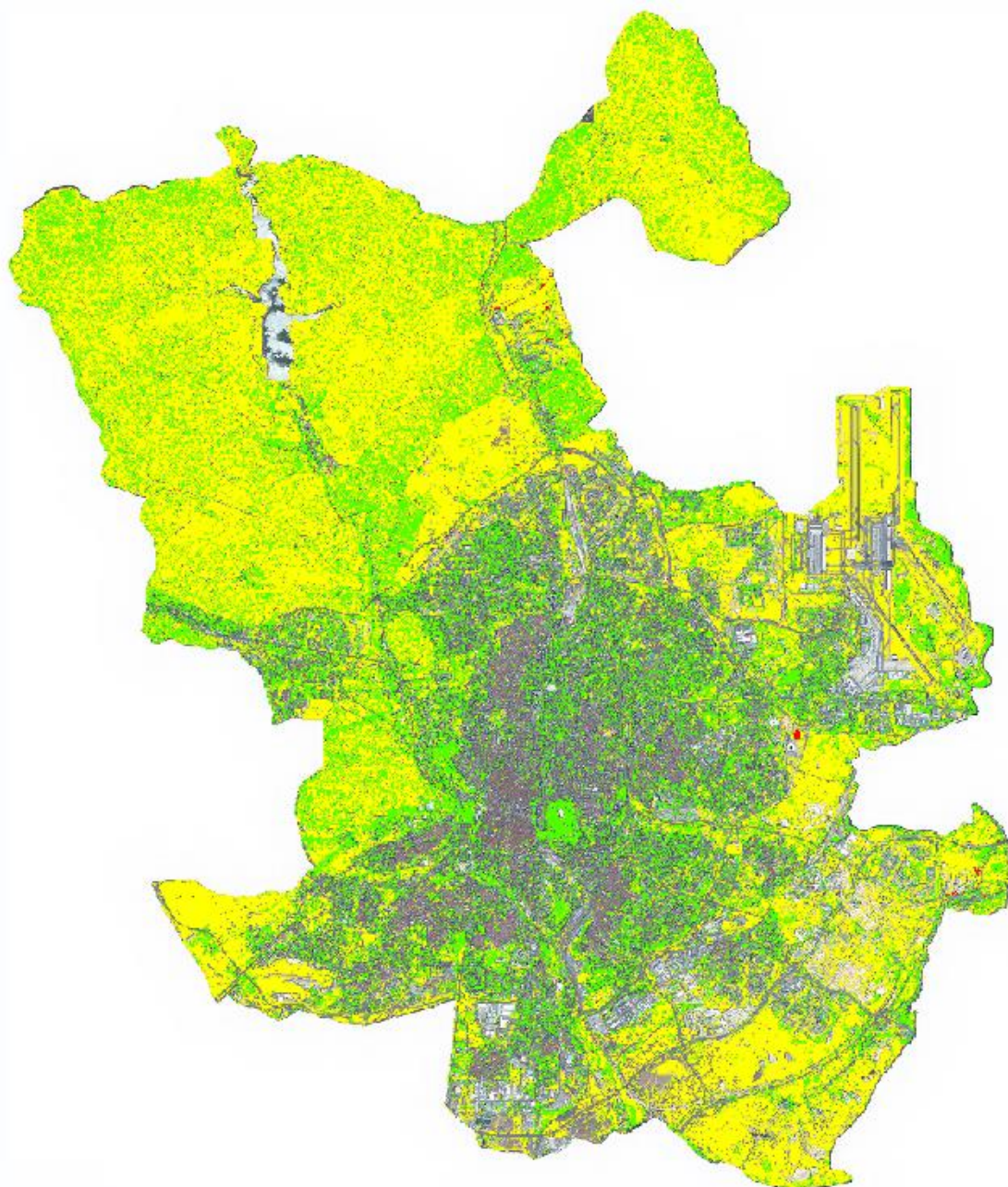


2024

Centro de Observación y
Teledetección Espacial,
S.A.U.

(COTESA – A47461066)



SERVICIO DE CARTOGRAFÍA DE ZONAS VERDES

[Metodología para la detección de zonas verdes en primavera de 2024 sobre el término municipal de Madrid más las pistas del aeropuerto Adolfo Suárez – Madrid Barajas]

Contenido

1.	Introducción.....	2
2.	Objetivo.....	2
3.	Metodología.....	2
3.1.	Cálculo de índices.....	2
3.2.	Selección de áreas de entrenamiento.....	3
3.3.	Clasificación por píxel.....	5
1.1.	Resultados.....	6
1.2.	Entregables.....	8

1. Introducción.

En este informe se muestra la metodología utilizada para la obtención de zonas verdes en Madrid para primavera del año 2024. Para ello se ha utilizado un algoritmo propio de COTESA que permite la automatización de cartografía de coberturas de suelo de alta resolución. Es una herramienta diseñada en software libre, que permite la semi- automatización de la cartografía de ocupación del suelo y con mayor precisión y recurrencia que la cartografía existente actual.

Mediante la clasificación supervisada basada en píxeles del mosaico obtenido anteriormente, se logra una cartografía de coberturas basadas en la detección de zonas verdes, que son arbolado, matorral y pastizal-zonas ajardinadas, correspondiendo cada polígono a una única cobertura. Se automatiza así el cartografiado de clases de vegetación con una alta resolución y precisión tanto geométrica como temática, reduciendo los tiempos y costes de producción según la metodología tradicional.

2. Objetivo.

Obtención de una capa raster donde queden reflejadas las zonas verdes para primavera de 2024 para el término municipal de Madrid (más pistas del Aeropuerto Adolfo Suárez - Madrid Barajas), con las siguientes clases:

- 1: Vegetación arbórea.
- 2: Prados-jardines.
- 3: Matorrales.

3. Metodología.

A continuación se describe cada una de las fases llevadas a cabo para la obtención de las zonas verdes, que radica en varios procesos.

3.1. Cálculo de índices.

El cálculo de índices corresponde a la preparación de las imágenes para poder realizar operaciones de segmentación y clasificación.

El uso de índices espectrales para facilitar la clasificación de la imagen procede de dos índices diferentes, NDVI y RVI. Con éstos índices adquirimos información explícita sobre vegetación y agua que favorecerá la clasificación de éstas zonas.

- NDVI: es un índice usado para estimar la cantidad, calidad y desarrollo de la vegetación con base a la medición (por medio de sensores remotos instalados comúnmente en una plataforma espacial) de la intensidad de la radiación de ciertas bandas del espectro electromagnético que la vegetación emite o refleja.

$$NDVI = \frac{(NIR - Red)}{(NIR + Red)}$$

- RVI: es una relación entre la reflectancia registrada en las bandas de infrarrojo cercano (NIR) y rojo. Esta es una forma rápida de distinguir las hojas verdes de otros objetos en la escena y estimar la biomasa relativa presente en la imagen. Además, este valor puede ser muy útil para distinguir la vegetación estresada de las áreas no estresadas.

$$RVI = \frac{NIR}{RED}$$

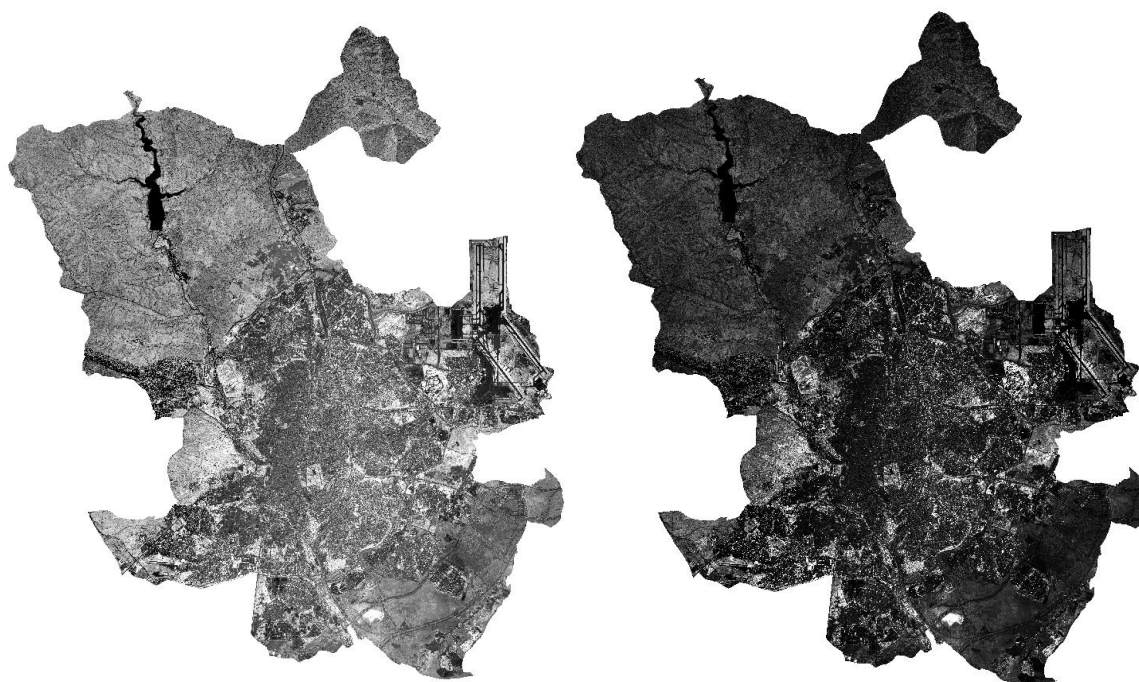


Imagen 1. NDVI a la izquierda, RVI a la derecha.

La realización de este proceso se elaboró utilizando los índices ya elaborados para el cálculo de índices, una vez obtenidos ambos, estos dos índices se unen como si de dos bandas más se tratasen al mosaico original, aportando así una información extra al mismo.

3.2. Selección de áreas de entrenamiento.

Las áreas de entrenamiento se obtienen mediante la fotointerpretación, realizando áreas de entrenamiento a lo largo de la imagen, como se puede observar en la Imagen 2.

En cada zona, se extraerán una serie de áreas de entrenamiento con el objetivo de optimizar la clasificación en los mejores escenarios. Estas áreas de entrenamiento han de cumplir los siguientes requisitos:

- Estar repartidos de manera equitativa por toda la superficie de la zona a calificar.

- La superficie total de cada clase, respecto a la superficie total de las áreas de entrenamiento, debe mantener los mismos criterios de proporcionalidad que en la extensión total de la zona a clasificar.
- Mantener un margen de seguridad en los límites y bordes de la cobertura sobre la que se está dibujando el área de entrenamiento.
- Las áreas de entrenamiento de la misma cobertura, deberán de mantener la mayor distancia posible entre ellas.

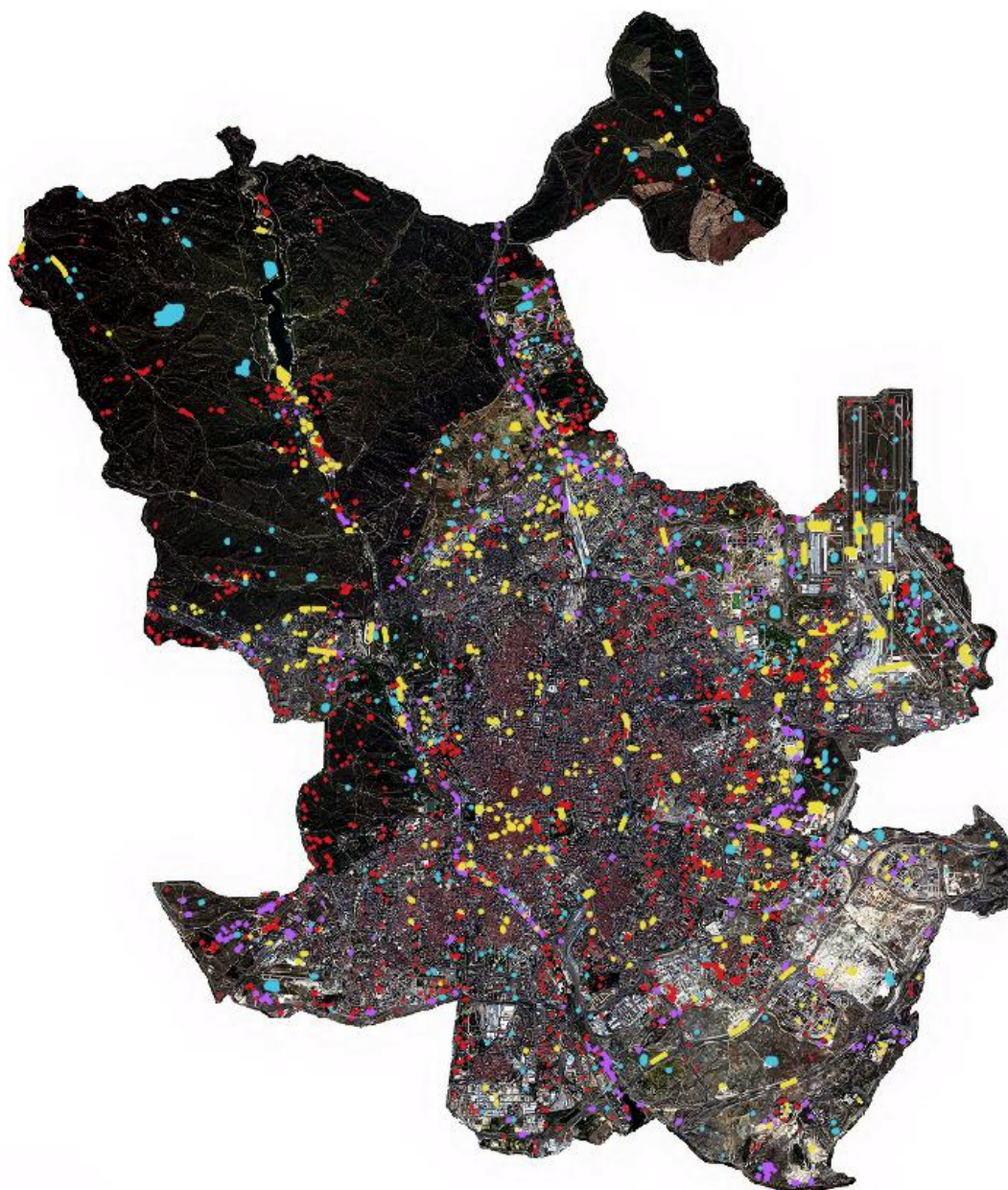


Imagen 2. Diferentes áreas de entrenamiento tomadas en el municipio de Madrid, en amarillo antrópico, en rojo vegetación arbórea, en azul prados-jardines y en morado matorrales.

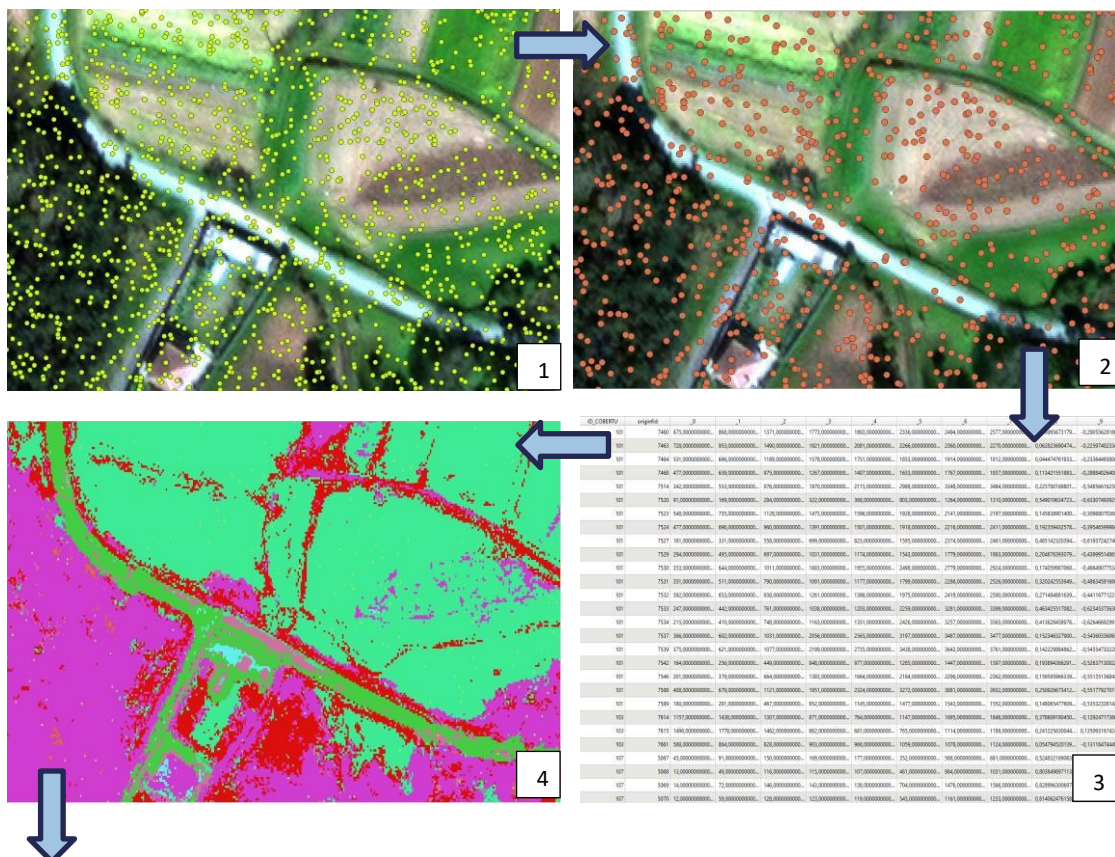
3.3. Clasificación por píxel.

La clasificación se realiza una vez obtenidas las áreas de entrenamiento, en la cual se crearán puntos aleatorios dentro de las mismas, con una estrategia de muestreo basada en la densidad de puntos. Estos puntos serán filtrados aleatoriamente y se extraerán estadísticas de las diferentes bandas del mosaico, siendo estos datos con los que el algoritmo entrenará obteniendo un modelo de clasificación. Con estos puntos se facilitará la ejecución de los algoritmos de clasificación.

Tras esto se utilizarán las siguientes herramientas:

- *PolygonClassStatistics*.
- *SampleSelection*.
- *SampleExtraction*.
- *ComputeImageStatistics*.
- *TrainVectorClassifier* donde se elegirá el método de entrenamiento de clasificación que es el *Random Forest*.
- *ImageClassifier* que utilizando las estadísticas previas clasificará la imagen.

Tras ello se realizan una serie de procesos para obtener los segmentos, siendo éstos un filtrado de los pixels para así poder depurar la imagen omitiendo el ruido que pueda haber en la clasificación mediante la herramienta de *Filtrado*, seguido de la utilización de la herramienta de ráster a polígono (*poligonizar*), y finalmente la herramienta *dissolve* de las clases obtenidas, como se puede observar en la imagen 3.



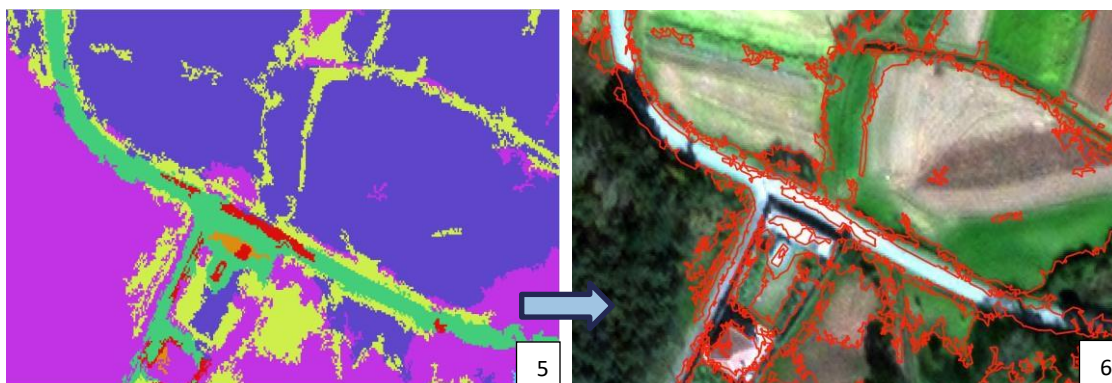


Imagen 3. Resultados de algunos de los diferentes procesos de la clasificación por píxel RF. (1) Random points; (2) sample selection; (3) sample straction; (4) image classifier; (5) filtrado; (6) poligonización.

1.1. Resultados.

En cuanto a los resultados obtenidos, se considera que son bastante satisfactorios para la dificultad que la imagen presentaba. Para la realización del QC se utilizará la matriz de confusión y el índice Kappa.

La matriz de confusión nos permite comparar el resultado de la clasificación con el conjunto de datos de referencia. Las muestras de referencia/validación están en filas mientras que los datos de clasificación están en columnas. Es equivalente a una matriz de contingencia. Los elementos diagonales corresponden a píxeles (o áreas) correctamente clasificados, mientras que todos los otros valores son clasificaciones erróneas, llamados confusiones.

	1	2	3	4
1	62578	716	1287	94
2	619	69542	549	101
3	794	654	79847	32
4	216	264	354	33756

Tabla 1. Matriz de confusión donde 1 representa la clase de árboles, 2 prados-jardines, 3 matorrales y 4 antrópico.

Con la realización de la matriz de confusión, se crean a su vez una serie de índices muy orientativos respecto a la clasificación. Como se puede observar en la tabla 1, aparecen representados todos los puntos extraídos con la herramienta *sample extraction*, y son estos los utilizados para realizar la clasificación, teniendo en las columnas y las filas las diferentes clases y cuánto ha variado la clasificación entre ellas, observando una amplia mayoría de puntos bien clasificados en función a la clasificación previa dada como área de entrenamiento.

Este primer índice ofrece una visión general de la precisión global de la clasificación. No refleja la diferencia del número de muestras entre clases. El índice Kappa se extraerá en función a las áreas de entrenamiento hechas previamente, ya que son de una alta fiabilidad porque fueron hechas a mano por un fotointérprete.

El índice de kappa responde más eficazmente a este desequilibrio frecuente, utilizando un ajuste estadístico. En ambos casos el resultado es un valor comprendido entre 0 y 1, cuanto más se acerca a 1 más eficiente es la clasificación.

El índice Kappa:

$$\kappa = \frac{N \sum_{i=1}^m x_{ii} - \sum_{i=1}^m x_{i\bar{i}} x_{\bar{i}i}}{N^2 - \sum_{i=1}^m x_{i\bar{i}} x_{\bar{i}i}}$$

m = número total de clases.
 N = número total de píxeles en las m clases de referencia.
 x_{ii} = elementos de la diagonal de la matriz de confusión.
 $x_{i\bar{i}}$ = suma de los píxeles de la clase i de referencia.
 $x_{\bar{i}i}$ = suma de los píxeles clasificados como la clase i .

Para ello, se utilizarán como referencia los puntos que se obtuvieron en el proceso de clasificación por píxel, en concreto tras el paso de la utilización de la herramienta de *sample extraction* con la que se cribaron los puntos obtenidos aleatoriamente dentro de las propias áreas de entrenamiento. Con ello veremos la bondad del algoritmo en esta imagen comparando las áreas de entrenamiento con los resultados obtenidos.

Para matorrales, encontramos que de 34590 puntos, 33756 puntos caen sobre un píxel clasificado como matorral, lo que representa una bondad del 97.59%; en cuanto a árboles, encontramos que de los 64967 puntos son detectados 62578, un 96.75%; y finalmente, de los 70811 puntos de prados, 69542 son detectados, un 98.2% del total.

Pueden parecer valores muy elevados, pero siempre hay que tener en cuenta, que este tipo de evaluación se hace siempre sobre los propios datos de entrada, es decir, de las áreas de entrenamiento, áreas que han sido escogidas y delimitadas de manera muy cautelosa, ya que sobre todo en esta fecha, tras el paso de Filomena, la vegetación quedó muy dañada, lo que se puede traducir en ocasiones en una sobrerrepresentación de la clase de prados-jardines.

Cabe destacar, que la clase antrópica desaparece del estudio, ya que no es interesante en la detección de zonas verdes, simplemente se utiliza para que sea detectada por el algoritmo, y es por ello que se elimina, dejando una capa con la clase arbórea, pastizal-jardín y matorral.

Así mismo hay que tener en cuenta varios aspectos relevantes, uno de ellos tiene que ver la detección de la clase de matorrales, debido a la dificultad de englobar éstos en una sola clase

por la gran variabilidad y poca distribución, una vez son detectados por el algoritmo son utilizados de nuevo para un segundo testeo que se realizará sobre las clases obtenidas a excepción de la clase antrópica, obteniendo así una segunda clasificación sobre las zonas verdes, y ahí donde coincida la clasificación de matorral se chequearán y serán sobre puestas sobre la detección primigenia, de ahí que “resalten” más que las otras clases.

Hay que tener en cuenta también que la detección de zonas verdes varía con cada imagen, por tanto, dependiendo del ángulo de captura de la imagen, la vegetación colindante a edificios puede no detectarse porque el edificio *ciegue* esa zona debido a la inclinación de este.

1.2. Entregables.

A continuación se muestran los entregables con una breve explicación de los mismos.

- *EsAytMadridPr2024OrtINVeg_ZonasVerdes.tif*; detección de árboles, matorrales y prados-jardines con el algortimo propio de COTESA.
- *EsAytMadridPr2024OrtINVeg_ZonasVerdes.xml*; metadatos del archivo raster.
- *EsAytMadridPr2024OrtINVeg_ZonasVerdes.pdf*; documentación con los procedimientos seguidos en la creación de la detección de zonas verdes.