



INSTITUTO TECNOLÓGICO SUPERIOR DE MISANTLA

MAESTRÍA EN SISTEMAS COMPUTACIONALES

ALGORITMOS ENSAMBLADOS Y APRENDIZAJE
PROFUNDO PARA CLASIFICACIÓN DE DAÑO FOLIAR
EN CÍTRICOS

T E S I S

QUE PARA OPTENER EL GRADO DE:

Maestro en Sistemas Computacionales

PRESENTA:

Juan Pablo Salazar

DIRECTOR:

Dr. Eddy Sánchez de la Cruz

CO-DIRECTOR:

Dr. Rajesh Roshan Biswal



Misantla, Veracruz, 2018



INSTITUTO TECNOLÓGICO SUPERIOR DE MISANTLA
DIVISIÓN DE ESTUDIOS PROFESIONALES
AUTORIZACIÓN DE IMPRESIÓN DE TRABAJO DE TITULACIÓN MAESTRÍA

FECHA: 29 de Noviembre de 2018.

ASUNTO: **AUTORIZACIÓN DE IMPRESIÓN**
DE TESIS.

A QUIEN CORRESPONDA:

Por medio de la presente se hace constar que el (la) C:

JUAN PABLO SALAZAR

estudiante de la maestría en SISTEMAS COMPUTACIONALES con No. de Control 162T0084 ha cumplido satisfactoriamente con lo estipulado por el **Lineamiento de Posgrado para la obtención del grado de Maestría** mediante **Tesis.**

Por tal motivo se **Autoriza** la impresión del **Tema** titulado:

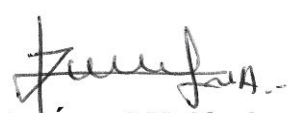
ALGORITMOS ENSAMBLADOS Y APRENDIZAJE PROFUNDO PARA
CLASIFICACIÓN DE DAÑO FOLIAR EN CÍTRICOS

Dándose un plazo no mayor de un mes de la expedición de la presente a la solicitud del examen para la obtención del grado de maestría.

ATENTAMENTE


M.S.C. Eddy Sánchez de la Cruz
Presidente




M.I.A. Roberto Ángel Meléndez Armenta
Secretario


M.C. Saúl Reyes Barajas
Vocal

Archivo.

Dedicatoria

A Mi madre:

Quien forjó en mí una visión a futuro y me enseñó que el estudio es la mejor herramienta para la superación personal.

A Mis hermanos:

Quienes han compartido sus experiencias conmigo, las cuales me han enseñado a enfrentar problemas, retos y riesgos para seguir adelante.

A mi esposa:

Que ha hecho mucho más fácil mi paso por la maestría y que más que mi esposa, se ha convertido en mi mejor amiga, mi compañera de desvelos, y mi columna de apoyo en momentos de presión para continuar con mi trabajo.

Vive como si fueses a morir mañana. Aprende como si fueses a vivir para siempre.
Mahatma Gandhi.

Agradecimientos

Mi gratitud, principalmente al Dios todopoderoso que lo hace todo posible, por darme la existencia, sabiduría y hacerme llegar al final del posgrado.

Mis más sincero agradecimiento a mi madre la Señora Luciana Salazar por brindarme su amor incondicional, enseñarme valores en todo momento y aconsejarme en las decisiones más difíciles de mi vida.

A mis hermanos, Eduardo Enrique Salazar y Victor Francisco Salazar por su apoyo incondicional.

A mi esposa Saraí Vazquez Montoya quien es mi motor y motivación.

Agradezco especialmente a mi director de tesis el Dr. Eddy Sanchez de la Cruz por guiarme constantemente y brindarme la oportunidad de trabajar bajo su asesoramiento.

A mi co-asesor Dr. Rajesh Roshan Biswal, Dr. Luis Alberto Morales, M.I.A. Roberto Ángel Meléndez Armenta, Dr. Simón Pedro Arguijo quienes formaron o forman parte del cuerpo académico de la Maestría en Sistemas Computacionales del ITSM y compartieron parte de su tiempo y conocimientos conmigo.

Al coordinador de M.S.C Galdino Matínez Flores por su aprobación en la conclusión de este trabajo de investigación.

Las personas que han sido mi ejemplo a seguir gracias a los conocimientos y experiencia de vida que han compartido conmigo Tamar Mizrahi y Jim Garofalo.

Al Consejo Nacional de Ciencia y Tecnología por el apoyo económico brindado para poder realizar mis estudios dentro del Programa Nacional de Posgrados de Calidad.

Resumen

El daño foliar es un indicador muy común de la presencia de enfermedades en plantas, las cuales pueden causar una disminución notable en la producción de cultivos, afectando la economía agrícola. Estas pueden ser detectadas por el cambio anormal en las características de las hojas de plantas, como la presencia de lesiones, decoloración, entre otras, lo que indica problemas relacionados con el crecimiento y un alto impacto en la productividad. La causa del daño en las hojas puede ser variable, como bacterias, virus o deficiencias nutricionales. Estas características sintomáticas de la hoja se pueden utilizar para clasificar enfermedades mediante procesamiento digital de imágenes y técnicas de aprendizaje automático. Esta investigación propone una alternativa para el mejoramiento de la precisión para la clasificación de enfermedades mediante el uso de características de síntomas foliares, la cual consiste en la implementación de aprendizaje automático. La base de datos fue creada realizando un proceso de extracción de características de las condiciones del daño foliar, aplicando métodos de procesamiento de imágenes. Después de realizar una clasificación binaria, la siguiente etapa consistió en añadir dos clases más de enfermedades para la ejecución de una clasificación multi-clase con la finalidad obtener las mejores combinaciones de algoritmos ensamblados con Deep Learning (*Multi-layer Perceptron*). Los resultados de combinar estos dos enfoques de clasificación, muestran un índice óptimo en conjuntos de datos binarios y porcentajes altamente competitivos en conjuntos multiclase. De este modo, se demuestra que la combinación de estas dos estrategias de clasificación es una alternativa muy fiable para la clasificación de enfermedades de naranja y otros tipos de cítricos.

Palabras clave: daño foliar, procesamiento de imágenes, algoritmos ensamblados, aprendizaje automático.

Abstract

Leaf damage is a very common indicator of plant diseases which can cause a remarkable decrease in the crop production thus affecting the agricultural economy. It can be seen as an abnormal change in the characteristics of plant leaves, such as the presence of lesions, discoloration, among others, indicating growth related problems and highly impacts the productivity. The cause of leaf damage can be variable, such as bacteria, virus or nutritional deficiencies. These symptomatic characteristics of the leaf can be used to classify diseases using image processing and automatic learning techniques. This research proposes an alternative for improving the accuracy of disease classification by the use of foliar symptom characteristics, which consists of the implementation of automatic learning. The database was created performing feature extraction from the characteristics of foliar damage using image processing methods. After performing a binary classification, the next step was to add two more classes of diseases for the execution of a multi-class classification to obtain the best combinations of assembled algorithms with Deep Learning (*Multi-layer Perceptron*). The results of combining these two classification approaches, show optimal rates in binary data sets and highly competitive percentages in multi-class sets. As a result it can be shown that the combination of these two classification strategies can be a highly reliable alternative for orange and other citrus plants.

Key words: leaf damage, image processing, assembled algorithms, deep learning.

Índice general

Resumen	I
Índice de figuras	V
Índice de tablas	VI
1. Generalidades	1
1.1. Introducción	2
1.2. Planteamiento del problema	3
1.2.1. Justificación	3
1.3. Objetivos	3
1.3.1. Objetivo General	3
1.3.2. Objetivos Específicos	4
1.4. Hipótesis	5
1.5. Propuesta de solución	5
1.6. Alcances y Limitaciones	5
1.6.1. Alcances	5
1.6.2. Limitaciones	6
1.7. Estructura de la tesis	6
2. Marco teórico	7
2.1. Patologías en plantas de cítricos causales de síntomas foliares	8
2.2. Procesamiento de Imágenes	10
2.2.1. Imágenes	10
2.2.1.1. Tipos de imágenes digitales	10
2.2.2. Pre-procesamiento digital de imágenes	12
2.2.2.1. Procesamiento puntual	12
2.2.2.2. Operaciones aritméticas	13
2.2.2.3. Histograma	13
2.2.2.4. Eliminación o reducción de ruido	14
2.2.2.5. Ecualización del histograma	14
2.2.3. Post-procesamiento digital de imágenes	14
2.2.3.1. Segmentación	15
2.2.3.2. Método de Otsu	15

2.3. Clasificación	16
2.3.0.1. Algoritmos ensamblados	17
2.3.0.2. Aprendizaje profundo	18
3. Estado del arte	21
4. Materiales y métodos	25
4.1. Materiales	26
4.2. Métodos	27
4.2.1. Adquisición de base de datos	27
4.2.2. Tratamiento de los datos	29
5. Experimentos y Resultados	33
6. Análisis de resultados	37
7. Conclusiones y trabajo futuro	41
7.1. Conclusiones	42
7.2. Trabajo futuro	42
Bibliografía	43

Índice de figuras

1.1. Método propuesto de solución.	5
2.1. Cubo de colores para el modelo de color RGB.	11
2.2. Representación gráfica de MLP.	19
4.1. Muestras. a) Saludable b) Deficiencia de Zinc c) Clorosis d) Mancha gracienta.	26
4.2. Procesamiento. a) Original b) Escala de grises c) Segmentada d) ROI. .	27
4.3. Captura de base de datos creada en CSV.	29
4.4. Diagrama de Validación cruzada(De Pablo Castellano)	30
5.1. Modelo de experimentación realizada de clasificación supervisada.	34
6.1. Criterio y algoritmos con mejores resultados.	38

Índice de tablas

2.1. Algoritmos ensamblados.	17
4.1. Resumen de características de la base de datos.	28
4.2. Conjuntos binarios y multiclase	28
4.3. Datos de entrenamiento y prueba para criterio 2/3-1/3.	30
4.4. Datos de entrenamiento y prueba para validación cruzada.	31
4.5. Datos de entrenamiento y prueba para muestra representativa.	32
5.1. Resultados del primer experimento con el conjunto 1 y clasificación binaria.	35
5.2. Mejores resultados de clasificación binaria y multiclase.	36
6.1. Matrices de confusión representativas de clasificación de los conjuntos binarios del 1 al 3	39
6.2. Matrices de confusión representativas de la clasificación de los conjuntos binarios del 4 al 6	39
6.3. Matrices de confusión representativas de la clasificación de los conjuntos multiclase 7 y 8	40
6.4. Comparativa con el estado del arte del % de clasificación de característi- cas foliares	40

Capítulo 1

Generalidades

1.1. Introducción

La detección oportuna y correcta de enfermedades en cultivos es determinante en una eficiente producción agrícola (18), por lo que es indispensable identificar de manera temprana sus causas. Los avances tecnológicos aplicados al área agrícola a nivel mundial crecen aceleradamente y recientemente están tomando un enfoque objetivo en la solución de problemas reales en cuanto a aplicaciones agrícolas se refiere. La implementación de herramientas computacionales para la clasificación de enfermedades de cítricos se ha enfocado con gran intensidad a la clasificación foliar derivado que el cambio de las características de las hojas es el principal indicador en estadios tempranos de alguna deficiencia nutrimental o presencia de enfermedades. Algunos síntomas visibles en las hojas de las plantas se caracterizan por la decoloración de las hojas, otros por el cambio a tonalidades oscuras, presencia de puntos de color marrón, presencia de hongos y parásitos de color blanco, estas características permiten la identificación y clasificación de los síntomas foliares para la determinación de acciones de prevención y erradicación de enfermedades. Hasta ahora se han implementado técnicas como análisis y diferenciación fitoquímica entre HLB y deficiencia de zinc(4), espectroscopia de infrarrojo medio para la detección de HLB (3), sensores de visión (16) y métodos de visión por computadora mediante método de co-ocurrencia para la clasificación de enfermedades de cítricos (17), algoritmo de vecinos más cercanos (KNN) (28), además de la aplicación de algoritmos de inteligencia artificial como redes neuronales probabilísticas (7) (33), redes neuronales convolucionales (14), discriminación varietal cuantitativa de folíolos de soya con redes neuronales(13) y transformaciones de color, segmentación con técnicas de clusterización usando k-means y clasificación estadística (2). Estas técnicas intentan resolver la clasificación de las enfermedades de las plantas identificando las características que conducen a diferentes tipos de daños en las hojas. Sin embargo, es posible mejorar la precisión de la clasificación utilizando combinaciones de algoritmos de filtrado y reducción de la dimensionalidad, minimizando así el margen de error de la clasificación.

Los algoritmos de inteligencia artificial para la clasificación de productos biológicos se utilizan cada vez más en la identificación de enfermedades. Debido a una amplia gama de algoritmos disponibles para los modelos de aprendizaje automático, es conveniente utilizar clasificadores que favorecen la detección de enfermedades tanto en aplicaciones médicas como agrícolas. Actualmente, se están buscando alternativas de bajo costo y fáciles de implementar en la agricultura, es aquí es donde se encuentra el nicho de oportunidad para el uso del reconocimiento de patrones y algoritmos de aprendizaje automático para clasificar las enfermedades de las plantas en base a las características de las hojas.

1.2. Planteamiento del problema

La producción de cítricos representa en México el 2.78 % del producto interno bruto (PIB) agrícola nacional (22); en donde el limón se posiciona con el 1.5 %, la naranja el 1.15 % y la toronja el 0.13 %; en donde la producción anual alcanzada para el año 2017 fue en limón de 2,520,797 toneladas, de naranja 4,670,259 toneladas y toronja de 451,814 toneladas(23). El estado de Veracruz es el mayor aportador a nivel nacional ingresando el 44.5 % del volumen total (21). La producción de este fruto se ve afectada por la diseminación de enfermedades en los campos de producción y es la principal causa de pérdidas económicas. Para solucionar esta problemática se han implementado técnicas con un enfoque fitoquímico con la finalidad de clasificar e identificar dichas enfermedades, aún así las soluciones no son completamente aptas para evitar la diseminación y daños en los cultivos por lo cual las investigaciones relativas a este tema están tomando un enfoque computacional para una mejor y eficiente determinación del tipo de enfermedad presente en los campos de producción.

1.2.1. Justificación

La investigación experimental en el ámbito de la informática es muy útil para estudiar problemas directamente relacionados con la sociedad, en particular con la agricultura. Ésta se puede beneficiar para mantener un nivel económico aceptable en las comunidades locales. La región de Misantla es reconocida como exportadora de cítricos, ocupando el undécimo lugar de producción en el Estado de Veracruz en México, sin embargo, la producción no mejora debido al bajo uso de herramientas de control de plagas y nutrición de las plantas. Un enfoque para resolver esto puede ser la implementación de algoritmos de procesamiento de imágenes e inteligencia artificial en imágenes de hojas enfermas, para identificar la naturaleza, tipo y extensión del daño en las hojas, permitiendo así, obtener una forma apropiada para el tratamiento de los cultivos y la generación de aplicaciones tecnológicas para la detección oportuna de enfermedades de cítricos.

1.3. Objetivos

1.3.1. Objetivo General

Alcanzar el valor óptimo de precisión, respecto al E. A., en la clasificación de enfermedades de cítricos aplicada al daño foliar mediante la combinación de algoritmos ensamblados y deep learning.

1.3.2. Objetivos Específicos

1. Recolectar hojas de árboles de naranja con presencia sintomática de diferentes enfermedades y visiblemente sanas .
2. Fotografiar cada una de las hojas recolectadas en un ambiente controlado.
3. Recopilar las imágenes obtenidas en una base de datos de imágenes de acuerdo el tipo de daño foliar.
4. Utilizar procesamiento digital de imágenes para la extracción de características.
5. Obtener el histograma de frecuencias de los niveles de gris de las imágenes.
6. Eliminar el brillo residual en las imágenes.
7. Realizar ecualización del histograma de frecuencias de las imágenes para mejorar el contraste.
8. Aplicar un método de segmentación para la obtención de la región de interés.
9. Selección de características estadísticas básicas de intensidad y los valores de niveles de gris del histograma de frecuencias.
10. Realizar la extracción de características.
11. Crear una base de datos general con la inclusión de todos los datos en base a las características.
12. Crear bases de datos binarias y multiclase.
13. Determinar que tipos de algoritmos de clasificación se emplearán.
14. Determinar que criterios de clasificación serán factibles.
15. Realizar pruebas de clasificación de daño foliar por medio de algoritmos ensamblados en combinación con deep learning para clasificación binaria empleando tres criterios de muestreo.
16. Determinar que combinaciones de algoritmos ensamblados y deep learning son las que obtienen mejores resultados de clasificación.
17. Emplear las mejores combinaciones de algoritmos de clasificación para procesar las bases de datos multiclase.
18. Analizar los resultados de las combinaciones de algoritmos de clasificación.
19. Realizar un análisis del prototipo para su validación con respecto al estado del arte.

1.4. Hipótesis

Es posible alcanzar el valor óptimo de precisión, respecto al E. A., en la clasificación de enfermedades de cítricos aplicada al daño foliar mediante la combinación de algoritmos ensamblados y deep learning.

1.5. Propuesta de solución

Esta investigación propone una alternativa para el mejoramiento de la precisión para la clasificación de enfermedades mediante síntomas foliares la cual consiste en (i) la creación de una base de datos de imágenes de varios tipos de daño foliar, (ii) el uso de métodos de procesamiento digital de imágenes para la extracción de características y (iii) la combinación de algoritmos ensamblados con Deep Learning (*Multi-layer perceptron*) para clasificar características foliares de hojas de árboles de naranja Valencia. Figura 1.1

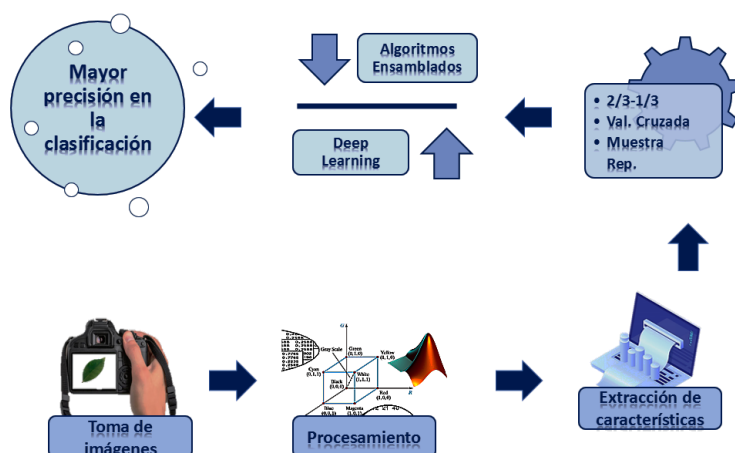


Figura 1.1: Método propuesto de solución.

1.6. Alcances y Limitaciones

1.6.1. Alcances

- Mejorar la precisión de clasificación y la identificación de los mejores algoritmos de clasificación mediante la experimentación.

- Clasificar anormalidades en imágenes digitales de hojas de árboles de naranja, utilizando algoritmos ensamblados y Deep Learning.
- Mejorar aplicaciones ya existentes en el diagnóstico de enfermedades de plantas.
- Combinar dos estrategias de clasificación empleando el enfoque de algoritmos ensamblados con Deep learning.

1.6.2. Limitaciones

- Dificil adquisición de bases de datos de imágenes de daño foliar de cítricos.
- Necesidad de reconocimiento de enfermedades en hojas por un experto.
- Baja disponibilidad de procesos de selección de hojas con diferentes tipos de daño foliar.

1.7. Estructura de la tesis

Este documento está estructurado en siete capítulos que se describen a continuación:

- El capítulo uno, explica las generalidades de este proyecto, en este se muestra el enfoque de la investigación como soporte a la problemática de diagnóstico de enfermedades de cítricos y las estrategias para su resolución; consta de la introducción al tema, se plantea el problema, el objetivo general, objetivos específicos, hipótesis, propuesta de solución, alcances y limitaciones de la investigación.
- El segundo capítulo describe el marco teórico, donde se encuentran las definiciones conceptuales de las técnicas implementadas.
- El estado del arte se encuentra en el capítulo tres en donde se presentan los trabajos referentes al tema de investigación los cuales sustentan la realización de la presente investigación.
- El cuarto capítulo presenta los materiales y métodos, se describen tanto el pre-procesamiento y procesamiento digital así como las combinaciones de algoritmos implementadas para la clasificación.
- Experimentos y resultados se encuentra en el quinto capítulo, donde se muestran los resultados de los experimentos realizados en la clasificación.
- El capítulo seis muestra el análisis de resultados, las matrices de confusión y tablas generadas de los resultados del capítulo anterior.
- Y por último las conclusiones y trabajo futuro estan en el capítulo siete, donde se culmina la investigación con las conclusiones basadas en los resultados, además, de establecer las acciones a realizar a corto y largo plazo.

Capítulo 2

Marco teórico

La disponibilidad de mayor información referente al área agrícola y la facilidad de su distribución a través de medios informáticos actualmente es un factor de soporte en el estudio de enfermedades de plantas y en específico las de cítricos, algunas de estas enfermedades se describen como punto de referencia para esta investigación.

2.1. Patologías en plantas de cítricos causales de síntomas foliares

En cítricos y la gran mayoría de plantas existen enfermedades y/o deficiencia de nutrientes que al afectar la planta derivan en la caracterización sintomática en las hojas, estas características foliares son variadas, desde decoloración hasta puntos que se pueden notar a simple vista. En México las enfermedades de cítricos en su mayoría están reglamentadas y controladas, en todas la característica principal es la presencia visible de diversos tipos de daño foliar (26).

Clorosis Variegada de los Cítricos (CVC): Esta enfermedad es causada por la bacteria *Xillela Fastidiosa*, la bacteria vive y se multiplica en la sabia de las hojas de las plantas de cítricos bloqueando el desplazamiento del agua. Los primeros síntomas de la hoja se parecen a la deficiencia de zinc con áreas cloróticas intervenales en la superficie superior, las hojas pueden ser más pequeñas de lo normal y los síntomas de la hoja como la decoloración de las áreas afectadas y las manchas oscuras o marrón con acumulación de gomosis en el envés se intensifican mucho más en hojas maduras, en la última etapa las manchas comienzan a extenderse hacia el borde y el tejido se seca (24).

Huanglongbing (HLB): HLB, citrus greening, amarillamiento ó enverdecimiento de los cítricos, es una enfermedad bacteriana causada por *Candidatus Liberibacter spp.*, y transmitida por dos especies de insectos psílidos. El principal síntoma es el aclaramiento de nervaduras de las hojas con presencia de manchas o moteados, asimétricas respecto al nervio central, el árbol presenta hojas verdes-claras o amarillas mezcladas con un verde normal sin una clara división entre ellas. Esta sintomatología está estrechamente asociada a la enfermedad. Una planta sin la enfermedad tiene los dos lados, derecho e izquierdo, cloróticos o verdes normales. En algunos casos se observa el engrosamiento de las nervaduras de las hojas. Las hojas jóvenes afectadas permanecen de tamaño pequeño, ocurriendo el proceso de forma más severa en su etapa mas avanzada. Los síntomas son parecidos a la deficiencia de minerales como: Zinc, Hierro, Manganese, Calcio, Azufre, Boro y Magnesio (25).

Virus psorosis de los cítricos (CPsV): En las hojas jóvenes de los árboles afectados por psorosis, se presentan flecos cloróticos que se observan a trasluz y que se sitúan normalmente entre los nervios laterales de las hojas y paralelas a ellos. En el caso de los árboles de campo afectados por psorosis B, algunas veces se observa la

aparición de necrosis en los brotes jóvenes (reacción de shock). Además suelen aparecer manchas amarillentas en el haz de las hojas adultas, que se corresponden con pequeñas lesiones de color pardo y aspecto gomoso en el envés. Además, la psorosis B puede mostrar variación en los síntomas de las hojas y frutos incluyendo los patrones de ringspot (manchas anulares) en las hojas (29).

Amarillamiento letal del limón persa: A partir de 1989 se detectó en el norte de Veracruz una enfermedad en árboles de limón Persa (*C. latifolia* Tanaka), injertados en naranjo agrio, al cual se le denominó Amarillamiento Letal. Esta enfermedad se presenta como un aclaramiento de las nervaduras del follaje, en algunas ramas y ocasionalmente moteados, el síntoma avanza lentamente al resto del árbol, el cual se torna completamente clorótico y bajo estrés hídrico se defolia. Los frutos son de color verde amarillento, en lugar de verde oscuro del fruto normal (32).

Cancro de los cítricos: causado por la bacteria *Xanthomonas citri* subsp. *citri*. Los síntomas foliares muestran características muy particulares como lesiones de color marrón, circulares, elevadas, acorchadas, con bordes húmedos y halo amarillo, con apariencia de cráter. La infección puede abarcar todo el grosor de la hoja y atravesar el haz y envés; las lesiones en el haz se observan generalmente más aplanadas y hundidas; mientras que en el envés se aprecian en forma de pequeñas ampollas cuando la lesión es joven, las lesiones miden aproximadamente de 2 a 10 mm. En lesiones maduras el tejido muerto corchoso puede desprenderse dejando huecos en las hojas (27).

Mancha grasienta de los cítricos: El agente causal es el hongo *Mycosphaerella citri* el primer síntoma en las hojas es una manchita verde amarillenta, la que se torna posteriormente amarillo naranja. Al área afectada se levanta a ambos lados de la hoja y su color se vuelve carmelita tostado, luego, carmelita oscura, tirando a negro y por último, parece como una masa traslúcida oscura de bajo de una epidermis semitransparente. Estas lesiones tienen un aspecto aceitoso. Las hojas fuertemente afectadas se ponen amarillas antes de caerse y son las primeras en hacerlo cuando el árbol recibe un shock, las pérdidas ocasionadas por la mancha grasienta son resultado de la defoliación prematura del árbol, lo que causa una reducción en el rendimiento y en el tamaño de la fruta. (31).

Deficiencia de Zinc: Los síntomas de deficiencia de Zinc se caracterizan por bandas verdes irregulares a lo largo de la nervadura central y las venas principales sobre un fondo amarillo. Las cantidades relativas de tejido verde y amarillo varían desde una condición de deficiencia leve de Zn en la que sólo hay pequeñas manchas amarillas entre las venas laterales más grandes hasta una condición en la que sólo una porción basal de la nervadura central es verde y el resto de la hoja es de color amarillo claro a blanco. En los estadios menos agudos, las hojas son de tamaño casi normal, mientras que en los casos más severos las hojas son puntiagudas, anormalmente estrechas con tendencia a estar erguidas, y de tamaño extremadamente reducido (34).

2.2. Procesamiento de Imágenes

El procesamiento digital de imágenes es un campo abierto de investigación. El progreso constante en esta área está creciendo en conjunción con otras áreas relacionadas como las matemáticas, la informática, y el creciente conocimiento de ciertos órganos del cuerpo humano involucrados en la percepción y manipulación de imágenes. Además de esto, la preocupación del hombre por imitar y utilizar ciertas características del ser humano como soporte en la solución de problemas, el avance del Procesamiento Digital de Imágenes se refleja en la medicina, la astronomía, la geología, la microscopía, entre otras aplicaciones. La información meteorológica, la transmisión y despliegue de imágenes y la visualización en internet están respaldadas por estos avances (15).

2.2.1. Imágenes

Una imagen digital difiere de una foto en que los valores x , y , y y $f(x, y)$ son todos discretos. Normalmente sólo toman valores enteros, por lo que una imagen tendrá un rango en x y y que va de 1 a 256 cada uno, es decir los valores de intensidad de brillo van de 0 (negro) a 255 (blanco). Una imagen digital, puede ser considerada como un gran arreglo de puntos muestreados de una imagen continua, cada una de las cuales tiene un brillo cuantizado; estos puntos son los píxeles que constituyen la imagen digital. Los píxeles que rodean a un píxel dado, constituyen su vecindad, la cual se caracteriza por su forma de la misma manera que una matriz la cual a su vez puede representar a una imagen digital.

2.2.1.1. Tipos de imágenes digitales

Se consideran cuatro tipos de imágenes digitales:

Binaria: Cada píxel es sólo blanco o negro, como sólo hay dos valores posibles para cada píxel, sólo es necesario un bit por píxel; por lo que, estas imágenes pueden ser muy eficaces en términos de almacenamiento. Las imágenes para las que puede ser adecuada una representación binaria incluyen texto ya sea impreso o manuscrito, huellas dactilares o planos arquitectónicos.

Escala de grises: Cada píxel es un tono de gris, normalmente de 0 (negro) a 255 (blanco). Este rango significa que cada píxel puede ser representado por ocho bits, o exactamente un byte, este es un rango muy natural para el manejo de archivos de imagen. Se utilizan otros rangos de escala de grises, pero generalmente son un exponente de 2. Algunos ejemplos de este tipo de imágenes surgen en la medicina como los rayos X e imágenes de obras impresas. Hasta ahora se ha demostrado que 256 niveles de gris diferentes son suficientes para el reconocimiento de la mayoría de los objetos naturales.

RGB ó color verdadero: Aquí cada píxel tiene un color en particular; ese color se describe por la cantidad de rojo, verde y azul que contiene. Si cada uno de estos componentes tiene un rango de 0 – 255, esto da un total de $255^3 = 16,777,216$ posibles colores diferentes en la imagen, suficiente información para cualquier imagen. Dado que el número total de bits necesarios para cada píxel es de 24, estas imágenes también se denominan imágenes en color de 24 bits y puede considerarse como una pila de tres matrices que representan los valores rojo, verde y azul de cada píxel, lo que significa que por cada píxel hay tres valores correspondientes.

De manera estándar los colores de una imagen, normalmente se definen mediante un subconjunto de un sistema de coordenadas tridimensional; dicho subconjunto se denomina modelo de color. De hecho, existen varios métodos diferentes para describir el color, pero para la visualización y almacenamiento de imágenes un modelo estándar es el RGB, para el cual podemos imaginar todos los colores que se encuentran dentro de un cubo de color de lado como se muestra en la Figura 2.1. Los colores a lo largo de la diagonal blanco-negro, mostrados en el diagrama como una línea de puntos, son los puntos del espacio donde todos los valores de R, G, B, son iguales representados por las diferentes intensidades de gris. También podemos pensar en los ejes del cubo de color como si fueran valores discretizados a números enteros en el rango 0-255.

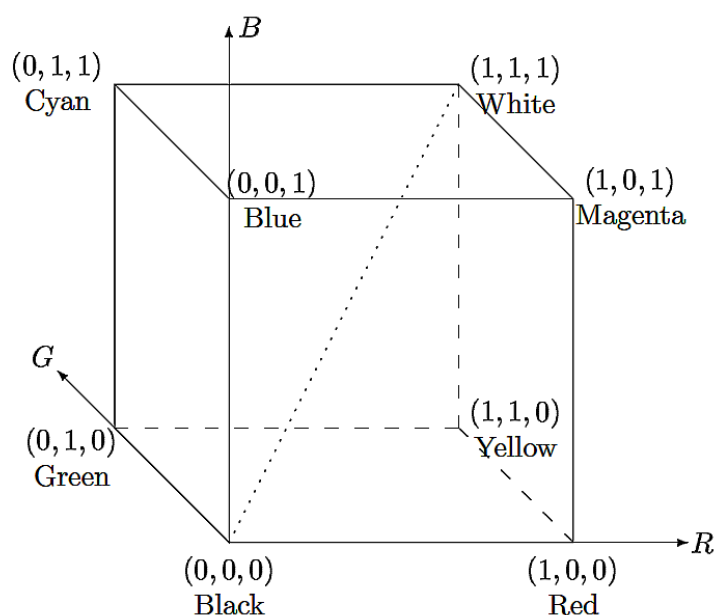


Figura 2.1: Cubo de colores para el modelo de color RGB.

Indexada: La mayoría de las imágenes a color solo tienen un pequeño subconjunto de los más de dieciséis millones de colores posibles. Para facilitar el almacenamiento y la gestión de archivos, la imagen tiene asociado un mapa de colores, o paleta de

colores, que es simplemente una lista de todos los colores utilizados en esa imagen. Cada píxel tiene un valor que no da su color, sino un índice del color en el mapa. Esto es conveniente si una imagen tiene 256 colores o menos, ya que entonces los valores del índice sólo requerirán un byte cada uno para almacenarlo. Algunos formatos de archivo de imagen, sólo permiten 256 colores o menos en cada imagen, precisamente por esta razón.

2.2.2. Pre-procesamiento digital de imágenes

El pre-procesamiento digital se lleva a cabo una vez cuantificados y codificados los valores de una imagen. Este consiste en obtener una nueva imagen que mejore su calidad o resalte algún atributo primario de los objetos capturados. En términos de calidad, tratará de corregir cualquier posible falta de iluminación, eliminar el ruido o aumentar el contraste en la imagen.

2.2.2.1. Procesamiento puntual

Cualquier operación de procesamiento de imágenes transforma los valores de gris de los píxeles. Sin embargo, las operaciones de procesamiento de imágenes pueden dividirse en tres clases en función de la información necesaria para realizar la transformación éstas son:

Transformación de imágenes: Una transformación representa los valores de los píxeles en alguna otra forma, pero equivalente. Las transformaciones permiten la implementación de algoritmos muy eficientes y poderosos, y al usarlas permiten que toda la imagen sea procesada como un solo gran bloque.

Procesamiento con píxeles vecinos: Para cambiar el nivel de gris de un píxel dado sólo es necesario conocer el valor de los niveles de gris en un pequeño vecindario de píxeles alrededor del píxel dado.

Operaciones puntuales: Otra forma de realizar transformaciones con procesamiento de imágenes es mediante operadores puntuales donde el valor de cada píxel de salida depende sólo del valor del píxel de entrada correspondiente es decir el valor de gris de un píxel se cambia sin tomar en cuenta sus píxeles vecinos; este tipo de operaciones se puede expresar matemáticamente con la relación (2.2.1).

$$q(x, y) = f(p(x, y)) \quad (2.2.1)$$

Donde $q(x, y)$ es el píxel de salida y $p(x, y)$ es el píxel de entrada. f especifica el mapeo del nivel de gris de la entrada al nivel de gris de la salida. La forma en que se transforma la imagen depende de la función ya definida.

Aunque las operaciones puntuales son las más simples, contienen algunas de las operaciones más potentes y ampliamente utilizadas en el procesamiento de imágenes,

éstas son especialmente útiles en el preprocesamiento de imágenes, donde se requiere modificar una imagen antes de realizar el trabajo principal.

2.2.2.2. Operaciones aritméticas

Estas operaciones actúan aplicando una función simple $y = f(x)$ a cada valor gris de la imagen. Por lo tanto $f(x)$ es una función que mapea el rango 0-255 sobre sí misma. Las funciones simples incluyen sumar o restar un valor constante a cada píxel; $y = x \pm C$ o multiplicando cada píxel por una constante; $y = C(x)$ con la finalidad de realizar el aclarado u oscurecimiento de una imagen. En cada caso se tiene que limitar ligeramente la salida para asegurar que los resultados son enteros en el rango 0-255. Esto se logra primero redondeando el resultado para obtener un número entero, y luego recortando los valores por ajuste como se muestra en la relación (2.2.2).

$$y \leftarrow \begin{cases} 255 & \text{si } y > 255, \\ 0 & \text{si } y < 0. \end{cases} \quad (2.2.2)$$

Existe una operación conocida como complemento de una imagen también llamado operador inverso o negativo fotográfico, mediante este se crea una imagen de salida que es la inversa de la imagen de entrada. Esta operación es útil en aplicaciones médicas. Su función de transformación se muestra en la ecuación número (2.2.3).

$$q(x, y) = 255 - p(x, y) \quad (2.2.3)$$

2.2.2.3. Histograma

Este es una función discreta que contabiliza el número de ocurrencias de cada uno de los niveles de gris contenidos en una imagen, el cual se puede representar como la siguiente función de probabilidad (2.2.4):

$$p(i) = \frac{h(i)}{(M)(N)} \quad (2.2.4)$$

Donde $h(i)$ = número de ocurrencia del nivel de gris i en la imagen.

2.2.2.4. Eliminación o reducción de ruido

El brillo residual en exceso en una imagen es considerado como un tipo de ruido, esta percepción visual de la luz reflejada si no se elimina puede alterar las características que se pretenden obtener de un objeto de estudio. El control del brillo en las imágenes se basa en la reducción ó aumento de la luminancia reflejada y el medio más simple es sumar o restar un valor constante a todos los valores de niveles de gris o también la multiplicación o división por un número constante, el computo de estas operaciones se puede representar con las siguientes operaciones aritméticas de los valores de intensidad en los pixeles de entrada de las imágenes (2.2.5):

$$\begin{aligned} g_{out}(i, j) &= k_1 g_{in}(i, j) + b_1 \\ \text{ó} \quad g_{out}(i, j) &= k[g_{in}(i, j) + b] \end{aligned} \quad (2.2.5)$$

Donde k_1, k y b_1, b son parámetros de ganancia y sesgo definidos por el usuario.

2.2.2.5. Ecualización del histograma

Este procedimiento se encarga de mejorar el contraste y consiste en la amplificación del rango de todos los tonos de gris por todo el histograma para resaltar detalles antes no visibles, la ecualización del histograma se establece igualando las funciones de distribución originales de la imagen con la función de distribución deseada esta se muestra en la relación (2.2.6)

$$\begin{aligned} F(r) &= \sum_{i=0}^r p(i) \Rightarrow F(r') = \\ (r' + 1) \frac{1}{I} &= F(r) \Rightarrow r' = F(r) * I - 1 \end{aligned} \quad (2.2.6)$$

2.2.3. Post-procesamiento digital de imágenes

Los procedimientos anteriores son parte del pre-procesamiento sin embargo para la extracción de catacterísticas de imágenes es necesaria la aplicación de métodos de procesamiento digital para lograr el realce de las imágenes, el objetivo principal es destacar los bordes de los objetos, segmentar para obtener regiones de interes, regularizar sus colores, acentuar sus texturas, entre otras acciones.

2.2.3.1. Segmentación

Esta técnica consiste en la agrupación de los píxeles mediante algún criterio de homogeneidad como la iluminación, color, bordes, texturas, y particionar las imágenes para obtener regiones de interés y además empleando principios como similitud, conectividad y discontinuidad entre los píxeles. Existen diversos métodos para segmentar imágenes pero uno de los más comunes es el de umbralización; el principio aplicable es la similitud entre los píxeles pertenecientes a la región de interés y sus diferencias respecto al resto de la imagen la cual corresponde al fondo. La umbralización arroja como resultado una imagen binarizada en donde se etiqueta con "1.^a los píxeles correspondientes a la región de interés y con "0.^a los pertenecientes al resto de la imagen, esta función se muestra en la relación (2.2.7).

$$g(x, y) = \begin{cases} 1 & f(x, y) < T \\ 0 & f(x, y) \geq T \end{cases} \quad (2.2.7)$$

Donde $f(x, y)$ es la función que retorna el nivel de gris del píxel (x, y) , $g(x, y)$ será la imagen binarizada y T el umbral.

2.2.3.2. Método de Otsu

Existen diferentes algoritmos para la selección óptima del umbral sin embargo uno de los más usados en aplicaciones de visión artificial es el método de Otsu (20). Este consiste en la suposición de que la función de densidad del fondo C^f y de los objetos C^o tienen modelo gaussiano, $N(\mu_f, \sigma_f^2)$ y $N(\mu_o, \sigma_o^2)$. Cada grupo estará formado por los niveles de gris fijados por el umbral T (2.2.8).

$$\begin{aligned} C_f &= 0, 1, 2, \dots, T \\ C_o &= T + 1, T + 2, \dots, I - 1 \end{aligned} \quad (2.2.8)$$

El umbral debe de minimizar la suma ponderada de cada una de las varianzas de las clases, C_f y C_o , ya que se supone que conforme la suma de las dos normales se aproxime más al histograma real, las desviaciones serán menores.

Para determinar los coeficientes se toman las probabilidades de cada una de las clases. Considerando un valor fijo de umbral, T , las probabilidades de cada categoría serán (2.2.9):

$$P_{C_f} = \sum_{i=0}^T p_i \quad P_{C_o} = \sum_{i=T+1}^{I-1} p_i \quad (2.2.9)$$

Donde p_i es la probabilidad de la intensidad i en la imagen. Las medias y varianzas de cada grupo corresponderán a (2.2.10):

$$\begin{aligned}\mu_{C_f} &= \frac{1}{P_{C_f}} \sum_{i=0}^T (i)(p_i) \\ \mu_{C_o} &= \frac{1}{P_{C_o}} \sum_{i=T+1}^{I-1} (i)(p_i) \\ \sigma_{C_f}^2 &= \frac{1}{P_{C_f}} \sum_{i=0}^T (i - \mu_{C_f})^2 (p_i) \\ \sigma_{C_o}^2 &= \frac{1}{P_{C_o}} \sum_{i=T+1}^{I-1} (i - \mu_{C_o})^2 (p_i)\end{aligned}\tag{2.2.10}$$

Siendo entonces la varianza ponderada (2.2.11):

$$\sigma_p^2 = (P_{C_f})(\sigma_{C_f}^2) + (P_{C_o})(\sigma_{C_o}^2)\tag{2.2.11}$$

Se recorre todo el rango de niveles de gris, de T igual a 0 a $I - 1$, calculándose la varianza ponderada y se elige el umbral, T , con la finalidad de minimizar el umbral a un valor óptimo.

2.3. Clasificación

La clasificación es un procedimiento que se emplea en muchas áreas de investigación, actualmente con los avances existentes en inteligencia artificial(IA) y desarrollo de nuevos algoritmos de aprendizaje automático los modelos de clasificación son mucho más eficientes y tienen la capacidad de aprender por si mismos. El procedimiento consiste en aceptar un conjunto de características y producir una etiqueta de clase para cada una de estas. Puede haber dos o muchas clases, aunque es habitual producir clasificadores multiclase de clasificadores de dos clases. Los clasificadores se construyen tomando un conjunto de ejemplos etiquetados y usarlos para crear una regla que asigne una etiqueta a cualquier nuevo ejemplo. En un problema general, se tiene un conjunto de datos de entrenamiento (x_i, y_i) ; cada uno de los vectores de características x_i consiste en la medición de las características de diferentes tipos de objetos, y los y_i son etiquetas que dan el tipo de objeto que generó el ejemplo (15). Se debe pensar en un clasificador como una regla por la cual se pasa un vector de características y esta devuelve una etiqueta de clase. Se conocen los costes relativos de etiquetar mal cada clase y se debe crear una regla que pueda tomar cualquier x plausible y asignarle una clase, de tal manera que el coste de etiquetado erróneo esperado sea lo más bajo posible, o al menos tolerable. La elección de la regla de clasificación debe depender del costo de cometer

Tabla 2.1: Algoritmos ensamblados.

Método Clasificador	Función
<i>AdaBoostM1</i>	Genera un clasificador nominal utilizando el método AdaBoostM1
<i>AttributeSelectedClassifier</i>	Reducir la dimensión de datos mediante selección de atributos.
<i>Bagging</i>	Clase para el ensacado de un clasificador para reducir la varianza, trabaja regresión.
<i>ClassificationViaClustering</i>	Un simple meta clasificador que utiliza una agrupación para la clasificación.
<i>ClasificaciónViaRegresión</i>	Realiza la clasificación usando un método de regresión, binarizando la clase y el edifica un modelo de regresión para cada valor.
<i>CostSensitiveClassifier</i>	Clasifica o predice la clase con el costo de clasificación erróneo mínimo esperado en lugar de la clase probablemente uno.
<i>CVParameterSelection</i>	Realiza la selección de parámetros de validación cruzada.
<i>Decorate</i>	Construye conjuntos de clasificadores mediante el uso de ejemplos de entrenamiento artificiales.
<i>FilteredClassifier</i>	Ejecuta un clasificador de los datos filtrados.
<i>IterativeClassifierOptimizer</i>	Optimiza el número de iteraciones realizadas por clasificadores iterativos, usando la validación cruzada en los datos de entrenamiento.
<i>LogitBoost</i>	Realiza una regresión logística aditiva.
<i>MultiClassClassifier</i>	Usa un clasificador de dos clases para un conjunto de datos multiclase.
<i>MultiScheme</i>	Usa validación cruzada para seleccionar un clasificador de varios candidatos.
<i>OrdinalClassClassifier</i>	Aplica algoritmos de clasificación estándar para problemas con valor de clase ordinal.
<i>RandomCommittee</i>	Crea un conjunto de clasificadores de base aleatoria.
<i>RandomSubSpace</i>	Hace un clasificador de árbol de decisión, donde mantiene mayor precisión en los datos de formación y mejor precisión de la generalización.
<i>Stacking</i>	Combina varios clasificadores utilizando el método de apilamiento.
<i>Vote</i>	Combina clasificadores utilizando el promedio de las estimaciones de probabilidad.

un error. Un clasificador de dos clases obtendrá cuatro resultados posibles en donde puede cometer dos tipos de errores el primero un falso positivo, el cual ocurre cuando un ejemplo negativo es clasificado como positivo y el segundo error; un falso negativo el cual ocurre cuando un ejemplo positivo es clasificado como negativo. Tomando en consideración este criterio los algoritmos de clasificación son los encargados de proporcionar dicha regla cada uno de forma diferente pero siempre con la finalidad de minimizar los errores de clasificación y alcanzar porcentajes de clasificación aceptables. En este contexto la idea de mejorar la precisión en la clasificación lleva a la implementación de combinaciones de algoritmos ya existentes, mejorando la estructura de las reglas de los modelos clasificadores.

2.3.0.1. Algoritmos ensamblados

Se categorizan dentro de los algoritmos que tienen la capacidad de aprender automáticamente. La función de estos consiste en mejorar el aprendizaje de los clasificadores convirtiéndolos en aprendices más poderosos (10), actualmente existen un gran número de éstos debido al amplio número de contribuciones en cuanto a desarrollo algorítmico. Para propósitos de conceptualización se proporciona una descripción breve en la Tabla 2.1 y se describe el funcionamiento más detallado de tres de ellos los cuales son de mayor relevancia en esta investigación.

AttributeSelectedClassifier: Este se encarga de seleccionar atributos reduciendo la dimensionalidad de los datos para eliminar datos redundantes y aplicar un clasificador al resultado del conjunto de datos obtenido; esto permite el uso de un método de selección de atributos y un algoritmo de aprendizaje específico como parte de un esquema de clasificación, asegurando que el conjunto de atributos elegidos es seleccionado solamente basado en los datos de entrenamiento y evaluado en los datos de prueba.

FilteredClassifier: Su función es aplicar un filtro al conjunto de datos antes de aplicar un algoritmo de aprendizaje, filtrando los datos de entrenamiento sin filtrar los de prueba, además este clasificador usa los datos previamente filtrados con la finalidad de reducir el conjunto de datos a solamente los de mayor importancia. Este algoritmo se utiliza para la evaluación equitativa de discretización supervisada esto lo realiza solo en los datos de prueba usando intervalos de discretización computados para los datos de entrenamiento que en la práctica equivale al procesamiento de datos completamente frescos.

RandomSubSpace: Construye un ensamble de clasificadores, cada uno entrenado y usando un subconjunto seleccionado aleatoriamente de atributos de entrada. Aparte del número de iteraciones y la inserción aleatoria usada, provee un parámetro de control del tamaño del subconjunto de atributos, reduciendo el tamaño de los conjuntos de datos para una eficiente clasificación.

2.3.0.2. Aprendizaje profundo

El aprendizaje profundo nace para resolver problemas como la comprensión de datos casi a nivel humano cuando es muy difícil extraer características abstractas de alto nivel a partir de datos brutos. El aprendizaje profundo resuelve este problema mediante la introducción de representaciones al aprendizaje de representaciones que se expresan en términos de otras representaciones más simples es decir permitiendo que la computadora construya conceptos complejos a partir de conceptos más simples.

Perceptrón Multicapa: También conocido como red neuronal multicapa con conexiones hacia adelante, este algoritmo pertenece a la categoría de funciones de clasificación, las cuales se desarrollan como ecuaciones matemáticas de una forma natural, tiene como base la función de regresión y clasificación, es un modelo de aprendizaje profundo por excelencia, el objetivo de este algoritmo es aproximar una función f^* . Un ejemplo claro para un clasificador $y = f^*(x)$ mapea un valor de x a una categoría y . Una red feedforward define un mapeo $y = f(x; \theta)$ y aprende los valores del parámetro θ el cual resulta en la mejor función de aproximación. En este algoritmo la información fluye a través de la función que ha sido evaluada desde x a través del computo intermediario para definir f hasta llegar a la salida y . Esta clase de red neuronal se compone de una capa de entrada, capas ocultas intermedias y una capa de salida, en donde cada conexión tiene una ponderación (peso) y cada nodo realiza una suma ponderada de sus entradas y umbralizando el resultado; no contiene ningún ciclo y la salida de la red depende solamente de las instancias actuales de entrada su ventaja es que tiene la capacidad de aprender a ignorar atributos irrelevantes y representa límites no lineales de decisión, su algoritmo [2.3.0.2](#) y su representación gráfica muestra el funcionamiento de un modelo básico para una sola unidad de salida y empleando una capa oculta en la [Figura 2.2](#).

```

1: Inicializar  $w$  y  $b$  a cero
2: repeat
3:   for  $l \in 1..L$  do
4:     if  $yl(w \times xl + b) \leq \beta$  then
5:       for  $i \in 1..N$  do
6:         if  $|vi \times xl + di| \leq 1$  then
7:            $vi \leftarrow vi + \lambda ylxl$ 
8:            $di \leftarrow di + \lambda yl$ 
9:         end if
10:      end for
11:       $b \leftarrow b + \lambda yl$ 
12:    end if
13:  end for
14: until criterio de finalización

```

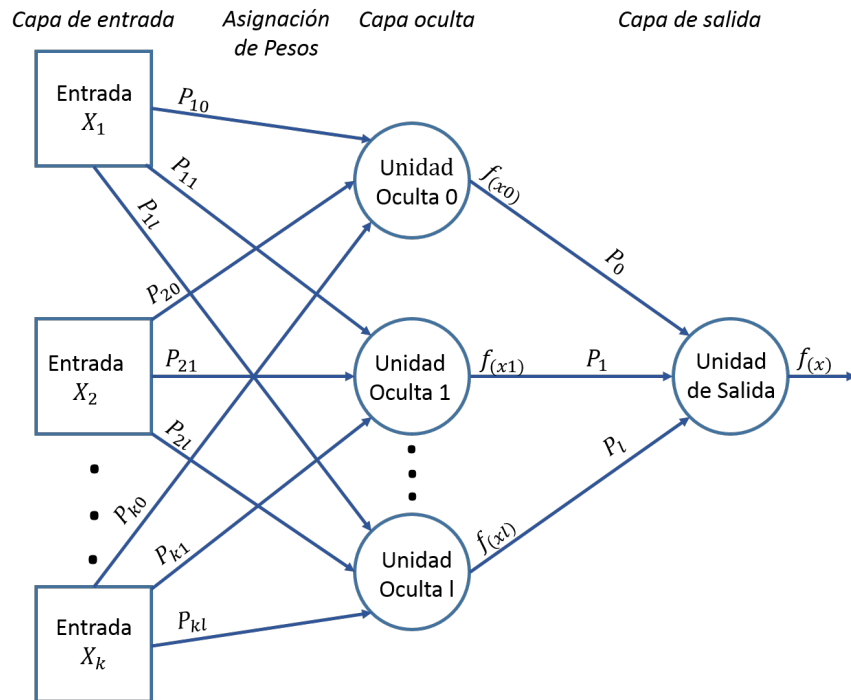


Figura 2.2: Representación gráfica de MLP.

El proceso aplicado con *Multi-layer Perceptron*, particularmente en este caso de estudio es el siguiente:

1. Entrada de datos: $D = \{x_0, x_1, x_2, \dots, x_k\}$ conjunto de cada columna, es decir perteneciente a cada uno de los atributos.
2. Asignación de pesos: Se aplica una inicialización normalizada mediante método de Xavier (8) la cual se basa en el siguiente modelo heurístico:

$$Var[W^L] = \frac{2}{n_{out} + n_{in}}$$

3. Se realiza la sumatoria de los pesos de los datos en la capa oculta:
$$x_i = \sum_j p_{ij} x_j$$
4. Se emplea la función de activación *Sigmoide* en la capa oculta en este caso es factible debido a que todos los valores de entrenamiento son positivos:
$$f(x) = \frac{1}{1+e^{-x}}$$
5. Se obtienen las funciones de salida, $f(x) = y_0, y_1, y_2, \dots, y_k$ las cuales pertenecen a las 4 clases de clasificación de daño foliar: saludable, clorosis, deficiencia de zinc y puntos grasos.

Capítulo 3

Estado del arte

Actualmente se han realizado investigaciones relativas a la identificación de enfermedades de cítricos y al análisis de defectos foliares, cada vez con un mayor enfoque en el área computacional, en esta sección se presentan los trabajos previos relacionados para sustentar esta proyecto; los artículos de éstas investigaciones previas se muestran a continuación:

Pourreza et al., en una de sus investigaciones desarrollaron un sensor de visión aprovechando la rotación del plano de polarización de la luz en algunas longitudes de onda para la identificación de síntomas de HLB; los resultados confirmaron que la iluminación polarizada de banda estrecha a 591 nm aumentó significativamente la precisión del diagnóstico en la clasificación y detección de HBL (16).

Camponetz et al., aplicaron un método de diagnóstico para la detección de presencia de HLB y CVC en hojas de árboles de naranja dulce mediante espectroscopia infraroja con reflectancia total atenuada de la transformada de Fourier y el clasificador inducido mediante regresión parcial de mínimos cuadrados. El modelo consideró cuatro clases de hojas, una con hojas saludables, con CVC sintomático, con HLB sintomático y con HBL asintomático en donde los resultados obtenidos arrojaron un porcentaje del 93.8 % en la correcta identificación de los cuatro tipos de hojas estudiados (3).

Pydipati et al., propusieron un modelo empleando un método de co-ocurrencia (CCM) para determinar si el tono de color, la saturación y la intensidad (HSI) basados en características de textura en conjunción con algoritmos de clasificación estadística podrían ser utilizados para identificar las hojas de cítricos enfermas y normales bajo condiciones de laboratorio. Se evaluaron muestras de hojas de cítricos normales y enfermas estas últimas con puntos grasos, melanosis y escamosis. El análisis discriminante de la muestra de hojas con características de textura CCM logró precisiones de clasificación de más del 95 % para todas las clases cuando se utilizaron la tonalidad de color y la saturación de las características de textura (17).

Ceballos et al., desarrollaron y una técnica optimizada basada en el método GC-MS para el análisis de metabolitos no alcanzables en hojas de cítricos. Las muestras de HLB y deficiencia de zinc de hojas de árboles de naranja dulce fueron sometidas al proceso de micro extracción en fase sólida y tratamientos de derivatización anteriores al análisis GC-MS. El análisis de componentes principales alcanzo una correcta clasificación de todos los derivados de extractos líquidos empleando biomarcadores para la rápida diferenciación de HLB y deficiencia de zinc (4).

Sankaran et al., mostraron en su investigación un enfoque en la aplicación de espectroscopia en el infrarojo medio para la detección de HLB. Se emplearon muestras de hojas saludables, con HLB y deficiencia de nutrientes en árboles mediante dos formas de procesamiento y analizandolas usando un espectrómetro portable y robusto. Se utilizaron datos de la absorvancia espectral en el rango de los 5.15–10.72nm y se

procesaron mediante corrección de alieamiento básica, corrección de datos atípicos y remoción de datos de las longitudes de banda de la absorvancia del agua. El primero y el segundo proceso derivativo fueron calculados mediante el método de Savitzky–Golay y analizados mediante análisis de componentes principales y los records de este fueron clasificados usando el algoritmo de analisis discriminante cuadrático (QDA) y el algoritmo de vecinos más cercanos (KNN), determinando que el algoritmo KNN obtuvo por arriba de 95 % en comparación con el algoritmo QDA, en este caso los datos analizados en bruto proporcionaron mejor precisión en la clasificación comparados con los datos derivativos del procesamiento uno y dos (28).

Mohanty et al., trabajaron en el desarrollo de un clasificador de imagenes de alta presición con la finalidad de diagnosticar enfermedades de plantas empleando una base de datos del proyecto de recolección de hojas con presencia de enfermedades y saludables llamado PlantVillage del cual se emplearon 54,306 imágenes en la clasificación de 26 enfermedades distintas de 16 especies de cultivos mediante el enfoque de una red neuronal convolucional. El desempeño de los modelos de clasificación fue medido basado en su habilidad de predecir los pares correctos cultivo-enfermedad, dadas 38 posibles clases. El modelo con mejor desempeño alcanzo un puntaje medio F1 de 0.9934 es decir una precisión del 99.35 % demostrando la factibilidad técnica de este enfoque (14).

Kadir et al., propusieron un método que captura información de color como un aspecto importante para la identificación de enfermedades en hojas. En esta investigación se incorporaron características de forma geométrica, color y textura de cada hoja; empleando una Red Neuronal Probabilística (PNN) como clasificador. Lo resultados experimentales mostraron que el método de clasificación elegido obtuvo una precisión media del 93.75 % cuando fue probado con la base de datos Flavia la cual contiene 32 tipos de hojas de plantas (1).

García-Cruz, et al. propusieron una investigación para analizar imágenes digitales de hojas de frijol "(Phaseolus vulgaris L.)" var. Cacahuete para la identificación de deficiencias de Fe y Mn mediante la aplicación de redes neuronales estadísticas, en etapas iniciales cuando aun es posible revertir los daños mediante fertilización. Las variables empleadas fueron los valores promedio de ocho variables de color y tres de textura, sobtenidas de seis muestras de imágenes digitales de 100X100 píxeles (360 muestras en total), de hojas de frijol obtenidas 74 dds las cuales se emplearon como variable de entrada para generar clasificadores con redes neuronales probabilísticas con el algoritmo de correlación en cascada de los tratamientos de deficiencias de Fe y Mn en donde los clasificadores que incluyeron características de textura y color, y seis clases de salida o tratamientos de deficiencias obtuvieron un porcentaje del 76.6 % comparado con la inclusión de características de textura que solo alcanzó un 44 % (7).

Dheep Al Bashish, et al., desarrollaron un software para la identificación y clasificación automática de enfermedades de hojas de plantas con la finalidad de proveer

una solución basada en procesamiento de imágenes, esta metodología se conformó de 4 fases principales, en la primera fase crearon una estructura de transformación de color para la imagen de hojas en RGB, después se aplicó una transformación en el espacio de color para la estructura de transformación de color. La segunda fase consistió en la segmentación mediante la técnica de clusterización k-means; en la tercera fase se calcularon las características de los objetos infectados segmentados. En la cuarta fase las características extraídas fueron pasadas a través de una red neuronal previamente entrenada. Como proceso de pruebas se empleó un conjunto de imágenes de hojas del área Al-Ghor de Jordania. El clasificador de red neuronal basado en clasificación estadística operó bien en todas las muestras de los tipos de enfermedades de hojas y logró la detección y clasificación con una precisión de alrededor del 93 % en las enfermedades examinadas (2).

Mari Oide, et al., propusieron un método cuantitativo para la discriminación práctica empleando redes neuronales en la cual utiliza como sus entradas la forma de hojas de soya. En este estudio se aplicó el modelo Hopfield y un perceptrón simple para discriminación varietal de las formas individuales de folíolos de 364 hojas de soya de 38 variedades. En la examinación de hasta 10 variedades el error de discriminación de la red neuronal con la imagen como entrada y mediante validación cruzada fue satisfactoriamente baja por lo cual concluyeron que el modelo funciona de manera adecuada para la discriminación varietal cuantitativa de folíolos de soya; la ventaja de no requerir ni la definición ni la extracción de ninguna característica de forma de las hojas determina que este modelo es ampliamente aplicable a otros casos (13).

WU, Stephen Gang, et al., realizaron un trabajo en el cual emplearon una red neuronal probabilística para el reconocimiento automatizado de hojas para la clasificación de plantas, aplicando técnicas de procesamiento de imágenes y de bases de datos. Se extrajeron 12 características foliares las cuales se ortogonalizaron en 5 variables principales que constituyeron el vector de entrada de la PNN. Esta PNN es entrenada fue entrenada con 1800 muestras de hojas para la clasificación de 32 clases de plantas con una precisión de más del 90 % (33).

Capítulo 4

Materiales y métodos

4.1. Materiales

Se llevó a cabo la recolección de hojas enfermas con presencia de clorosis, deficiencia de Zinc, mancha gracieta y hojas saludables de árboles de naranja dulce Valencia (*Citrus Sinensis*), se empleó una cámara con calidad de 5 megapíxeles y una resolución de 2592 x 1456 para tomar fotografías de cada una de ellas en un ambiente controlado (Manteniendo la misma cantidad de iluminación en la toma de fotos) con la finalidad de reducir el brillo en las imágenes, las fotografías obtenidas se recopilaron en una base de datos, 60 imágenes de hojas saludables, 60 con clorosis sintomática, 60 con deficiencia de Zinc y 60 con mancha gracieta, las muestras de estas hojas se muestran en la Figura 4.1).

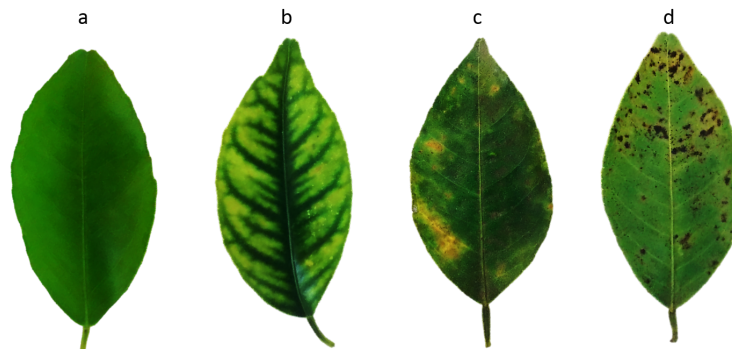


Figura 4.1: Muestras. a) Saludable b) Deficiencia de Zinc c) Clorosis d) Mancha gracieta.

4.2. Métodos

4.2.1. Adquisición de base de datos

Una vez recopiladas las imágenes fue necesaria la aplicación de pre-procesamiento digital de imágenes para el mejoramiento de las imágenes antes de realizar la extracción de características. Para cumplir este objetivo se empleó el histograma de frecuencia de niveles de gris, eliminación de brillo residual, ecualización del histograma para mejora del contraste, una vez realizado estos procedimientos el siguiente fue el post-procesamiento el cual consistió en la segmentación de la imagen mediante método de Otsu con la finalidad de obtener la región de interés (ROI) como se muestra en la Figura 4.2.

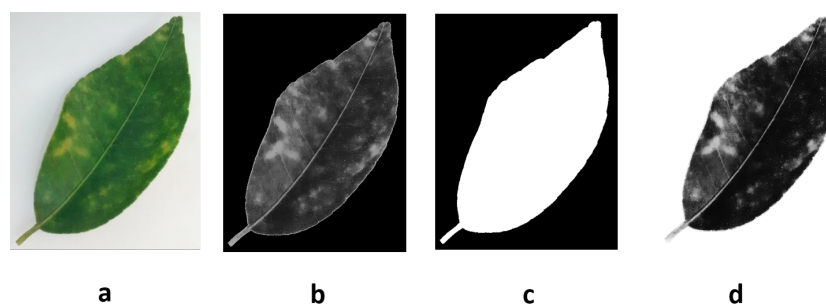


Figura 4.2: Procesamiento. a) Original b) Escala de grises c) Segmentada d) ROI.

La creación de la base de datos tiene la finalidad de recopilar los atributos necesarios para la clasificación de daño foliar, primero; se creó una base de datos general en la cual se adjuntaron 62880 datos totales; de esta se obtuvieron 6 bases de datos binarias para realizar clasificación supervisada mediante la herramienta Weka; las cuales se conformaron con los pares: saludables-clorosis, saludables-def. de zinc, saludables-puntos grasos, clorosis-def. de zinc, clorosis-puntos grasos y def. de zinc-puntos grasos. Después se crearon dos bases de datos multiclase primero considerando solo las 3 enfermedades para un séptimo conjunto y para el octavo se agregó la clase de hojas saludables a las 3 clases de hojas enfermas. Las columnas de la base de datos contiene los atributos de cada hoja iniciando con los valores de cada pixel del histograma de frecuencia de niveles de gris desde 0 hasta 255, es decir ocupa las primeras 256 columnas; las últimas 6 columnas son ocupadas por características estadísticas de intensidad del histograma de frecuencia incluyendo la media, desviación estándar, curtosis, índice de asimetría, media laplaciana y gradiente medio. Por último los nombres de las clases se ubican en la última columna con el propósito de evitar errores en su procesamiento. En total se construyeron 8 bases de datos de las cuales las Tablas (4.1) y (4.2) brindan una descripción de las características de las bases de datos generadas y sus clases respectivamente.

Tabla 4.1: Resumen de características de la base de datos.

Características del Histograma de frecuencia	Numero de Características
Valores de pixeles	61440
Media	240
Desviación Estándar	240
Curtosis	240
Índice de Asimetría	240
Media Laplaciana	240
Gradiente medio	240
Total	62880

Tabla 4.2: Conjuntos binarios y multiclase

Combinaciones	Tipo de conjunto	Descripción
{he,chl}	Binario	Healthy and Chlorosis
{he,zd}	Binario	Healthy and Zinc deficiency
{he,gs}	Binario	Healthy and Greasy Spot
{chl,zd}	Binario	Chlorosis and Zinc deficiency
{chl,gs}	Binario	Chlorosis and Greasy Spot
{zd,gs}	Binario	Zinc deficiency and Greasy Spot
{chl,zd,gs}	Multiclase	Chlorosis, Zinc deficiency and Greasy Spot
{he,chl,zd,gs}	Multiclase	Healthy, Chlorosis, Zinc deficiency and Greasy Spot

La base de datos se generó en un archivo CSV para su manipulación como requisito en el procesamiento de los datos su captura se muestra en la Figura 4.3.

IS	IT	IU	IV	IW	IX	IY	IZ	JA	JB	JC
pixel252	pixel253	pixel254	pixel255	media	desvstd	oblicuidad	curtosis	medialap	gradiente	estado
0	0	0	0	101.25	12.806	5.8553	0.54007	0.10854	15.783	healthy
0	0	0	0	81.039	15.249	8.3193	1.566	0.17158	16.744	healthy
0	0	0	0	91.229	15.009	7.1562	1.1015	0.18427	15.638	healthy
0	0	0	0	87.682	13.987	5.8147	0.85637	0.22006	17.68	healthy
0	0	0	0	120.46	12.183	4.8868	0.54016	0.075764	9.1562	healthy
0	0	0	0	93.358	15.144	6.6381	0.74349	0.10688	17.589	healthy
0	0	0	0	43.558	13.451	11.752	2.3081	0.12917	11.139	healthy
0	0	0	0	88.607	12.16	8.2593	1.235	0.10162	12.374	healthy
0	0	0	0	88.399	14.52	9.6434	1.7538	0.1415	15.383	healthy
0	0	0	0	103.77	15.449	8.9751	1.9633	0.14534	11.122	healthy
0	0	0	0	80.366	12.643	9.7998	1.4324	0.16115	12.765	healthy
0	0	0	0	96.444	11.188	7.1004	0.45039	0.14274	17.43	healthy
0	0	0	0	92.632	15.603	6.9893	1.1532	0.15519	15.312	healthy
n	n	n	n	83.542	13.496	7.7703	1.2305	0.10217	13.301	healthy

Figura 4.3: Captura de base de datos creada en CSV.

4.2.2. Tratamiento de los datos

La finalidad de este procedimiento es obtener el patrón de comportamiento de los algoritmos en el proceso de clasificación y el porcentaje de precisión que obtiene cada combinación. El proceso de clasificación de los datos se realizó empleando la combinación de algoritmos ensamblados con el algoritmo deep learning (*multy-layer perceptron*) empleando 3 criterios de evaluación: Validación cruzada con 10 iteraciones, división de porcentaje 2/3 para entrenamiento y 1/3 para pruebas, y por último muestra representativa.

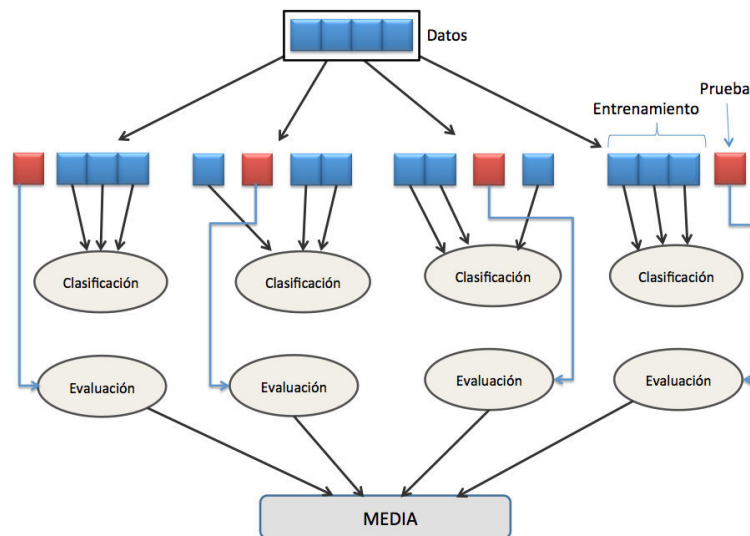
Para el tratamiento de todos los conjuntos se emplearon datos balanceados para evitar métricas sesgadas a la clase con mayor número de datos en el proceso de clasificación, es decir los algoritmos de clasificación tienden a favorecer a la clase con la mayor proporción de observaciones en caso de emplear datos desbalanceados por lo que en todos los criterios de evaluación se empleo el mismo número de instancias en cada clase.

División de Porcentaje: También conocido como método de conjunto de prueba, este elige aleatoriamente el 30 % de los datos para que este pertenezca a un conjunto de pruebas, el resto forma parte de un conjunto de entrenamiento. Para este caso en específico se empleó un porcentaje de 66.5 % para el segmento de datos de entrenamiento y el resto para los segmentos de prueba, para cada uno de los conjuntos a clasificar como se muestra en la Tabla 4.3.

Tabla 4.3: Datos de entrenamiento y prueba para criterio 2/3-1/3.

Conjunto	Entrenamiento	Prueba	Total
Binario	80	40	120
Multiclase(3 clases)	120	60	180
Multiclase(4 clases)	160	80	240

Validación Cruzada: La validación cruzada es una importante técnica estadística que funciona como una simple variante iterativa la cual divide aleatoriamente el conjunto de datos en k particiones equitativas. En la validación cruzada al decidir un número fijo de divisiones o particiones de los datos, cada una de estas partes se utiliza para realizar pruebas y el resto para entrenamiento como se muestra en la Figura 4.4.

**Figura 4.4:** Diagrama de Validación cruzada(De Pablo Castellano)

Para este caso de estudio se usó validación cruzada con 10 iteraciones, realizando la partición en 10 segmentos del total de datos, por lo que nueve décimos de los datos se destinaron para el entrenamiento y un décimo para las pruebas, el proceso consiste en repetir el mismo procedimiento diez veces para que al final cada segmento de datos sea utilizado exactamente una vez en la realización de pruebas. 4.4.

Tabla 4.4: Datos de entrenamiento y prueba para validación cruzada.

Conjunto	Entrenamiento	Prueba	Total
Binario	108	12	120
Multiclase(3 clases)	168	18	180
Multiclase(4 clases)	216	24	240

Muestra representativa: La muestra representativa es la encargada de representar parte de un conjunto total de datos, esta se calcula en base al nivel de confianza para acotar un valor con una alta probabilidad y al margen de error permitido con el cual se determina el número de veces esperado en un rango específico dentro de un conjunto total de datos. Derivado que en nuestro caso de estudio se conoce el total de la población de datos el calculo se realiza con la relación matemática (4.2.1).

$$n = \frac{k^2 \times p \times q \times N}{e^2 \times (N - 1) + k^2 \times p \times q} \quad (4.2.1)$$

Donde:

N : es el tamaño total de la población (número total de datos).

k : es una constante que depende del nivel de confianza asignado. El nivel de confianza indica la probabilidad mínima deseada de que los resultados sean ciertos; un 95.5 % de confianza significa que puede haber un margen de error con una probabilidad del 4,5 %.

e : precisión (error máximo admisible en términos de proporción).

p : es la proporción de la población con la característica objetivo de estudio. Este dato es generalmente desconocido y se suele suponer que $p=q=0.5$ que es la opción más segura.

q : es la proporción de la población que no posee esa característica, es decir, es $1-p$.

n : es el tamaño de la muestra representativa.

Para realizar el cálculo de la muestra representativa se consideró un nivel de confianza de 95 % y un margen de error del 5 %. Para la clasificación binaria se obtuvo un total de 92 instancias que en porcentaje equivale al de 23.33 % destinado para la realización de prueba. Para los casos multiclase en el conjunto 7 se obtuvo un 31.66 % y en el conjunto 8 un 38.3 %. El número de datos para cada conjunto se puede apreciar en la Tabla 4.5.

Tabla 4.5: Datos de entrenamiento y prueba para muestra representativa.

Conjunto	Entrenamiento	Prueba	Total
Binario	92	28	120
Multiclase(3 clases)	123	57	180
Multiclase(4 clases)	148	92	240

Capítulo 5

Experimentos y Resultados

El procedimiento de experimentación consistió en combinar cada uno de los algoritmos ensamblados (metaclasificadores) con el algoritmo de deep learning D14jMlp (*Multi-layer Perceptron*) y evaluar el porcentaje de precisión que obtiene cada combinación mediante el uso de la herramienta WEKA, el modelo de experimentación se muestra en la Figura 5.1.

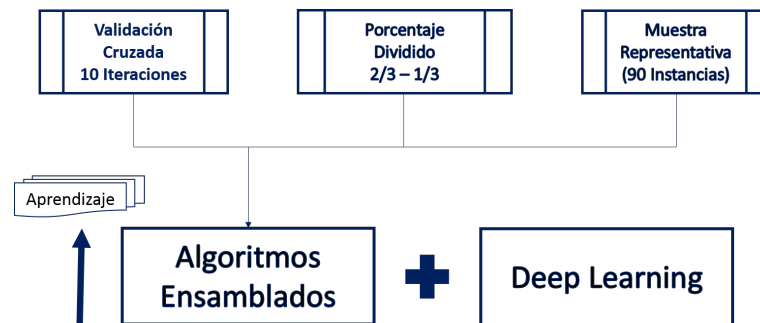


Figura 5.1: Modelo de experimentación realizada de clasificación supervisada.

El primer experimento consistió en emplear cada uno de los algoritmos ensamblados y combinarlos con el algoritmo Deep Learning *Multi-Layer Perceptron* para realizar una clasificación binaria. En esta fase se encontró que los mejores resultados se obtuvieron con el criterio de porcentaje dividido 2/3-1/3 con tres metaclasificadores, el primero AttributeSelectedClassifier, el segundo FilteredClassifier y el tercero RandomSubSpace todos con un porcentaje del 100 % de instancias clasificadas correctamente. Los experimentos de las combinaciones y los resultados se muestran en la Tabla 5.1.

Tabla 5.1: Resultados del primer experimento con el conjunto 1 y clasificación binaria.

Meta-clasificador	Aprendizaje	10- folds	2/3 - 1/3	Muestra
	Profundo	Cross-Validation		representativa
	Dl4jMlp	97.5	97.5	92.2222
AttributeSelectedClassifier	Dl4jMlp	99.1667	100	98.8889
Bagging	Dl4jMlp	50	42.5	44.4444
ClassificationViaClustering	Dl4jMlp	62.5	67.5	
ClassificationViaRegression	Dl4jMlp	50	42.5	44.4444
CVParameterelection	Dl4jMlp	50	42.5	44.4444
FilteredClassifier	Dl4jMlp	99.1667	100	96.6667
LogitBoost	Dl4jMlp	50	57.5	55.5556
MultiClassClasifier	Dl4jMlp	50	42.5	44.4444
MultiScheme		50	42.5	
MultiSearch	Dl4jMlp	50	42.5	44.4444
OrdinalClassClassifier	Dl4jMlp	97.5	97.5	92.2222
RandomCommite	Dl4jMlp	50	42.5	44.4444
RandomizableFilteredClassif	Dl4jMlp	54.1667	60	51.1111
RandomSubSpace	Dl4jMlp	96.6667	100	98.8889
Stacking	Dl4jMlp	50	57.5	44.4444
ThresholdSelector	Dl4jMlp	54	42.5	55.5556
Vote	Dl4jMlp	50		
WeightedInstancesHandlerWrap	Dl4jMlp	50	42.5	44.4444

En la experimentación realizada para los conjuntos multiclase la clasificación alcanzó un porcentaje máximo de 96.66 %. La compilación de los mejores resultados de los experimentos tanto para la clasificación binaria como para las combinaciones multiclase se muestran en la Tabla 5.2.

Tabla 5.2: Mejores resultados de clasificación binaria y multiclase.

Clasificadores	Deep	Criterio de	he-chl	he-zd	he-gs	chl-zd	chl-gs	zd-gs	chl-zd-gs	he-chl-zd-gs
Ensamblados	Learning	Evaluación								
<i>AttSelClassifier</i>	Dl4jMlp	Val. Cruzada			96.66		96.66	99.16	94.44	
<i>FilteredClass</i>	Dl4jMlp	2/3-1/3	100	100				100	96.66	
<i>RamSupSpace</i>	Dl4jMlp	Muestra Rep.								
<i>AttSelClassifier</i>	Dl4jMlp	2/3-1/3	100	100			100			
<i>FilteredClass</i>	Dl4jMlp	Muestra Rep.								93.91
<i>RamSupSpace</i>	Dl4jMlp	Val. Cruzada			96.66	96.66	95.83	98.33	95.55	
<i>AttSelClassifier</i>	Dl4jMlp	Muestra Rep.		100						
<i>FilteredClass</i>	Dl4jMlp	Val. Cruzada		100	95.83		95.83	99.16		93.75
<i>RamSupSpace</i>	Dl4jMlp	2/3-1/3	100	100		100				96.25
<i>OrdinalClassifier</i>	Dl4jMlp	2/3-1/3				100				

Los experimentos se realizaron en un ordenador con: OS Windows 8, Intel(R) Core(TM) i3-2330M CPU 2.20 GHz, RAM 8.00 GB, HD 320 GB, sistema operativo de 64 bits y procesador x64. Además, para el procesamiento de imágenes y la ejecución de los algoritmos descritos en esta sección, se empleó la herramienta de software matemático y entorno de desarrollo integrado (IDE) MATLAB® y el software de minería de datos WEKA v.3.8. respectivamente.

Capítulo 6

Análisis de resultados

Las combinaciones de algoritmos ensamblados con Deep Learning (*Multi-layer perceptron*) fueron determinantes en los resultados de clasificación, debido a que tienen la capacidad de mejorar el aprendizaje, además de realizar un tratamiento previo de los datos antes de realizar el proceso de clasificación. Cada uno tiene una característica única, para este caso de estudio, donde tres fueron los clasificadores de mayor eficiencia y los resultados obtenidos se deben a:

- Que el algoritmo *AttributeSelectedClassifier* realiza una reducción de la dimensionalidad de los datos de entrenamiento al seleccionar los atributos antes de pasarlos al clasificador.
- El algoritmo clasificador *FilteredClassifier* pasa previamente los datos por un filtro arbitrario antes de ser clasificados y por último.
- El algoritmo *RandomSubSpace* construye una clasificación basada en un árbol de decisión la cual logra mantener una alta precisión en el entrenamiento de datos, mejorando la precisión cuanto mas crece la complejidad de los datos. La idea general obtenida de estos resultados se muestra en la Figura (6.1).

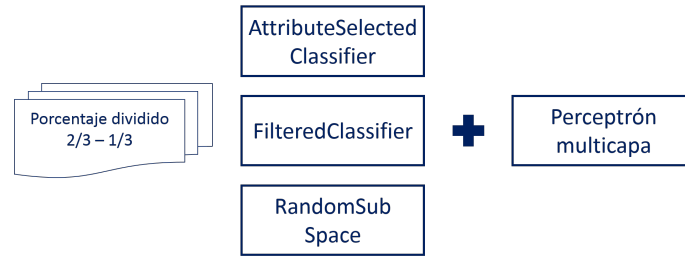


Figura 6.1: Criterio y algoritmos con mejores resultados.

Los resultados obtenidos de las tres combinaciones se representan en una matriz de confusión para cada conjunto, cabe aclarar que para los casos de clasificación binaria de los conjuntos del 1 al 6 las matrices de confusión obtenidas muestran una ligera variación en el comportamiento y criterio de evaluación por lo que solo se muestran las representativas para cada conjunto binario, en los casos de clasificación multiclase para los conjuntos 7 y 8 los algoritmos siguen el mismo patron que en la clasificación binaria pero con un porcentaje ligeramente más bajo. Para visualizar el comportamiento de las instancias correctamente clasificadas se muestran las matrices de confusión en las Tablas 6.1, 6.2, 6.3.

En la clasificación binaria las matrices de confusión de los algoritmos *AttributeSelectedClassifier*, *FilteredClassifier* así como *RandomSubSpace* en combinación con *Multilayer Perceptron* muestran el número total de instancias para clasificar de las cuales se puede observar no se obtuvo ninguna clasificada erróneamente. Es decir, los valores 23 y 17 son los casos de hojas saludables y enfermas pertenecientes al total de instancias de entrada. Por otro lado en la clasificación multiclase mientras que el número de instancias clasificadas incorrectamente fue mínimo se puede apreciar que del

total de instancias si hubo algunas clasificadas erroneamente, indicadas por el valor 1 en las matrices de confusión tanto para el caso de tres clases en donde se clasificaron erroneamente 2 instancias y para el caso de 4 clases en donde el error fue de 3 instancias, aún así se puede apreciar que persiste una reducción en la confusión en la clasificación de dichas instancias debido a la reducción de dimensionalidad y el filtrado previo de los datos, aplicado por los algoritmos ensamblados y la eficiencia de *Deep Learning*.

Tabla 6.1: Matrices de confusión representativas de clasificación de los conjuntos binarios del 1 al 3

Clasificación Binaria								
1. he-chl			2. he-zd			3. he-gs		
% Div. 2/3-1/3			% Div 2/3-1/3			Val. Cruzada 10 Iteraciones		
a	b	Clasificadas como	a	b	Clasificadas como	a	b	Clasificadas como
23	0	a = healthy	23	0	a = healthy	59	1	a = healthy
0	17	b = chlorosis	0	17	b = zinc def.	3	57	b = greasy spot

Tabla 6.2: Matrices de confusión representativas de la clasificación de los conjuntos binarios del 4 al 6

Clasificación Binaria								
4. chl-zd			5. chl-gs			6. zd-gs		
% Div. 2/3-1/3			% Div. 2/3-1/3			% Div. 2/3-1/3		
a	b	Clasificadas como	a	b	Clasificadas como	a	b	Clasificadas como
23	0	a = chlorosis	23	0	a = chlorosis	59	1	a = zinc def.
0	17	b = zinc def.	0	17	b = greasy spot	3	57	b = greasy spot

Los porcentajes de clasificación en comparación con otras investigaciones previas son favorables resaltando que el método de Redes Neuronales Convolucionales (CNN) obtiene resultados mucho más elevados que los resultados obtenidos en esta investigación, sin embargo se debe tomar en consideración que el costo computacional y el consumo de recursos para el caso de CNN es muy elevado (19); mientras que los métodos empleados en esta investigación reducen el costo computacional y el consumo de recursos se mantiene dentro de los estandares básicos. Los porcentajes de cada investigación y el método se pueden apreciar en la Tabla 6.4

Tabla 6.3: Matrices de confusión representativas de la clasificación de los conjuntos multiclase 7 y 8

Clasificación Multiclase									
7. chl-zd-gs					8. he-chl-zd-gs				
% Div. 2/3-1/3					% Div. 2/3-1/3				
a	b	c	Clasificadas como		a	b	c	d	Clasificadas como
24	0	0	a = chlorosis		20	0	0	1	a = healthy
0	18	1	b = zinc def.		0	20	1	1	b = chlorosis
1	0	16	c = greasy spot		0	0	17	0	c = zinc def.
					0	0	0	20	d = greasy spot

Tabla 6.4: Comparativa con el estado del arte del % de clasificación de características foliares

Autor	Método de Clasificación	Porcentaje máximo obtenido
<i>WU, Stephen Gang, et al.</i>	Red Neuronal Probabilística	90 %
<i>Dheep Al Bashish, et al.</i>	Red Neuronal Probabilística	93 %
<i>Kadir et al.</i>	Red Neuronal Probabilística	93.75 %
<i>Camponoz et al.</i>	Transformada de Fourier	93.8 %
<i>Pydipati et al.</i>	Análisis discriminante	95 %
<i>Sankaran et al.</i>	KNN	95 %
<i>Esta investigación</i>	FilteredClassifier + Deep Learning	96.66 %
<i>Mohanty et al.</i>	Red Neuronal Convolutacional	99.35 %

Capítulo 7

Conclusiones y trabajo futuro

7.1. Conclusiones

La obtención de porcentajes de clasificación en su mayoría de 100 % en la clasificación de conjuntos binarios nos demostró que este enfoque es muy eficiente para conjuntos con diferenciación linear. En cuanto a la clasificación multiclase, el enfoque muestra una mejora satisfactoria con respecto a los porcentajes obtenidos en investigaciones previas, en donde se emplearon diferentes técnicas como: método de vecinos más cercanos, transformada de Fourier y redes neuronales probabilísticas en su mayoría. Los algoritmos *filteredclassifier*, *attributeselecterclassifier*, *randomsubspace* y en ultima instancia un nuevo candidato el algoritmo *ordinalclassifier* siguen un patrón de mejora en la precisión de clasificación cuando se combinan con el algoritmo Deep Learning *Multy-Layer Perceptron*. El porcentaje máximo de clasificación en conjuntos multiclase del 96.66 % obtenido en esta investigación solamente fue superado por el uso de Redes Neuronales Convolucionales pero a costa de un alto consumo de recursos y coste computacional por lo que el enfoque de combinar algoritmos para maximizar la precisión de clasificación es más factible que invertir en equipos costosos.

La mejora en general del método aquí aplicado se debió a la combinación de factores como: las técnicas de procesamiento digital de imágenes, la combinación de algoritmos de aprendizaje automático y el tamaño ligero de la base de datos, además de reducción de dimensionalidad y filtrado de datos aplicado a la base de datos por los algoritmos ensamblados por lo que se considera que debido al comportamiento de estos algoritmos, estos son lo indicados para la clasificación de características de daño foliar de cítricos y que las técnicas pueden ser aplicadas en la clasificación de cualquier conjunto de características que cumplan con el mismo enfoque.

7.2. Trabajo futuro

Para la continuación futura de esta investigación se propone realizar las siguientes acciones:

- Ampliar la base de datos de imágenes agregando más enfermedades de cítricos.
- Seleccionar los mejores algoritmos de clasificación para realizar pruebas con bases de datos más complejas.
- Aumentar e implementar las características usadas en esta investigación para realizar una clasificación no supervisada.
- En base a los algoritmos de reconocimiento de patrones empleados en el procesamiento digital de imágenes para la extracción de características y a los algoritmos de aprendizaje automático para la clasificación; construir un sistema experto.
- Desarrollar un sistema experto, para el usuario final, en una segunda fase de investigación.

Bibliografía

- [1] Kadir Abdul, Nugroho Lukito Edi, Susanto Adhi, and Santosa Paulus Insap. Leaf classification using shape color and texture features. *arXiv preprint arXiv:1401.4447*, 2013. [23](#)
- [2] Dheeb Al Bashish, Malik Braik, and Sulieman Bani-Ahmad. Detection and classification of leaf diseases using k-means-based segmentation and. *Information Technology Journal*, 10(2):267–275, 2011. [2](#), [24](#)
- [3] Marcelo Camponez, Paulino Ribeiro Villas Boasa, Débora Marcondes Bastos Pereira Miloria, Ednaldo José Ferreira, Marina Franca e Silva, Marcos Antonio Machadoc, Barbara Sayuri Belleted, and Maria Fatima das Gracias Fernandes da Silva. Infrared spectroscopy: A potential tool in huanglongbing and citrus variegated chlorosis diagnosis. *Talanta*, 91(1):1–6, 2012. [2](#), [22](#)
- [4] Juan Manuel Cevallos-Cevallos, Rosalía García-Torres, Edgardo Etxeberria, and José Ignacio Reyes-De-Corcuera. Gc-ms analysis of headspace and liquid extracts for metabolomic differentiation of citrus huanglongbing and zinc deficiency in leaves of ‘valencia’ sweet orange from commercial groves. *Phytochemical Analysis*, 22(3):236–246, 2010. [2](#), [22](#)
- [5] Ronan Collobert and Samy Bengio. Links between perceptrons, mlps and svms. In *Proceedings of the twenty-first international conference on Machine learning*, page 23. ACM, 2004.
- [6] David Serrano, Esther Serrano, Megan Dewdney, Christina Southwick. Citrus variegated chlorosis (cvc). <http://idtools.org/id/citrus/diseases/factsheet.php?name=Citrus+variegated+chlorosis+>
- [7] Edgar García-Cruz, Manuel Sandoval-Villa, José A Carrillo-Salazar, Jorge M Valdéz-Carrasco, and Paulina H González-Fierro. Identificación con redes neuronales probabilísticas de las deficiencias de hierro y manganeso, usando imágenes digitales de hojas de frijol (*phaseolus vulgaris* l.). *Agrociencia*, 49(4):395–409, 2015. [2](#), [23](#)

- [8] Xavier Glorot and Yoshua Bengio. Understanding the difficulty of training deep feedforward neural networks. In *Proceedings of the thirteenth international conference on artificial intelligence and statistics*, pages 249–256, 2010. [20](#)
- [9] Aaron Courville Ian Goodfellow, Yoshua Bengio. *Deep Feedforward Networks*. MIT Press, 2016.
- [10] Mark A. Hall Ian H. Witten, Eibe Frank. *Data Mining*. Morgan Kaufmann, 2011. [17](#)
- [11] Liberato JR, Queiroz-Voltan, RB, Matsuoka K, Laranjeira FF, Miles AK. Citrus variegated chlorosis (xylella fastidiosa). <http://www.padil.gov.au/pests-and-diseases/pest/main/136652>, 2018.
- [12] Pedro Huamaní Navarrete. Using multiple thresholding otsu method to recognize the red light on traffic lights. *Scientia*, 17(17):247–242, 2015.
- [13] Mari Oide and Seishi Ninomiya. Discrimination of soybean leaflet shape by neural networks with image input. *Computers and electronics in agriculture*, 29(1-2):59–72, 2000. [2](#), [24](#)
- [14] Mohanty Sharada P, Hughes David P, and Salathé Marcel. Using deep learning for image-based plant disease detection. *Frontiers in plant science*, 7:1419, 2016. [2](#), [23](#)
- [15] Jean Ponce, David Forsyth, Equipe-projet Willow, S Antipolis-Méditerranée, R d’activité RAweb, L Inria, and I Alumni. Computer vision: a modern approach. *Computer*, 16(11), 2011. [10](#), [16](#)
- [16] Alireza Pourreza, Won Suk Lee, Ed Etxeberria, and Arunava Banerjee. An evaluation of a vision based sensor performance in huanglongbing disease identification. *Biosystems Engineering*, 130(1):13–22, 2015. [2](#), [22](#)
- [17] R. Pydipati, T.F. Burks, and W.S. Lee. Identification of citrus disease using color texture features and discriminant analysis. *Computers and Electronics in Agriculture*, 52(2):49–59, 2006. [2](#), [22](#)
- [18] M. B. Riley, M. R. Williamson and O. Maloy, Plant disease diagnosis, *The plant health instructor* (2002). [2](#)
- [19] Pablo Rodríguez-Sahagún Alesanco. Aplicación de redes neuronales convolucionales y recurrentes al diagnóstico de autismo a partir de resonancias magnéticas funcionales. 2018. [39](#)
- [20] Prasanna K Sahoo, SAKC Soltani, and Andrew KC Wong. A survey of thresholding techniques. *Computer vision, graphics, and image processing*, 41(2):233–260, 1988. [15](#)

- [21] Secretaría de Agricultura, Ganadería, Desarrollo Rural, Pesca y Alimentación (SAGARPA). Se consolida México como quinto productor mundial de naranja. <https://www.gob.mx/sagarpa/prensa/se-consolida-mexico-como-quinto-productor-mundial-de-naranja?idiom=es>, 2017. 3
- [22] Servicio de Información Agroalimentaria y Pesquera (SIAP). Cítricos: Limón, naranja y toronja mexicanos. planeación agrícola nacional 2016-2030. [https://www.gob.mx/cms/uploads/attachment/file/257073/Potencial-Citricos – parteuno.pdf](https://www.gob.mx/cms/uploads/attachment/file/257073/Potencial-Citricos--parteuno.pdf), 2016. 3
- [23] Servicio de Información Agroalimentaria y Pesquera (SIAP). Avance de siembras y cosechas resumen nacional por cultivo. [http://infosiap.siap.gob.mx:8080/agricola_siap_gobmx /AvanceNacionalSinPrograma.do](http://infosiap.siap.gob.mx:8080/agricola_siap_gobmx/AvanceNacionalSinPrograma.do), 2018. 3
- [24] Servicio Nacional de Sanidad, Inocuidad y Calidad Agroalimentaria. Ficha técnica: Clorosis variegada de los cítricos. https://www.gob.mx/cms/uploads/attachment/file/215850/04_Ficha_Tecnica-_Clorosis_Variegada_de_los_Citricos.pdf, 2017. 8
- [25] Servicio Nacional de Sanidad, Inocuidad y Calidad Agroalimentaria. Huanglongbing de los cítricos. <https://www.gob.mx/senasica/documentos/huanglongbing-de-los-citricos-110925>, 2017. 8
- [26] Servicio Nacional de Sanidad Inocuidad y Calidad Agroalimentaria. Plagas reglamentadas de los cítricos (x. fastidiosa subsp. pauca, x. citri subsp. citri, g. citricarpa y clv-c). <http://sinavef.senasica.gob.mx/SIRVEF/GuiasSintomas.aspxs>, 2018. 8
- [27] Servicio Nacional de Sanidad, Inocuidad y calidad Agroalimentaria (SENASICA). Cancro de los cítricos. https://www.gob.mx/cms/uploads/attachment/file/215782/07_Ficha_Tecnica-_Cancro_de_los_citricos.pdf, 2016. 9
- [28] Sankaran Sindhuja, Ehsani Reza, and Etxeberria Edgardo. Mid-infrared spectroscopy for detection of huanglongbing in citrus leaves. *Talanta*, 83(2):574–581, 2010. 2, 23
- [29] Sistema Nacional Argentino de Vigilancia y Monitoreo de plagas. Citrus psorosis virus (cpsv). <https://www.sinavimo.gov.ar/plaga/citrus-psorosis-virus-cpsv>, 2018. 9
- [30] Da-Wen Sun. *Image Processing*. Nikki Levy, 2016.
- [31] Suárez Pérez Rosendo, Hernández Alvides Jorge, Serrano Romero Jorge, de Armas Arredondo Georrgina. Mancha grasienta del cítrico. [https://www.ecured.cu/Mancha-grasienta-del.c](https://www.ecured.cu/Mancha-grasienta-del-c)9
- [32] Alberto Uc-Vázquez, Daniel Leobardo Ochoa-Martínez, Elizabeth Cárdenas-Soriano, and Gustavo Mora-Aguilera. Sintomatología e histopatología del amarillamiento letal

de la lima persa citrus latifolia tanaka. *Revista Mexicana de Fitopatología*, 23(2), 2005.

[9](#)

- [33] Stephen Gang Wu, Forrest Sheng Bao, Eric You Xu, Yu-Xuan Wang, Yi-Fan Chang, and Qiao-Liang Xiang. A leaf recognition algorithm for plant classification using probabilistic neural network. In *Signal Processing and Information Technology, 2007 IEEE International Symposium on*, pages 11–16. IEEE, 2007. [2](#), [24](#)

- [34] Mongi Zekri and Tom Obreza. Manganese (mn) and zinc (zn) for citrus trees. *Citrus Tree Nutrient series*, 23(SL403), 2016. [9](#)