



CIENCIA ergo-sum, Revista Científica
Multidisciplinaria de Prospectiva
ISSN: 1405-0269
ISSN: 2395-8782
ergosum@uaemex.mx
Universidad Autónoma del Estado de México
México

Identificación de las principales enfermedades de la planta del café (*Coffea arabica*) a través de visión artificial

 Flores Colorado, Oscar Eder

 Cervantes Canales, Jair

 García-Lamont, Farid

 Ruiz Castilla, José Sergio

Identificación de las principales enfermedades de la planta del café (*Coffea arabica*) a través de visión artificial

CIENCIA ergo-sum, Revista Científica Multidisciplinaria de Prospectiva, vol. 30, núm. 3, 2023

Universidad Autónoma del Estado de México

Disponible en: <https://www.redalyc.org/articulo.oa?id=10475688011>

DOI: <https://doi.org/10.30878/ces.v30n3a8>

Atribución — Usted debe dar crédito de manera adecuada, brindar un enlace a la licencia, e indicar si se han realizado cambios. Puede hacerlo en cualquier forma razonable, pero no de forma tal que sugiera que usted o su uso tienen el apoyo de la licenciante. NoComercial — Usted no puede hacer uso del material con propósitos comerciales. SinDerivadas — Si remezcla, transforma o crea a partir del material, no podrá distribuir el material modificado. No hay restricciones adicionales — No puede aplicar términos legales ni medidas tecnológicas que restrinjan legalmente a otras a hacer cualquier uso permitido por la licencia.



Esta obra está bajo una Licencia Creative Commons Atribución-NoComercial-SinDerivar 4.0 Internacional.

Ciencias Exactas y Aplicadas

Identificación de las principales enfermedades de la planta del café (*Coffea arabica*) a través de visión artificial

Identification of the main diseases of the coffee plant (*Coffea arabica*) through artificial vision

Oscar Eder Flores Colorado
Universidad Autónoma del Estado de México, México
flores.eder.93@gmail.com

 <https://orcid.org/0000-0002-5913-0526>

Jair Cervantes Canales
Universidad Autónoma del Estado de México, México
jcervantesc@uaemex.mx

 <https://orcid.org/0000-0003-2012-8151>

Farid García-Lamont
Universidad Autónoma del Estado de México, México
fgarcial@uaemex.mx

 <https://orcid.org/0000-0002-9739-3802>

José Sergio Ruiz Castilla
Universidad Autónoma del Estado de México, México
jsruizc@uaemex.mx

 <https://orcid.org/0000-0001-7821-4912>

DOI: <https://doi.org/10.30878/ces.v30n3a8>

Recepción: 15 Julio 2021

Aprobación: 01 Febrero 2022



Acceso abierto diamante

Resumen

Se utilizan técnicas de reconocimiento de patrones para identificar hojas sanas y cuatro enfermedades de la planta del café *Coffea arabica*. Las enfermedades son la roya del café, el minador de la hoja, phoma quema y *Cercospora coffeicola*. Para lograrlo, se ocuparon diferentes técnicas de segmentación, entre ellas Otsu, PCA y método de frontera global. Con el fin de obtener el vector de características, las imágenes se procesaron para extraer las características cromáticas, geométricas y textuales. Por último, se implementaron cuatro algoritmos de clasificación entre los que se encuentran *support vector machine*, *random forest*, Naive Bayes y redes neuronales artificiales *backpropagation*. La mejor precisión obtenida es del 83% con segmentación Otsu y clasificación con redes neuronales artificiales *backpropagation*.

Palabras clave: reconocimiento de patrones, visión artificial, enfermedades y plagas del café.

Abstract

Pattern recognition techniques in digital images are applied to identify healthy leaves and four important diseases of the coffee plant *Coffea arabica*. The diseases are coffee rust, leaf miner, *Phoma* leaf spot and *Cercospora coffeicola*. To achieve this, different segmentation techniques were used, among them: Otsu, PCA and Global Border Method. To get the features vector, the images were processed to extract the Chromatic, geometric and textural features. Finally, four classification algorithms were implemented, including support vector machine, *random forest*, Naive Bayes and Artificial neural networks backpropagation. The best accuracy obtained is 83% with Otsu segmentation and classification with backpropagation Artificial neural networks.

Keywords: pattern recognition, artificial vision, coffee plant diseases.

Notas de autor

flores.eder.93@gmail.com

Introducción

La planta del café es un cultivo muy importante en países como Brasil, Colombia y Vietnam. En México, este cultivo representa alrededor de 0.66% del PIB agrícola del país. Se estima que hay alrededor de 500 000 productores que se dedican a la producción de este grano en 15 estados del país con un aproximado de 675 000 ha de superficie plantada (Vichi, 2015), donde Chiapas, Veracruz y Puebla se posicionan como los principales estados productores.

Entre las plagas que más daño producen a los cultivos del café se encuentra la roya del café. Esta plaga redujo entre 30%-50% la producción de este cultivo durante el brote epidémico de la roya del café en 2012 y 2013 en México (Hubert y Torres, 2017). Esta situación trajo consigo diversos problemas en las familias productoras como la emigración y el abandono definitivo de la cafecultura (Henderson, 2019). Otra enfermedad que se estudia en este artículo es el minador de la hoja (*Leucoptera coffeella*) que, en caso de infestación abundante, provoca defoliación en la planta del café (Ramírez y Jiménez, 2021). La phoma quema (*Phyllosticta coffeicola*) es un hongo que provoca manchas necróticas de tamaño variable en la hoja como si se tratara de quemaduras. Los frutos afectados se desprenden de la planta, lo cual repercute en la productividad de la planta (Echandi, 1957). La infección causada por *Cercospora coffeicola* es visible por la presencia de manchas de color marrón en las hojas y frutos de la planta de café; al respecto, se ha reportado que es capaz de disminuir entre 15% y 30% la productividad de una planta y afectar la calidad final del producto (Paula *et al.*, 2019).

Este trabajo pretende aportar información sobre el desempeño de cuatro algoritmos de clasificación: *Random forest*, Naive Bayes, *support vector machine* y redes neuronales artificiales *backpropagation* (ANN por sus siglas en inglés) y, de esta manera, coadyuvar en la labor científica de investigadores en fitopatología para ayudar a mantener la sanidad de los cultivos de café de agricultores y productores. El uso de la visión artificial representa nuevas formas de abordar la problemática de la identificación de enfermedades en las plantas y su uso se ha venido ampliando gracias a la masificación de teléfonos móviles. Entre los principales usos están a) la identificación de diferentes especies y tipos de hojas de plantas, es decir, hojas de diferentes cultivos (Sampallo, 2003), b) la clasificación de enfermedades en diferentes cultivos, esto es, identificar la presencia de lesiones en las hojas (Barbedo, 2019) y c) la estimación de la severidad o cuantificación del daño provocado por enfermedades en hojas y frutos, dicho de otro modo, el porcentaje de la hoja afectada por alguna enfermedad (Manso *et al.*, 2019). Como se aprecia, la visión artificial puede aprovecharse con diferentes propósitos en la agricultura y pretende ser una herramienta de soporte para investigadores y agricultores para obtener un diagnóstico más confiable y preciso a la hora de identificar, clasificar y cuantificar diferentes plagas y enfermedades del café. El objetivo que se persigue es actuar de manera más oportuna para minimizar los impactos de cada enfermedad evitando la dispersión regional de plagas y enfermedades de la planta del café.

Entre los algoritmos de clasificación de enfermedades en plantas más aplicados se encuentran las SVM, ANN, clasificadores bayesianos, árboles de decisión (*Decision Trees*, en inglés) y clasificadores del vecino más cercano (KNN por sus siglas en inglés). De los primeros trabajos de visión artificial aplicados a la planta del café está el de Price *et al.* (1993), quienes tomaron fotografías en Papúa Nueva Guinea de hojas de café dañadas por la roya del café. Se compararon tres formas de medición: a) estimación visual, b) medición con planímetro^[1] y c) análisis digital de imágenes donde se binarizaron las imágenes para entonces producir mediciones lineales y axiales de las lesiones conectadas. Los resultados mostraron que las estimaciones con el análisis digital de imágenes se aproximaron en mayor medida a las mediciones con el planímetro cuando el área afectada es >20% y con un mayor número de imágenes (en buenas condiciones de captura) el modelo podría entregar hasta una $R^2 = 93\%$, lo cual demuestra el buen desempeño del sistema de análisis desarrollado por los autores (Gavhale *et al.*, 2014). Trabajaron en la identificación de enfermedades de las plantas de los cítricos que incluyen toronja, limón, lima y naranjas atacadas por el cancro y antracnosis. El modelo propuesto incluye como preprocesamiento la conversión de RGB a diferentes espacios de color, la segmentación se hizo con agrupamiento por *K-means* y se extrajeron

características texturales de Haralick y características de color. Los resultados muestran una tasa del 95% de precisión con SVM (Bhange y Hingoliwala, 2015). Subiendo imágenes a un sistema web para extraer características, como el color, morfología y el vector de coherencia de color, estos autores experimentaron con cámaras de 10, 5 y 3 megapíxeles para averiguar si una granada está infectada o no.

Obtuvieron una precisión de 85% con SVM para imágenes de 10 megapíxeles, 82% para las de 5 megapíxeles y 79% para las de 3 megapíxeles. Entre los autores que han trabajado con los algoritmos de clasificación SVM se encuentran Qin *et al.* (2016), quienes intentaron diagnosticar y clasificar cuatro enfermedades de la alfalfa, además extrajeron 129 características texturales, cromáticas y geométricas. Compararon los algoritmos SVM, KNN y bosque aleatorio. El que mejor resultados obtuvo fue el algoritmo SVM con 94.74% de exactitud. También, en ese mismo año, Mengistu *et al.* (2016) publicaron un trabajo donde identifican tres enfermedades del café por medio de los algoritmos de redes neuronales artificiales, vecino más cercano, Naive Bayes y un híbrido de mapa autoorganizable y función de base radial. El conjunto de datos utilizado consta de 9 100 datos de los cuales el 70% sirvió para entrenamiento y el resto para *testing*. Los resultados de precisión entregados por los algoritmos fueron de 58.16% para el clasificador del vecino más cercano; para las redes neuronales, 79.04%; para Naive Bayes, 53.47%; para híbrido de mapa autoorganizable y función de base radial, 90.07%. Otro trabajo interesante es el de Esgario *et al.* (2020), quienes clasificaron la severidad de las hojas del café afectadas por cuatro enfermedades (roya del café, minador de la hoja, phoma quema y cercospora) por medio de *deep learning* y utilizaron diferentes arquitecturas de redes neuronales convolucionales. AlexNet tuvo 61 millones de parámetros y 8 capas; GoogleNet; 6.9 millones de parámetros y 22 capas; VGG19, 138 millones de parámetros y 19 capas; por último, ResNet50, 25 millones de parámetros con 50 capas. Encontraron que ResNet50 (Multi-task) con 95.24% de exactitud, 95.29% de precisión y 91.14% de sensibilidad fue la que mejores resultados entregó para la clasificación de enfermedades y para el cálculo de la severidad de las enfermedades fue VGG16, con 86.51% de exactitud, 82.49% de precisión y 80.89% de sensibilidad. Montalbo y Hernández (2020) identificaron tres enfermedades del café gracias a redes neuronales convolucionales con el modelo VGG16. Para entrenar este modelo se valieron de un conjunto de datos de 3 958 imágenes de hojas de café con diferentes enfermedades y para la validación, un 20% del conjunto de datos total. Como es usual en este tipo de arquitecturas, el modelo VGG16 consume mucho tiempo para ser entrenado, pero consigue obtener resultados de hasta el 100%. En la actualidad, la visión artificial no sólo se aplica en el ámbito agronómico para detectar enfermedades sino también es posible detectar deficiencias nutrimentales de las plantas a través del análisis de las hojas. Lewis y Espineli (2020) implementaron un dispositivo que con ayuda de una Raspberry Pi 4 ofrece un diagnóstico de las deficiencias nutrimentales de la planta del café como falta de fósforo, nitrógeno, calcio, potasio, zinc, entre otros. El algoritmo que eligieron es una red neuronal convolucional que obtiene una exactitud general de 91.49%. La exactitud más alta alcanzada es al detectar deficiencias de fósforo con 93%, mientras que la exactitud más baja es al detectar deficiencias de potasio con 90%. Yebasse *et al.* (2021) clasifican hojas enfermas y hojas sanas gracias al método llamado *enfoque guiado* (*Guided Approach*), que es una serie de pasos que incluyen la eliminación de ruido, la segmentación del área de interés y la clasificación con el algoritmo Naive Bayes de donde se obtiene 71% con 4 ciclos y hasta 98% con 14 ciclos. Los resultados de este método los compararon con los del el algoritmo Naive Bayes sin el preprocesamiento del enfoque guiado que fueron de máximo 75% en ciclo 14. Esto comprueba la efectividad de su método para mejorar la clasificación de enfermedades del café.

En este trabajo se realizó la identificación hojas sanas de la planta de café y de cuatro enfermedades diferentes del café: la roya del café, el minador de la hoja del café, la phoma quema y la mancha de hierro. Se segmentaron las regiones de interés con diferentes métodos como PCA (*Principal Component Analysis* en inglés) (Hotelling, 1933), el método Otsu (1979) y el método de frontera global. También, se extrajeron características texturales, geométricas y cromáticas. De acuerdo con los resultados obtenidos, la propuesta presentada en este trabajo es competitiva, ya que se obtuvieron porcentajes de hasta el 83% de precisión.

1. Materiales y métodos

En esta sección se describe cada uno de los pasos de la metodología, cómo se constituyó el conjunto de datos, el preprocesamiento de las imágenes, las técnicas de segmentación empleadas, los tipos de características obtenidas y los parámetros usados para la optimización de los algoritmos de clasificación de este estudio.

La figura 1 ilustra el flujo de la metodología utilizada.

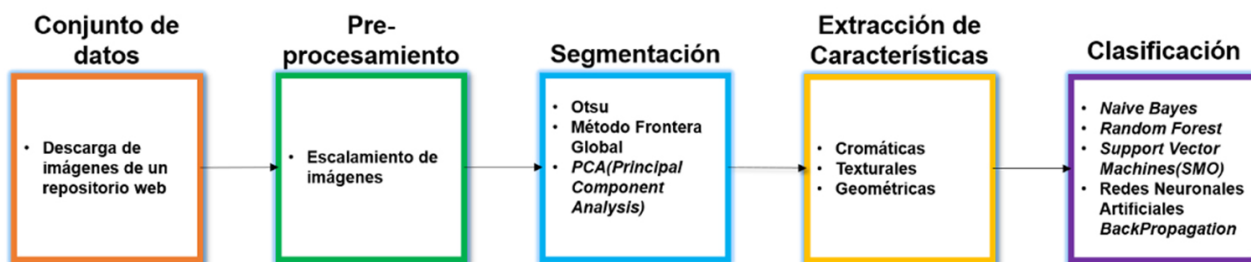


FIGURA 1

Pasos de la metodología propuesta

Fuente: elaboración propia.

1. 1. Características del conjunto de datos

El conjunto de datos empleado para este trabajo se obtuvo de Mendeley Data.^[2] Los autores lo nombraron “Bracol - A brazilian Arabica Coffee Leaf images dataset to identification and quantification of coffee diseases and pests” (Krohling *et al.*, 2019). A partir de imágenes que se tomaron de cultivos de café en Santa Maria de Marechal Floreano en Brasil, este trabajo contiene imágenes de hojas sanas y de cuatro enfermedades comunes en la planta del café de la especie *Coffea arabica*, especie que cuenta con amplia distribución en México.

Las fotos fueron tomadas con teléfonos de la marca ASUS Zenfone 2, Xiaomi Redmi 5A, Xiaomi S2, Galaxy S8, and iPhone 6S. Todas las hojas están etiquetadas con la enfermedad y la severidad de la enfermedad a la que pertenecen. Las imágenes tienen una resolución de 2 048 pixeles de ancho $\times$ 1 024 pixeles de alto.

Para el propósito de esta investigación se crearon cinco clases; la hoja sana y las hojas enfermas con alguna de las cuatro enfermedades sin importar la intensidad de su severidad, esto es, la incidencia de alguna enfermedad. Del total de las imágenes del conjunto de datos original, se hizo una preselección para obtener las más representativas de cada clase o enfermedad. Cabe destacar que en cada hoja se presenta una sola enfermedad, por lo que esta investigación se limita a reconocer una enfermedad por hoja; la detección múltiple de enfermedades está fuera del alcance de este trabajo debido a la complejidad para el sistema de detectar la enfermedad predominante. La figura 2 presenta las imágenes de las diferentes clases del conjunto de datos y la cantidad de imágenes por cada clase.



FIGURA 2

Imágenes de las clases del conjunto de datos

Fuente: elaboración propia.

Se utilizaron 739 fotos para este trabajo. Cabe mencionar que el conjunto de datos aplicado presenta un desbalance de datos, en consecuencia, el número de imágenes de hojas no es el mismo para todas las clases. Cabe remarcar que estas imágenes se tomaron en un ambiente controlado, puesto que las hojas de café

fueron cortadas y puestas sobre un fondo blanco con el fin de diferenciar mejor la región de interés del fondo y eliminar el ruido de otro tipo de vegetación presente en la plantación de café.

1. 2. Preprocesamiento

Por medio de MatLab R2017B, se redujeron las imágenes a un tamaño de 410 píxeles de ancho x 205 píxeles de alto, que representa sólo el 20% de su tamaño original. Al reducir el tamaño de las imágenes, la pérdida de información para la extracción de características fue mínima, por lo que ese proceso se hace en menor tiempo.

1. 3. Segmentación

La segmentación es el proceso donde se busca separar una región de interés (ROI por sus siglas en inglés) de su fondo con el objetivo de aprovechar y mejorar el análisis sobre el área que interesa. Las técnicas que se utilizaron para este trabajo fueron análisis de componentes principales PCA, método de frontera global y Otsu, de las cuales esta última fue la que mejor resultados generó, ya que el sistema de reconocimiento de patrones se enfoca únicamente en la región de interés que se determina por su contorno y propiedades que se usarán en el siguiente paso de la metodología, que es la extracción de características.

En la figura 3 se muestran hojas afectadas por las cuatro diferentes enfermedades del café. Cada hoja se segmentó por tres métodos diferentes. Se puede observar que el método de frontera global tiene problemas para segmentar algunas hojas. Otsu y PCA, por su parte, obtienen muy buenos resultados, pero el hecho de que algunas imágenes contengan sombras derivadas de los ángulos y hora del día en que se tomaron las fotos hacen que estas sombras confundan a los algoritmos de segmentación implementados y terminan por segmentar esas pequeñas áreas como si fueran parte de la hoja. No obstante, los algoritmos de segmentación toman el área más grande segmentada, que en este caso es el área de la hoja, e ignoran las sombras aisladas e irrelevantes. Los resultados entregados por el algoritmo Otsu fueron mejores que los del algoritmo que utiliza PCA, por lo que se seleccionó Otsu para segmentar todas las imágenes.

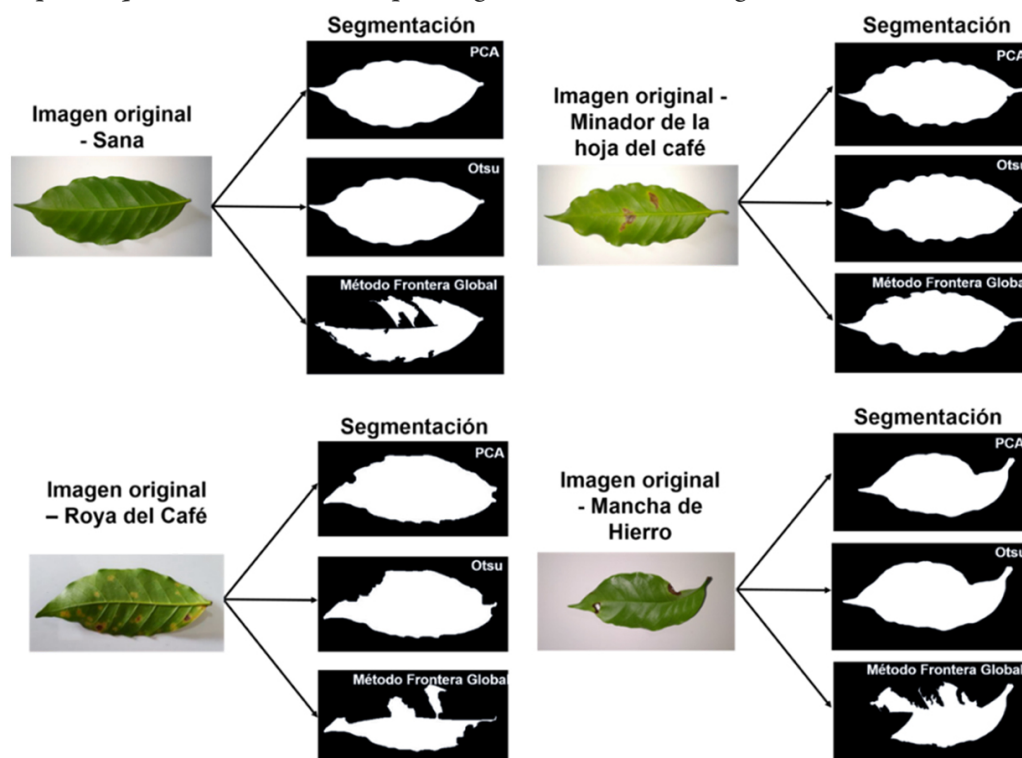


FIGURA 3

Diferentes técnicas de segmentación aplicadas a las hojas de café

Fuente: elaboración propia.

1. 4. Extracción de características

La extracción de características es una de las tareas determinantes en el procesamiento digital de imágenes. En este paso se obtienen las formas, colores, y relieves ocasionados por los síntomas de las plagas y enfermedades en las hojas del café. Este proceso genera vectores de características que sirven como parámetros para los algoritmos de clasificación. El total de las características geométricas es de 78, para las características cromáticas es de 273 y para las textuales 84 características, que da como resultado un vector de 436 características por cada imagen cuando se combinan todos los tipos de características para procesar por los algoritmos de clasificación en Weka 3.8. Lo anterior se aprecia en la figura 4.

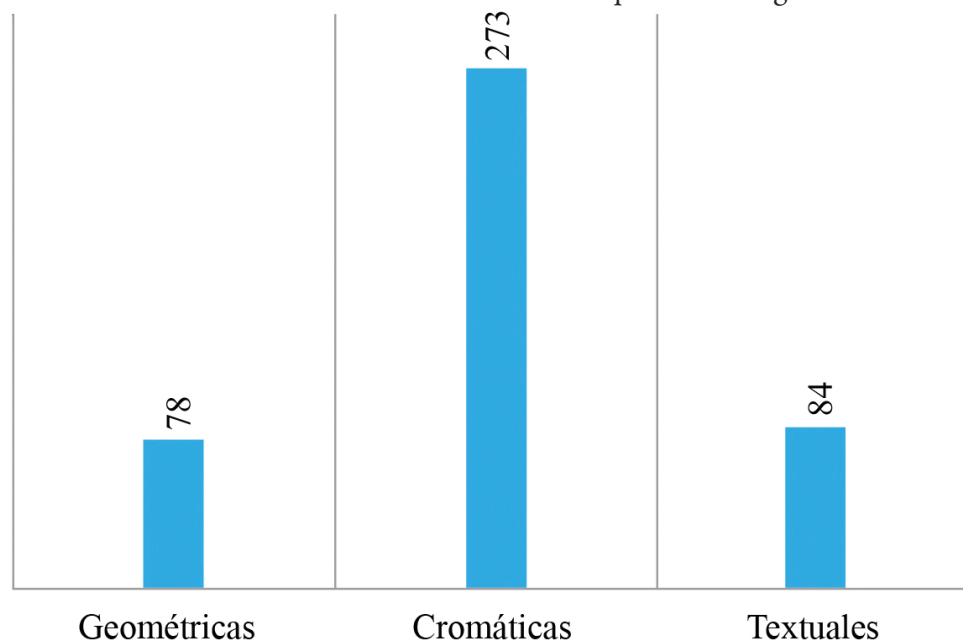


FIGURA 4

Cantidad de atributos por tipo de características

Fuente: elaboración propia.

1. 5. Características geométricas

Para el cálculo de características geométricas se utilizaron características geométricas básicas como altura, ancho, área, perímetro, redondez, diámetro equivalente, número de Euler, excentricidad, centro de gravedad, rectangularidad, factor de Danielson, área convexa y área rellena, entre otros, además de momentos geométricos como los momentos centrales, que son funciones que ayudan a reconocer una imagen dentro de un eje de coordenadas independientemente de su transformación geométrica como la traslación, escalado y rotación. De igual manera, se extrajeron los momentos de Hu (1962), que consisten en siete momentos que ayudan también a precisar si una hoja es la misma sin importar su orientación, localización o tamaño.

Se utilizaron los primeros ocho descriptores de Fourier para el cálculo del contorno de la hoja, así como también descriptores de Flusser, los cuales permiten saber sobre la difuminación de la imagen (Flusser *et al.*, 2009). Los momentos de Hu pueden ayudar a diferenciar una hoja de otra en cuestión de forma, pero si son hojas muy similares los momentos de hu geométricos, no serán suficientes para discriminar las hojas y se tendrán que utilizar otro tipo de características. Los momentos de Hu se obtienen con las siguientes fórmulas:

$$\emptyset 1 = \mu 20 + \mu 02$$

(1)

$$\emptyset 2 = (\mu_{20} - \mu_{02})^2 + 4(\mu_{11})^2 \quad (2)$$

$$\emptyset 3 = (\mu_{30} - 3\mu_{12})^2 + (\mu_{03} - 3\mu_{21})^2 \quad (3)$$

$$\emptyset 4 = (\mu_{30} + \mu_{12})^2 + (\mu_{03} + \mu_{21})^2 \quad (4)$$

$$\emptyset 5 = (\mu_{30} - 3\mu_{12})(\mu_{30} + \mu_{12})((\mu_{30} + \mu_{12})^2 - 3(\mu_{21} + \mu_{03})^2) + (3\mu_{21} - \mu_{03})(\mu_{21} + \mu_{03})(3(\mu_{30} + \mu_{12})^2 - (\mu_{03} + \mu_{21})^2) \quad (5)$$

$$\emptyset 6 = (\mu_{20} - \mu_{02})((\mu_{30} + \mu_{12})^2 - (\mu_{21} + \mu_{03})^2) + 4\mu_{11}(\mu_{30} + \mu_{12})(\mu_{21} + \mu_{03}) \quad (6)$$

$$\emptyset 7 = (3\mu_{21} - \mu_{03})(\mu_{30} + \mu_{12})((\mu_{30} + \mu_{12})^2 - 3(\mu_{21} + \mu_{03})^2) + (\mu_{30} - 3\mu_{12})(\mu_{21} + \mu_{03})(3(\mu_{30} + \mu_{12})^2 - (\mu_{03} + \mu_{21})^2) \quad (7)$$

Una descripción más detallada de los siete momentos de Hu está disponible en Hu (1962). La figura 5 presenta algunas de las características geométricas obtenidas, así como los momentos de Hu geométricos.

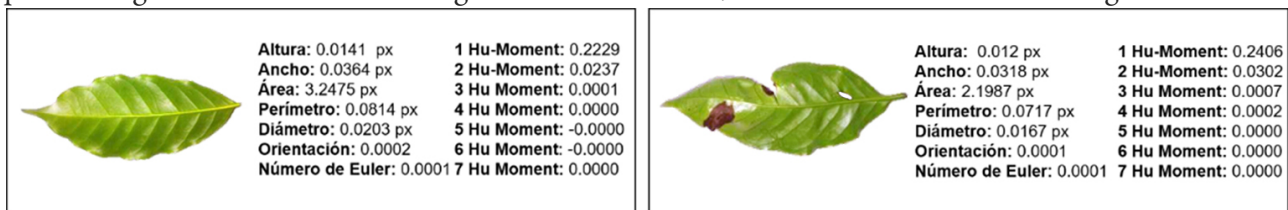


FIGURA 5

Comparación de características geométricas de dos hojas: una hoja sana y otra afectada por phoma quemado
Fuente: elaboración propia.

1. 6. Características cromáticas

Las características cromáticas permiten saber información sobre la intensidad del color en una región de interés y son una de las características que más información proveen. En este caso se utilizó el espacio de color RGB, que son los tres colores primarios rojo, verde y azul, y también una máscara de 5 x 5 para obtener el contraste, lo cual permite obtener la diferencia entre el tono más negro y el tono más claro de la imagen. En la figura 6 se observan algunas de las características cromáticas obtenidas.

	R	G	B
1 Hu Momento:	2.3301	1.3727	0.2229
2 Hu Momento:	2.5922	0.8996	0.0237
3 Hu Momento:	0.0245	0.0065	0.0001
4 Hu Momento:	0.0016	0.0004	0.0000
5 Hu Momento:	-0.0000	-0.0000	-0.0000
6 Hu Momento:	-0.0024	-0.0004	-0.0000
7 Hu Momento:	0.0000	0.0000	-0.0000
DCT I:	0.3279	0.4025	0.1191
DCT II:	0.0505	0.0484	0.0425
DCT III:	0.0598	0.0673	0.0301
DCT IV:	0.0492	0.0570	0.0251

	R	G	B
1 Hu Momento:	2.5932	1.6789	0.2406
2 Hu Momento:	3.5124	1.4724	0.0302
3 Hu Momento:	0.2556	0.0862	0.0007
4 Hu Momento:	0.0683	0.0230	0.0002
5 Hu Momento:	0.0084	0.0010	0.0000
6 Hu Momento:	0.0988	0.0216	0.0000
7 Hu Momento:	0.0032	0.0004	0.0000
DCT I:	0.2528	0.3194	0.1188
DCT II:	0.0425	0.0502	0.0480
DCT III:	0.0500	0.0543	0.0419
DCT IV:	0.0511	0.0538	0.0413

FIGURA 6

Comparación de características geométricas de dos hojas: una hoja sana y otra afectada por phoma quemada
Fuente: elaboración propia.

1. 7. Características texturales

La textura provee información valiosa de la superficie de algún objeto como el arreglo estructural de la intensidad de los colores. Para esto, se recurrió a las características de Haralick, las cuales se obtienen de la matriz de concurrencia del nivel de gris. Se calcularon 14 características de Haralick, entre las que se encuentran el segundo momento angular, contraste, correlación, suma de cuadrados, momento de diferencia inversa, suma promedio, suma entropía, suma de varianzas, entropía, diferencia de varianzas, diferencia de entropías, información de medidas de correlación y coeficiente máxima de correlación. A continuación, se muestran las fórmulas para obtener las 14 características de Haralick.

$$f_1 = \sum_i \sum_j [p(i,j)^2]$$

(8)

$$f_2 = \sum_{n=0}^{N_g-1} n^2 \left\{ \sum_{i=1}^{N_g} \sum_{j=1}^{N_g} p(i,j) \right\}$$

(9)

$$f_3 = \frac{\sum_{i=1}^{N_g} \sum_{j=1}^{N_g} [ijp(i,j) - \mu_x \mu_y]}{\sigma_x \sigma_y}$$

(10)

$$f_4 = \sum_i \sum_j (i - \mu_x)^2 p(i, j) \quad (11)$$

$$f_5 = \sum_i \sum_j \frac{1}{1 + (i - j)^2} p(i, j) \quad (12)$$

$$f_6 = \sum_{i=2}^{2N_g} iP_{x+y}(i) \quad (13)$$

$$f_7 = \sum_{i=2}^{2N_g} (i - f_8)^2 p_{x+y}(i) \quad (14)$$

$$f_8 = -\sum_{i=2}^{2N_g} p_{x+y}(i) \log\{p_{x+y}(i)\} \quad (15)$$

$$f_9 = -\sum_i \sum_j p(i, j) \log\{p(i, j)\} \quad (16)$$

$$f_{10} = \sum_{i=0}^{N_R-1} (i - f_8)^2 p_{x-y}(i) \quad (17)$$

$$f_{11} = -\sum_{i=0}^{N_g-1} P_{x-y}(i) \log\{p_{x-y}(i)\} \quad (18)$$

$$f_{12} = \frac{HXY - HXY1}{\max\{HX, HY\}} \quad (19)$$

$$f_{13} = (1 - e^{[-2(HXY2 - HXY)])}^{\frac{1}{2}} \quad (20)$$

Una descripción más detallada de los descriptores de Haralick se encuentra en Haralick *et al.* (1973). La figura 7 muestra las características texturales obtenidas.

	Segundo momento angular:	R	G	B
	Contraste:	0.3319	0.3170	0.4411
	Correlación:	0.2455	0.4247	0.3952
	Suma de cuadrados:	242.9466	244.4889	242.3573
	Momento de diferencia inversa:	0.2455	0.4426	0.3952
	Suma promedio:	0.8968	0.8584	0.8824
	Suma de entropía:	6.9767	6.8685	8.1021
	Suma de varianza:	5.6858	5.5498	6.8544
	Entropía:	1.2909	1.3203	1.2477
	Diferencia de Varianzas:	1.4868	1.6086	1.5066
	Diferencia de Entropías:	0.0790	0.0676	0.0749
	Medidas de Correlación:	0.5496	0.6820	0.6278
	Máximo Coeficiente de Correlación:	-0.3696, 0.4884	-0.2262, 0.3349	-0.2821, 0.3895
		0.5494	0.3443	0.5309
	Segundo momento angular:	R	G	B
	Contraste:	0.2262	0.4411	0.2109
	Correlación:	0.8021	0.3952	0.6481
	Suma de cuadrados:	59.6198	242.3573	55.9104
	Momento de diferencia inversa:	0.8021	0.3952	0.6481
	Suma promedio:	0.8290	0.8824	0.7995
	Suma de entropía:	3.4084	8.1021	3.3864
	Suma de varianza:	1.8607	6.8544	1.8439
	Entropía:	1.5498	1.2477	1.5426
	Diferencia de Varianzas:	1.9488	1.5066	1.9112
	Diferencia de Entropías:	0.0602	0.0749	0.0546
	Medidas de Correlación:	0.8603	0.6278	0.8161
	Máximo Coeficiente de Correlación:	0.2428, 0.4006	-0.2821, 0.3895	-0.1711, 0.2981
		0.4512	0.5309	0.3646

FIGURA 7

Comparación de características texturales de dos hojas: una hoja sana y otra afectada por phoma quemada
Fuente: elaboración propia.

1. 8. Experimentos

Con el programa Weka 3.8 se seleccionaron los algoritmos de clasificación y se optimizaron parámetros para obtener los mejores resultados.

1. 8. 1. *Support vector machine*

Las SVM (*support vector machine*) son técnicas muy utilizadas para la clasificación de enfermedades debido a que entregan resultados muy competitivos y con un menor consumo de recursos en comparación con otros algoritmos como las redes neuronales artificiales cuando se trata de conjuntos de datos pequeños (Pujari *et al.*, 2016). Una de las características fundamentales de las SVM es que hacen uso de *kernels* cuando la clasificación binaria del conjunto de datos no está separada linealmente, por lo que se recurre a los *kernels* para transformar los espacios de entrada en un espacio multidimensional donde los conjuntos pueden ser separados de ese modo una vez transformados.

Después de realizar pruebas con los diferentes *kernels* de SVM disponibles en Weka 3.8, se decidió utilizar el PolyKernel con el valor de la variable compleja de 2.0; para la validación cruzada, 10 *folds*; como calibrador, *Logistic*. El número de iteraciones se dejó por defecto (hasta que converja).

1. 8. 2. Redes neuronales artificiales *backpropagation*

Las redes neuronales artificiales son conjuntos de neuronas artificiales interconectadas que con algoritmos y modelos matemáticos procesan conjuntos de datos para aprender de ellos y ofrecer una clasificación de acuerdo al objetivo con el que fueron desarrolladas. Una red neuronal artificial comprende a menudo de tres tipos de capa de neuronas: capa de entrada, capa oculta y capa de salida.

Uno de los algoritmos más populares de redes neuronales es el de *backpropagation*, el cual cambia de manera iterativa los pesos de las neuronas mientras minimiza el error cuadrático de la salida que es requerida y la que es obtenida con los actuales pesos. Para realizar eso, se utilizan los ejemplos del conjunto de entrenamiento que se definen como: $\{(x_1, y_1), (x_2, y_2), \dots, (X_n, y_n)\}$.

La red neuronal es una red *backpropagation* ajustada a 500 ciclos, una tasa de aprendizaje α que es igual a 0.1, además de la variable *Momentum* que su valor es de 0.1.

1. 8. 3. *Random forest*

Es un algoritmo que construye múltiples árboles de decisión; cada árbol depende de los valores de un vector aleatorio. El error de generalización converge a un límite a medida que el número de árboles en el bosque se incrementa. Tiene mejor rendimiento que los árboles individuales y es más eficiente que los algoritmos de clasificación tradicionales como las redes neuronales artificiales, en especial cuando se trata de conjuntos de datos grandes XinHai (2013). *Random forest* maneja miles de variables explicativas o independientes y puede modelar interacciones complejas entre variables.

En este estudio, *Random forest* utilizó 200 iteraciones. El valor de la variable *depth* se fijó en 0, lo cual quiere decir que el algoritmo controlará en automático la profundidad del árbol.

1. 8. 4. Naive Bayes

El clasificador Naive Bayes está basado en la suposición de que los valores de los atributos son condicionalmente independientes dado el valor objetivo. Naive Bayes ignora las posibles dependencias o correlaciones entre las entradas y reduce un problema multivariables a un grupo de problemas univariable.

1. 9. Métricas de desempeño

Evaluar el desempeño del modelo con distintas métricas es esencial para tener una idea más clara de su desempeño real. En este artículo se recurre a distintas métricas que permiten tener un juicio más acertado sobre el comportamiento del modelo propuesto.

a) *Accuracy*: es usual que se conozca como la proporción de predicciones correctas sobre el total de ejemplos.

$$Accuracy = \frac{\text{Número de predicciones correctas}}{\text{Número total de predicciones realizadas}} \quad (21)$$

b) Matriz de confusión: la matriz de confusión nos permite visualizar el desempeño del modelo. Esta métrica es construida a partir de diferentes medidas.

c) Verdaderos positivos (VP): son los casos en los que la predicción alcanzada es positiva (+) y la clase real es positiva (+).

d) Verdaderos negativos (VN): son los casos en los que la predicción alcanzada es negativa (-) y la clase real es negativa (-).

e) Falsos positivos (FP): son los casos en los que la predicción alcanzada es positiva (+) y la clase real es negativa (-).

f) Falsos negativos (FN): son los casos en los que la predicción alcanzada es negativa (-) y la clase real es positiva (+).

g) Precisión: esta métrica es calculada como la proporción de verdaderos positivos sobre la suma de verdaderos positivos y falsos positivos.

$$Precisión = \frac{VP}{VP + FP} \quad (22)$$

h) *TP Rate (Sensitividad o Recall)* se define mediante la siguiente proporción.

$$Sensitividad = \frac{TP}{FN + TP}$$

(23)

i) Especificidad: se define mediante la siguiente proporción.

$$Especificidad = \frac{TN}{TN + FP}$$

(24)

j) *FP Rate*: se define mediante la siguiente proporción.

$$Sensitividad = \frac{FP}{TN + FP}$$

(25)

k) Área Bajo la Curva (AUC): esta métrica se evalúa para diferentes umbrales de FPR (abscisas) contra TPR (ordenadas). El área calculada debajo de la curva obtenida es la que se reporta, mientras el área se acerque más a 1 el desempeño del modelo es mejor, y mientras más cercano a 0 sea el desempeño del modelo es peor.

k) *F-measure*: es una métrica muy utilizada en conjuntos de datos con desbalance y, combinando la esencia de Precisión y *Recall*, captura sus métricas de desempeño.

$$F - measure = \frac{2 * precisión * Recall}{FN + TP precisión + Recall}$$

(26)

2. Resultados

En la figura 8 están exhibidas las matrices de confusión de las características cromáticas, donde se aprecia que la mayoría de las instancias se clasifican correctamente en cuanto a su clase, exceptuando a la mancha de hierro o cercospora que tiene pocas instancias que son clasificadas correctamente. En el recuadro de mancha de hierro se observa que sólo 22 instancias de 77 posibles son clasificadas correctamente. En este caso, la mancha de hierro es confundida por el clasificador SVM en su mayoría con el minador de la hoja del café, pero también con la roya del café y la phoma quema.

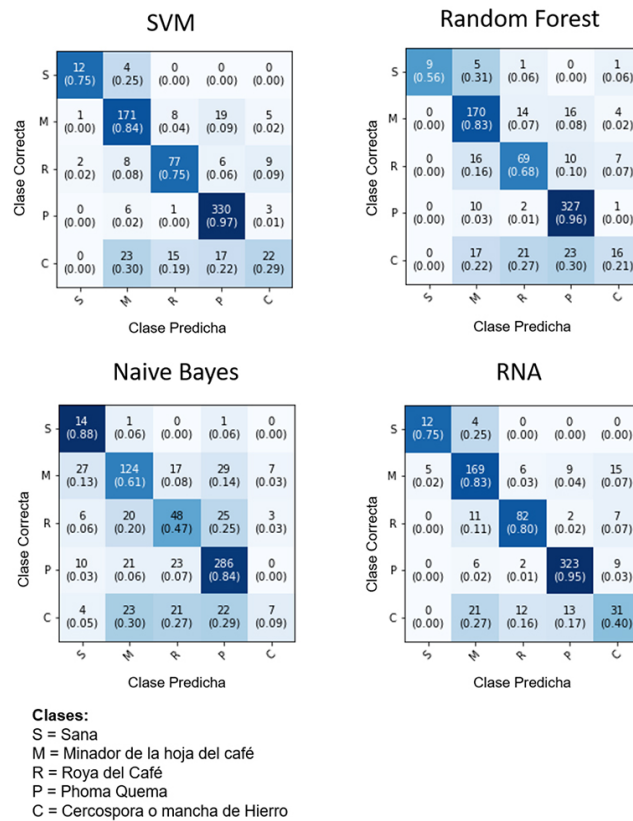


FIGURA 8

Matriz de confusión de los resultados obtenidos en las características cromáticas

Fuente: elaboración propia.

El cuadro 1 condensa los resultados que obtuvieron los cuatro diferentes algoritmos para las características cromáticas. De este modo, observa que los algoritmos entregaron alta sensibilidad para este tipo de características; no obstante, la sensibilidad reportada por Naive Bayes (64%) fue baja si se compara con los demás algoritmos. Todos los algoritmos reportaron baja tasa de falsos positivos, excepto Naive Bayes que presenta una tasa más alta respecto a los demás algoritmos.

Cuadro 1

Métricas obtenidas usando características cromáticas

Parámetros/algoritmos	Redes neuronales	Naive Bayes	Random forest	SVM
Sensitividad	0.83	0.648	0.794	0.828
FP Rate	0.059	0.139	0.092	0.078
Precisión	0.82	0.652	0.775	0.812
Valor F	0.831	0.635	0.773	0.814
Roc Area	0.94	0.839	0.936	0.895

Fuente: elaboración propia.

En el cuadro 2 se muestran las métricas de sensibilidad y precisión con diferentes valores de K -folds para las características cromáticas. K -folds es un parámetro en la validación cruzada que indica el número de grupos para validación en un conjunto de datos limitado. Si se tiene por ejemplo $k = 10$, habrá 10 iteraciones para validación.

Cuadro 2Sensitividad y precisiones obtenidas a partir de características cromáticas con diferentes valores *K-folds*

# <i>K-folds</i> / algoritmos	Redes neuronales- sensibilidad/ precisión	Naive Bayes sensibilidad/ precisión	<i>Random forests</i> sensitividad/ precisión	SVM sensibilidad/ precisión
2	0.819/0.817	0.660/0.649	0.775/0.758	0.817/0.803
4	0.808/0.794	0.639/0.635	0.770/0.739	0.828/0.816
6	0.840/0.836	0.637/0.624	0.779/0.752	0.831/0.819
8	0.831/0.827	0.640/0.641	0.773/0.749	0.834/0.822
12	0.828/0.816	0.636/0.622	0.778/0.753	0.832/0.819
14	0.827/0.817	0.640/0.636	0.783/0.761	0.840/0.819

Fuente: elaboración propia.

De acuerdo con el cuadro 3, se determina que los cuatro algoritmos utilizados en este artículo tienen un desempeño más modesto con las características geométricas que con otro tipo de características, ya que la sensibilidad más alta con las características geométricas es de 73% por medio de *Random forest*.

Cuadro 3

Métricas obtenidas usando características geométricas

Parámetros/algoritmos	Redes neuronales	Naive Bayes	<i>Random forest</i>	SVM
Sensibilidad	0.708	0.43	0.732	0.714
<i>FP Rate</i>	0.123	0.107	0.128	0.151
Precisión	0.685	0.639	0.703	0.683
Valor F	0.694	0.461	0.707	0.68
<i>Roc Area</i>	0.872	0.788	0.898	0.799

Fuente: elaboración propia.

En el cuadro 4 se registran las métricas de sensibilidad y precisión con diferentes valores de *K-folds* para las características geométricas.

Cuadro 4

Cuadro 4 Sensitividad y precisiones obtenidas usando características geométricas con diferentes valores *K-folds*

# <i>K-folds</i> / algoritmos	Redes neuronales- sensitividad/ precisión	Naive Bayes sensitividad/ precisión	<i>Random forest sensitividad</i> / precisión	SVM sensitividad/ precisión
2	0.529/0.513	0.396/0.536	0.345/0.498	0.817/0.803
4	0.564/0.516	0.415/0.580	0.332/0.464	0.828/0.816
6	0.562/0.534	0.392/0.633	0.323/0.466	0.831/0.819
8	0.537/0.504	0.417/0.620	0.326/0.475	0.834/ 0.822
12	0.537/0.513	0.421/0.634	0.306/0.463	0.832/0.819
14	0.559/0.532	0.418/0.643	0.317/0.471	0.840/0.819

Fuente: elaboración propia.

En el cuadro 5, a partir de las características texturales, se tienen resultados de sensibilidad que promedian el 73% para todos los algoritmos. Las redes neuronales artificiales, *random forest* y SVM consiguen resultados de sensibilidad parecidos. Al algoritmo Naive Bayes le sucede lo contrario, ya que termina con 62% de sensibilidad. En general, los algoritmos tienen bajas tasas de falsos positivos y las curvas ROC indican un buen desempeño al clasificar.

Cuadro 5

Métricas obtenidas a partir de características texturales

Parámetros/algoritmos	Redes neuronales	Naive Bayes	<i>Random Forest</i>	SVM
Sensitividad	0.774	0.622	0.752	0.76
<i>FP Rate</i>	0.078	0.1	0.099	0.102
Precisión	0.763	0.655	0.736	0.728
Valor F	0.766	0.632	0.735	0.724
<i>Roc Area</i>	0.916	0.862	0.924	0.857

Fuente: elaboración propia.

El cuadro 6 presenta las métricas de sensibilidad y precisión con diferentes valores de *K-folds* para las características texturales.

Cuadro 6Sensitividad y precisión obtenida usando características texturales con diferentes valores *K-folds*

# <i>K-folds</i> / algoritmos	Redes neuronales- sensibilidad/ precisión	Naive Bayes sensibilidad/ precisión	<i>Random Forest</i> sensibilidad/ precisión	<i>svm</i> sensibilidad/ precisión
2	0.740/0.733	0.633/0.662	0.750/0.729	0.752/0.699
4	0.766/0.757	0.639/0.664	0.760/0.745	0.755/0.698
6	0.749/0.743	0.639/0.664	0.752/0.734	0.747/0.716
8	0.764/0.759	0.631/0.663	0.760/0.740	0.753/0.699
12	0.763/0.757	0.629/0.655	0.747/0.732	0.749/0.692
14	0.769/0.760	0.631/0.660	0.755/0.741	0.755/0.723

Fuente: elaboración propia.

En los resultados mostrados en Esgario *et al.* (2020) los autores muestran los parámetros que la red ha aprendido. Es claro que las posibilidades de sobreajuste de un conjunto de datos aumentan con la cantidad de parámetros. La cantidad de parámetros que se pueden aprender en una red convolucional es definida por la cantidad de canales que se están utilizando en las imágenes, el número de filtros, los tamaños de los filtros y el aplanamiento de la salida convolucional. En el caso de este artículo, el número de parámetros aprendidos por las técnicas clásicas sólo depende del número de características extraídas. Los dos enfoques son opuestos: mientras que las técnicas clásicas extraen características a partir de la imagen con el propósito de que cada una de esas características sean fundamentales para detectar cierta clase, en el caso de las redes convolucionales extrae las características con la certeza de que estas influyen en la identificación de la clase. No obstante, cada técnica tiene sus propias ventajas y desventajas que las vuelven fundamentales. Una ventaja significativa de la metodología propuesta es el tiempo de entrenamiento relativamente bajo debido a que los modelos se entrenan con un número de características fijas, además es posible que mediante métodos de selección de características el desempeño de los modelos propuestos se vea incrementado al eliminar aquellas características que introducen ruido al clasificador.

Conclusiones

Este trabajo implementó cuatro algoritmos de clasificación: Naive Bayes, *Random forest*, redes neuronales artificiales *backpropagation* y *support vector machine* con SMO; estos dos últimos fueron los que mejor resultados lograron con porcentajes que rondan el 80% de sensibilidad. Las redes neuronales artificiales *backpropagation* obtuvieron el mejor porcentaje de sensibilidad (83%), pero también es el algoritmo que más tiempo consume para ser entrenado. Por su parte, *support vector machine* destaca en rapidez para clasificar y buena precisión.

Asimismo, se presentan resultados competitivos, pues se obtiene un promedio de sensibilidad de todos los algoritmos de 72%. Esto quiere decir que la mayoría de los algoritmos entregaron resultados satisfactorios, a excepción de Naive Bayes que obtuvo resultados deficientes en las características geométricas y mejorables en el resto de características. Como menciona Barbedo (2019), la mayoría de las veces, la precisión de los sistemas de reconocimiento patrones en plantas está muy ligada a conjuntos de datos variados y con alta representatividad. El reto es obtener o mejorar el conjunto de datos para dar mayor representatividad a cada una de las clases y, por ende, mejorar la clasificación de las enfermedades y plagas del café.

Análisis prospectivo

Este estudio contribuye a la obtención de parámetros interesantes como rendimiento, sensibilidad, exactitud y curva ROC de los diferentes algoritmos implementados, los cuales puede aportar información valiosa a investigadores al momento de seleccionar algoritmos más eficientes para desarrollar una herramienta como una *app* que sea más práctica y accesible para agrónomos e investigadores en fitopatología y ser utilizada *in situ* para poder identificar enfermedades y plagas de la planta del café. Lo que se persigue es ofrecer un diagnóstico rápido y certero para que se tomen las acciones pertinentes de control y contrarrestar la dispersión de plagas y enfermedades.

Así también, gracias a este trabajo es posible identificar una enfermedad o plaga de entre cuatro que se estudiaron, es decir, se obtiene la incidencia de alguna enfermedad o plaga en la hoja, la cual es una variable cualitativa, que permite realizar análisis sobre la dispersión de una plaga en un cultivo. Lo ideal sería cuantificar o medir el daño causado por alguna enfermedad o plaga, por ejemplo, si la hoja está afectada en un 20% o 80%. Por tal motivo, como trabajo futuro se tiene pensado desarrollar un sistema para la cuantificación del daño causado específicamente por alguna enfermedad, es decir, el porcentaje de la hoja que se encuentra dañada. Esta variable cuantitativa resulta de mayor interés para agrónomos e investigadores en fitopatología, dado que esta variable serviría para determinar en qué momento del desarrollo de la enfermedad es más pertinente aplicar mecanismos de control (productos químicos para

atacar a una enfermedad en particular) en una plantación de café para mejorar su condición general fitosanitaria y evitar así, la resistencia de plagas y enfermedades a productos químicos.

Referencias

- Barbedo, J. G. A. (2019). Plant disease identification from individual lesions and spots using deep learning. *Biosystems Engineering*, 180, 96-107.
- Bhange, M. & Hingoliwala, H. (2015). Smart Farming: Pomegranate Disease Detection Using Image Processing. *Procedia Computer Science*, 58, 280-288
- Echandi, E. (1957). La quema de los cafetos causada por *Phoma costarricensis* n. sp. *Revista de Biología Tropical*, 5(1), 81-102.
- Esgario, J. G., Krohling, R. A., & Ventura, J. A. (2020). Deep learning for classification and severity estimation of coffee leaf biotic stress. *Computers and Electronics in Agriculture*, 169, 105162.
- Flusser, J., Zitova, B., & Suk, T. (2009). *Moments and moment invariants in pattern recognition*. Wiley Publishing.
- Gabor, D. (1946). Theory of communication. Part 1: The analysis of information. *Journal of the Institution of Electrical Engineers-Part III: Radio and Communication Engineering*, 93(26), 429-441
- Gavhale, K. R., Gawande, U., & Hajari, K. O. (2014). Unhealthy region of citrus leaf detection using image processing techniques. *International Conference for Convergence for Technology*. IEEE.
- Haralick, R. M., Shanmugam, K., & Dinstein, I. H. (1973). Textural features for image classification. *IEEE Transactions on Systems, Man, and Cybernetics*, 6, 610-621.
- Henderson, T. P. (2019). La roya y el futuro del café en Chiapas. *Revista Mexicana de Sociología*, 81(2), 389-416.
- Hotelling, H. (1933). Analysis of a complex of statistical variables into principal components. *Journal of Educational Psychology*, 24(6), 417-441.
- Hu, M. K. (1962). Visual pattern recognition by moment invariants. *IRE Transactions on Information Theory*, 8(2), 179-187.
- Hubert, M.-C. R., & Torres, R. M. (2017). Política pública y sustentabilidad de los territorios cafetaleros en tiempos de roya: Chiapas y Veracruz. *Estudios Latinoamericanos*. Universidad Nacional Autónoma de México.
- Krohling, R. A., Esgario, G. J. M., & Ventura, J. A. (2019). BRACOL - A Brazilian Arabica coffee leaf images dataset to identification and quantification of coffee diseases and pests. *Mendeley Data*, V1. <https://doi.org/10.17632/yy2k5y8mxg.1>
- Lewis, K. P., & Espineli, J. D. (2020). Classification and detection of nutritional deficiencies in coffee plants using image processing and convolutional neural network (Cnn). *International Journal of Scientific & Technology Research*, 9(4), 2076-2081.
- Manso, G. L., Knidel, H., Krohling, R. A., & Ventura, J. A. (2019). A smartphone application to detection and classification of coffee leaf miner and coffee leaf rust. *Computer Vision and Pattern Recognition*.
- Mengistu, A. D., Alemayehu, D. M., & Mengistu, S. G. (2016). Ethiopian coffee plant diseases recognition based on imaging and machine learning techniques. *International Journal of Database Theory and Application*, 9(4), 79-88.
- Montalbo, F. J. P., & Hernandez, A. A. (2020). An optimized classification model for coffea liberica disease using deep convolutional neural networks. *16th IEEE International Colloquium on Signal Processing & Its Applications (CSPA)* (pp. 213-218). IEEE.
- Otsu, N. (1979). A threshold selection method from gray-level histograms. *IEEE Transactions on Systems, Man, and Cybernetics*, 9(1), 62-66.

- Paula, P. V. A. A. D., Pozza, E. A., Alves, E., Moreira, S. I., Paula, J. C. A., & Santos, L. A. (2019). *Infection process of Cercospora coffeicola in immature coffee fruits*. <http://www.sbicafe.ufv.br/handle/123456789/12068>
- Price, T. V., Gross, R., Ho, W. J., & Osborne, C. F. (1993). A comparison of visual and digital image-processing methods in quantifying the severity of coffee leaf rust (*Hemileia vastatrix*). *Australian Journal of Experimental Agriculture*, 33(1), 97-101.
- Pujari, D., Yakkundimath, R., & Byadgi, A. S. (2016). SVM and ANN based classification of plant diseases using feature reduction technique. *International Journal of Interactive Multimedia and Artificial Intelligence*, 3(7), 6-14.
- Qin, F., Liu, D., Sun, B., Ruan, L., Ma, Z., & Wang, H. Luvisi, A. (Ed.) (2016). Identification of Alfalfa Leaf Diseases Using Image Recognition Technology. *PLOS ONE*, 11, e0168274.
- Ramírez, D., & Jiménez, F. G. (2021). Manejo del minador de la hoja (*Leucoptera coffeella*) en el cultivo de café en Costa Rica. *Agronomía costarricense: Revista de ciencias agrícolas*, 45(2), 143-153.
- Sampallo, G. (2003). Reconocimiento de tipos de hojas. Inteligencia Artificial. *Revista Iberoamericana de Inteligencia Artificial*, 7(21), 55-62.
- Vichi, F. F. (2015). La producción de café en México: ventana de oportunidad para el sector agrícola de Chiapas. *Espacio I+ D: Innovación más Desarrollo*, 4(7).
- XinHai, L. (2013). Using random forest for classification and regression. *Chinese Journal of Applied Entomology*, 50, 1190-1197
- Yebasse, M., Shimelis, B., Warku, H., Ko, J., & Cheoi, K. J. (2021) Coffee Disease Visualization and Classification. *Plants*, 10, 1257. <https://doi.org/10.3390/plants10061257>

Notas

- [1] Herramienta para medir y calcular áreas irregulares.
[2] Disponible en <https://data.mendeley.com>

Enlace alternativo

<https://cienciaergosum.uaemex.mx/article/view/16715> (html)