mejores puntuaciones. También se observó una reducción en las diferencias en los scores entre los datos de entrenamiento y validación, donde el modelo KNN mostró diferencias cercanas a cero (Figura 35A).

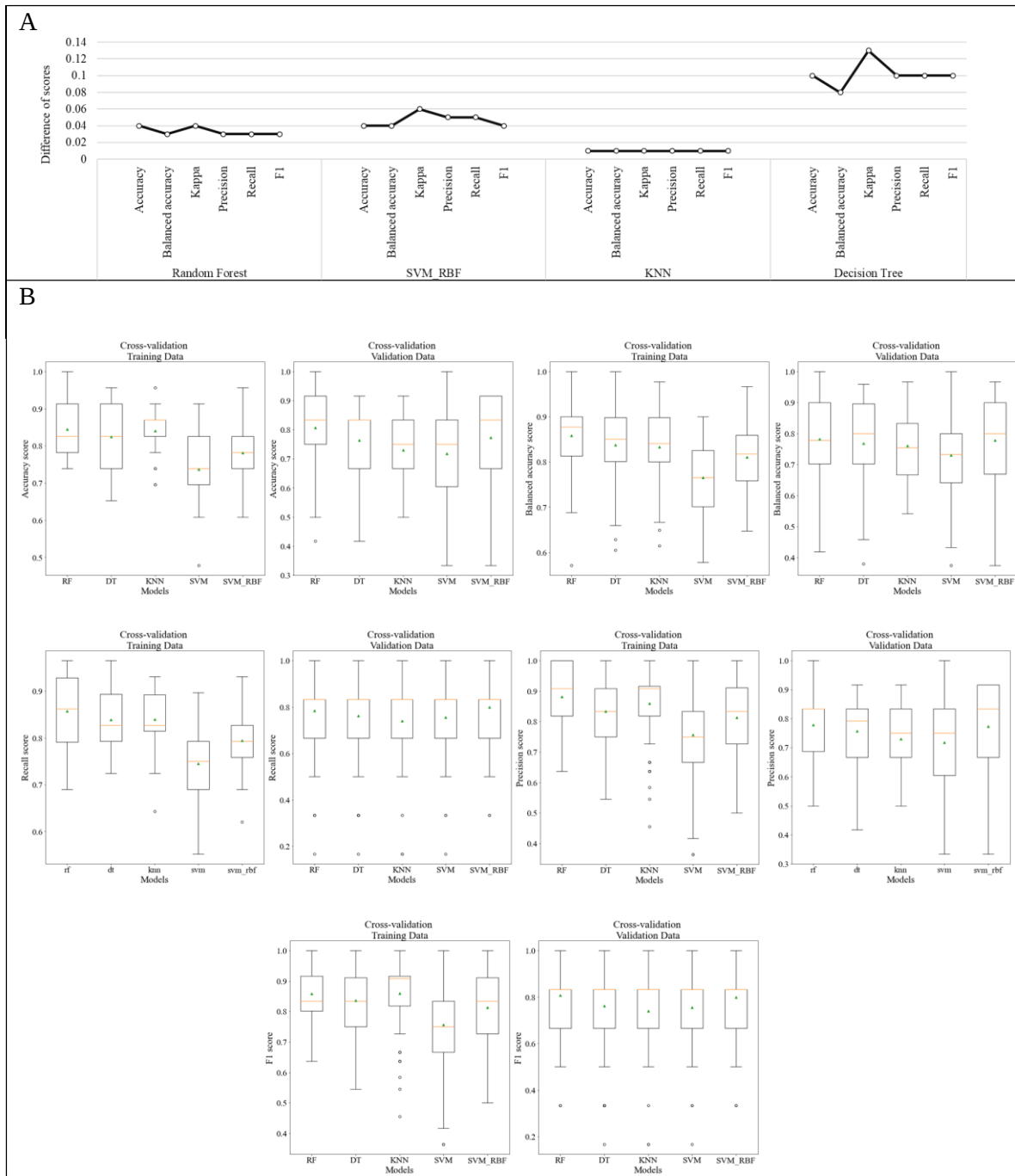


Figura 35. A) Diferencia de métricas para los modelos Random Forest, SVM RBF, KNN y Decision Tree. El modelo con las menores diferencias fue KNN, con diferencias de 0.1 en todas las métricas calculadas. B) Análisis de validación cruzada para las métricas utilizadas. Este análisis se realizó para los conjuntos de entrenamiento y validación. Observamos que el modelo KNN presenta una menor variación en los scores obtenidos mediante este método.

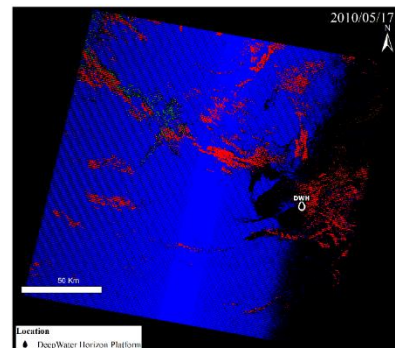
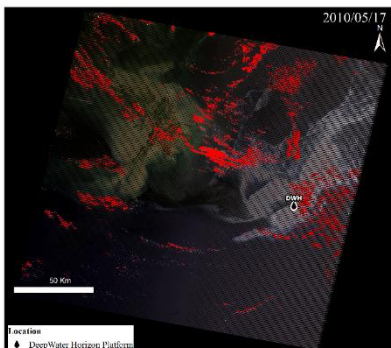
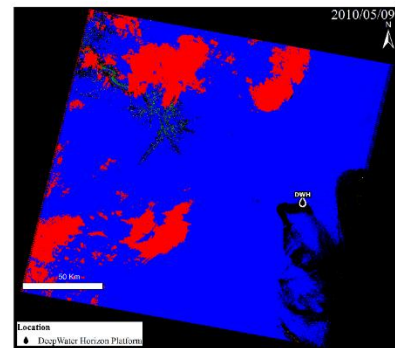
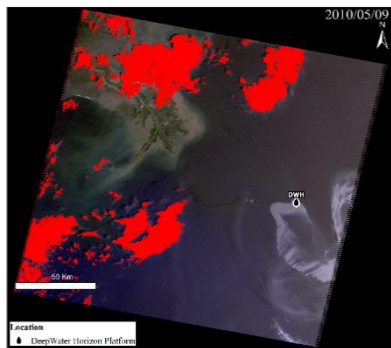
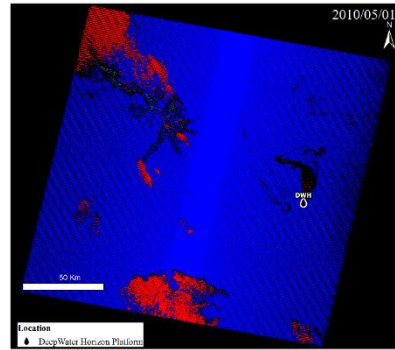
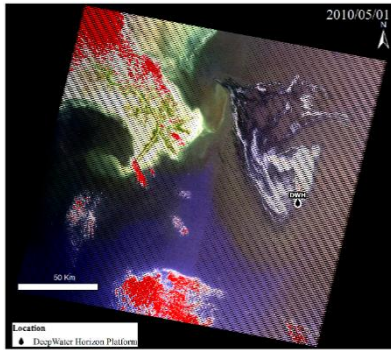
Tras probar los modelos, se observó que el modelo KNN fue el que presentó la mejor aproximación en la detección de aceite. Las diferencias cercanas a cero entre los conjuntos de entrenamiento y validación podrían explicar este resultado (Figura 35A). Además, al realizar el análisis de validación cruzada, este comportamiento es más evidente, ya que la variación de las métricas del modelo KNN es menor en comparación con los otros (Figura 35B). Estas comparaciones y los resultados de la subsección de detección de superficies (Figura 19), nos indican que el mejor modelo para detectar derrames de petróleo, cuantitativa y cualitativamente, es el modelo KNN.

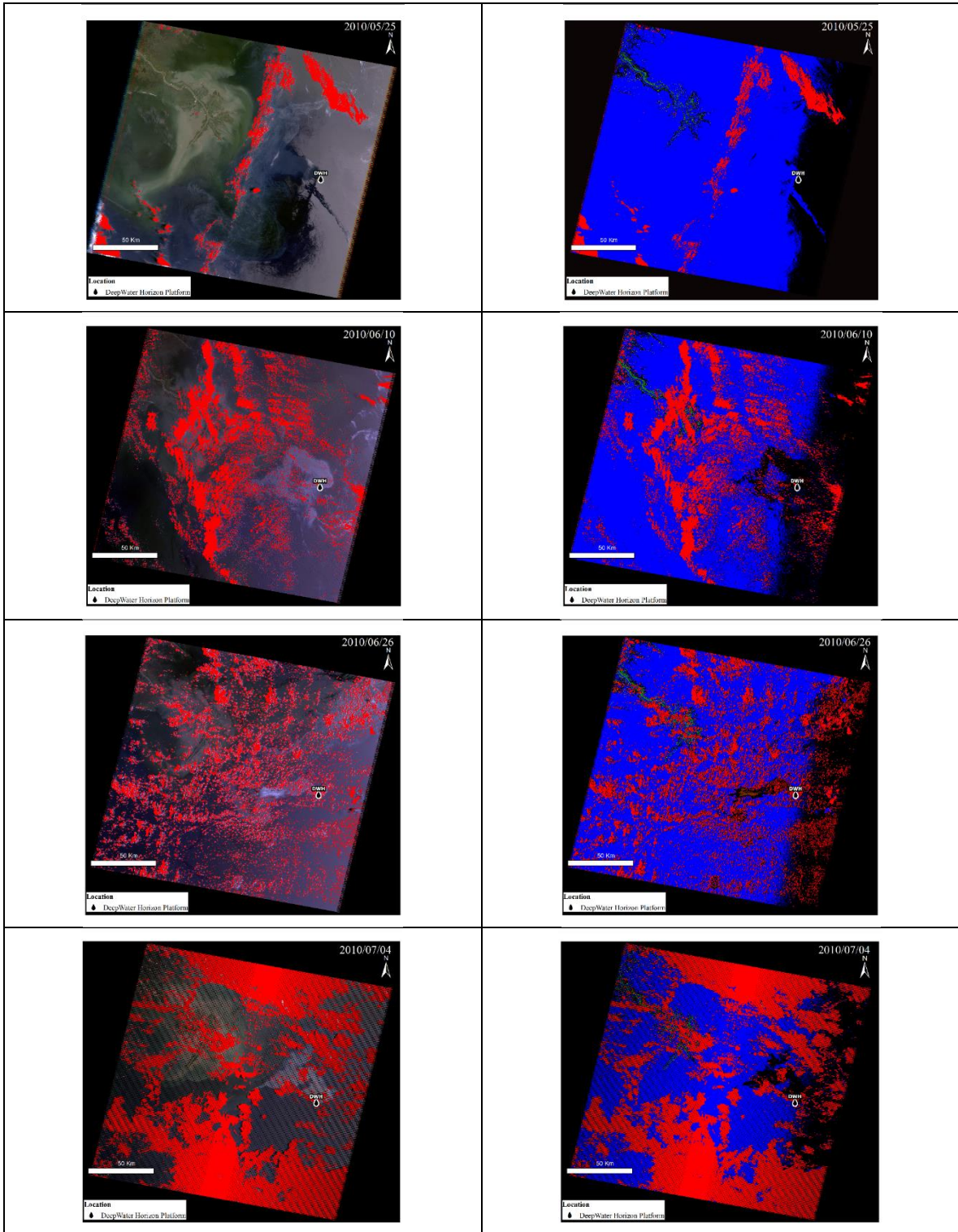
Además, se probó el modelo KNN en ocho imágenes con petróleo (Figura 36) y 10 imágenes sin petróleo (Figura 37), que corresponden al derrame ocurrido en el año 2010. Se obtuvo que, en las ocho imágenes, el modelo pudo clasificar la mancha del derrame en la clase de aceite, y en las 10 imágenes sin aceite, el modelo efectivamente no identificó aceite, es decir, obtuvo un 100% de precisión en la identificación del aceite en las imágenes (Tabla 8).

Tabla 8. Matriz de confusión con los resultados de clasificación del modelo KNN para las imágenes con aceite y sin aceite. El modelo obtiene una precisión del 100% para ambos casos.

		Predicción	
		Imágenes con aceite detectado	Imágenes en las que no se detecta aceite
Observación	Imágenes con aceite	8 / 8	0
	Imágenes sin aceite	0	10 / 10

Imágenes con aceite Derrame DeepWater Horizon





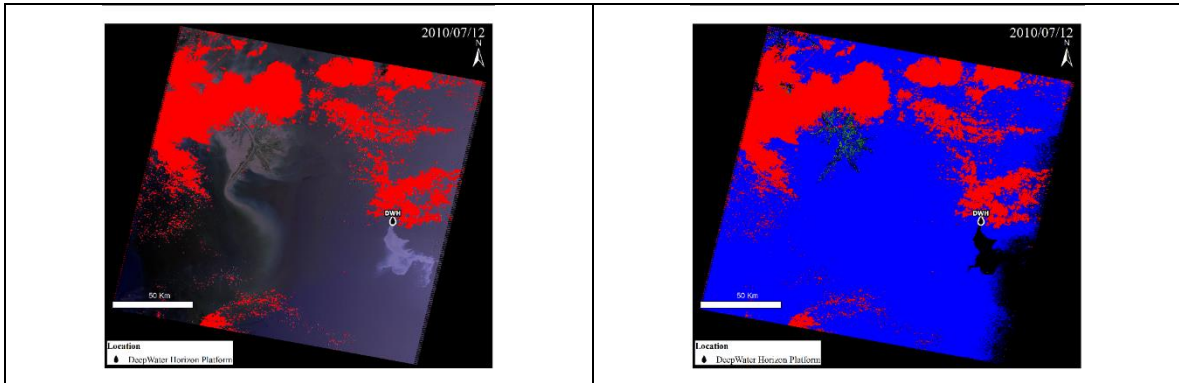
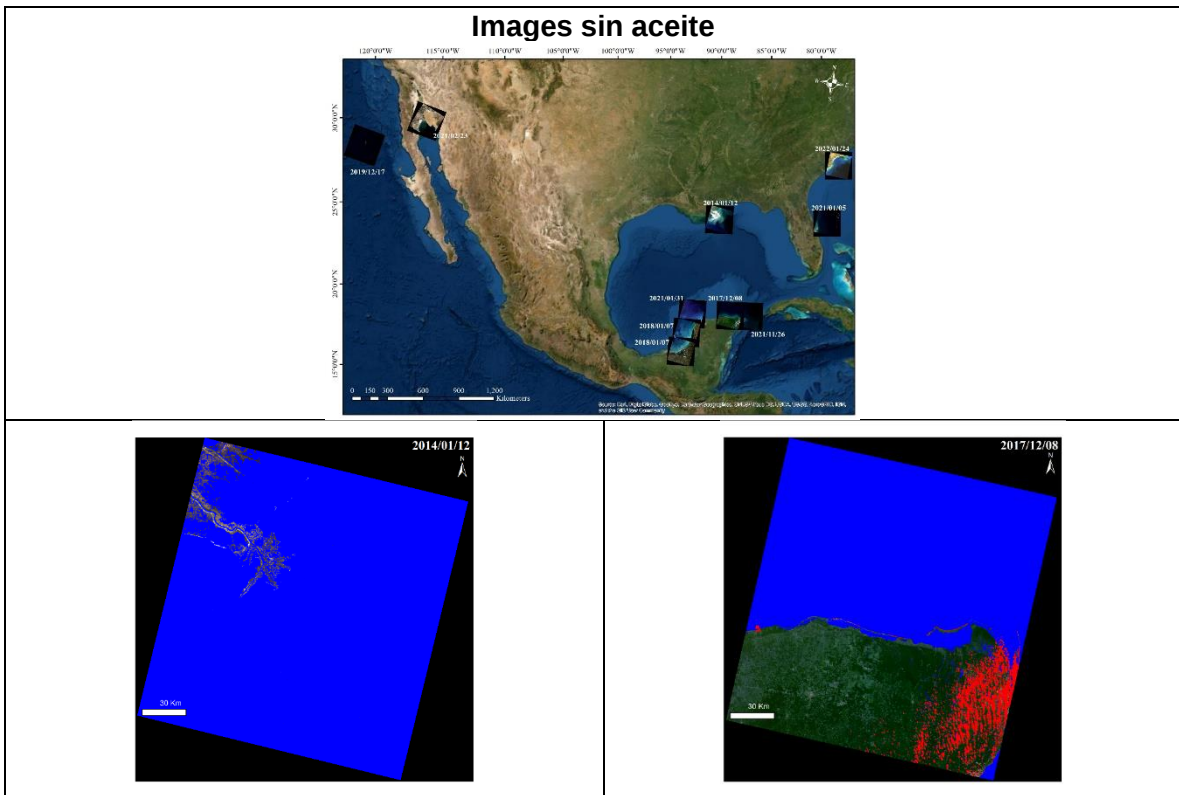


Figura 36. Resultados del modelo KNN de clasificación de variables importantes para imágenes con petróleo. La primera imagen es la localización de las imágenes, que corresponden a la zona de la plataforma DeepWater Horizon. A continuación, la columna de la izquierda muestra las imágenes RGB obtenidas, y la columna de la derecha muestra los resultados de clasificación del modelo. El color azul corresponde al agua, negro al petróleo, marrón al suelo, verde a la vegetación y amarillo al plástico. Las nubes aparecen enmascaradas en rojo en todas las imágenes.



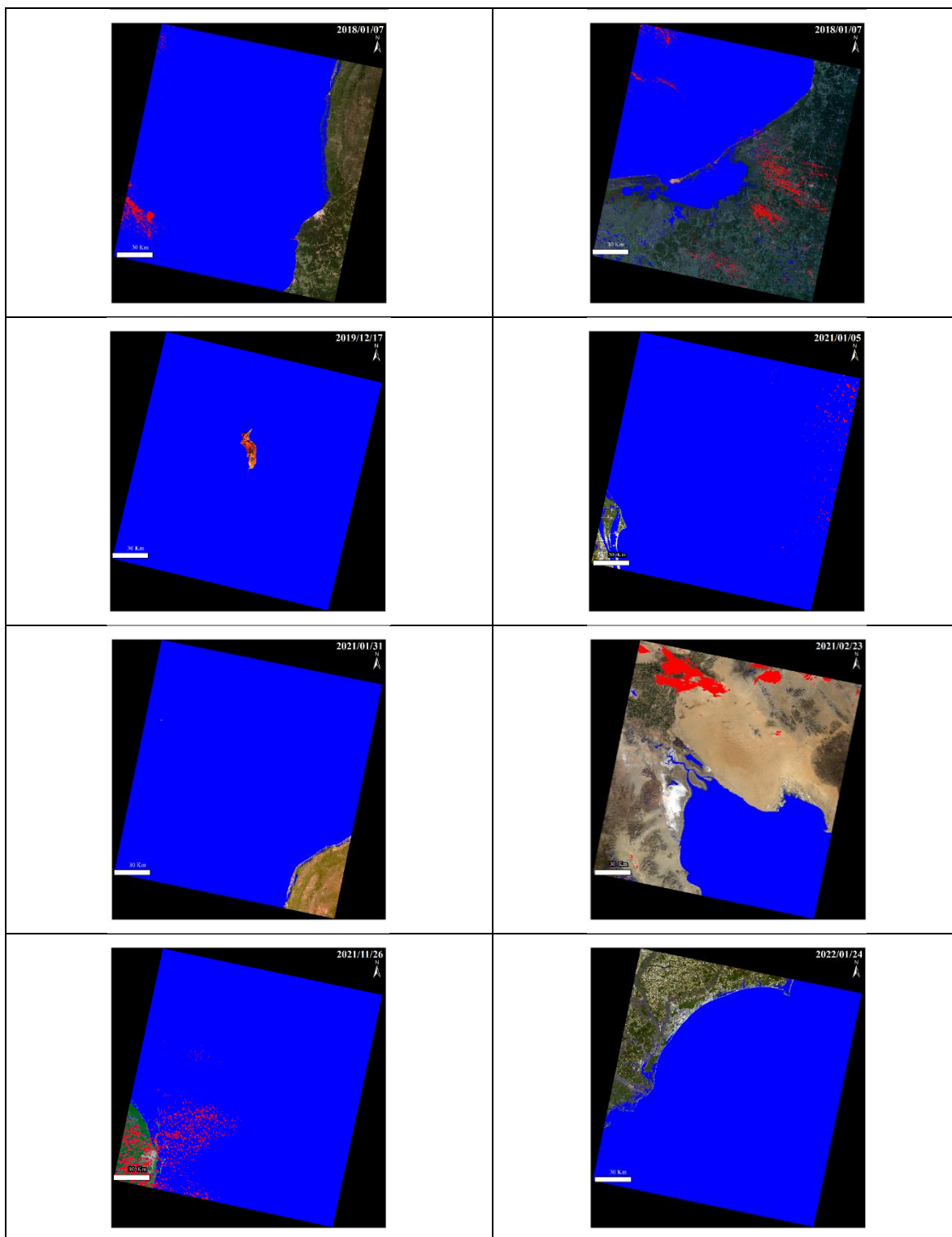


Figura 37. Resultado del modelo KNN de variables importantes. La primera imagen corresponde a la localización de todas las imágenes sin aceite. A continuación, se muestran los resultados de la clasificación. La parte de agua sólo se asignó a la clase de agua, y no se observó petróleo. La Tierra se enmascaró y se observa en color verdadero, y las nubes se enmascaran en rojo en todas las imágenes.

En las imágenes del derrame de 2010 también se puede observar que algunas zonas costeras fueron clasificadas como petróleo, esto podría deberse a que, durante el evento, el petróleo estuvo presente en la costa de Luisiana y los deltas del Mississippi (Beyer *et al.*, 2016; Garcia-Pineda *et al.*, 2017; Thyng, 2019).

Potencial de detección de aceite en imágenes satelitales ópticas

En los derrames de petróleo, los factores ambientales afectan su espesor y concentración, lo que hace que la detección sea una tarea difícil (Clark *et al.*, 2010; Leifer *et al.*, 2012; Fingas and Brown, 2018). Estos factores pueden provocar un proceso de meteorización constante, y ocasionar que petróleo pueda encontrarse en emulsiones de agua en petróleo (WO por sus siglas en inglés) y petróleo en agua (OW por sus siglas en inglés) en las imágenes (Figura 38A, Figura 38B), las cuales en ocasiones no generan contrastes positivos con el agua, y por tanto pueden no ser detectables (Leifer *et al.*, 2012; Lu *et al.*, 2020).

En el estudio de Lu *et al.*, 2020, para la imagen del 25 de mayo, se identificaron algunas zonas con emulsiones WO y OW mediante una combinación de bandas de falso color (SWIR1, NIR, RED) (Figura 38C). En nuestro estudio, el modelo de clasificación KNN, de bandas importantes, detectó ambas emulsiones y les asignó la clase de aceite (Figura 38D), sin embargo, en ambos casos no se detecta completamente la mancha observada en la imagen de falso color. Cabe señalar que, en nuestro caso, no se distinguió entre emulsiones porque el objetivo principal era detectar el aceite y diferenciarlo de otros materiales.

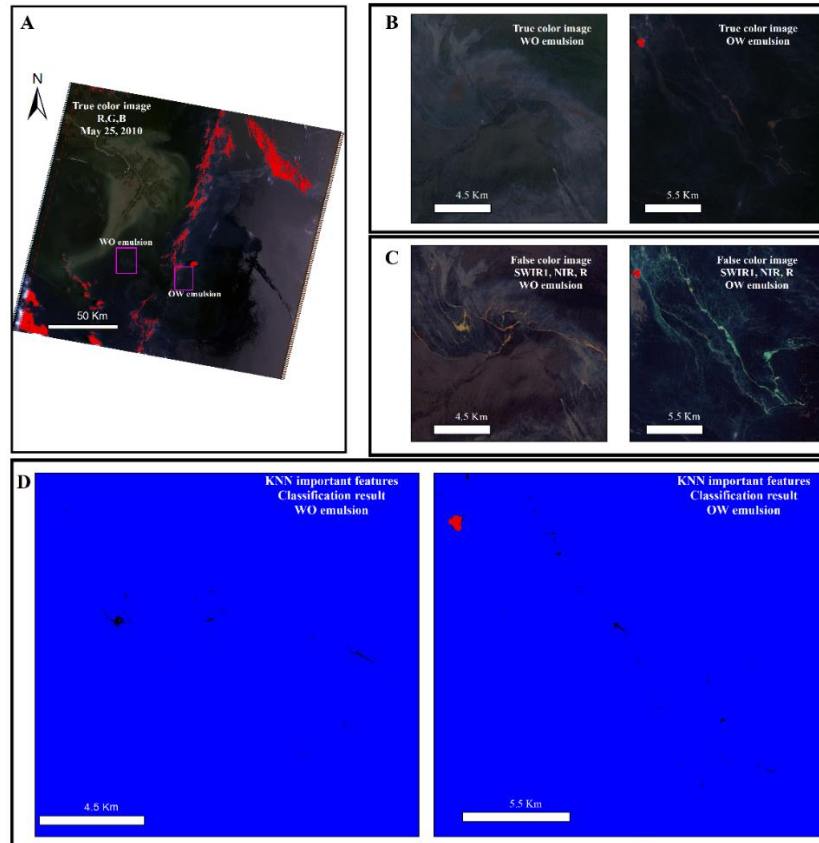


Figura 38. A) Imagen RGB del 25 de mayo de 2010, los recuadros corresponden a las zonas donde hay emulsiones (Lu et al., 2020). B) Imágenes RGB del acercamiento a las emulsiones WO y OW. C) Imágenes en falso color (SWIR1, NIR, R) de las emulsiones. D) Resultado de la clasificación con KNN para bandas importantes, ambas emulsiones se clasificaron parcialmente como aceite.

Esta identificación parcial puede deberse a la heterogeneidad del petróleo causada por los procesos de meteorización, lo que provoca que el modelo confunda el petrolero con agua, materiales orgánicos (Fingas and Brown, 2017; Leifer et al., 2012; Sun et al., 2015), o incluso al plástico (Hörig et al., 2001). Asimismo, en la detección de derrames en imágenes ópticas, es importante considerar que zonas o superficies con alto contraste (nubosidad, o zonas de bajo viento) pueden ser confundidas con manchas de hidrocarburos (Bonn Agreement Accord de Bonn, 2007; Leifer et al., 2012; Lu et al., 2019, 2020).

Datos espectrales e imágenes ópticas para la estimación de la concentración

Se desarrollaron modelos lineales, generales (Figura 13A) y por banda (Figura 13B), para analizar la relación entre reflectancia y concentración, mostrando una relación inversa. Estudios previos han mencionado que, al tener datos de muestras de aceite con diferentes concentraciones, se pueden desarrollar modelos lineales y no lineales para relacionar la reflectancia con la concentración, y así predecir el proceso de degradación del hidrocarburo (Daling and Strøm, 1999; Daling *et al.*, 2014; Lu *et al.*, 2019). En nuestro caso, utilizar ML para estimar la concentración, en lugar de un modelo lineal convencional, nos permitió aumentar la precisión.

En los resultados se observó que la concentración estimada para las cuatro imágenes fluctuaba entre 40% y 60% (Figura 22). Al realizar los perfiles, se puede observar esta fluctuación en la concentración estimada a lo largo de las dos líneas trazadas en la imagen (Figura 23B y Figura 23C). Esta fluctuación es un comportamiento esperado, ya que, a mayor concentración, menor reflectancia y viceversa, es decir, reflectancia y concentración tienen una relación inversa, concordando con Lu *et al.*, 2019.

En las gráficas (Figura 39) del transecto 1 (Figura 23A), se observan valores bajos de reflectancia al inicio debido a la ausencia de petróleo en esa región. Este comportamiento espectral característico del agua se hace más evidente en las bandas SWIR1 y SWIR2 con valores cercanos a cero, dado que el agua tiende a absorber mayor radiación electromagnética en estas zonas del espectro. Y cuando se detecta petróleo, se observa un aumento brusco de los valores de reflectancia.

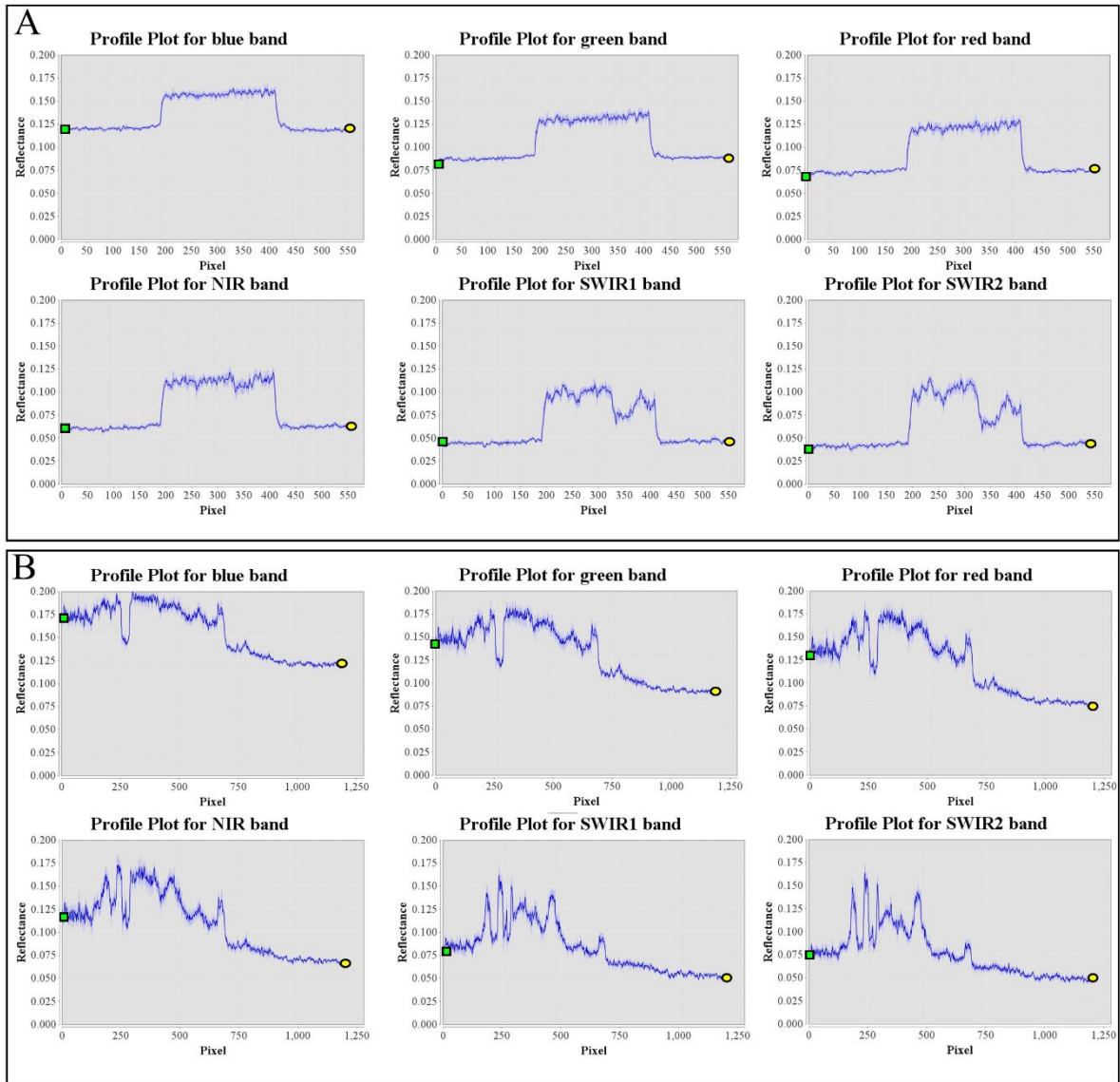


Figura 39. Gráficos de los valores de reflectancia de los perfiles de la Figura 23 para cada banda utilizada. Sólo se utiliza la imagen de mayo de 09. El cuadrado verde representa el inicio, y la elipse amarilla representa el final del perfil en la imagen A) Corresponde al transecto 1 en la Figura 23A. B) Corresponde al transecto 2 de la Figura 23A.

Por lo tanto, con el modelo ML, entrenado utilizando datos de aceite con diferentes concentraciones, se genera una buena aproximación en la estimación de la concentración. Sin embargo, se trata de una primera aproximación, lo que indica que este modelo puede mejorarse.

Sensores activos

Importancia del contexto y forma en la discriminación de aceite

Para la clasificación y discriminación de imágenes con aceite se desarrolló un modelo de DL, específicamente un modelo CNN, el cual logró obtener buenas aproximaciones en la clasificación de imágenes con presencia de petróleo, alcanzando valores de Exactitud de 0.99. Este resultado se debe a que este tipo de modelos son capaces de extraer la información contextual y características específicas de las imágenes, por medio de las operaciones de convolución (LeCun, Bengio and Hinton, 2015; Yamashita *et al.*, 2018).

Esta extracción de características las podemos observar en cada una de las convoluciones del modelo CNN. En el caso de los derrames de petróleo (Figura 40), muestran una progresión desde características muy simples como bordes e intensidad en las primeras convoluciones, hasta características más complejas a medida que se profundiza en el modelo. En general, se observa que el modelo aprendió un patrón lineal regular (Figura 40), esta característica de linealidad la presentan muchos de los derrames que provienen de embarcaciones, ya que el movimiento constante de estas fuentes evita que el petróleo se acumule en un punto (Solberg, 2012). Además, este patrón de linealidad puede presentarse en los derrames de plataforma, ya que a pesar de ser fuentes fijas, existe el viento y los movimientos del agua que ocasiona que el petróleo tome una forma lineal (Solberg, Brekke and Husoy, 2007; Solberg, 2012)

Además, existen otros aspectos que se consideran y pueden observarse en las imágenes SAR, como barcos y plataformas (Solberg, Brekke and Husoy, 2007; Solberg, 2012). Los cuales pueden observarse como puntos aislados al inicio del derrame en algunas convoluciones (Figura 40).

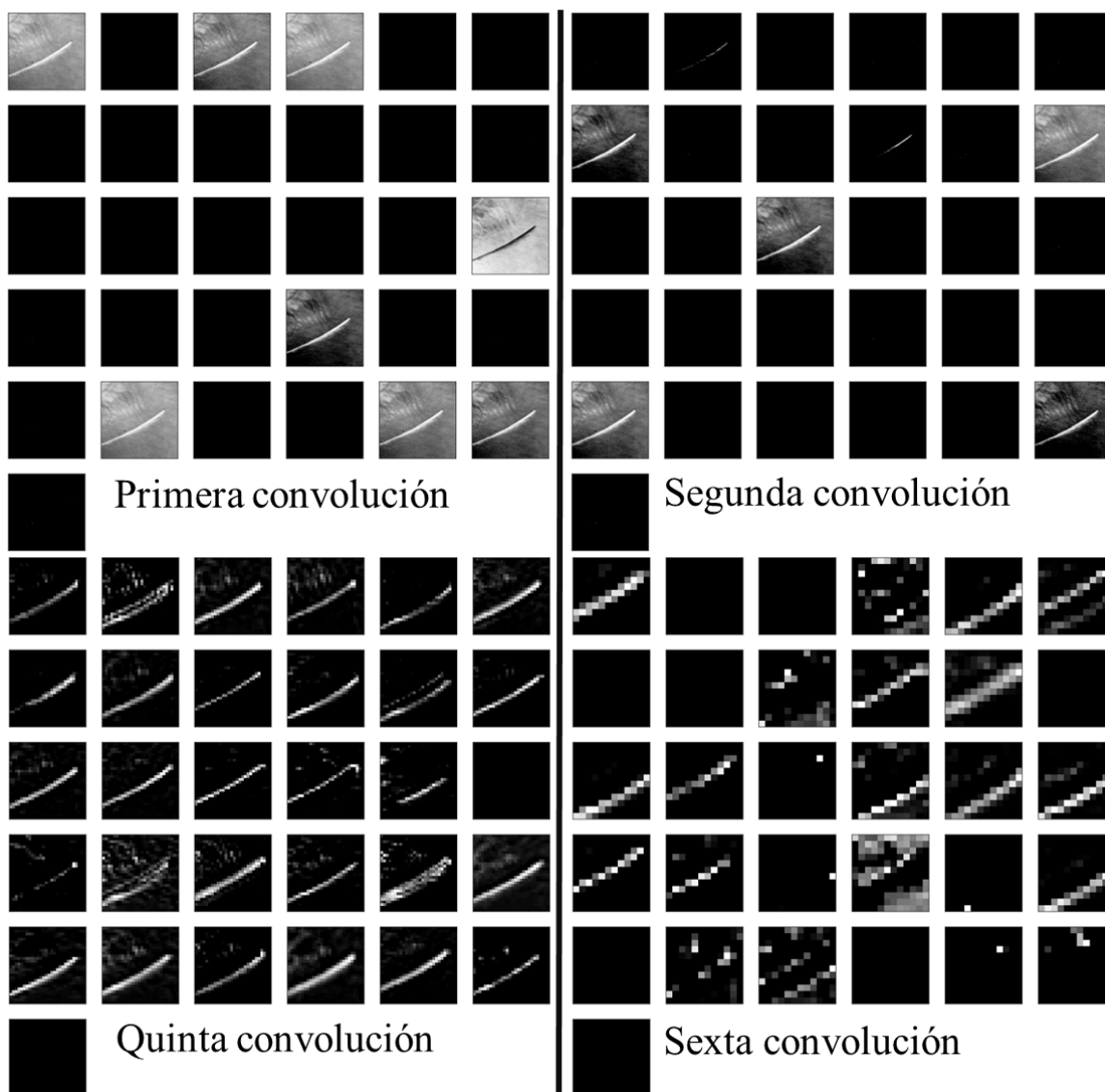


Figura 40. Extracción de características en algunas convoluciones del modelo CNN para la imagen Sentinel-1 SAR con presencia de aceite. Se observa que en la primera convolución se extraen características sencillas como el borde e intensidad; y en las otras convoluciones son más complejas. Se observa que, para el aceite, el modelo logra extrae y aprender un patrón regular y lineal correspondiente a la mancha.

Para el caso donde no se encuentra la presencia de aceite dentro de las imágenes, las características que aprende el modelo son patrones irregulares y, hasta cierto punto, caóticos, lo cual representa la naturaleza misma del movimiento y rugosidad del agua (Kim, Moon and Kim, 2010; Li *et al.*, 2018) (Figura 41).

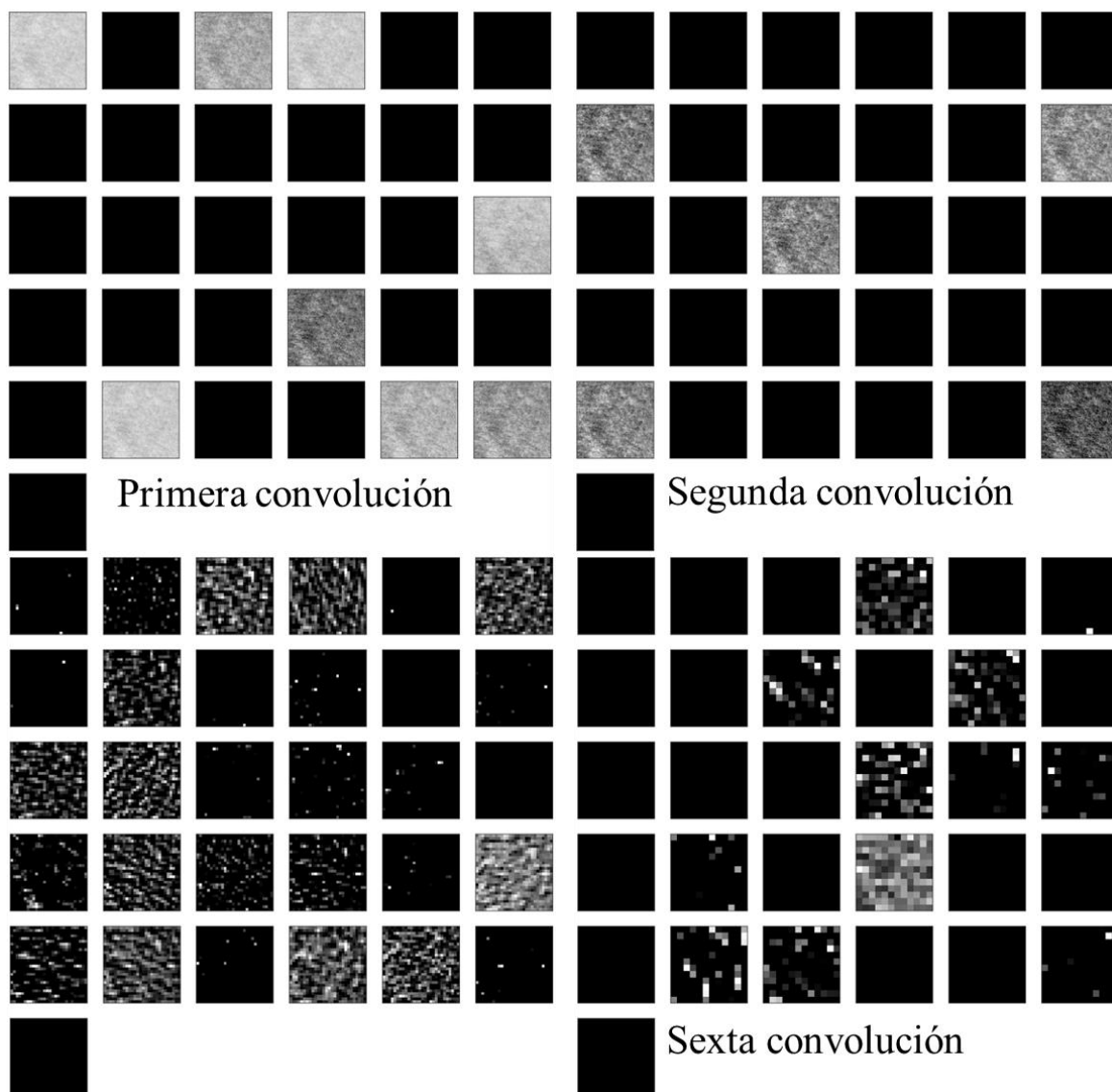


Figura 41. Extracción de características en algunas convoluciones del modelo CNN para la imagen Sentinel-1 SAR sin aceite. En general se observa que el modelo aprende aspectos y características sin un patrón aparente.

Sin embargo, el aprender a diferenciar entre agua y petróleo es considerada una tarea relativamente fácil, ya que como se observó (Figura 40, Figura 41) los patrones, y formas son completamente diferentes, encontrándose únicamente un patrón regular en el derrame (Solberg, Brekke and Husoy, 2007; Solberg, 2012).

La tarea difícil es discriminar el aceite de *look-alikes*, debido a que existen fenómenos que generan manchas similares a las que presenta el aceite dentro de las imágenes SAR (Brekke and Solberg, 2005; Solberg, 2012; Fingas and Brown, 2018; Alpers, Liu and Wu, 2019). Estas pueden ser películas/manchas naturales (florecimientos algales), hielo, áreas con viento menor a 3 m/s), viento protegido por tierra, celdas de lluvia, zonas de cizalla, entre otros (Solberg, 2012). Por lo tanto, el realizar la identificación de derrames de petróleo se convierte en una tarea complicada en zonas donde es común la presencia de estos fenómenos (Brekke and Solberg, 2005; Solberg, 2012).

Sin embargo, al observar las características extraídas y aprendidas por el modelo en cada una de las convoluciones para la imagen del *look-alike* (Figura 42), se aprecia que los patrones al inicio se basan en aspectos de intensidad y borde, pero en la quinta convolución dichos patrones son más caóticos, lo cual difiere de las características visuales correspondientes al derrame (Figura 41).

En este sentido, la forma de la mancha es importante en términos de discriminación entre derrames de petróleo y *look-alike*, además, la homogeneidad de la mancha oscura y de su entorno también es útil, ya que es más probable que una mancha

aislada en un entorno homogéneo sea petróleo que la misma mancha en un entorno heterogéneo (Solberg, 2012).

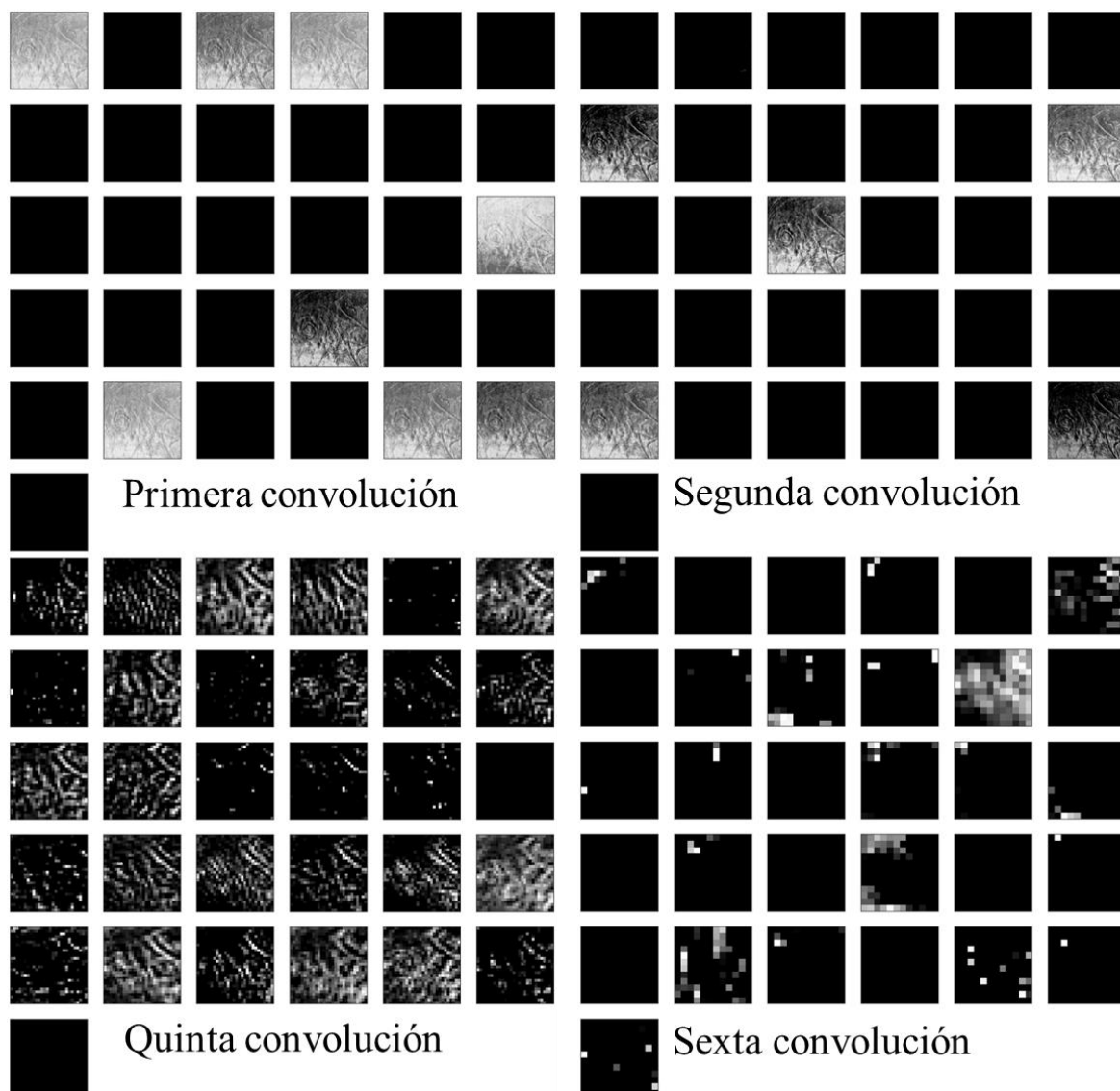


Figura 42. Extracción de características en cada una de las convoluciones del modelo CNN para la imagen Sentinel-1 de look-alike. Para este tipo de superficie se observa que el modelo extrae y aprende características que no presentan patrones muy definidos, es decir son muy irregulares.

Al analizar este comportamiento se logró observar, que, a pesar de que los look-alikes y los derrames puedan presentar características similares, si es posible

discriminarlos al presentar patrones visuales distintos, además, de que las manchas, en su mayoría se presentan aisladas, en comparación con los look-alikes.

Comparación de modelos de segmentación

El incorporar la información contextual y espacial de las imágenes, no solo se considera para la parte de clasificación, sino también en la parte de segmentación, ya que el entrenar un modelo con únicamente los valores de los píxeles, como el caso de los modelos de MLP, generan de falsos positivos (Figura 43). Esto se debe a que los derrames de petróleo amortiguan el ruido del mar circundante entre 0.6 y 13 dB, y a las películas naturales entre 0.8 dB y 11.3 dB, teniendo valores con los mismos rangos, lo que genera los errores (Brekke and Solberg, 2005; Solberg, Brekke and Husoy, 2007).

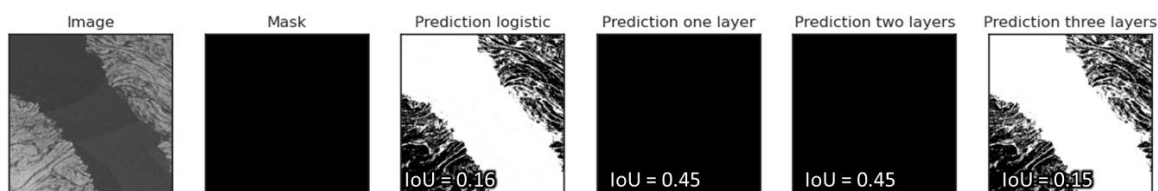


Figura 43. Resultado de la segmentación realizada por los modelos MLP para una imagen de look-alike. Se observa que los modelos de una y dos capas ocultas no segmentan nada, pero el modelo logístico y de tres capas ocultas segmentan la superficie como aceite, generando un falso positivo.

Estos errores en los modelos MLP para la segmentación, podrían minimizarse con el modelo CNN al realizar un primer filtrado y evitar la segmentación de imágenes con look-alikes. Sin embargo, en escenarios reales, pueden existir casos en los cuales dentro de la misma imagen se encuentren tanto aceite como look-alike. En estos casos el modelo CNN clasificaría a la imagen con presencia de aceite, y si se

utilizara el modelo MLP para segmentación se podrían generar falsos positivos debido a que podrían segmentar los look-alikes como aceite como se observó en la Figura 43.

Por esto, para la parte de segmentación se tienen dos perspectivas, una mediante el uso de un MLP, y otro con el modelo U-net que considera el contexto al igual que el modelo CNN. Se realizó la comparación de ambas estrategias utilizando tres casos donde se encontró la presencia de aceite y look-alike (Figura 44). Para estas imágenes el resultado del modelo CNN ubicó a las imágenes en la clase aceite, por lo que el siguiente paso es la segmentación del derrame. Se observa que el modelo de MLP obtuvo valores de IoU de 0.36 para el primer caso, 0.58 para el segundo y 0.61 para el tercero. Por otra parte, el modelo Unet2D dio valores más altos, de 0.81, 0.82, y 0.96.

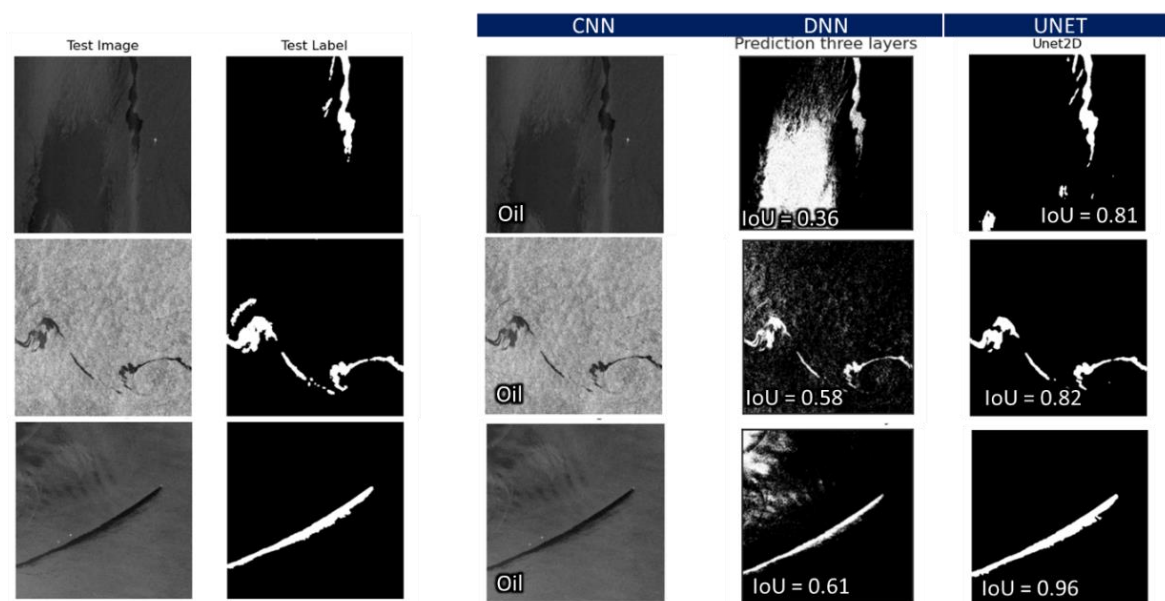


Figura 44. Comparación de las estrategias de segmentación posterior a la identificación de aceite dentro de las imágenes. En la primera columna se encuentra la imagen utilizada para prueba, en la segunda columna se encuentra la máscara o ground truth, en la tercera columna se encuentra la predicción realizada mediante el

modelo CNN, en el cual para los ejemplos utilizados se detectó la presencia de aceite. La cuarta columna corresponde a la predicción realizada por el modelo MLP de tres capas ocultas, donde se observa que los valores de IoU se encuentran entre 0.36 y 0.61. La última columna es la predicción realizada por el modelo Unet2D, dando valores de IoU de 0.81, 0.82, y 0.96 para el tercer caso.

Con esta comparación, se destaca la importancia del contexto en las tareas de segmentación de derrames de petróleo, donde el modelo Unet2D logra una segmentación precisa del aceite, mientras que el modelo MLP tiende a generar falsos positivos al clasificar erróneamente zonas de look-alikes como aceite. Esto es crucial para un sistema de alerta o monitoreo, evitando gastos innecesarios y movilizaciones indebidas.

Modelo U-net con pérdida focal como alternativa para segmentación de aceite

Tras comparar los resultados, el modelo Unet2D con mayor número de capas convolucionales y filtros ($16 \rightarrow 32 \rightarrow 64 \rightarrow 128 \rightarrow 256 \rightarrow 512 \rightarrow 1024$), utilizando *Pérdida Focal* como función de pérdida e IoU como métrica para monitorear el entrenamiento, además de Upsampling2D en la parte de expansión del modelo (Figura 45), demostró tener el score IoU más alto. Alcanzando un valor de IoU de 0.96 en entrenamiento y 0.90 en validación y pruebas (Figura 30, Figura 33).

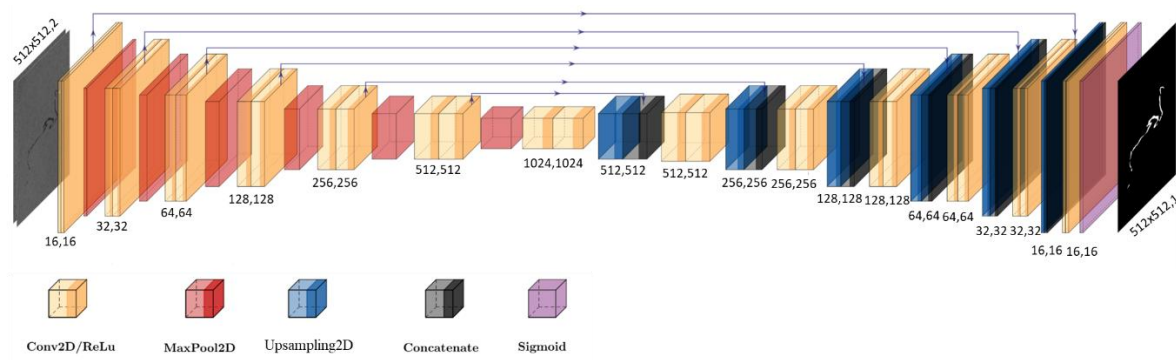


Figura 45. Esquema de la arquitectura del modelo U-net con las mejores aproximaciones para la segmentación de aceite.

En el modelo, se destaca el uso de *Pérdida Focal (Focal Loss)* en lugar de *Entropía Cruzada* para abordar el desequilibrio de clases (Lin et al., 2017), lo cual es crucial, ya que en las imágenes pueden encontrarse derrames pequeños o microderrames, que generan manchas aisladas difíciles de diferenciar y segmentar (Fingas, 2011; Schmidt-Etkin, 2015), y que provocan un desbalance de clases entre el *fondo* y el *plano principal (foreground)*. Al usar *Pérdida Focal*, se penalizan los ejemplos fáciles de clasificar, como el fondo de la imagen, evitando que abrumen el detector durante el entrenamiento (Lin et al., 2017; Basit et al., 2022).

Existen trabajos que utilizan *Entropía Cruzada* y *Exactitud* (Tabla 9), pero estas métricas pueden generar sesgo y errores en las tareas de segmentación al enfocarse en la clase mayoritaria para dar un score (Lin et al., 2017; Basit et al., 2022). Esto se refleja en valores de exactitud de entre 0.85 y 0.94, scores considerados altos, pero al utilizar IoU para evaluar la precisión de la segmentación, los scores que reportan, en algunos casos, disminuyen drásticamente dando valores de entre 0.43 y 0.83 (Tabla 9).

En contraste, nuestro modelo presenta métricas más consistentes, con una *Exactitud* de 0.99 y un *IoU* de 0.90, lo que podría atribuirse al uso de *Pérdida Focal* como función de pérdida. Esta elección resulta beneficiosa para problemas con un gran desbalance de clases, como los derrames de petróleo, ya que ha demostrado funcionar mejor que otras funciones de pérdida en este tipo de casos (Lin *et al.*, 2017; Basit *et al.*, 2022).

Tabla 9. Trabajos realizados para la segmentación de derrames de petróleo con imágenes SAR. Se observa el modelo propuesto, así como algunas métricas que utilizaron.

Modelo	Exactitud	IoU	F1	Autores
ANN	91.6			Singha et al., 2013
Mask R-CNN	96.6			Temitope Yekeen et al., 2020
U-net con atención	91.2	83.4		Zhu et al., 2022
U-Net	85.61	81.46		Zhu et al., 2022
Conditional Advesarial Networks	99			Li et al., 2021
AlexNet	97-98			Wang et al., 2021
Faster -CNN	89.23			Huang et al., 2022
Unet		53.79		Krestenitis et al., 2019
SVM-ANN	98.13			Dasari et al., 2022
DeepLabV3+			75.08	Dehghani-Dehcheshmeh et al., 2023
FC-DenseNet			73.94	Dehghani-Dehcheshmeh et al., 2023
U-Net			60.85	Dehghani-Dehcheshmeh et al., 2023
CNN y ViT			78.48	Dehghani-Dehcheshmeh et al., 2023
DeepLabv3+ VisionTransformers	94.45	64.59		Dehghani-Dehcheshmeh et al., 2023
U-Net	94.17	43.73		Dehghani-Dehcheshmeh et al., 2023
Feature Merging Network		46.49		Fan et al., 2021
Dual Attention Encoding Network		85.3		Zhai et al., 2022
Ensemble model con UNET		71.12		Rousso et al., 2022
Two-stages CNN		66.67		Nieto-Hidalgo et al., 2018
Unet Efficient-net-B3		62.19		de Moura et al., 2022
DeepLabv3		96.08		Ma et al., 2023

Bilateral Segmentation Network (BiSeNetV2)		87.6	Chen et al., 2022
Multichannel Deep neural network	98.56	97.17	Hasimoto-Beltran et al., 2023
Nuestro modelo Unet profunda con pérdida focal	99	96	95

El trabajo de Hasimoto-Beltran et al., 2023, destaca por obtener valores altos en las métricas de *Exactitud* (98.56) e *IoU* (97.17), ya que utilizan *IoU* para monitorear el entrenamiento, lo cual es considerado alternativa efectiva en modelos de segmentación semántica (Ni wattanakul et al., 2013; Rahman and Wang, 2016). En nuestro estudio, también hemos adoptado el uso de *IoU* como métrica en el entrenamiento de nuestro modelo de segmentación, corroborando su efectividad en la segmentación de derrames de petróleo.

Sin embargo, es importante mencionar que en su enfoque utilizan solo 16 imágenes, lo cual puede generar un sesgo y limitar la generalización del modelo a otros eventos, como se evidencia en su análisis comparativo con imágenes de otros sitios (Krestenitis et al., 2019), donde obtuvieron scores de *IoU* más bajos (54.27 y 57.13).

En cambio, en nuestro estudio, se ha utilizado un enfoque más amplio y diverso en la selección de imágenes para el entrenamiento (Figura 4). Esto nos ha permitido abarcar una gran parte de la variabilidad de los distintos tipos de aceite y los patrones que toma el derrame debido a su interacción con los aspectos físicos, químicos y ambientales en cada sitio. Al hacerlo, hemos minimizado el sesgo del modelo al aprender datos de un solo sitio, lo que es crucial para lograr una mejor capacidad de generalización en diferentes escenarios de derrames de petróleo. Con esto, también se consigue minimizar o evitar el "data leakage" o fuga de datos,

común en conjuntos de datos complejos (Kaufman *et al.*, 2012). Así hemos obtenido métricas más consistentes y resultados más confiables en la segmentación de derrames.

Necesidad de un esquema clasificación y segmentación de derrames en dos etapas

El modelo Unet2D desarrollado para la segmentación de derrames de petróleo demostró tener un buen desempeño, en entrenamiento y validación (IoU de 95 y 90), así como en las pruebas (IoU 0.90). Sin embargo, el modelo no alcanza una precisión del 100%. Si se utilizara únicamente el modelo U-net para segmentación, se enfrentaría a desafíos en imágenes con presencia de look-alikes, lo que afectaría su desempeño. En otras palabras, el modelo generaría errores en la segmentación en presencia de *look-alikes* (Figura 46), lo que podría llevar a falsos positivos. Esto puede deberse a la complejidad y heterogeneidad de los look-alikes (Brekke and Solberg, 2005; Solberg, Brekke and Husoy, 2007).

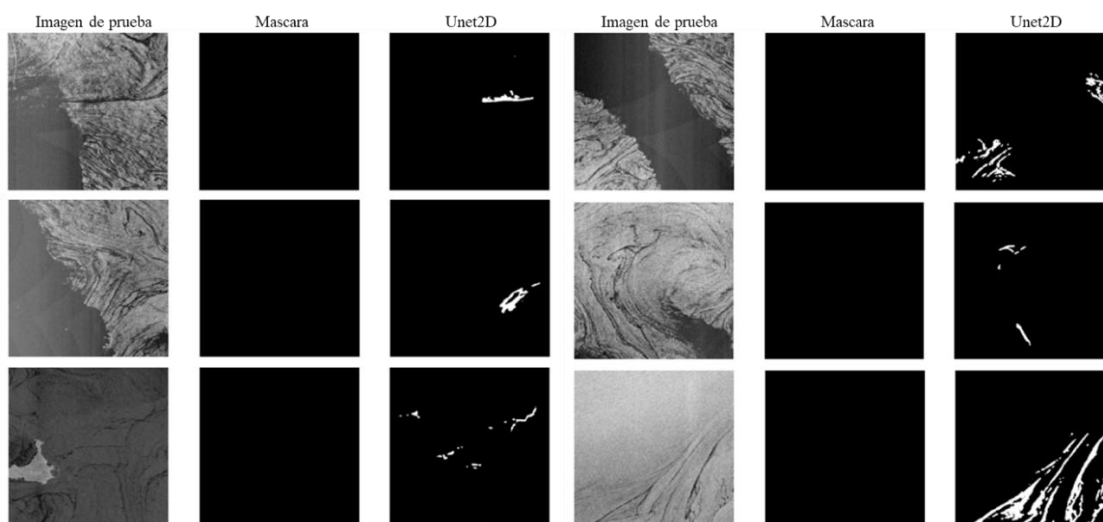


Figura 46. Ejemplos de imágenes con presencia de look-alikes. Se observa la imagen con su correspondiente mascara, así como la predicción realizada por el modelo Unet2D, la cual muestra errores en la segmentación.

Al realizar pruebas con el modelo CNN utilizando las mismas imágenes (Figura 46), los resultados muestran que el modelo clasifica estas imágenes en la clase "No Aceite", es decir, libres de petróleo (Figura 47). Esta observación destaca la necesidad del esquema en dos etapas, jerárquico, para la detección y posterior segmentación de derrames de petróleo, ya que este enfoque ayuda a minimizar la posibilidad de generar falsos positivos o falsas alarmas. Con esto, se observa que el uso jerárquico de ambos modelos, la detección inicial con CNN y la posterior segmentación con Unet2D, proporciona una estrategia más robusta para la identificación confiable de derrames de petróleo en imágenes SAR.

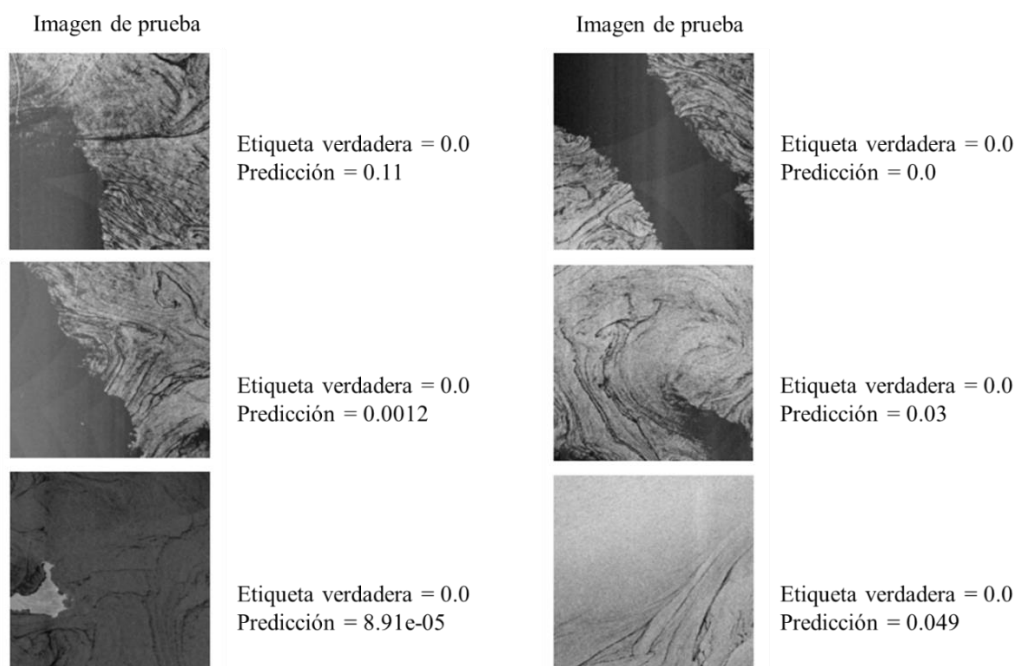


Figura 47. Resultados del modelo CNN. Las imágenes utilizadas fueron las mismas que en la Figura 46. Se observa que la etiqueta verdadera tiene valor de 0, ya que son imágenes que pertenecen a la clase No aceite, y el resultado de la predicción hecha por el modelo CNN ubica a las imágenes dentro de dicha clase al dar valores de probabilidad de 0 o cercanos a 0.

Por tanto, el esquema propuesto para las imágenes Sentinel-1 incluye las siguientes etapas (Figura 48): obtención y preprocesamiento de la imagen, división en sub-imágenes, detección con el modelo CNN para identificar la presencia de aceite, y segmentación con el modelo Unet2D para una delimitación más precisa. Si se detecta aceite en múltiples sub-imágenes cercanas, se realiza un mosaico para estimar el área total del derrame, y por último se obtiene la ubicación del derrame.

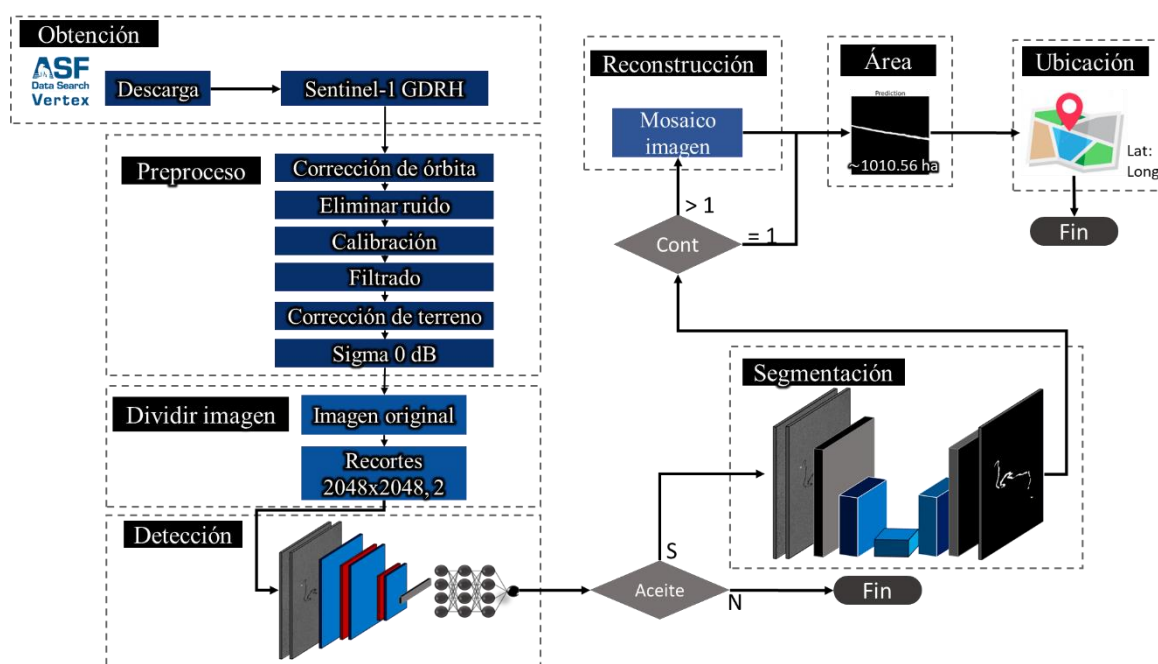


Figura 48. Esquema general de clasificación, segmentación y estimación de área para imágenes Sentinel-1 SAR. El esquema abarca desde la obtención, el preproceso y división de la imagen, posteriormente se encuentra la detección por el modelo CNN, y en dado caso de encontrarse la presencia de aceite pasara al modelo de segmentación. Por último, se realiza la estimación del área y localización de la mancha de aceite.

Utilizar un mecanismo de dos etapas ya ha sido probado por otros autores con buenos resultados, aunque aún se reportaban errores graves en la segmentación (Nieto-Hidalgo *et al.*, 2018; Shaban *et al.*, 2021). Sin embargo, nuestro modelo, al utilizar la función *Pérdida Focal*, ha demostrado obtener mejoras significativas en la

segmentación, lo que lleva a una identificación más precisa y reducción de errores en la delimitación de los derrames de petróleo en las imágenes Sentinel-1. Por esto, nuestro enfoque proporciona una herramienta más efectiva para la detección y cuantificación de derrames de petróleo en aplicaciones de monitoreo y respuesta ante estos desastres.

Un ejemplo de su implementación se muestra en la Figura 49 de dos imágenes Sentinel-1 del 17 de Julio del 2023, seleccionadas dada la presencia de un supuesto derrame de petróleo reportado por diversos medios de comunicación, el cual se cree que fue originado por la explosión de una de las plataformas petroleras cercanas a Campeche, México, el día 07 de Julio del 2023. Para ambas imágenes, se utilizó el esquema completo, preproceso, división, detección, segmentación, reconstrucción, estimación de área y localización. De acuerdo con nuestro enfoque, se detectó la presencia de un derrame de petróleo con coordenadas UTM WGS84 20.74285 (Superior), -95.15959 (Izquierda), 18.52329 (Inferior), -92.62886 (Derecha), y se estimó su extensión en ~38,215.9 hectáreas (~382.159 km²), con cercanía a las costas de Tabasco, México.

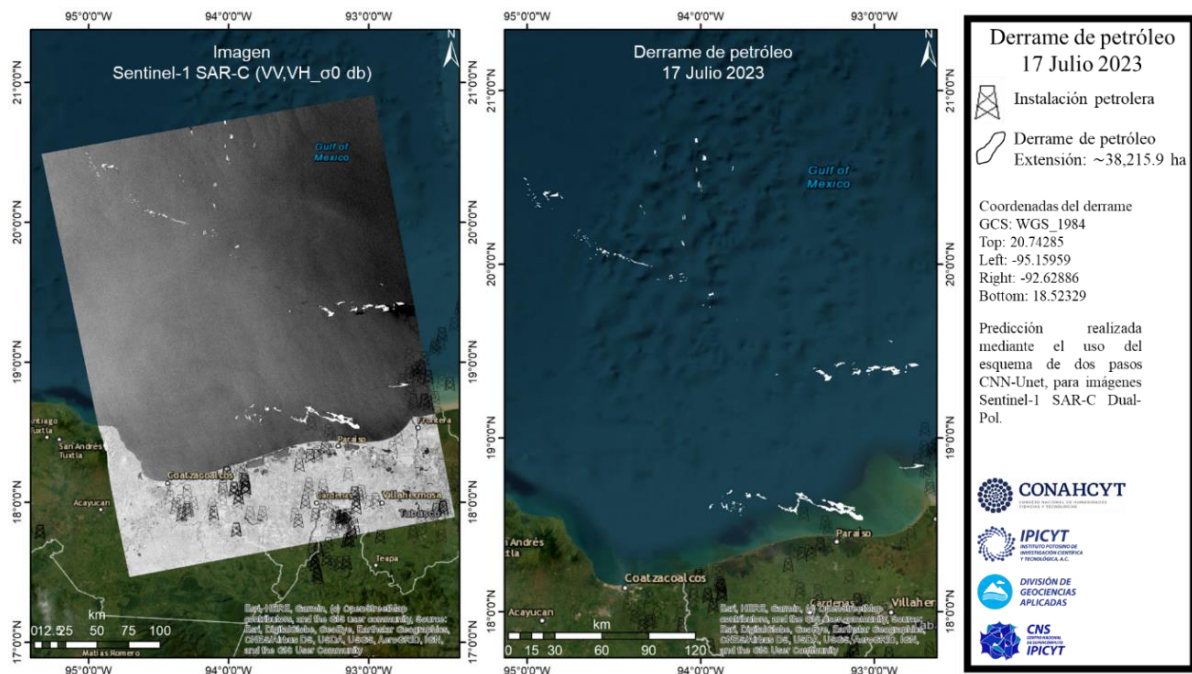


Figura 49. Mapa realizado a partir del esquema de clasificación por CNN y segmentación por el modelo U-net. Asimismo, se observa en la leyenda las coordenadas del derrame de petróleo y la extensión estimada en el recuadro de la derecha.

Consideraciones

Sensores pasivos

Aunque este estudio demuestra la posibilidad de detectar manchas de petróleo en imágenes ópticas utilizando datos de respuesta espectral y algoritmos ML, y la implementación de las características más importantes para diferenciar los materiales, hay ciertos aspectos que deben tenerse en cuenta para implementar un modelo capaz de diferenciar correctamente el petróleo de otros materiales. Hay que tener en cuenta que el petróleo está constantemente expuesto a factores ambientales que favorecen su degradación (Daling and Strøm, 1999; Daling *et al.*, 2014). Asimismo, la proximidad a las costas, y los flujos de sedimentos o lodos cercanos a la fuente del vertido, hacen que el petróleo pueda combinarse con otros materiales, generando en algunos casos agregados minerales, o el hundimiento de partes pesadas de petróleo (Daly *et al.*, 2016). Estos fenómenos de degradación y combinación generan una gran heterogeneidad en los valores de respuesta espectral del petróleo, haciendo que diferenciar la mancha de petróleo de otros materiales sea una tarea complicada. Esto puede remediarse aumentando el número de observaciones de petróleo y otros materiales, para lograr una mayor separabilidad. También puede mejorarse añadiendo diferentes tipos de petróleo y disponiendo de imágenes de satélite debidamente segmentadas, por clase de material. Y si es posible, disponer de mediciones in situ que permitan la verificación de los resultados.

En la parte de estimación de la concentración, se debe considerar la degradación constante del petróleo ya que los fenómenos pueden ocurrir de un día para otro, generando variaciones a nivel subpíxel, lo que complica el seguimiento de la evolución de la mancha de petróleo (Svejkovsky *et al.*, 2016), además los satélites tienen la desventaja de proporcionar información desactualizada debido a su baja resolución temporal (Oscar Garcia-Pineda *et al.*, 2013; Garcia-Pineda *et al.*, 2017). Sin embargo, se ha señalado que si el hidrocarburo está bien caracterizado y se conocen las condiciones ambientales de velocidad del viento, estado del mar, corrientes, salinidad, temperatura y radiación solar, debería ser posible calcular las tasas de muchos de estos procesos, y así establecer cómo cambia el estado del hidrocarburo con el tiempo (Daling *et al.*, 1990). Además, hay que tener en cuenta que, al igual que en los métodos de clasificación, puede haber bandas que generen ruido u otras que aporten información relevante para mejorar las estimaciones del modelo.

Estas cuestiones hacen que la detección de hidrocarburos, así como su seguimiento a lo largo del tiempo, sea una tarea compleja. Esta complejidad en la detección de derrames de petróleo es un campo que necesita ser estudiado y ampliado, enfocándose en incorporar nuevas tecnologías de teledetección y análisis de datos. Sin embargo, al igual que otros estudios, este trabajo ofrece una perspectiva sobre la aplicación de algoritmos ML entrenados con datos de respuesta espectral en la detección de petróleo en imágenes ópticas, cubriendo el espectro visible e infrarrojo, que pueden ser utilizados para apoyar la respuesta y atención a futuros derrames de petróleo.

Sensores activos

El desarrollo e implementación de modelos de ML y DL para la detección y segmentación de derrames de petróleo con imágenes SAR ha demostrado un gran potencial. Sin embargo, existen diversos aspectos que deben considerarse para incrementar la precisión de los modelos. Uno de estos es considerar que los modelos pueden tender a fugas de conocimiento, sesgándose a problemas particulares, e imposibilitando su generalización, lo cual es clave al tratarse de un problema que es de índole mundial (Kaufman *et al.*, 2012; Fingas and Brown, 2018). Siempre se debe buscar ampliar la cantidad de datos, y que sean datos de distintas partes del mundo con la finalidad de incrementar la variabilidad que puede aprender el modelo.

Al igual que en los sensores pasivos, es fundamental tener en cuenta la degradación constante del petróleo desde el momento de su liberación, ya que esto puede dificultar su caracterización, y que con el tiempo pueda parecerse más a superficies de origen biológico (Solberg, 2012). Es importante considerar esta degradación, ya que los modelos pueden aprender características presentes en los bordes del derrame, donde el petróleo puede variar en concentración, volumen e incluso combinarse para formar emulsiones, además, todas las variaciones que pueda tener van a depender de la fuente y el lugar donde se encuentre el derrame (Gade *et al.*, 1998; Kim, Moon and Kim, 2010; Solberg, 2012).

En este sentido, para incrementar la precisión de los modelos, en un esquema más completo, debe considerarse la ubicación de las fuentes principales, como pozos,

plataformas y embarcaciones, y con ello tener un sistema de monitoreo centrado en dichos sitios, aunque el realizar esto limita un poco la parte de exploración en zonas potenciales por presencia de derrames naturales (Solberg, Brekke and Husoy, 2007; Solberg, 2012; Topouzelis *et al.*, 2015).

Otro aspecto que considerar es la resolución temporal de los satélites, ya que para desarrollar un mecanismo de monitoreo se requiere que los datos se encuentren disponibles de manea frecuente, lo cual es una desventaja ya que en el caso de Sentinel-1 SAR su tiempo de revisita es de cinco días, lo que limita un monitoreo en tiempo real, sin embargo, se aprovecha su disponibilidad en todo el globo.

Por último, una parte a considerar para el desarrollo del modelo es la ventaja y desventaja que pueda tener cada uno de ellos, ya sea de ML o DL, y tratar de realizar búsquedas exhaustivas de parámetros, métricas y profundidad en las arquitecturas, que puedan ser de utilidad para la resolución del problema a tratar, ya que en ocasiones los problemas pueden ser específicos y a la vez complejos que requieren el desarrollo de un modelo o perspectiva única.

Esquema de monitoreo conjunto

Como se ha descrito a lo largo del documento, los derrames de petróleo son un problema grave de índole internacional, por lo que se busca un sistema de monitoreo en tiempo real que genere resultados precisos, que ayude a eficientizar los tiempos de respuesta y que minimice los costos, además de mejorar la toma de decisión.

Sin embargo, un sistema de monitoreo en tiempo real solo sería posible con la existencia de un satélite geoestacionario, es decir, que se encuentre de manera permanente sobre un determinado punto de la superficie terrestre (Chuvieco, 1995), por lo que se debe buscar la manera de generar un sistema integral que utilice información de diversos satélites y sensores para subsanar la resolución temporal.

En este sentido se debe buscar un esquema integral de monitoreo, el cual utilice lo desarrollado para los sensores ópticos y activos, en donde, dicho sistema trabaje de manera paralela para obtener imágenes de ambos satélites y sensores, y con ello lograr reducir los tiempos de revisita, hasta en 3 días. Esto obteniendo información de los datos de los satélites con sensores ópticos Landsat y Sentinel-2, y del satélite con sensores activos Sentinel-1.

Con esto se podría incrementar el tiempo de respuesta, aunado a que, de acuerdo con los tratados internacionales, la disponibilidad de datos espaciales se encontraría disponible en caso de grandes derrames, por lo que tener un sistema para sensores pasivos y activos ayudaría a la toma de decisión y vigilancia.

Conclusiones

Sensores pasivos

Los datos de respuesta espectral (procedentes de técnicas de laboratorio, espectroscopia de campo y datos de satélite) son útiles para detectar superficies en imágenes ópticas. Sin embargo, diversos procesos de meteorización y mezcla, así como las propiedades físicas y químicas, favorecen que algunos datos se parezcan a otras superficies. El uso del análisis estadístico para estudiar los datos espectrales de las superficies proporcionó una gran perspectiva para comprender y explorar su comportamiento. Debido a su naturaleza, los materiales mostraron un comportamiento característico en determinados rangos del espectro electromagnético, lo que ayudó a discriminarlos. Los modelos ML generados, para todas las bandas obtuvieron buenas puntuaciones; sin embargo, no pudieron discriminar completamente las superficies en imágenes de satélite. Seleccionando las bandas Azul, NIR, SWIR1 y SWIR2, las diferencias entre los valores de entrenamiento y validación disminuyeron, mejorando así los modelos. En este estudio, el modelo con las mejores aproximaciones para la detección de petróleo, cuantitativa y cualitativamente, fue el KNN para las bandas Azul, NIR, SWIR1 y SWIR2. Este modelo mostró una buena precisión en la detección de aceite, ya que sólo identificó aceite en imágenes con presencia de aceite, es decir, verdaderos positivos, y en imágenes con ausencia de aceite el modelo no lo identificó. En el caso de la estimación de la concentración de petróleo, utilizando los datos de respuesta espectral de petróleo con diferentes concentraciones es posible generar

estimaciones de la concentración de petróleo en imágenes de satélite. Existe un cierto grado de incertidumbre en las estimaciones debido a la existencia de muchos fenómenos físicos y químicos que alteran la composición del petróleo en el entorno real. Para los dos modelos de aprendizaje supervisado generados, detección de petróleo (clasificación de superficies) y estimación de la concentración (regresión), es conveniente aumentar el número de datos de ejemplo para obtener mejores modelos, y lograr generalizarlos a otros derrames en el mundo. El uso de datos de respuesta espectral, la selección de variables importantes y técnicas de ML, y los resultados de nuestra metodología aumentan la aplicabilidad de las imágenes ópticas de satélite en la detección de vertidos de petróleo.

Sensores activos

Es posible crear un modelo sencillo, como el modelo logístico, a partir de los datos extraídos de los píxeles, teniendo resultados similares a un modelo MLP con mayor complejidad. El modelo entrenado únicamente con los valores de los píxeles se confunde cuando se presentan imágenes con *look-alikes*, como sargazo, zonas de bajo o nulo viento, o celdas microconvectivas. El aplanar una imagen para extraer los valores y realizar segmentaciones, ocasiona que se pierda mucha información contextual, la cual es importante en la detección y segmentación de superficies. Es necesario considerar el contexto y las formas dentro de las imágenes para lograr diferenciar a al aceite de otros materiales. Se necesita tener un primer filtro para eliminar la posibilidad de falsos positivos, por lo que un modelo de clasificación de imágenes, CNN, fue de utilidad para este fin. Considerar un modelo de segmentación que utilice la información contextual de las imágenes, incrementa la

certidumbre del modelo. En principio se tiene un problema de clases desbalanceadas, por lo que las métricas y funciones de pérdida convencionales, como *exactitud* y entropía cruzada, no pueden utilizarse como en cualquier otro caso. Al utilizar *Pérdida Focal* como función de pérdida en el entrenamiento del modelo, se aborda el desequilibrio de clases, se penalizan los ejemplos fáciles de clasificar, y se pone énfasis en ejemplos difíciles mal clasificados. Utilizar la métrica *IoU* mejora el entrenamiento del modelo, ya que al igual que la *Pérdida Focal*, se penaliza los ejemplos mal clasificados e incrementa la precisión del modelo. El utilizar una UNET con solo dos entradas, y más profunda, en comparación con la UNET original, incrementa la certidumbre del modelo. El comparar diferentes tipos de UNET nos da un panorama de las ventajas y desventajas de cada una, con ello saber cuál es la que mejores resultados da para resolver el problema. Para nuestro caso el mejor modelo de segmentación fue la UNET “original”.

Perspectivas del proyecto

Durante este proyecto doctoral, se han abordado diversos aspectos para crear una herramienta de detección y segmentación de derrames de petróleo, utilizando datos satelitales y técnicas de inteligencia artificial. Sin embargo, se han identificado varias áreas de oportunidad y futuras líneas de investigación que podrían ampliar y mejorar este enfoque.

En el caso de los sensores pasivos, se puede explorar el uso de modelos de aprendizaje profundo para la detección y segmentación de derrames de petróleo. Estos modelos podrían analizar directamente los datos de imágenes provenientes de sensores como Landsat y Sentinel-2, e incluso considerar el uso de sensores con menor resolución espacial pero mayor resolución temporal, como MODIS. Además, se podría investigar la estimación de la concentración del derrame a partir de estas imágenes, lo que permitiría obtener información más detallada sobre la extensión y gravedad del derrame.

En cuanto a los sensores activos, se pueden explorar enfoques que proporcionen una aproximación de la concentración y degradación del aceite. Por ejemplo, se podría utilizar algún método no supervisado para generar máscaras de las áreas afectadas y utilizar algoritmos de *machine learning* o *deep learning* para estimar la concentración a partir de estas máscaras. Esto permitiría obtener información sobre la distribución y evolución del derrame a lo largo del tiempo.

Otra perspectiva interesante sería la fusión de los enfoques de detección y segmentación de los sensores pasivos y activos. Esto permitiría combinar la información de ambos tipos de sensores, y utilizarla de manera conjunta para mejorar la precisión y la capacidad de monitoreo de los derrames. Esta integración, podría llevar a la creación de una herramienta de monitoreo en tiempo casi real, generando predicciones actualizadas en pocos días. Esto implicaría una conexión constante con los servidores de descarga de imágenes satelitales y un procesamiento eficiente de los datos para proporcionar alertas tempranas sobre la presencia de derrames.

Además, se ha evaluado la importancia de desarrollar herramientas para pronosticar la dispersión y movilidad del contaminante. Esto permitiría generar escenarios de los posibles movimientos y áreas de mayor riesgo de dispersión del petróleo, lo que ayudaría a enfocar los esfuerzos de respuesta, así como minimizar el impacto ambiental y socioeconómico de los derrames. Integrar esta herramienta de pronóstico en el sistema de monitoreo completo, que incluya detección, segmentación, estimación de área y concentración, y pronóstico de dispersión, sería de gran utilidad para la toma de decisiones informadas y la implementación de medidas preventivas y de mitigación. Además, esta aproximación sería de bajo costo, ya que todos los activos a utilizar serían de acceso libre.

Es importante destacar que el conocimiento adquirido durante este proyecto no se limita únicamente a la detección y segmentación de derrames de petróleo, sino que puede extrapolarse a otras situaciones y áreas de monitoreo ambiental, como el seguimiento del sargazo en las costas del Caribe mexicano o el monitoreo y

cambios en la vegetación. Incluso llevarlo a diversas áreas de conocimiento en ciencias de la Tierra y del medio ambiente.

Todo esto, nos abre nuevas posibilidades de aplicación y desarrollo en el campo de la teledetección y la inteligencia artificial. Estas tecnologías nos brindan herramientas poderosas para la toma de decisiones informadas y la implementación de estrategias de manejo ambiental sostenible.

Referencias

(ESA), E. S. A. (2018) ‘SNAP - ESA Sentinel Application Platform’. Available at: <http://step.esa.int>.

Abou Samra, R. M. and Ali, R. R. (2022) ‘Monitoring of oil spill in the offshore zone of the Nile Delta using Sentinel data’, *Marine Pollution Bulletin*, 179, p. 113718. doi: <https://doi.org/10.1016/j.marpolbul.2022.113718>.

Adamo, M. *et al.* (2009) ‘Detection and tracking of oil slicks on sun-glittered visible and near infrared satellite imagery’, *International Journal of Remote Sensing*. Taylor & Francis, 30(24), pp. 6403–6427. doi: 10.1080/01431160902865772.

Al-Ruzouq, R. *et al.* (2020) ‘Sensors, Features, and Machine Learning for Oil Spill Detection and Monitoring: A Review’, *Remote Sensing* . doi: 10.3390/rs12203338.

Albawi, S., Mohammed, T. A. and Al-Zawi, S. (2017) ‘Understanding of a convolutional neural network’, in *2017 International Conference on Engineering and Technology (ICET)*, pp. 1–6. doi: 10.1109/ICEngTechnol.2017.8308186.

Alpers, W., Liu, A. K. and Wu, S. Y. (2019) ‘Satellite remote sensing SAR’, in *Encyclopedia of Ocean Sciences*. Academic Press, pp. 429–442. doi: 10.1016/B978-0-12-409548-9.11615-X.

Alzubaidi, L. *et al.* (2021) ‘Review of deep learning: concepts, CNN architectures, challenges, applications, future directions’, *Journal of Big Data*, 8(1), p. 53. doi: 10.1186/s40537-021-00444-8.

Arslan, N. (2018) ‘Assessment of oil spills using Sentinel 1 C-band SAR and Landsat 8 multispectral sensors’, *Environmental Monitoring and Assessment*, 190(11), p. 637. doi: 10.1007/s10661-018-7017-4.

Arzt, M. *et al.* (2022) ‘LABKIT: Labeling and Segmentation Toolkit for Big Image Data’, *Frontiers in Computer Science*, 4. doi: 10.3389/fcomp.2022.777728.

Azzari, G. and Lobell, D. B. (2017) 'Landsat-based classification in the cloud: An opportunity for a paradigm shift in land cover monitoring', *Remote Sensing of Environment*, 202, pp. 64–74. doi: <https://doi.org/10.1016/j.rse.2017.05.025>.

Babagolimatikolaie, J. (2022) 'Monitoring of oil slicks in the Persian Gulf using Sentinel 1 images', *Journal of Ocean Engineering and Science*. doi: <https://doi.org/10.1016/j.joes.2022.05.029>.

Baldrige, A. M. *et al.* (2009) 'The ASTER spectral library version 2.0', *Remote Sensing of Environment*, 113(4), pp. 711–715. doi: <https://doi.org/10.1016/j.rse.2008.11.007>.

Balogun, A.-L. *et al.* (2020) 'Spatio-Temporal Analysis of Oil Spill Impact and Recovery Pattern of Coastal Vegetation and Wetland Using Multispectral Satellite Landsat 8-OLI Imagery and Machine Learning Models', *Remote Sensing*. doi: 10.3390/rs12071225.

Bar, S., Parida, B. R. and Pandey, A. C. (2020) 'Landsat-8 and Sentinel-2 based Forest fire burn area mapping using machine learning algorithms on GEE cloud platform over Uttarakhand, Western Himalaya', *Remote Sensing Applications: Society and Environment*, 18, p. 100324. doi: <https://doi.org/10.1016/j.rsase.2020.100324>.

Basit, A. *et al.* (2022) 'Comparison of CNNs and Vision Transformers-Based Hybrid Models Using Gradient Profile Loss for Classification of Oil Spills in SAR Images', *Remote Sensing*, 14(9). doi: 10.3390/rs14092085.

Bessis, J.-L. (2003) 'International Charter "Space and Major Disasters" Evolution', in. doi: 10.2514/6.iac-03-c.2.01.

Beyer, J. *et al.* (2016) 'Environmental effects of the Deepwater Horizon oil spill: A review', *Marine Pollution Bulletin*, 110(1), pp. 28–51. doi: <https://doi.org/10.1016/j.marpolbul.2016.06.027>.

Bonn Agreement Accord de Bonn, (BAOAC) (2007) 'Bonn Agreement Oil Appearance Code'. BAOAC. Available at: <https://www.bonnagreement.org/publications>.

Brekke, C. and Solberg, A. H. S. (2005) 'Oil spill detection by satellite remote sensing',

Remote Sensing of Environment, 95(1), pp. 1–13. doi: <https://doi.org/10.1016/j.rse.2004.11.015>.

Briggs, I. L. and Briggs, B. C. (2018) ‘Petroleum Industry Activities and Human Health’, in *The Political Ecology of Oil and Gas Activities in the Nigerian Aquatic Ecosystem*. Academic Press, pp. 143–147. doi: 10.1016/B978-0-12-809399-3.00010-0.

Bulgarelli, B. and Djavidnia, S. (2012) ‘On MODIS Retrieval of Oil Spill Spectral Properties in the Marine Environment’, *IEEE Geoscience and Remote Sensing Letters*, 9(3), pp. 398–402. doi: 10.1109/LGRS.2011.2169647.

Burgherr, P. (2007) ‘In-depth analysis of accidental oil spills from tankers in the context of global spill trends from all sources’, *Journal of Hazardous Materials*, 140(1), pp. 245–256. doi: <https://doi.org/10.1016/j.jhazmat.2006.07.030>.

Cantorna, D. *et al.* (2019) ‘Oil spill segmentation in SAR images using convolutional neural networks. A comparative analysis with clustering and logistic regression algorithms’, *Applied Soft Computing*, 84, p. 105716. doi: <https://doi.org/10.1016/j.asoc.2019.105716>.

Cao, Z. *et al.* (2020) ‘A machine learning approach to estimate chlorophyll-a from Landsat-8 measurements in inland lakes’, *Remote Sensing of Environment*, 248, p. 111974. doi: <https://doi.org/10.1016/j.rse.2020.111974>.

Chander, G. *et al.* (2009) ‘Remote Sensing of Environment Summary of current radiometric calibration coefficients for Landsat MSS, TM, ETM+, and EO-1 ALI sensors’, *Remote Sensing of Environment*. Elsevier Inc., 113(5), pp. 893–903. doi: 10.1016/j.rse.2009.01.007.

Chaturvedi, S. K., Banerjee, S. and Lele, S. (2020) ‘An assessment of oil spill detection using Sentinel 1 SAR-C images’, *Journal of Ocean Engineering and Science*, 5(2), pp. 116–135. doi: <https://doi.org/10.1016/j.joes.2019.09.004>.

Chen, Y. *et al.* (2022) ‘A novel lightweight bilateral segmentation network for detecting oil spills on the sea surface’, *Marine Pollution Bulletin*, 175, p. 113343. doi: <https://doi.org/10.1016/j.marpolbul.2022.113343>.

Chollet, F. and others (2015) 'Keras'. GitHub. Available at: <https://keras.io>.

Chowdhury, S., Evans, C. and Shipman, T. C. (2021) 'The Role of Landsat-8 Multispectral Data in Spill Response: Three Case Studies BT - Advances in Remote Sensing for Infrastructure Monitoring', in Singhroy, V. (ed.). Cham: Springer International Publishing, pp. 291–305. doi: 10.1007/978-3-030-59109-0_13.

Chuvieco, E. (1995) *Fundamentos De Teledeteccion Espacial*. Segunda. Madrid, España: Ediciones RIALP S.A.

Clark, R. N. *et al.* (2010) *A method for quantitative mapping of thick oil spills using imaging spectroscopy, Open-File Report*. doi: 10.3133/ofr20101167.

Cococcioni, M. *et al.* (2012) 'SVME: an ensemble of support vector machines for detecting oil spills from full resolution MODIS images', *Ocean Dynamics*, 62(3), pp. 449–467. doi: 10.1007/s10236-011-0510-8.

CONAE, (Comision Nacional de Actividades Espaciales) (no date) 'Biblioteca de Firmas Espectrales de CONAE'. Argentina: Comisión Nacional de Actividades Espaciales (CONAE). Available at: <https://catalogos5.conae.gov.ar/firmasespectrales/documentacion.aspx>.

Daling, P. S. *et al.* (1990) 'Characterization of crude oils for environmental purposes', *Oil and Chemical Pollution*, 7(3), pp. 199–224. doi: [https://doi.org/10.1016/S0269-8579\(05\)80027-9](https://doi.org/10.1016/S0269-8579(05)80027-9).

Daling, P. S. *et al.* (2014) 'Surface weathering and dispersibility of MC252 crude oil', *Marine Pollution Bulletin*, 87(1), pp. 300–310. doi: <https://doi.org/10.1016/j.marpolbul.2014.07.005>.

Daling, P. S. and Strøm, T. (1999) 'Weathering of Oils at Sea: Model/Field Data Comparisons', *Spill Science & Technology Bulletin*, 5(1), pp. 63–74. doi: [https://doi.org/10.1016/S1353-2561\(98\)00051-6](https://doi.org/10.1016/S1353-2561(98)00051-6).

Daly, K. L. *et al.* (2016) 'Assessing the impacts of oil-associated marine snow formation and sedimentation during and after the Deepwater Horizon oil spill', *Anthropocene*, 13, pp. 18–

33. doi: <https://doi.org/10.1016/j.ancene.2016.01.006>.

Dasari, K., Anjaneyulu, L. and Nadimikeri, J. (2022) ‘Application of C-band sentinel-1A SAR data as proxies for detecting oil spills of Chennai, East Coast of India’, *Marine Pollution Bulletin*, 174, p. 113182. doi: <https://doi.org/10.1016/j.marpolbul.2021.113182>.

Dehghani-Dehcheshmeh, S., Akhoondzadeh, M. and Homayouni, S. (2023) ‘Oil spills detection from SAR Earth observations based on a hybrid CNN transformer networks’, *Marine Pollution Bulletin*, 190, p. 114834. doi: <https://doi.org/10.1016/j.marpolbul.2023.114834>.

Distribution, A. S. (2020) ‘Anaconda Software Distribution’, *Anaconda Documentation*. Anaconda Inc. Available at: <https://docs.anaconda.com/>.

Elkan, C. (2012) ‘Evaluating Classifiers’, *University of San Diego, California*, retrieved [01-11-2012] from <http://cseweb.ucsd.edu/~elkan> B. Citeseer, 250, pp. 1–11. Available at: <http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.233.2746&rep=rep1&type=pdf%5Cnhttp://cseweb.ucsd.edu/~elkan/250Bwinter2012/classifiereval.pdf>.

Everitt, B. (1992) ‘Book reviews : Chambers JM, Hastie TJ eds 1992: Statistical models in S. California: Wadsworth and Brooks/Cole. ISBN 0 534 16765-9’, *Statistical Methods in Medical Research*, 1(2), pp. 220–221. doi: 10.1177/096228029200100208.

Fan, Y. *et al.* (2021) ‘Feature Merged Network for Oil Spill Detection Using SAR Images’, *Remote Sensing*, 13(16). doi: 10.3390/rs13163174.

Feinauer, D. M. *et al.* (2022) ‘Oil Spill Identification using Deep Convolutional Neural Networks’, in 2022 14th International Conference on Computational Intelligence and Communication Networks (CICN), pp. 240–245. doi: 10.1109/CICN56167.2022.10008373.

Filipponi, F. (2019) ‘Sentinel-1 GRD Preprocessing Workflow’, *Proceedings*, 18(1). doi: 10.3390/ECRS-3-06201.

Fingas, M. (2011) ‘Introduction’, in Fingas, M. (ed.) *Oil spill science and technology*. Oxford, UK: Elsevier Inc.

Fingas, M. and Brown, C. E. (2017) 'Chapter 5 - Oil Spill Remote Sensing', in Fingas, M. B. T.-O. S. S. and T. (Second E. (ed.). Boston: Gulf Professional Publishing, pp. 305–385. doi: <https://doi.org/10.1016/B978-0-12-809413-6.00005-9>.

Fingas, M. and Brown, C. E. (2018) 'A Review of Oil Spill Remote Sensing', *Sensors* . doi: 10.3390/s18010091.

Fingas, M. and Brown E., C. (2015) 'Oil spill remote sensing', in Fingas, M. (ed.) *Handbook of oil spill science and technology*. First. John Wiley & Sons, Inc.

Fiscella, B. *et al.* (2000) 'Oil spill detection using marine SAR images', *International Journal of Remote Sensing*. Taylor & Francis, 21(18), pp. 3561–3566. doi: 10.1080/014311600750037589.

Fisher, R. A. (1992) 'Statistical methods for research workers', in *Breakthroughs in statistics*. Springer, pp. 66–70.

Del Frate, F. *et al.* (2000) 'Neural networks for oil spill detection using ERS-SAR data', *IEEE Transactions on Geoscience and Remote Sensing*, 38(5), pp. 2282–2287. doi: 10.1109/36.868885.

Freedman, D., Pisani, R. and Purves, R. (2007) 'Statistics (international student edition)', *Pisani, R. Purves, 4th edn. WW Norton & Company, New York*.

Gade, M. *et al.* (1998) 'Imaging of biogenic and anthropogenic ocean surface films by the multifrequency/multipolarization SIR-C/X-SAR', *Journal of Geophysical Research: Oceans*, 103(C9), pp. 18851–18866. doi: <https://doi.org/10.1029/97JC01915>.

Gao, S., Li, S. and Liu, H. (2022) 'Oil Spill Detection by CP SAR Based on the Power Entropy Decomposition', *Remote Sensing*, 14(19). doi: 10.3390/rs14195030.

Garain, S., Mitra, D. and Das, P. (2021) 'Detection of hydrocarbon microseepage prospects using Landsat 8-based vegetation stress analysis in part of Assam-Arakan Fold Belt, NE India', *Arabian Journal of Geosciences*, 14(19), p. 1984. doi: 10.1007/s12517-021-08376-6.

Garcia-Pineda, Oscar *et al.* (2013) ‘Detection of Floating Oil Anomalies From the Deepwater Horizon Oil Spill With Synthetic Aperture Radar’, *Oceanography*, 26(2). doi: 10.5670/oceanog.2013.38.

Garcia-Pineda, O *et al.* (2013) ‘Oil Spill Mapping and Measurement in the Gulf of Mexico With Textural Classifier Neural Network Algorithm (TCNNA)’, *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, 6(6), pp. 2517–2525. doi: 10.1109/JSTARS.2013.2244061.

Garcia-Pineda, O. *et al.* (2017) ‘Detection of Oil near Shorelines during the Deepwater Horizon Oil Spill Using Synthetic Aperture Radar (SAR)’, *Remote Sensing* . doi: 10.3390/rs9060567.

Gauthier, M.-F. *et al.* (2007) ‘Integrated satellite tracking of pollution: A new operational program’, in *2007 IEEE International Geoscience and Remote Sensing Symposium*, pp. 967–970. doi: 10.1109/IGARSS.2007.4422960.

GDAL/OGR contributors (2021) ‘{GDAL/OGR} Geospatial Data Abstraction software Library’. doi: 10.5281/zenodo.5884351.

Ghorbanian, A. *et al.* (2020) ‘Improved land cover map of Iran using Sentinel imagery within Google Earth Engine and a novel automatic workflow for land cover classification using migrated training samples’, *ISPRS Journal of Photogrammetry and Remote Sensing*, 167, pp. 276–288. doi: <https://doi.org/10.1016/j.isprsjprs.2020.07.013>.

Goutte, C. and Gaussier, E. (2005) ‘A Probabilistic Interpretation of Precision, Recall and F-Score, with Implication for Evaluation’, in Losada, D. E. and Fernández-Luna, J. M. (eds) *Advances in Information Retrieval*. Berlin, Heidelberg: Springer Berlin Heidelberg, pp. 345–359.

Guo, Y. and Zhang, H. Z. (2014) ‘Oil spill detection using synthetic aperture radar images and feature selection in shape space’, *International Journal of Applied Earth Observation and Geoinformation*, 30, pp. 146–157. doi: <https://doi.org/10.1016/j.jag.2014.01.011>.

Harris, C. R. *et al.* (2020) ‘Array programming with {NumPy}’, *Nature*. Springer Science

and Business Media {LLC}, 585(7825), pp. 357–362. doi: 10.1038/s41586-020-2649-2.

Hasimoto-Beltran, R. *et al.* (2023) ‘Ocean oil spill detection from SAR images based on multi-channel deep learning semantic segmentation’, *Marine Pollution Bulletin*, 188, p. 114651. doi: <https://doi.org/10.1016/j.marpolbul.2023.114651>.

Hooper, C. H. (1982) ‘The IXTOC I oil spill : the Federal scientific response’. Edited by N. O. S. United States Office of Marine Pollution Assessment, and N. H. M. R. P. (U.S.). (NOAA special report). Available at: <https://repository.library.noaa.gov/view/noaa/13576>.

Hörig, B. *et al.* (2001) ‘HyMap hyperspectral remote sensing to detect hydrocarbons’, *International Journal of Remote Sensing*. Taylor & Francis, 22(8), pp. 1413–1422. doi: 10.1080/01431160120909.

Hu, C. *et al.* (2009) ‘Detection of natural oil slicks in the NW Gulf of Mexico using MODIS imagery’, *Geophysical Research Letters*. John Wiley & Sons, Ltd, 36(1). doi: <https://doi.org/10.1029/2008GL036119>.

Hu, C. *et al.* (2018) ‘Remote sensing estimation of surface oil volume during the 2010 Deepwater Horizon oil blowout in the Gulf of Mexico: scaling up AVIRIS observations with MODIS measurements’, *Journal of Applied Remote Sensing*, 12(02), p. 1. doi: 10.1117/1.JRS.12.026008.

Huang, X. *et al.* (2022) ‘A novel deep learning method for marine oil spill detection from satellite synthetic aperture radar imagery’, *Marine Pollution Bulletin*, 179, p. 113666. doi: <https://doi.org/10.1016/j.marpolbul.2022.113666>.

Hunter, J. D. (2007) ‘Matplotlib: A 2D graphics environment’, *Computing in Science & Engineering*. IEEE COMPUTER SOC, 9(3), pp. 90–95. doi: 10.1109/MCSE.2007.55.

Huz, R., Lastra, M. and López, J. (2019) ‘Other environmental health issues: Oil spill’, in *Encyclopedia of Environmental Health*. Elsevier, pp. 792–796. doi: 10.1016/B978-0-12-409548-9.11156-X.

Indolia, S. *et al.* (2018) ‘Conceptual Understanding of Convolutional Neural Network- A

Deep Learning Approach’, *Procedia Computer Science*, 132, pp. 679–688. doi: <https://doi.org/10.1016/j.procs.2018.05.069>.

Ivshina, I. B. *et al.* (2015) ‘Oil spill problems and sustainable response strategies through new technologies’, *Environmental Science: Processes & Impacts*. The Royal Society of Chemistry, 17(7), pp. 1201–1219. doi: 10.1039/C5EM00070J.

Jones, C. E. (2023) ‘An automated algorithm for calculating the ocean contrast in support of oil spill response’, *Marine Pollution Bulletin*, 191, p. 114952. doi: <https://doi.org/10.1016/j.marpolbul.2023.114952>.

Kaufman, S. *et al.* (2012) ‘Leakage in data mining’, *ACM Transactions on Knowledge Discovery from Data*, 6(4), pp. 1–21. doi: 10.1145/2382577.2382579.

De Kerf, T. *et al.* (2020) ‘Oil Spill Detection Using Machine Learning and Infrared Images’, *Remote Sensing* . doi: 10.3390/rs12244090.

Kim, D., Moon, W. M. and Kim, Y.-S. (2010) ‘Application of TerraSAR-X Data for Emergent Oil-Spill Monitoring’, *IEEE Transactions on Geoscience and Remote Sensing*, 48(2), pp. 852–863. doi: 10.1109/TGRS.2009.2036253.

Kingma, D. P. and Ba, J. (2014) ‘Adam: A Method for Stochastic Optimization’, *3rd International Conference on Learning Representations, ICLR 2015 - Conference Track Proceedings*. doi: <https://doi.org/10.48550/arXiv.1412.6980>.

Kokaly, R. F. *et al.* (2017) *USGS Spectral Library Version 7, Data Series*. Reston, VA. doi: 10.3133/ds1035.

Kolokoussis, P. and Karathanassi, V. (2018) ‘Oil Spill Detection and Mapping Using Sentinel 2 Imagery’, *Journal of Marine Science and Engineering* . doi: 10.3390/jmse6010004.

Kourdi, E., Salem, M. and Qiblawi, T. (2021) ‘Syrian oil spill spreads across the Mediterranean and could reach Cyprus on Wednesday’. CNN. Available at: <https://edition.cnn.com/2021/08/31/middleeast/syria-cyprus-oil-spill-intl/index.html>.

Krestenitis, M. *et al.* (2019) ‘Oil Spill Identification from Satellite Images Using Deep Neural Networks’, *Remote Sensing*, 11(15). doi: 10.3390/rs11151762.

Krijthe, J. H. (2015) ‘{Rtsne}: T-Distributed Stochastic Neighbor Embedding using Barnes-Hut Implementation’. Available at: <https://github.com/jkrijthe/Rtsne>.

Kubat, M., Holte, R. C. and Matwin, S. (1998) ‘Machine Learning for the Detection of Oil Spills in Satellite Radar Images’, *Machine Learning*, 30(2), pp. 195–215. doi: 10.1023/A:1007452223027.

de la Huz, R., Lastra, M. and López, J. (2011) ‘Oil Spills’, in *Encyclopedia of Environmental Health*. Elsevier, pp. 251–255. doi: 10.1016/B978-0-444-52272-6.00568-7.

Lary, D. J. *et al.* (2018) ‘Machine Learning Applications for Earth Observation BT - Earth Observation Open Science and Innovation’, in Mathieu, P.-P. and Aubrecht, C. (eds). Cham: Springer International Publishing, pp. 165–218. doi: 10.1007/978-3-319-65633-5_8.

LeCun, Y., Bengio, Y. and Hinton, G. (2015) ‘Deep learning’, *Nature*, 521(7553), pp. 436–444. doi: 10.1038/nature14539.

Lee, J. S. H. *et al.* (2016) ‘Detecting industrial oil palm plantations on Landsat images with Google Earth Engine’, *Remote Sensing Applications: Society and Environment*, 4, pp. 219–224. doi: <https://doi.org/10.1016/j.rsase.2016.11.003>.

Leifer, I. *et al.* (2012) ‘State of the art satellite and airborne marine oil spill remote sensing: Application to the BP Deepwater Horizon oil spill’, *Remote Sensing of Environment*, 124, pp. 185–209. doi: <https://doi.org/10.1016/j.rse.2012.03.024>.

Li, P. *et al.* (2016) ‘Offshore oil spill response practices and emerging challenges’, *Marine Pollution Bulletin*, 110(1), pp. 6–27. doi: <https://doi.org/10.1016/j.marpolbul.2016.06.020>.

Li, Y. *et al.* (2017) ‘Detection and Monitoring of Oil Spills Using Moderate/High-Resolution Remote Sensing Images’, *Archives of Environmental Contamination and Toxicology*, 73(1), pp. 154–169. doi: 10.1007/s00244-016-0358-5.

Li, Y. *et al.* (2018) ‘Detection of Oil Spill Through Fully Convolutional Network BT - Geo-Spatial Knowledge and Intelligence’, in Yuan, H. *et al.* (eds). Singapore: Springer Singapore, pp. 353–362.

Li, Y. *et al.* (2021) ‘Oil Spill Detection with Multiscale Conditional Adversarial Networks with Small-Data Training’, *Remote Sensing*, 13(12). doi: 10.3390/rs13122378.

Lin, T.-Y. *et al.* (2017) ‘Focal Loss for Dense Object Detection’, *CoRR*, abs/1708.0. Available at: <http://arxiv.org/abs/1708.02002>.

Lin, Y. *et al.* (2021) ‘An optimized machine learning approach to water pollution variation monitoring with time-series Landsat images’, *International Journal of Applied Earth Observation and Geoinformation*, 102, p. 102370. doi: <https://doi.org/10.1016/j.jag.2021.102370>.

Liu, B. *et al.* (2016) ‘Extraction of Oil Spill Information Using Decision Tree Based Minimum Noise Fraction Transform’, *Journal of the Indian Society of Remote Sensing*, 44(3), pp. 421–426. doi: 10.1007/s12524-015-0499-4.

Lu, Y. *et al.* (2019) ‘Optical interpretation of oil emulsions in the ocean – Part I: Laboratory measurements and proof-of-concept with AVIRIS observations’, *Remote Sensing of Environment*, 230, p. 111183. doi: <https://doi.org/10.1016/j.rse.2019.05.002>.

Lu, Y. *et al.* (2020) ‘Optical interpretation of oil emulsions in the ocean – Part II: Applications to multi-band coarse-resolution imagery’, *Remote Sensing of Environment*, 242, p. 111778. doi: <https://doi.org/10.1016/j.rse.2020.111778>.

Luciani 1, R. and LANEVE 1, G. (2018) ‘Oil Spill Detection Using Optical Sensors: A Multi-Temporal Approach’, *Satellite Oceanography and Meteorology; Pre-published article* DO - 10.18063/som.v0i0.816 . Available at: <http://ojs.whioce.com/index.php/som/article/view/816>.

Ma, X. *et al.* (2023) ‘Detection of marine oil spills from radar satellite images for the coastal ecological risk assessment’, *Journal of Environmental Management*, 325, p. 116637. doi: <https://doi.org/10.1016/j.jenvman.2022.116637>.

Ma, Y. *et al.* (2014) ‘Feature selection and classification of oil spills in SAR image based on statistics and artificial neural network’, in *2014 IEEE Geoscience and Remote Sensing Symposium*, pp. 569–571. doi: 10.1109/IGARSS.2014.6946486.

Maechler, M. *et al.* (2015) ‘Package “cluster”: Cluster Analysis Basics and Extensions’, *R topics Documented*, p. 79. Available at: <https://cran.r-project.org/web/packages/cluster/cluster.pdf>.

Martín Abadi *et al.* (2022) ‘TensorFlow: Large-scale machine learning on heterogeneous systems’. Available at: <https://zenodo.org/record/6574269>.

Mathias Davidson, Michael Askne, L. M. H. U. (2012) *Sentinel-1: ESA’s radar observatory mission for GMES operational services*, *ESA Special Publication*. Available at: https://sentinel.esa.int/documents/247904/349449/S1_SP-1322_1.pdf.

McKinney, W. (2010) ‘Data Structures for Statistical Computing in Python’, in van der Walt, S. and Millman, J. (eds) *Proceedings of the 9th Python in Science Conference*, pp. 56–61. doi: 10.25080/Majora-92bf1922-00a.

Meerdink, S. K. *et al.* (2019) ‘The ECOSTRESS spectral library version 1.0’, *Remote Sensing of Environment*, 230, p. 111196. doi: <https://doi.org/10.1016/j.rse.2019.05.015>.

Mega, E. R. (2022) ‘Unprecedented oil spill catches researchers in Peru off guard’, *Nature*. doi: 10.1038/d41586-022-00333-x.

Mehlig, B. (2019) ‘Artificial Neural Networks’, *CoRR*, abs/1901.0. Available at: <http://arxiv.org/abs/1901.05639>.

de Mendiburu, F. (2020) ‘agricolae: Statistical Procedures for Agricultural Research’. Available at: <https://cran.r-project.org/package=agricolae>.

Mera, D. *et al.* (2017) ‘On the use of feature selection to improve the detection of sea oil spills in SAR images’, *Computers & Geosciences*, 100, pp. 166–178. doi: <https://doi.org/10.1016/j.cageo.2016.12.013>.

Mitchell, T. (1997) *Machine Learning, Machine Learning*. McGraw-Hill Science/Engineering/Math. doi: 10.1007/BF00116892.

Mo, Y. *et al.* (2022) ‘Review the state-of-the-art technologies of semantic segmentation based on deep learning’, *Neurocomputing*, 493, pp. 626–646. doi: <https://doi.org/10.1016/j.neucom.2022.01.005>.

Mohammadi, M. *et al.* (2021) ‘Detection of Oil Pollution Using SAR and Optical Remote Sensing Imagery: A Case Study of the Persian Gulf’, *Journal of the Indian Society of Remote Sensing*. doi: 10.1007/s12524-021-01399-2.

Mohammadiun, S. *et al.* (2021) ‘Intelligent computational techniques in marine oil spill management: A critical review’, *Journal of Hazardous Materials*, 419, p. 126425. doi: <https://doi.org/10.1016/j.jhazmat.2021.126425>.

de Moura, N. V. A. *et al.* (2022) ‘Deep-water oil-spill monitoring and recurrence analysis in the Brazilian territory using Sentinel-1 time series and deep learning’, *International Journal of Applied Earth Observation and Geoinformation*, 107, p. 102695. doi: <https://doi.org/10.1016/j.jag.2022.102695>.

Ng, H.-W. *et al.* (2015) ‘Deep Learning for Emotion Recognition on Small Datasets Using Transfer Learning’, in *Proceedings of the 2015 ACM on International Conference on Multimodal Interaction*. New York, NY, USA: Association for Computing Machinery (ICMI ’15), pp. 443–449. doi: 10.1145/2818346.2830593.

Nieto-Hidalgo, M. *et al.* (2018) ‘Two-Stage Convolutional Neural Network for Ship and Spill Detection Using SLAR Images’, *IEEE Transactions on Geoscience and Remote Sensing*, 56(9), pp. 5217–5230. doi: 10.1109/TGRS.2018.2812619.

Niwattanakul, S. *et al.* (2013) ‘Using of Jaccard coefficient for keywords similarity’, in *Proceedings of the international multiconference of engineers and computer scientists*, pp. 380–384.

O’Shea, K. and Nash, R. (2015) ‘An Introduction to Convolutional Neural Networks’, *CoRR*, abs/1511.0. Available at: <http://arxiv.org/abs/1511.08458>.

Oktaý, O. *et al.* (2018) ‘Attention U-Net: Learning Where to Look for the Pancreas’, *CoRR*, abs/1804.0. Available at: <http://arxiv.org/abs/1804.03999>.

Ozigis, M. S., Kaduk, J. D. and Jarvis, C. H. (2019) ‘Mapping terrestrial oil spill impact using machine learning random forest and Landsat 8 OLI imagery: a case site within the Niger Delta region of Nigeria’, *Environmental Science and Pollution Research*, 26(4), pp. 3621–3635. doi: 10.1007/s11356-018-3824-y.

De Padova, D. *et al.* (2017) ‘Synergistic use of an oil drift model and remote sensing observations for oil spill monitoring’, *Environmental Science and Pollution Research*, 24(6), pp. 5530–5543. doi: 10.1007/s11356-016-8214-8.

Pearson, K. (1901) ‘LIII. On lines and planes of closest fit to systems of points in space’, *The London, Edinburgh, and Dublin philosophical magazine and journal of science*. Taylor & Francis, 2(11), pp. 559–572.

Pedregosa, F. *et al.* (2011) ‘Scikit-learn: Machine Learning in Python’, *Journal of Machine Learning Research*, 12, pp. 2825–2830.

Polychronis, K. and Vassilia, K. (2013) ‘Detection of Oil Spills and Underwater Natural Oil Outflow Using Multispectral Satellite Imagery’, *International Journal of Remote Sensing Applications*, 3(3).

QGIS Development Team, O. (2018) ‘QGIS Geographic Information System’. Available at: <http://qgis.osgeo.org>.

R Core Team (2021) ‘R: A Language and Environment for Statistical Computing’. Vienna, Austria. Available at: <https://www.r-project.org/>.

Rahman, M. A. and Wang, Y. (2016) ‘Optimizing Intersection-Over-Union in Deep Neural Networks for Image Segmentation’, in Bebis, G. *et al.* (eds) *Advances in Visual Computing*. Cham: Springer International Publishing, pp. 234–244.

Ronneberger, O., Fischer, P. and Brox, T. (2015) ‘U-Net: Convolutional Networks for Biomedical Image Segmentation’, in Navab, N. *et al.* (eds) *Medical Image Computing and*

Computer-Assisted Intervention -- MICCAI 2015. Cham: Springer International Publishing, pp. 234–241.

Van Rossum, G. and Drake, F. L. (2009) *Python 3 Reference Manual*. Scotts Valley, CA: CreateSpace.

Rousseeuw, P. J. (1987) ‘Silhouettes: A graphical aid to the interpretation and validation of cluster analysis’, *Journal of Computational and Applied Mathematics*, 20, pp. 53–65. doi: [https://doi.org/10.1016/0377-0427\(87\)90125-7](https://doi.org/10.1016/0377-0427(87)90125-7).

Rousso, R. *et al.* (2022) ‘Automatic Recognition of Oil Spills Using Neural Networks and Classic Image Processing’, *Water*, 14(7). doi: 10.3390/w14071127.

Ruby, U. and Yendapalli, V. (2020) ‘Binary cross entropy with deep learning technique for Image classification’, *International Journal of Advanced Trends in Computer Science and Engineering*, 9. doi: 10.30534/ijatcse/2020/175942020.

Rumelhart, D. E., Hinton, G. E. and Williams, R. J. (1986) ‘Learning representations by back-propagating errors’, *Nature*, 323(6088), pp. 533–536. doi: 10.1038/323533a0.

Schmidt-Etkin, D. (2015) ‘Risk Analysis and Prevention’, in Fingas, M. (ed.) *Handbook of oil spill science and technology*. First. John Wiley & Sons, Inc.

Shaban, M. *et al.* (2021) ‘A Deep-Learning Framework for the Detection of Oil Spills from SAR Data’, *Sensors*, 21(7). doi: 10.3390/s21072351.

Simon, H. A. (1996) *The Sciences of the Artificial*. Third Edit, Cambridge, MA. Third Edit. London, England: MIT Press. doi: 10.1016/S0898-1221(97)82941-0.

Singh, R. K. *et al.* (2021) ‘A machine learning-based classification of LANDSAT images to map land use and land cover of India’, *Remote Sensing Applications: Society and Environment*, 24, p. 100624. doi: <https://doi.org/10.1016/j.rsase.2021.100624>.

Singha, S., Bellerby, T. J. and Trieschmann, O. (2012) ‘Detection and classification of oil spill and look-alike spots from SAR imagery using an Artificial Neural Network’, in *2012*

IEEE International Geoscience and Remote Sensing Symposium, pp. 5630–5633. doi: 10.1109/IGARSS.2012.6352042.

Singha, S., Bellerby, T. J. and Trieschmann, O. (2013) ‘Satellite Oil Spill Detection Using Artificial Neural Networks’, *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, 6(6), pp. 2355–2363. doi: 10.1109/JSTARS.2013.2251864.

Solberg, A. H. S. (2012) ‘Remote Sensing of Ocean Oil-Spill Pollution’, *Proceedings of the IEEE*, 100(10), pp. 2931–2945. doi: 10.1109/JPROC.2012.2196250.

Solberg, A. H. S., Brekke, C. and Husoy, P. O. (2007) ‘Oil Spill Detection in Radarsat and Envisat SAR Images’, *IEEE Transactions on Geoscience and Remote Sensing*, 45(3), pp. 746–755. doi: 10.1109/TGRS.2006.887019.

Srivastava, H. and Singh, T. P. (2010) ‘Assessment and development of algorithms to detection of oil spills using MODIS data’, *Journal of the Indian Society of Remote Sensing*, 38(1), pp. 161–167. doi: 10.1007/s12524-010-0007-9.

Sudha, V. and Vijendran, A. S. (2021) ‘Oil Spill Discrimination of SAR Satellite Images Using Deep Learning Based Semantic Segmentation BT - Computing Science, Communication and Security’, in Chaubey, N., Parikh, S., and Amin, K. (eds). Cham: Springer International Publishing, pp. 127–139.

Sun, S., Hu, C. and Tunnell, J. W. (2015) ‘Surface oil footprint and trajectory of the Ixtoc-I oil spill determined from Landsat/MSS and CZCS observations’, *Marine Pollution Bulletin*, 101(2), pp. 632–641. doi: <https://doi.org/10.1016/j.marpolbul.2015.10.036>.

Svejkovsky, J. *et al.* (2016) ‘Characterization of surface oil thickness distribution patterns observed during the Deepwater Horizon (MC-252) oil spill with aerial and satellite remote sensing’, *Marine Pollution Bulletin*, 110(1), pp. 162–176. doi: <https://doi.org/10.1016/j.marpolbul.2016.06.066>.

Taravat, A. and Del Frate, F. (2012) ‘Development of band ratioing algorithms and neural networks to detection of oil spills using Landsat ETM+ data’, *EURASIP Journal on Advances in Signal Processing*, 2012(1), p. 107. doi: 10.1186/1687-6180-2012-107.

Taye, M. M. (2023) 'Theoretical Understanding of Convolutional Neural Network: Concepts, Architectures, Applications, Future Directions', *Computation*, 11(3). doi: 10.3390/computation11030052.

Teluguntla, P. *et al.* (2018) 'A 30-m landsat-derived cropland extent product of Australia and China using random forest machine learning algorithm on Google Earth Engine cloud computing platform', *ISPRS Journal of Photogrammetry and Remote Sensing*, 144, pp. 325–340. doi: <https://doi.org/10.1016/j.isprsjprs.2018.07.017>.

Temitope Yekeen, S., Balogun, A. and Wan Yusof, K. B. (2020) 'A novel deep learning instance segmentation model for automated marine oil spill detection', *ISPRS Journal of Photogrammetry and Remote Sensing*, 167, pp. 190–200. doi: <https://doi.org/10.1016/j.isprsjprs.2020.07.011>.

Thyng, K. M. (2019) 'Deepwater Horizon Oil could have naturally reached Texas beaches', *Marine Pollution Bulletin*, 149, p. 110527. doi: <https://doi.org/10.1016/j.marpolbul.2019.110527>.

Topouzelis, K. *et al.* (2007) 'Detection and discrimination between oil spills and look-alike phenomena through neural networks', *ISPRS Journal of Photogrammetry and Remote Sensing*, 62(4), pp. 264–270. doi: <https://doi.org/10.1016/j.isprsjprs.2007.05.003>.

Topouzelis, K. *et al.* (2015) 'Detection, tracking, and remote sensing: Satellites and image processing (Spaceborne oil spill detection)', in Fingas, M. (ed.) *Handbook of oil spill science and technology*. First. John Wiley & Sons, Inc.

Topouzelis, K. and Psyllos, A. (2012) 'Oil spill feature selection and classification using decision tree forest on SAR image data', *ISPRS Journal of Photogrammetry and Remote Sensing*, 68, pp. 135–143. doi: <https://doi.org/10.1016/j.isprsjprs.2012.01.005>.

Trujillo-Acatitla, R., Tuxpan-Vargas, J. and Ovando-Vázquez, C. (2022) 'Oil spills: Detection and concentration estimation in satellite imagery, a machine learning approach', *Marine Pollution Bulletin*, 184, p. 114132. doi: <https://doi.org/10.1016/j.marpolbul.2022.114132>.

- Tukey, J. W. (1977) *Exploratory data analysis*. Reading, MA.
- van der Maaten, L. and Hinton, G. (2012) ‘Visualizing non-metric similarities in multiple maps’, *Machine Learning*, 87(1), pp. 33–55. doi: 10.1007/s10994-011-5273-4.
- Vasudevan, R. K. *et al.* (2021) ‘Off-the-shelf deep learning is not enough, and requires parsimony, Bayesianity, and causality’, *npj Computational Materials*, 7(1), p. 16. doi: 10.1038/s41524-020-00487-0.
- Verrelst, J. *et al.* (2012) ‘Machine learning regression algorithms for biophysical parameter retrieval: Opportunities for Sentinel-2 and -3’, *Remote Sensing of Environment*, 118, pp. 127–139. doi: <https://doi.org/10.1016/j.rse.2011.11.002>.
- Wan, J. and Cheng, Y. (2013) ‘Remote sensing monitoring of Gulf of Mexico oil spill using ENVISAT ASAR images’, in *2013 21st International Conference on Geoinformatics*, pp. 1–5. doi: 10.1109/Geoinformatics.2013.6626165.
- Wang, D. *et al.* (2022) ‘BO-DRNet: An Improved Deep Learning Model for Oil Spill Detection by Polarimetric Features from SAR Images’, *Remote Sensing*, 14(2). doi: 10.3390/rs14020264.
- Wang, D. *et al.* (2023) ‘An improved semantic segmentation model based on SVM for marine oil spill detection using SAR image’, *Marine Pollution Bulletin*, 192, p. 114981. doi: <https://doi.org/10.1016/j.marpolbul.2023.114981>.
- Wang, X. *et al.* (2021) ‘Detection of Oil Spill Using SAR Imagery Based on AlexNet Model’, *Computational Intelligence and Neuroscience*. Edited by N. Zhang. Hindawi, 2021, p. 4812979. doi: 10.1155/2021/4812979.
- Warnes, G. R. *et al.* (2020) ‘gplots: Various R Programming Tools for Plotting Data’. Available at: <https://cran.r-project.org/package=gplots>.
- White, H. K. *et al.* (2012) ‘Impact of the *Deepwater Horizon* oil spill on a deep-water coral community in the Gulf of Mexico’, *Proceedings of the National Academy of Sciences*, 109(50), pp. 20303–20308. doi: 10.1073/pnas.1118029109.

Winkelmann, K. H. (2005) *On the applicability of imaging spectrometry for the detection and investigation of contaminated sites with particular consideration given to the detection of fuel hydrocarbon contaminants in soil*. Available at: <https://opus4.kobv.de/opus4-btu/frontdoor/index/index/docId/132>.

Wolanin, A. *et al.* (2019) ‘Estimating crop primary productivity with Sentinel-2 and Landsat 8 using machine learning methods trained with radiative transfer simulations’, *Remote Sensing of Environment*, 225, pp. 441–457. doi: <https://doi.org/10.1016/j.rse.2019.03.002>.

Xing, Q. *et al.* (2015) ‘Observation of Oil Spills through Landsat Thermal Infrared Imagery: A Case of Deepwater Horizon’, *Aquatic Procedia*. Elsevier B.V., 3, pp. 151–156. doi: [10.1016/j.aqpro.2015.02.205](https://doi.org/10.1016/j.aqpro.2015.02.205).

Xu, J. *et al.* (2017) ‘Research on marine radar oil spill network monitoring technology’, in *2017 IEEE International Geoscience and Remote Sensing Symposium (IGARSS)*, pp. 1868–1871. doi: [10.1109/IGARSS.2017.8127341](https://doi.org/10.1109/IGARSS.2017.8127341).

Yamashita, R. *et al.* (2018) ‘Convolutional neural networks: an overview and application in radiology’, *Insights into Imaging*, 9(4), pp. 611–629. doi: [10.1007/s13244-018-0639-9](https://doi.org/10.1007/s13244-018-0639-9).

Yang, Z. *et al.* (2021) ‘Decision support tools for oil spill response (OSR-DSTs): Approaches, challenges, and future research perspectives’, *Marine Pollution Bulletin*, 167, p. 112313. doi: <https://doi.org/10.1016/j.marpolbul.2021.112313>.

Ye, X. *et al.* (2019) ‘A simulation-based multi-agent particle swarm optimization approach for supporting dynamic decision making in marine oil spill responses’, *Ocean & Coastal Management*, 172, pp. 128–136. doi: <https://doi.org/10.1016/j.ocecoaman.2019.02.003>.

Zeng, K. and Wang, Y. (2020) ‘A Deep Convolutional Neural Network for Oil Spill Detection from Spaceborne SAR Images’, *Remote Sensing*, 12(6). doi: [10.3390/rs12061015](https://doi.org/10.3390/rs12061015).

Zhai, J. *et al.* (2022) ‘A Dual Attention Encoding Network Using Gradient Profile Loss for Oil Spill Detection Based on SAR Images’, *Entropy*, 24(10). doi: [10.3390/e24101453](https://doi.org/10.3390/e24101453).

Zhang, J. *et al.* (2020) ‘Oil Spill Detection in Quad-Polarimetric SAR Images Using an

Advanced Convolutional Neural Network Based on SuperPixel Model’, *Remote Sensing*, 12(6). doi: 10.3390/rs12060944.

Zhou, T. *et al.* (2021) ‘Prediction of soil organic carbon and the C:N ratio on a national scale using machine learning and satellite data: A comparison between Sentinel-2, Sentinel-3 and Landsat-8 images’, *Science of The Total Environment*, 755, p. 142661. doi: <https://doi.org/10.1016/j.scitotenv.2020.142661>.

Zhou, Z. *et al.* (2018) ‘UNet++: {A} Nested U-Net Architecture for Medical Image Segmentation’, *CoRR*, abs/1807.1. Available at: <http://arxiv.org/abs/1807.10165>.

Zhu, Q. *et al.* (2022) ‘Oil Spill Contextual and Boundary-Supervised Detection Network Based on Marine SAR Images’, *IEEE Transactions on Geoscience and Remote Sensing*, 60, pp. 1–10. doi: 10.1109/TGRS.2021.3115492.

Zhu, Z., Wang, S. and Woodcock, C. E. (2015) ‘Improvement and expansion of the Fmask algorithm: cloud, cloud shadow, and snow detection for Landsats 4–7, 8, and Sentinel 2 images’, *Remote Sensing of Environment*, 159, pp. 269–277. doi: <https://doi.org/10.1016/j.rse.2014.12.014>.

Zhu, Z. and Woodcock, C. E. (2012) ‘Object-based cloud and cloud shadow detection in Landsat imagery’, *Remote Sensing of Environment*, 118, pp. 83–94. doi: <https://doi.org/10.1016/j.rse.2011.10.028>.

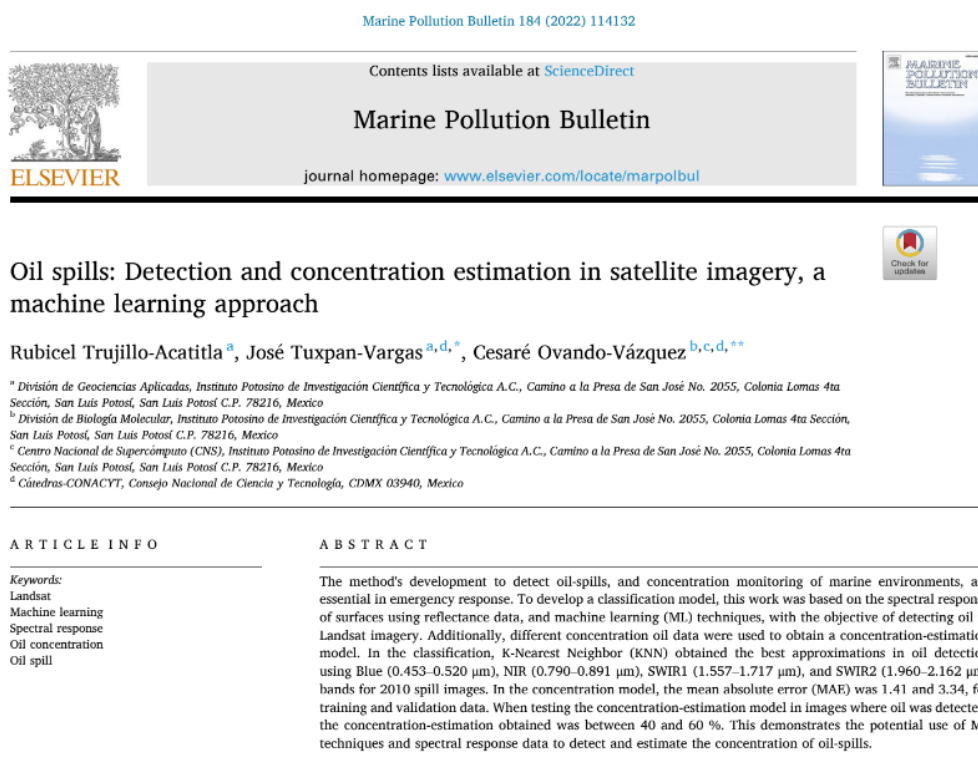
Zhuang, F. *et al.* (2019) ‘A Comprehensive Survey on Transfer Learning’, *CoRR*, abs/1911.0. Available at: <http://arxiv.org/abs/1911.02685>.

Zou, Y. *et al.* (2016) ‘Oil spill detection by a support vector machine based on polarization decomposition characteristics’, *Acta Oceanologica Sinica*, 35(9), pp. 86–90. doi: 10.1007/s13131-016-0935-5.

Anexos

Artículo científico derivado de este proyecto doctoral:

Trujillo-Acatitla, R., Tuxpan-Vargas, J., Ovando-Vázquez, C., 2022. Oil spills: Detection and concentration estimation in satellite imagery, a machine learning approach. *Mar. Pollut. Bull.* 184, 114132. [https://doi.org/https://doi.org/10.1016/j.marpolbul.2022.114132](https://doi.org/10.1016/j.marpolbul.2022.114132)



1. Introduction

The oil industry has experienced spills of great magnitude throughout its history, causing damage and losses to environmental, social, economic, and even political sectors. These spills occur at different stages of the industry process chain (exploration, production, refining, storage, and/or distribution), either in marine or terrestrial environments (Burgherr, 2007; Daly et al., 2016; Ivshina et al., 2015; Li et al., 2016; Yang et al., 2021).

The marine environment has the greatest impact, with an estimation of two million tons of oil released annually due to ship strikes, oil well blowouts, pipeline ruptures, and explosions at storage facilities; as a result, there is a constant damage to the ecosystem, and economic losses (Burgherr, 2007; Daly et al., 2016; Ivshina et al., 2015). Since the 1980s, the use of remote sensing tools in the detection and monitoring of oil

ABSTRACT

The method's development to detect oil-spills, and concentration monitoring of marine environments, are essential in emergency response. To develop a classification model, this work was based on the spectral response of surfaces using reflectance data, and machine learning (ML) techniques, with the objective of detecting oil in Landsat imagery. Additionally, different concentration oil data were used to obtain a concentration-estimation model. In the classification, K-Nearest Neighbor (KNN) obtained the best approximations in oil detection using Blue (0.453–0.520 μm), NIR (0.790–0.891 μm), SWIR1 (1.557–1.717 μm), and SWIR2 (1.960–2.162 μm) bands for 2010 spill images. In the concentration model, the mean absolute error (MAE) was 1.41 and 3.34, for training and validation data. When testing the concentration-estimation model in images where oil was detected, the concentration-estimation obtained was between 40 and 60 %. This demonstrates the potential use of ML techniques and spectral response data to detect and estimate the concentration of oil-spills.

spills began (Hooper, 1982), but it was not until one of the worst spills (over 4.9 million barrels of crude oil) that has been documented to date occurred in April 2010 in the Deepwater Horizon platform in the Gulf of Mexico that its use increased exponentially (García-Pineda et al., 2017). During this event, human observations from the air were used to detect and monitor oil trajectories (García-Pineda et al., 2017, 2013; Leifer et al., 2012), which ended up in high costs due to the extent of the slick (~180,000 km^2) and the duration of the spill (Clark et al., 2010; Leifer et al., 2012; Topouzelis et al., 2015).

Therefore, better methods were sought for spill detection and monitoring, retaking the use of remote sensing tools that are low-cost, cover large areas and are more efficient, making them useful for an emergency response. With which it has started a new era in the application of remote sensing to oil spill detection, in which various agencies, institutions, and governments began contributing their spatial assets

* Correspondence to: J. Tuxpan-Vargas, División de Geociencias Aplicadas, Instituto Potosino de Investigación Científica y Tecnológica A.C., Camino a la Presa de San José No. 2055, Colonia Lomas 4ta Sección, San Luis Potosí, San Luis Potosí C.P. 78216, Mexico.

** Correspondence to: C. Ovando-Vázquez, División de Biología Molecular, Instituto Potosino de Investigación Científica y Tecnológica A.C., Camino a la Presa de San José No. 2055, Colonia Lomas 4ta Sección, San Luis Potosí, San Luis Potosí C.P. 78216, Mexico.

E-mail addresses: jose.tuxpan@ipicyt.edu.mx (J. Tuxpan-Vargas), cesare.ovando@ipicyt.edu.mx (C. Ovando-Vázquez).

<https://doi.org/10.1016/j.marpolbul.2022.114132>

Received 18 April 2022; Received in revised form 8 September 2022; Accepted 9 September 2022

Available online 26 September 2022

0025-326X/© 2022 Elsevier Ltd. All rights reserved.