

**Centro de Investigación Científica y de Educación
Superior de Ensenada, Baja California**



**Maestría en Ciencias
en Electrónica y Telecomunicaciones con orientación
en Telecomunicaciones**

**Desarrollo de algoritmos para el reconocimiento de patrones
en imágenes oftálmicas.**

Tesis

para cubrir parcialmente los requisitos necesarios para obtener el grado de
Maestro en Ciencias

Presenta:

Jesús Ismael Casarrubias Garín

Ensenada, Baja California, México

2017

Tesis defendida por

Jesús Ismael Casarrubias Garín

y aprobada por el siguiente Comité

Dr. Roberto Conte Galván

Codirector del Comité

Dr. Gabriel Alejandro Galaviz Mosqueda

Codirector del Comité

Dr. Miguel Ángel Alonso Arévalo

Dr. Salvador Villarreal Reyes

Dr. Jorge Torres Rodríguez



Dr. Miguel Ángel Alonso Arévalo

Coordinador del Programa de Posgrado en Electrónica y Telecomunicaciones

Dra. Rufina Hernández Martínez

Director de Estudios de Posgrado

Jesús Ismael Casarrubias Garín © 2017

Queda prohibida la reproducción parcial o total de esta obra sin el permiso formal y explícito del autor

Resumen de la tesis que presenta **Jesús Ismael Casarrubias Garín**, como requisito parcial para la obtención del grado de Maestro en Ciencias en Electrónica y Telecomunicaciones con orientación en Telecomunicaciones.

Desarrollo de algoritmos para el reconocimiento de patrones en imágenes oftálmicas.

Resumen aprobado por:

Dr. Roberto Conte Galván

Codirector de Tesis

Dr. Gabriel Alejandro Galaviz Mosqueda

Codirector de Tesis

La diabetes es una de las enfermedades que ha visto aumentada su prevalencia durante los últimos años; es una enfermedad crónica-degenerativa que con el paso del tiempo causa daños a los tejidos y desencadena complicaciones potencialmente letales. La retinopatía diabética (RD) es una de las complicaciones más frecuentes, el 95 % de las personas que padecen diabetes eventualmente sufrirá de ésta. Los microaneurismas son lesiones en los vasos sanguíneos, su presencia en la retina es el primer signo que evidencia la RD. El método de diagnóstico es la exploración del fondo de ojo mediante la técnica de oftalmoscopia directa e indirecta; con las cámaras de fondo de ojo, es posible obtener fotografías del fondo de ojo (retinografías). Con el objetivo de reducir la carga de trabajo de los oftalmólogos se han propuesto sistemas de diagnóstico asistido por computadora y de detección automática. Esta tesis propone una metodología para la detección de microaneurismas en imágenes del fondo de ojo siguiendo el esquema general de la detección de objetos; en la metodología se incluye una etapa de optimización basada en el cómputo evolutivo (Algoritmos Genéticos "AG") para mejorar la etapa de generación de candidatos (Segmentación) asumiendo que entre más completas aparezcan las lesiones (microaneurismas), más exacta y precisa será su clasificación y por ende se mejorará la tasa de detección. La base de datos de imágenes del fondo de ojo fue conformada con 58 imágenes de la base de datos DiaretDB1. El objetivo principal es demostrar que la sintonización de parámetros utilizando un algoritmo genético, efectivamente le da mejores oportunidades de detección al sistema. La validación consistió en comparar el número de lesiones presentes en la imagen resultado de aplicar la segmentación y umbralización utilizando 1) parámetros de procesamiento aleatorios (los parámetros con mejor aptitud en las primeras generaciones del AG), 2) los parámetros con mejor aptitud de la última generación del AG. Los métodos fueron implementados en el lenguaje "Python 2.7" a través del entorno de trabajo "Spyder 2"; se aplicaron técnicas de procesamiento digital de imágenes (mejora de contraste, reducción de ruido, segmentación, umbralización) y la parte de la clasificación se desarrolló en el software "Weka 3.6.13".

Palabras Clave: Retinopatía Diabética, Detección de Microaneurismas, Segmentación, Optimización, Algoritmo Genético

Abstract of the thesis presented by **Jesús Ismael Casarrubias Garín**, as a partial requirement to obtain the Master of Science degree in Electronics and Telecommunications with orientation in Telecommunications.

Algorithms design for pattern recognition in ophthalmic images.

Abstract approved by:

Dr. Roberto Conte Galván

Thesis Co-Director

Dr. Gabriel Alejandro Galaviz Mosqueda

Thesis Co-Director

Diabetes is one of the diseases with the highest prevalence among the world population; it's a chronic-degenerative disease that causes severe damage to all body tissues over time and triggers potentially lethal complications. Diabetic retinopathy is one of the most common complications, and 95 % of diabetics will eventually develop it. Microaneurysms are small lesions in the blood vessels, and their presence in the retina is the first clinical sign that evidences diabetic retinopathy. The diagnose method consists in exploring the eye fundus through direct and indirect ophthalmoscopy; thanks to the fundus eye cameras it's possible to obtain photographs of the eye fundus (retinographies). In order to reduce the workload of ophthalmologists, computer assisted diagnosis and automatic detection systems have been proposed. This research proposes a methodology for microaneurysms detection in fundus eye images according to the general scheme for objects detection; in this methodology an optimization stage is included based on evolutionary computation (Genetic Algorithms "GA") to improve the candidate generation stage (Segmentation) under the belief that the better the lesions (microaneurysms) appear in the resulting image, the more accurate and precise will be their classification and thus the detection rate will be improved. DiaretDB1 was used as the fundus eye image database, and 58 images were selected. Gathered results demonstrate that the optimization technique effectively gives better detection opportunities to the system. Validation consisted in comparing the number of lesions appearing in the resulting image after segmentation and thresholding stages using 1) random processing parameters (best fitness among the first generations of the GA), and 2) processing parameters with the highest fitness at the GA last generation. All digital image processing techniques (contrast enhancement, noise reduction, segmentation, thresholding) were implemented using "Python 2.7" through the "Spyder 2" work environment; the classification stage was developed by means of the "Weka 3.6.13" software.

Keywords: Diabetic Retinopathy, Microaneurysm Detection, Segmentation, Optimization, Genetic algorithm

Dedicatoria

Provehito in altum:

*Por ser fruto de su esfuerzo, con amor para mis padres
por sus consejos, enseñanzas, valores y su apoyo
constante en pro de alcanzar mis anhelos, sueños y*

metas:

Marilú Garín e

Ismael Casarrubias

Agradecimientos

Lo que hagas con tu vida será insignificante, pero es muy importante que lo hagas porque nadie más lo hará. -*Gandhi*.

Al Centro de Investigación Científica y de Educación Superior de Ensenada, Baja California por darme la oportunidad de crecer académicamente y permitirme enriquecerme de los conocimientos y experiencias de sus investigadores.

Al Consejo Nacional de Ciencia y Tecnología (CONACyT) por brindarme el apoyo económico con número de becario 338064 para realizar mis estudios de maestría.

A mis padres y hermano por su apoyo incondicional, por acompañarme siempre en el camino de la búsqueda de mi crecimiento y desarrollo pleno.

A mis codirectores de tesis, el Dr. Roberto Conte Galván por creer en este trabajo de investigación, por su confianza, su apoyo constante, por sus lecciones dentro y fuera del aula, sus enseñanzas académicas y lecciones de vida, por sus anécdotas que colaboraron tanto en mi crecimiento y formación profesional como personal; y al Dr. Gabriel Alejandro Galaviz Mosqueda por creer en este trabajo de investigación, sus valiosas contribuciones, su guía para mejorarlo, su compromiso, tiempo y conocimientos compartidos, y por enseñarme a considerar siempre un cambio de variable.

A los miembros del comité de tesis, Dr. Miguel Ángel Alonso Arévalo, Dr. Salvador Villarreal Reyes y Dr. Jorge Torres Rodríguez por sus comentarios y aportaciones para enriquecer este trabajo de investigación y crecer profesionalmente.

A mi novia, Linda, por su apoyo, cariño, amor y compañía constante. Por creer en mí cuando ni yo mismo creí en mí.

Al Dr. Jorge Torres Rodríguez por aceptarme como oyente en sus clases que me permitieron obtener los conceptos y bases fundamentales para el desarrollo de este trabajo de investigación.

A la Dra. Diana Tentori por aceptarme como oyente en su clase de Óptica geométrica y radiometría, y al Dr. Víctor Ruíz por su guía.

A los miembros del grupo ARTS por sus enseñanzas durante el año de cursos y su apoyo y recomendaciones durante el año de tesis.

A mis amigos Telecos, Memo, Andrea, Ramón, Jairo, Javi, Ernesto, Héctor, compañeros de clase, de crossfit, de cubo, de futbol, de salidas, de comida; gracias por su amistad y los buenos momentos que hicieron menos pesado el camino. Sin su amistad y apoyo no veo de que manera hubiera sido posible llegar hasta la recta final.

A mis compañeros de generación y amigos, Alex, Eduardo, Héctor Cuesta, Marco, Alicia, Isaac, Wilbert, Cetina, Marlon, Alberto, y demás amigos del CICESE, Roilhi, Roger, Calixto, Giovanni, Alejandro, Sergio, Víktor, por su amistad, sus eventuales aportaciones a mi formación profesional, los debates, discusiones, por las lecturas recomendadas, por enseñarme a ver la otra cara de la moneda y saber que existe más de una manera de hacer una misma cosa; por las enseñanzas que cada uno dejó en mí, gracias.

Al oftalmólogo Dr. Jairo Enoc Rodríguez Lizondro por su tiempo, asesoría y guía.

Al personal del CICESE (Secretarias, Personal Médico, Investigadores, Técnicos, Guardias) por su apoyo y guía en trámites, orientación y atención, por colaborar a sentirme cobijado y permitir desarrollarme plenamente durante mi estancia en el CICESE.

Tabla de contenido

	Página
Resumen en español	v
Resumen en inglés	vii
Dedicatoria	ix
Agradecimientos	xi
Lista de figuras	xvii
Lista de tablas	xxi
Capítulo 1. Introducción	1
1.1 Motivación	2
1.2 Planteamiento del problema	5
1.3 Objetivos	5
1.3.1 Objetivo general	5
1.3.2 Objetivos específicos	5
1.4 Organización de la tesis	6
Capítulo 2. Análisis de imágenes oftálmicas	8
2.1 Introducción	8
2.2 Antecedentes	8
2.3 Fundamentos	23
2.4 Procesamiento digital de imágenes	30
2.4.1 Adquisición de la imagen	30
2.4.2 Canales y espacios de color	32
2.4.3 Técnicas de reducción de ruido	35
2.4.4 Modificaciones al histograma	41
2.4.5 Segmentación	44
2.4.5.1 Umbralización	52
2.5 Reconocimiento de patrones	52
2.5.1 Algoritmo de etiquetaje	53
2.5.2 Extracción de características	55
2.5.3 Clasificación	56
2.5.3.1 Validación	57
2.5.3.2 Métricas	58

Capítulo 3.	Algoritmos genéticos	62
3.1	Introducción	62
3.2	Técnicas de búsqueda	62
3.2.1	Cómputo evolutivo	63
3.2.2	Algoritmos genéticos	64
3.3	Definiciones y operadores genéticos	65
3.3.1	Individuos	65
3.3.2	Genes	67
3.3.3	Población	69
3.3.4	Aptitud	72
3.3.5	Codificación	73
3.3.6	Reproducción	76
3.3.6.1	Selección	76
3.3.6.2	Cruzamiento	80
3.3.6.3	Mutación e inversión	83
3.3.6.4	Reemplazo	86
3.3.7	Criterio de convergencia	87
Capítulo 4.	Metodología	89
4.1	Introducción	89
4.2	Banco de imágenes	90
4.3	Procesamiento digital de imágenes oftálmicas	92
4.3.1	Preprocesamiento	93
4.3.1.1	Adquisición de imágenes	93
4.3.1.2	Separación de canales	94
4.3.1.3	Reducción de ruido	96
4.3.1.4	Mejora de contraste	97
4.3.2	Generación de regiones candidatas a MA	97
4.3.2.1	Segmentación	97
4.3.2.2	Umbralización	98
4.4	Búsqueda de soluciones cuasi-óptimas	99
4.4.1	Propuesta de algoritmo genético	101
4.4.1.1	Codificación (Representación)	103
4.4.1.2	Selección	106
4.4.1.3	Cruzamiento	107
4.4.1.4	Mutación	108
4.4.1.5	Reemplazo	110
4.4.1.6	Convergencia	111
4.4.2	Función objetivo	111
4.4.3	Clasificación	112
4.4.4	Extracción de características (descriptores)	113
4.4.5	Clasificación de candidatos	114
Capítulo 5.	Experimentación y resultados	116
5.1	Introducción	116
5.2	Experimento 1	116
5.2.1	Descripción	116

5.2.2	Resultados	117
5.3	Experimento 2	118
5.3.1	Descripción	118
5.3.2	Resultados	119
5.4	Experimento 3	121
5.4.1	Descripción	121
5.4.2	Resultados	122
5.5	Experimento 4: Selección de características	124
5.5.1	Descripción	124
5.5.2	Resultados	125
5.6	Experimento 5: Clasificación de candidatos	127
5.6.1	Descripción	127
5.6.2	Resultados	127
5.7	Conclusiones	134
Capítulo 6.	Conclusiones y trabajo a futuro	136
6.1	Introducción	136
6.2	Conclusiones	136
6.3	Aportaciones	139
6.4	Trabajo a futuro	140
Literatura citada		142
Anexos		148

Lista de figuras

Figura		Página
1	a) Se observan microaneurismas (MAs), exudados (EXs) y hemorragias (HMs), conformando la retinopatía diabética no proliferativa. b) Se observa crecimiento anormal de los vasos sanguíneos (neovascularización), indicando la presencia de retinopatía diabética proliferativa. Adaptación de Amin <i>et al.</i> (2016)	11
2	Esquema del fondo de ojo. Imagen tomada de Riordan-Eva (2012).	13
3	Configuración óptica de la cámara de fondo de ojo de bajo costo propuesta por Tran <i>et al.</i> (2012). Imagen tomada de Tran <i>et al.</i> (2012).	14
4	Ejemplo de un arreglo de sensores.	24
5	Imagen digital de $m \times n$, $m=8$ y $n=6$	26
6	a) Escena real b) Digitalización de la escena. Imagen tomada de Gonzalez y Woods (2011).	26
7	Tratamiento de imágenes.	28
8	Esquema general de la detección de objetos.	29
9	Captura de una escena real. Imagen tomada de Gonzalez y Woods (2011).	31
10	Modelo de color RGB.	33
11	Modelo de color HSV.	34
12	Modelo de color HSL.	35
13	Modelo de ruido en una imagen.	36
14	Filtrado espacial.	38
15	Ejemplo de máscaras para filtro de suavizado.	39
16	Ejemplo de máscaras para filtro gaussiano.	40
17	Ejemplo del histograma de una imagen.	42
18	Ejemplo de mejora de contraste (Algoritmo CLAHE).	42
19	Ejemplo de erosión de una imagen binaria.	46
20	Ejemplo de dilatación de una imagen binaria.	47
21	Operador Hit-Miss.	48
22	Ejemplo de operadores morfológicos aplicados a una imagen de niveles de gris.	50
23	a) Perfil de intensidad b) Apertura c) Resultado de la apertura d) Cierre e) Resultado del cierre.	51
24	Algoritmo de etiquetaje: Casos posibles.	54

Figura	Página
25 Ejemplo del algoritmo de etiquetaje.	55
26 Ejemplo de validación cruzada.	58
27 Matriz de confusión.	59
28 Ejemplo de curva ROC.	61
29 Algoritmo genético simple.	66
30 Representación del genotipo y fenotipo (adaptación de Sivanandam y Deepa (2008)).	66
31 Métodos de codificación de cromosomas.	67
32 Fenotipo de los cromosomas propuestos en la figura 31.	67
33 Estructura de un cromosoma.	68
34 Fenotipo del cromosoma propuesto en la Figura 33.	68
35 Población de 10 individuos.	70
36 Diversidad en la población. Cruces azules: Alta diversidad. Puntos rojos: Baja diversidad. Adaptación de MathWorks (2005).	71
37 Paisaje de aptitud.	73
38 Mapeo del espacio de los cromosomas al espacio de soluciones (adaptación de Gen y Cheng (2000)).	74
39 Mapeo de cromosomas a soluciones (adaptación de Gen y Cheng (2000)).	75
40 Selección por torneo con diferentes valores de S (tamaño de torneo). Imagen tomada de Goldberg y Deb (1991).	78
41 Ejemplo de muestreo estocástico universal.	79
42 Cruzamiento "scanning".	82
43 Esquema general del sistema propuesto.	90
45 Esquema general de la detección de objetos.	92
44 a) Imagen 27 (original), b) Imagen ground truth exudados suaves, c) Imagen ground truth hemorragias, d) Imagen ground truth microaneurismas, e) Imagen ground truth exudados duros.	92
46 Penetración retiniana de tres longitudes de onda distintas. Imagen tomada de Manzano (2004).	95
47 Canales RGB de la imagen 017.	95
48 Canales HSV de la imagen 017.	96
49 Ejemplo de reducción de ruido.	96

Figura	Página
50 Ejemplo de CLAHE.	97
51 Segmentación mediante la transformada de bottom-hat.	98
52 Umbralización de la imagen resultado de la transformada de bottom-hat. . .	99
53 Importancia de la optimización de parámetros: a) Imagen 17 sin preproce- samiento, b) Imagen 17 procesada con parámetros sub-óptimos, c) Imagen 17 procesada con parámetros cuasi-óptimos.	100
54 Evaluación de la generación de candidatos: a) Objeto a detectar (cuadro negro) y detección (azul), b) Métricas de evaluación de la segmentación. .	101
55 Esquema general del algoritmo genético simple.	102
56 a) Matriz multidimensional genotipo, b) Matriz multidimensional fenotipo. . .	103
57 Cromosoma propuesto.	104
58 Diccionario para el tamaño de la ventana del filtro mediano.	105
59 Diccionario para el elemento estructural.	105
60 Cruzamiento en un punto.	107
61 Cruzamiento en dos puntos.	108
62 Mutación de tipo "flipping".	109
63 Mutación de tipo "interchanging".	110
64 Mutación de tipo "reversing".	110
65 Número de microneurismas detectados por cada solución.	117
66 a) Evolución de la población, b) Paisaje de aptitud.	120
67 Curva ROC para el modelo "Näive Bayes Multinomial Updateable" con 28 descriptores.	128
A.1 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	162
A.2 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	163
A.3 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	164
A.4 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	165
A.5 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	166

Figura	Página
A.6 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	167
A.7 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	168
A.8 Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.	169

Lista de tablas

Tabla	Página
1 Escala clínica internacional de gravedad de la retinopatía diabética.	12
2 Escala clínica internacional de gravedad de edema macular diabético.	13
3 Comparación de los países con mayor prevalencia de diabetes (FID: Federación Internacional de la Diabetes (2011)) contra su respectivo número de oftalmólogos (ICO: International Council of Ophtalmology (2012)).	16
4 Comparación de metodologías para la detección de microaneurismas.	23
5 Comparación entre conceptos de evolución natural y la terminología de los algoritmos genéticos.	65
6 Soluciones propuestas.	117
7 Porcentaje de área de microaneurismas reales detectada.	118
8 Soluciones propuestas.	120
9 Número de microaneurismas reales pre-detectados.	121
10 Evaluación de las soluciones encontradas por el Algoritmo Genético.	123
11 Evaluación de 22 modelos de clasificación con 28 descriptores.	129
12 Evaluación de 22 modelos de clasificación con 9 descriptores.	130
13 Evaluación de 22 modelos de clasificación con 8 descriptores.	131
14 Evaluación de 22 modelos de clasificación con 13 descriptores.	132
15 Comparación de la metodología propuesta con metodologías en el estado del arte.	134
A.1 Imágenes utilizadas para los experimentos con su respectivo número de microaneurismas de acuerdo a los datos ground truth.	148
A.2 Imágenes de control de la segmentación del experimento 1.	149
A.3 Imágenes de control de la segmentación del experimento 2.	155

Capítulo 1. Introducción

En el presente trabajo de investigación, se presenta una metodología para la detección de retinopatía diabética (RD) en etapa no proliferativa, lo que se logra detectando pequeñas lesiones que aparecen en los vasos sanguíneos de retina, i.e. microaneurismas. Para lograr esto se propuso una metodología que combina el análisis de imágenes con técnicas de optimización guiadas y aleatorias basadas en poblaciones de soluciones, i.e. cómputo evolutivo.

El enfoque utilizado para analizar las imágenes se basa en el esquema de la detección de objetos, lo que consiste en conformar un banco de imágenes, dividido en dos conjuntos: de entrenamiento y de prueba. Se implementan técnicas para mejorar las imágenes de estos dos conjuntos, y técnicas para aislar los microaneurismas de las demás estructuras del ojo que no interesan, lo que nos genera regiones que pueden ser microaneurismas; a esta etapa se le considera una predetección, las regiones son candidatos y no implica que todos sean una lesión. Para discernir si las regiones se tratan o no de microaneurismas, se utiliza un modelo de clasificación basado en vectores de características, donde de cada región se extraen características o descriptores (e.g. intensidad promedio en canal verde, intensidad promedio en canal de brillo, área de la región, perímetro de la región, circularidad, etc.), el modelo genera una superficie (hiperplano) de decisión que divide las regiones que son un microaneurisma de las que no lo son.

El problema que se encontró, es que una vez implementadas las técnicas de procesamiento digital de imágenes, se deben encontrar los parámetros de estas técnicas para que en todas las imágenes del banco de datos se puedan identificar sino todos, la mayoría de los microaneurismas. Se tiene un problema de optimización, por lo que en esta tesis, se hace uso del cómputo evolutivo bajo el enfoque de los algoritmos genéticos para optimizar los parámetros de procesamiento.

La selección de características, es una tarea que no se debe trivializar y requiere de un análisis estadístico para seleccionar las que en conjunto con el modelo de clasificación generan mejores resultados. La etapa de selección es de gran importancia, debido a que hay características que no permiten discernir de sí una región es o no es un micro-

aneurisma; estos descriptores, es mejor desecharlos para reducir la dimensionalidad y disminuir el coste computacional y tiempo de procesamiento en la etapa de extracción de características de los candidatos y clasificación de éstos.

La estrategia para seleccionar las características se basó en los métodos para evaluación de atributos (descriptores) implementados en el software WEKA; una vez que se encontraron los mejores descriptores se procedió a clasificar cada región utilizando el mismo software y probando distintos modelos de clasificación, e.g. KNN¹ (K- vecinos más cercanos), Nāive-Bayes, adaboost, perceptrón multicapa, regresión logística, etc; con el objetivo de obtener una comparación del rendimiento de los modelos de clasificación implementados en WEKA para la detección de microaneurismas.

1.1. Motivación

La motivación del presente trabajo de investigación radica en el contexto de la telemedicina en México y a nivel mundial, cuyo perfil epidemiológico se encuentra en una transición en la cual los índices de mortalidad y morbilidad se han tornado hacia enfermedades cardiovasculares y crónico-degenerativas.

El problema se agrava por la alta prevalencia de la diabetes, lo que ha aumentado la carga de trabajo de los oftalmólogos, debido a que de acuerdo con las guías de práctica clínica, una vez que se identifica la presencia de algún tipo de diabetes, al menos una vez al año se debe acudir al oftalmólogo para descartar indicios de retinopatía diabética. Según estimaciones de la FID: Federación Internacional de la Diabetes (2013), el número de diabéticos a nivel mundial aumentará en 205 millones de personas para el año 2035.

La diabetes es una de las enfermedades de mayor relevancia en los últimos años debido a su alto índice de prevalencia en la población mexicana y en general, en la población mundial. La diabetes es una de las causas más importantes de mortalidad en el país (México), y es la primera causa de ceguera adquirida (SSA: Secretaría de Salud, 2013, 2014).

En los últimos años, se ha observado una creciente adopción en el uso de las tecno-

¹KNN: Del inglés K-Nearest Neighbor

logías de la información y comunicación (TIC's) en diferentes aspectos de la vida cotidiana con el objetivo de facilitar múltiples tareas al ser humano y mejorar su calidad de vida. Los grandes avances tecnológicos en la electrónica y las TIC's han llevado a que los sistemas de salud adopten el uso de éstas tecnologías para la optimización de sus procesos; han surgido líneas de investigación (Telemedicina) que buscan brindar servicios de salud a distancia y la implementación de sistemas de diagnóstico asistido por computadora, para aumentar la eficiencia y el alcance de los servicios de salud.

La Retinopatía Diabética (RD) es una de las complicaciones más frecuentes de la diabetes; la oftalmología es la rama de la medicina que estudia y trata clínicamente esta complicación; la tele-oftalmología por su parte, es la rama de la telemedicina que pretende reducir la brecha del sin número de personas sin acceso a un médico especialista en oftalmología. La Teleoftalmología pretende implementar sistemas que permitan evaluar clínicamente a pacientes de localidades remotas sin necesidad de un oftalmólogo presente, utilizando las TIC's para transmitir y recibir información clínica válida y videoconferencia para su diagnóstico. En Reino Unido se ha configurado un sitio de Internet dedicado en el "Moorfields Eye Hospital" con el objetivo de ofrecer servicios de tele-oftalmología a países africanos como Ghana, Gambia y Tanzania, entre otros (Yogesan *et al.*, 2008).

En la India se ha puesto en marcha un camión (unidad de telemedicina) configurado especialmente para llevar atención oftálmica a regiones denominadas inalcanzables. El camión está seccionado en tres partes, la primera consiste en la cabina, donde se encuentra el chofer y viaja el equipo médico y técnico. La segunda sección es el área de atención y consulta, que cuenta con los dispositivos requeridos y una sala de espera. En la tercera sección se encuentra un almacén, las fuentes de energía y demás equipo. El camión cuenta además con una antena satelital, debido a que el camión se encuentra conectado con los especialistas a través de una red VSAT (Yogesan *et al.*, 2008).

Otro ejemplo de sistema de tele-oftalmología se encuentra en Australia, un país con una gran extensión territorial y con una densidad de población relativamente baja. Esto genera el problema de que la población se encuentra dispersa, contando así con localidades remotas donde no se cuentan con especialistas tanto en oftalmología como de

otras ramas de la medicina. Para subsanar esta problemática se han planteado servicios oftálmicos basados en Internet, para que de esta manera una mayor parte de la población tenga acceso a estos especialistas. Se desarrolló el servicio de tele-oftalmología en el “Center of Excellence in e-Medicine” en conjunto con “Lions Eye Institute” en Perth. El sistema almacena y transmite los datos multimedia de los pacientes a una base de datos segura (Yogesani *et al.*, 2008).

El estado del arte de la teleoftalmología tiene otro enfoque que está orientado hacia el procesamiento digital de imágenes para el desarrollo e implementación de sistemas de diagnóstico asistido por computadora; de manera que se han diseñado códigos, algoritmos y metodologías que permiten mejorar las imágenes, observar regiones de interés en las fotografías e incluso detectarlas automáticamente (detección de objetos) (Fleming *et al.*, 2007; Niemeijer *et al.*, 2007; Sopharak *et al.*, 2013; Walter *et al.*, 2007; Arenas-Cavalli *et al.*, 2015). Este enfoque pretende colaborar con el problema de la detección de retinopatía diabética, que debido a la alta prevalencia de la diabetes, la carga de trabajo a los oftalmólogos se ha visto aumentada y para el año 2035 el panorama empeorará según la estimación de la Federación Internacional de la Diabetes.

El objetivo de la adopción de técnicas de detección de objetos como herramienta para detectar retinopatía diabética, se centra principalmente en detectar esta complicación en etapa temprana, de manera que las técnicas de detección se enfocan en la detección de microaneurismas (primer signo de retinopatía diabética). Para poder implementar estas técnicas, el primer paso es la exploración del fondo del ojo y la obtención de fotografías del mismo por medio de cámaras de fondo de ojo.

Las cámaras de fondo de ojo (fundus camera) son básicamente dispositivos que se componen de un microscopio especializado de baja potencia con una cámara. Este instrumento permite registrar el fondo de ojo obteniendo una serie de fotografías de la retina para ser utilizadas para generar datos diagnósticos y documentar la patología ocular del fondo de ojo (Tran *et al.*, 2012).

En conclusión, la motivación del presente trabajo de investigación es diseñar una metodología capaz de detectar microaneurismas utilizando fotografías del fondo de ojo para,

1) reducir la carga de trabajo de los oftalmólogos, 2) aumentar la capacidad de atención de pacientes en riesgo de desarrollar retinopatía diabética, 3) diagnosticar y tratar oportunamente la RD, 4) evitar la ceguera adquirida de diabéticos.

1.2. Planteamiento del problema

El problema en específico al que se pretende dar solución, es al que dado un conjunto de imágenes de fondo de ojo, se optimicen los parámetros de procesamiento para generar los candidatos en los que la mayor cantidad de éstos sean microaneurismas reales (verdaderos positivos).

Los algoritmos genéticos han sido utilizados con anterioridad para optimizar la detección de objetos en imágenes. Se plantea la hipótesis de que utilizando algoritmos genéticos es posible sintonizar los parámetros de procesamiento digital de imágenes, y que esto permitirá mejorar la etapa de detección de regiones de interés; encontrando así un mayor número de microaneurismas reales (verdaderos positivos) y reduciendo a la par el número de falsos positivos y falsos negativos.

1.3. Objetivos

1.3.1. Objetivo general

El objetivo de este proyecto de investigación consiste en el análisis de imágenes de la retina para mejorar y resaltar zonas de interés con el objetivo de facilitar la identificación de patrones que deriven en la identificación de indicios de retinopatía diabética en su etapa no proliferativa, donde específicamente se desea detectar microaneurismas.

1.3.2. Objetivos específicos

Los objetivos específicos de este proyecto se describen a continuación.

- Estudio de: la estructura y funcionamiento del ojo humano; los conceptos básicos de oftalmología y las enfermedades oftálmicas; los algoritmos para el procesamiento de imágenes oftálmicas; los fundamentos de reconocimiento de patrones.

- Conformación de un banco de datos de imágenes de la retina a partir de la base de datos DIARETDB1 (Kauppi *et al.*, 2007).
- Generar imágenes binarias de cada una de las imágenes en el banco de datos, resaltando únicamente los microaneurismas de acuerdo a los datos de referencia disponibles en la base DIARETDB1.
- Programación de técnicas de procesamiento digital de imágenes para mejorar las imágenes oftálmicas y facilitar la detección de microaneurismas.
- Implementación de la técnica de cómputo evolutivo de algoritmos genéticos como herramienta para la optimización de los parámetros de procesamiento.
- Extraer descriptores y seleccionar los que ofrecen mayor capacidad de discriminación utilizando WEKA.
- Obtener una comparación de rendimiento de múltiples modelos de clasificación utilizando el software WEKA.

1.4. Organización de la tesis

La estructura del presente trabajo de investigación se describe en los siguientes párrafos:

- Capítulo I: Introducción. Se presenta una pequeña introducción con el objetivo de establecer qué se va a hacer, para qué y cómo.
- Capítulo II: Retinopatía diabética. Se define la enfermedad de retinopatía diabética, para lo que en primera instancia se define lo que es la diabetes y se describe la estructura del ojo humano para poder identificar las regiones que afecta y el alcance de la enfermedad. Brevemente se describen los métodos comúnmente utilizados para su detección y algunos de los antecedentes de los Sistemas de ayuda computarizada para la detección de retinopatía diabética.
- Capítulo III: Análisis de imágenes oftálmicas. En este capítulo se describen los conceptos de análisis de imágenes, las dos etapas en las que se subdivide; se detalla

el esquema general de la detección de objetos y algunas de las técnicas utilizadas comúnmente en cada fase con el respectivo objetivo de cada fase. Respecto a la etapa de reconocimiento de patrones, se describen algunas métricas que se suelen utilizar para evaluar el rendimiento del sistema de detección implementado.

- **Capítulo IV: Algoritmos genéticos.** En este capítulo se aborda la técnica de búsqueda de algoritmos genéticos. Se describen los fundamentos, características y cada una de las etapas. En cada etapa se presentan una o más opciones que pueden ser adoptadas.
- **Capítulo V: Metodología.** Se describe la metodología del trabajo de investigación, el objetivo es que éste sirva de guía para trabajos futuros, para que pueda ser reproducido total o parcialmente y tener un marco de referencia sobre el cuál comparar ambos trabajos.
- **Capítulo VI: Experimentos y resultados.** Se describen los experimentos realizados y su respectivo objetivo. Así mismo, se presentan los resultados de dicho experimento y se realiza un breve análisis de éstos.
- **Capítulo VII: Conclusiones y trabajo a futuro.** En última instancia se presentan las conclusiones de este trabajo de investigación, de igual manera se destacan las principales aportaciones y se establece un marco de tareas identificadas con potencial para explorarse en mayor detalle y mejorar trabajos de investigación en un tema relacionado al presente trabajo de investigación.

Capítulo 2. Análisis de imágenes oftálmicas

2.1. Introducción

El presente trabajo de investigación tiene como objetivo detectar microaneurismas en imágenes de fondo de ojo. Mediante el uso de herramientas para el tratamiento de imágenes; por lo que en este capítulo se describe el concepto de análisis de imágenes en sus dos etapas principales; de las cuales se detallan sus conceptos básicos y fundamentos en el presente capítulo con el objetivo de facilitar la comprensión de la etapa de análisis de imágenes.

El capítulo se divide en tres secciones: la primera describe conceptos básicos, tales como la definición de imagen, pixel, bandas espectrales y análisis de imágenes; la segunda sección describe el procesamiento digital de imágenes y el esquema general de procesamiento; finalmente, la tercera etapa describe el reconocimiento de patrones y sus etapas.

2.2. Antecedentes

La diabetes mellitus es un grupo heterogéneo de trastornos que se caracterizan por concentraciones elevadas de glucosa en sangre (IMSS: Instituto Mexicano del Seguro Social, 2014). Es un trastorno crónico que se produce ya sea cuando el páncreas no produce suficiente insulina (hormona reguladora de glucosa en sangre) o cuando el cuerpo no puede utilizar eficazmente la insulina producida. Es una de las principales enfermedades no transmisibles a nivel mundial. El número de casos y la prevalencia de la diabetes han aumentado de forma constante durante las últimas décadas. Se ha pronosticado que para el año 2035 el número de adultos con diabetes se duplicará en todo el mundo, desde 177 millones en 2000 a 370 millones (WHO: World Health Organization, 2016).

De acuerdo a la ADA: American Diabetes Association (2015), la diabetes se clasifica en las siguientes categorías:

1. Diabetes tipo 1, es un trastorno autoinmune, el cual consiste en la destrucción de células del páncreas con déficit absoluto de insulina, lo que requiere la administración de insulina exógena para regular los niveles de glucosa.

2. Diabetes tipo 2, constituye una pérdida progresiva de la secreción de insulina, lo que conlleva a una menor sensibilidad de los tejidos a la acción metabólica de la hormona, haciendo una resistencia a la insulina. Representa 90 % del total de casos de diabetes mellitus.
3. Diabetes Mellitus Gestacional (DMG), diabetes que se diagnostica en el segundo o tercer trimestre del embarazo.
4. Diabetes específica por otras causas (por ejemplo: MODY (Diabetes de comportamiento maduro en jóvenes), fibrosis quística, diabetes inducida por medicamentos).

Los criterios de diagnóstico de la diabetes tipo 2 y la diabetes tipo 1 de la ADA son los mismos. Los criterios diagnósticos consisten en identificar cualquiera de los siguientes:

1. Glucosa en ayuno $> 126\text{mg/dL}$ (sin haber tenido ingesta calórica en las últimas 8 horas)
2. Glucosa plasmática a las 2 horas $> 200\text{mg/dL}$ durante una prueba oral de tolerancia a la glucosa. La prueba debe ser realizada con una carga de 75 gramos de glucosa anhidra disuelta en agua.
3. Hemoglobina glicosilada $> 6.5\%$.
4. Paciente con síntomas clásicos de hiperglucemia o crisis hiperglucémica con una glucosa al azar $> 200\text{mg/dL}$.

La diabetes es una enfermedad crónica que requiere atención médica continua y educación al paciente para su propio manejo; para prevenir complicaciones agudas y reducir el riesgo de complicaciones a largo plazo (ADA: American Diabetes Association, 2015). La retinopatía diabética es la complicación microvascular de la diabetes mellitus más común y una de las principales causas de ceguera y discapacidad visual en personas en edad económicamente activa.

La retinopatía diabética es una causa importante de pérdida visual evitable en países desarrollados y una de las principales causas de ceguera en países de ingresos medios

(Arroyo, 2015). Los reportes indican que de los 37 millones de personas ciegas en el mundo, el 5% corresponde a retinopatía diabética y se estima que hasta el 39% de los recién diagnosticados diabéticos ya presentan algún grado de retinopatía. Se presenta en el 27% de los pacientes que tiene entre los 5 y 10 años de evolución de la enfermedad, en el 71 % a 90 % de aquellos con más de 10 años y en el 95% después de 20 años; de estos entre el 30% – 50% desarrollan una etapa proliferativa (SSA: Secretaría de Salud, 2014; Dr. Ariel Prado-Serrano, 2009).

En México, los resultados de la Encuesta Nacional de Salud y Nutrición del año 2012, muestran que 3 de cada 4 diabéticos, requieren de un mayor control del padecimiento de diabetes mellitus que permita reducir las complicaciones que se presentan a largo plazo. Las más frecuentes son baja visual con 3 millones (47.6%) y daños a la retina con 889 mil (13.9%). Es la principal causa de años de vida saludables perdidos (AVISA) y contribuye con 13% del total de AVISA por la población derechohabiente (SSA: Secretaría de Salud, 2014; Gutiérrez *et al.*, 2012).

El suministro de sangre hacia todas las capas de la retina se lleva a cabo a través de los vasos sanguíneos microvasculares, que son susceptibles a cambios en los niveles de glucosa. Cuando una gran cantidad de glucosa o fructosa (hiperglucemia) se reúne en la sangre, se produce una distribución insuficiente de oxígeno a las células de la retina y los vasos sanguíneos se vasoconstruyen (Amin *et al.*, 2016).

Los microaneurismas indican que existe cierre capilar; la pared debilitada de los vasos sanguíneos puede romperse, lo cual puede generar hemorragias, o permitir la fuga de líquido intravascular, que ocasiona un edema en el humor vítreo alrededor del vaso sanguíneo roto; el edema macular es la causa más frecuente de pérdida visual en los pacientes diabéticos con retinopatía (Sorrentino *et al.*, 2016; Coleman *et al.*, 2016).

El diagnóstico de retinopatía diabética incluye microaneurismas (MAs), exudados (EXs), y hemorragias (HMs), así como el crecimiento anormal de los vasos sanguíneos (Figura 1). Se realiza por medio de exploración clínica, mediante oftalmoscopia indirecta, para tener una visión de fondo de ojo. Normalmente tiene dos etapas diferentes nombradas como retinopatía diabética proliferativa y no proliferativa (Amin *et al.*, 2016).

Las mejores estrategias para identificar a los pacientes con retinopatía diabética que amenaza la visión, son las que califican la gravedad de las lesiones de acuerdo con criterios específicos. A continuación se presenta la escala clínica internacional de gravedad de la retinopatía diabética (Tabla 1), observada a través de oftalmoscopia con midriasis farmacológica (dilatación de pupila), ya que sin la misma la detección disminuye hasta en un 50% (Lima Gómez, 2006).

En cualquier grado de retinopatía diabética puede presentarse edema que altere la función visual. Cuando ese edema se localiza en la mácula requiere un tratamiento intensivo (por tratarse del área de visión fina), así como una calificación adicional a la del grado de la retinopatía diabética (Tabla 2).

El método de diagnóstico de retinopatía diabética consiste en evaluar el fondo de ojo y se basa en los exámenes oftalmológicos básicos de oftalmoscopia directa e indirecta para evaluar la fisiología y anatomía de éste (ojo derecho e izquierdo); ambos exámenes básicos tienen el objetivo de evaluar la retina y detectar anomalías, e.g. microaneurismas, exudados, hemorragias, desprendimiento de retina, crecimiento anormal de vasos sanguíneos.

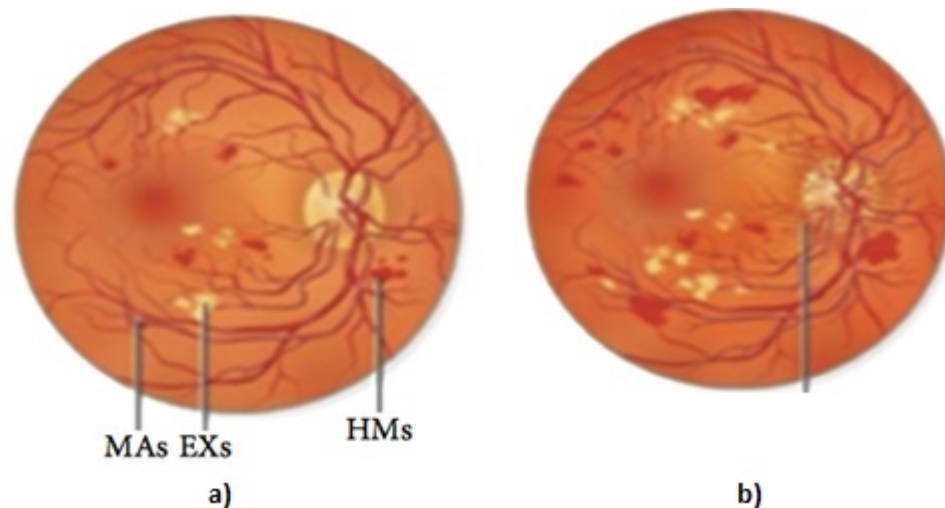


Figura 1: a) Se observan microaneurismas (MAs), exudados (EXs) y hemorragias (HMs), conformando la retinopatía diabética no proliferativa. b) Se observa crecimiento anormal de los vasos sanguíneos (neovascularización), indicando la presencia de retinopatía diabética proliferativa. Adaptación de Amin *et al.* (2016)

Tabla 1: Escala clínica internacional de gravedad de la retinopatía diabética.

Nivel de severidad propuesto	Hallazgos en oftalmoscopia con dilatación
Sin retinopatía aparente	Sin alteraciones
Retinopatía diabética no proliferativa leve	Sólo microaneurismas
Retinopatía diabética no proliferativa moderada	Más que sólo microaneurismas pero menos que retinopatía diabética no proliferativa severa.
Retinopatía diabética no proliferativa severa	Cualquiera de los siguientes: ■ Más de 20 hemorragias retinianas en cada uno de los cuatro cuadrantes. ■ Tortuosidad (arrosariamiento) venosa en dos o más cuadrantes ■ Anormalidades microvasculares intrarretinianas en uno o más cuadrantes y sin signos de retinopatía proliferativa
Retinopatía diabética proliferativa	Uno o más de los siguientes: ■ Neovascularización. ■ Hemorragia vítrea o prerretiniana .

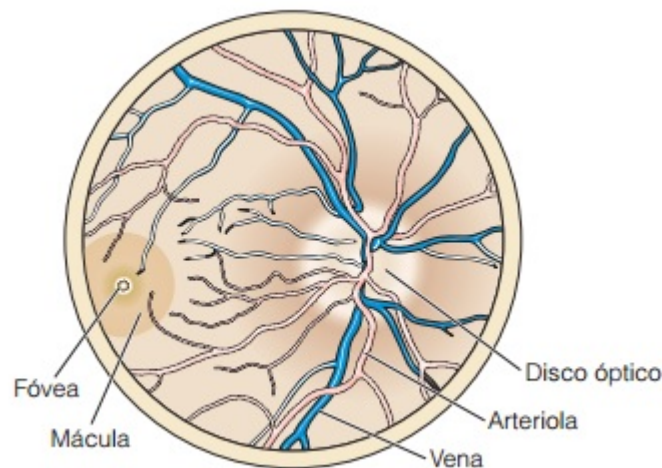
Muchos pacientes no reciben el diagnóstico oportuno ni el tratamiento adecuado de la enfermedad, por lo que se ha considerado la evaluación con apoyo de sistemas de telemedicina orientados a la atención oftalmológica para identificar la retinopatía diabética a distancia, teniendo el potencial de mejorar el acceso a la atención y preservar la visión, sin embargo no hay normas universales respecto a la elección de la cámara, ni un protocolo para la utilización e implementación de sistemas de tele-oftalmología. Por otro lado, se han planteado sistemas de detección automática de retinopatía diabética pero no existe una metodología lo suficientemente robusta de acuerdo a las normas internacionales (Horton *et al.*, 2016).

La exploración del fondo de ojo (ver Figura 2) busca servir como herramienta de exploración de la retina para la detección de anormalidades (e.g. retinopatía diabética) y valoración del grado de avance de éstas. Las técnicas basadas en oftalmoscopia se basan en métodos complicados, por lo que es necesario que sea realizada por un of-

Tabla 2: Escala clínica internacional de gravedad de edema macular diabético.

Nivel de severidad propuesto	Hallazgos en oftalmoscopia con dilatación
Edema macular aparentemente ausente	No hay engrosamiento de retina ni exudados en el polo posterior
Edema macular aparentemente presente	Aparente engrosamiento de la retina y exudados en el polo posterior
Si existe Edema macular presente se clasifica como:	■ Leve: engrosamiento de la retina o exudados en el polo posterior pero alejados del centro de la mácula. ■ Moderado: engrosamiento de la retina o exudados en el polo posterior cercanos al centro de la mácula sin compromiso del centro. ■ Severo: engrosamiento de la retina o exudados en el polo posterior con compromiso del centro de la mácula.

talmólogo experto (Sender Palacios, 2007). De acuerdo a Chia y Yap (2004), la fotografía de la retina logra una alta sensibilidad en la captura de lesiones en la retina características de retinopatía diabética y es comparable a un examen clínico por parte de un oftalmólogo.

**Figura 2:** Esquema del fondo de ojo. Imagen tomada de Riordan-Eva (2012).

La fotografía de fondo de ojo ha surgido como una alternativa para la evaluación en busca de retinopatía diabética principalmente porque además de registrar los detalles del

fondo de ojo, permite documentarlos para el estudio futuro y para tener una comparación y control de la evolución de la enfermedad (Riordan-Eva, 2012). De acuerdo a Sender Palacios (2007), las principales ventajas del método son: es fácil de utilizar, eficaz, barato y accesible para el paciente.

La fotografía del fondo de ojo permite al médico-oftalmólogo observar y documentar el fondo de ojo del paciente, i.e. las estructuras más importantes de la parte posterior del globo ocular; se obtiene con una cámara especial conocida como cámara fundus o cámara de fondo de ojo.

La cámara de fondo de ojo es un arreglo óptico (ver Figura 3) que se basa en dos principales procedimientos: la formación de imágenes y la radiometría. Para la captura de las imágenes puede requerir la midriasis del ojo del paciente, de ahí la existencia de cámaras midriáticas y no midriáticas.

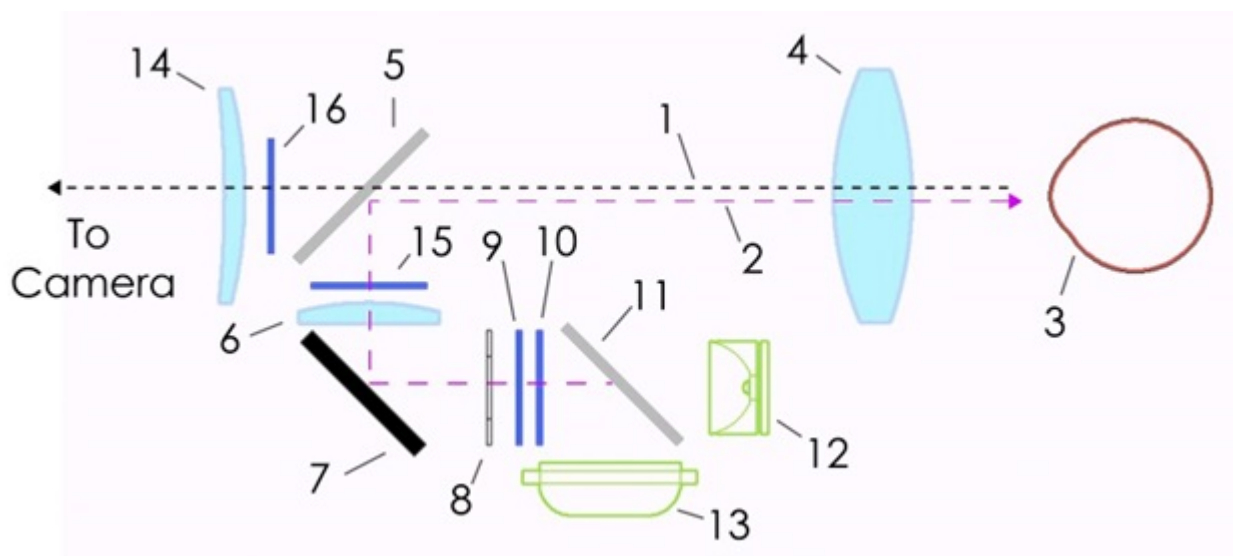


Figura 3: Configuración óptica de la cámara de fondo de ojo de bajo costo propuesta por Tran *et al.* (2012). Imagen tomada de Tran *et al.* (2012).

Con la evolución de las cámaras de fondo de ojo y la facilidad de la técnica de fotografía de fondo de ojo en comparación con la oftalmoscopia indirecta, la exploración del fondo de ojo mediante ésta se ha convertido en el estándar de oro; es el método predilecto de acuerdo a las guías de práctica clínica para la detección de retinopatía diabética (SSA: Secretaría de Salud, 2014) puesto que además de explorar el fondo de ojo permite

documentar la evolución de esta complicación de la diabetes.

Un método de diagnóstico que ha llamado el interés de investigadores alrededor del mundo es la implementación de sistemas de diagnóstico automatizado por computadora utilizando técnicas de análisis de imágenes, la alta disponibilidad de las imágenes de fondo de ojo ha permitido implementar y probar algoritmos con éste objetivo.

La evaluación de las imágenes de fondo de ojo en busca de signos de retinopatía diabética es una tarea compleja que requiere de entrenamiento, sin embargo el sistema visual humano tiene una alta capacidad de discriminación e interpretación, lo que pone en duda el por qué de un sistema automatizado, y si éste puede ser tan exacto y preciso como un oftalmólogo calificado.

La problemática radica en la alta prevalencia de la diabetes combinada con el limitado número de oftalmólogos a nivel mundial, en conjunto generan el problema de una sobrecarga de trabajo a los oftalmólogos en cuanto a análisis de fotografías de fondo de ojo se refiere. De acuerdo a la Academia Americana de Oftalmología (AAO: American Academy of Ophthalmology, 2015), el número de oftalmólogos en 2012 era igual a 210,730 y la prevalencia de la diabetes en 2011 de acuerdo a la FID: Federación Internacional de la Diabetes (2011) era de 366'200,000. De acuerdo a las estimaciones anteriores, si ninguno de los diabéticos tuviera indicios de retinopatía diabética, ni ninguna mujer diabética estuviera embarazada y suponiendo que los oftalmólogos están distribuidos uniformemente, cada oftalmólogo debería evaluar alrededor de 1,737.76 imágenes de fondo de ojo anualmente.

El panorama planteado anteriormente es ideal y aún así se aprecia un problema de alta carga de trabajo, además de que sólo el 27 % de los pacientes que tienen entre 5 y 10 años con diabetes desarrollan retinopatía diabética, lo que deriva en un alto porcentaje de imágenes de fondo de ojo sin ningún signo, resultando en una inversión de tiempo de análisis infactible. El segundo gran problema es la distribución de los oftalmólogos a nivel mundial (Tabla 3), los cuáles tienen mayor presencia en países desarrollados y además, suelen aglomerarse en "clústers médicos". En el estudio de ICO: International Council of Ophthalmology (2012), se encontró que hay 23 países en los que existe menos

de un oftalmólogo por millón de habitantes, 30 países con 1-4 oftalmólogos por millón de habitantes, 48 países con 4-25, 74 países con 25-100 y 18 países con más de 100 por millón de habitantes.

Tabla 3: Comparación de los países con mayor prevalencia de diabetes (FID: Federación Internacional de la Diabetes (2011)) contra su respectivo número de oftalmólogos (ICO: International Council of Ophtalmology (2012)).

País	Número de habitantes en millones	Prevalencia de diabetes en millones (2011)	Porcentaje de población con diabetes	Número de oftalmólogos (2012)	Tasa de habitantes por oftalmólogo	Tasa de habitantes diabéticos por oftalmólogo
China	1376.049	90	6.54	28338	48558.43	3175.94
India	1311.05	61.3	4.67	15000	87403.4	4086.66
Estados Unidos de América	321.774	23.7	7.36	18805	17111.08	1260.30
Rusia	143.45	12.6	8.78	14600	9825.82	863.01
Brasil	207.84	12.4	5.96	14679	14159.54	844.74
Japón	126.57	10.7	8.45	13911	9098.77	769.17
México	127.017	10.3	8.10	5459	23267.44	1886.79
Bangladesh	160.99	8.4	5.21	610	263927.86	13770.49
Egipto	91.50	7.3	7.97	2400	38128.33	3041.66
Indonesia	257.56	7.3	2.83	4814	53503.11	1516.41

En base a las estadísticas expuestas anteriormente, los sistemas de detección automática despiertan un gran interés para solucionar el problema de la sobrecarga de trabajo a los oftalmólogos. Los métodos de diagnóstico automatizado se basan en la detección de diversas lesiones, tales como exudados, hemorragias, crecimiento anormal de vasos sanguíneos, microaneurismas, desprendimiento de retina, y en base a la detección de alguna de estas lesiones determinan si es una imagen de fondo normal o anormal. Existen otras metodologías en el estado del arte que detectan más de un tipo de lesión y son capaces de diagnosticar el grado de retinopatía diabética de la fotografía en base a la escala clínica internacional de gravedad de la retinopatía diabética.

Aunque existen metodologías como: la propuesta de Fleming *et al.* (2007) que tiene el objetivo de detectar exudados, que se basó en operadores morfológicos, propiedades locales y un clasificador binario basado en umbralización. El sistema fue probado con

13,219 imágenes y obtuvo una sensibilidad del 95 % y un 84.6 % de especificidad; la propuesta de Niemeijer *et al.* (2007) en la que se evaluaron 300 imágenes (430 total, 130 de entrenamiento) en busca de exudados con un algoritmo de clasificación de píxeles en base a la respuesta del píxel a 14 filtros gaussianos combinada con descriptores espaciales y un clasificador de k vecinos más cercanos (k-NN). La sensibilidad del sistema fue de 95 % y un 86 % de especificidad. Y la propuesta de Bae *et al.* (2011) para la detección de hemorragias mediante la técnica de “template matching” que alcanzó una sensibilidad del 85 %. La mayoría de las metodologías propuestas para detectar retinopatía diabética se centran en la detección de microaneurismas. El argumento para enfocarse en este tipo de lesiones es que son las primeras lesiones en manifestarse y observarse en la retina, i.e. son un signo de que existe riesgo o se encuentra en una fase temprana de retinopatía diabética.

Sopharak *et al.* (2013) presentó una metodología para la detección de microaneurismas en 45 imágenes (Thammasat University Hospital) y clasificación del grado de severidad de retinopatía diabética de cada imagen en base al número de microaneurismas en la imagen. La metodología opera sobre el canal verde de una imagen RGB, se utilizó un filtro mediano para disminuir el ruido en la imagen, se aplicó la técnica de CLAHE para mejorar el contraste en la imagen, posteriormente para reducir la iluminación no uniforme se le restó a la imagen mejorada con CLAHE el resultado de filtrar esa misma imagen con un filtro mediano de tamaño 35×35 . Posteriormente elimina los exudados y vasos sanguíneos de la imagen utilizando procesamiento morfológico. Se extrajeron 15 descriptores de cada pixel candidato a microaneurisma y en el software Weka se clasificó utilizando un modelo de clasificación de Nāive-Bayes. La metodología propuesta obtuvo una sensibilidad del 85.68 %, 99.99 % de especificidad, 83.34 % de precisión y 99.99 % de exactitud. En cuanto a la determinación del grado de retinopatía diabética de cada imagen, se obtuvo una exactitud promedio del 88.31 %. Se comparó con un clasificador de vecino más cercano, el cual obtuvo 87.15 % de sensibilidad con 96.5 % de especificidad.

Walter *et al.* (2007) presentó una metodología para la detección de microaneurismas basada en la mejora de contraste del canal verde de imágenes RGB utilizando una trans-

formación de punto polinomial y una normalización para disminuir la iluminación no uniforme del fondo, posteriormente se utilizó un filtro gaussiano para la disminución de ruido en las imágenes (5×5 píxeles con $\sigma = 2$). La generación de candidatos se basó en la operación morfológica de cierre mejorada seguida de una binarización. Se extrajeron 15 características espaciales, de intensidad y de color para clasificar cada candidato como microaneurisma o no microaneurisma utilizando una estimación de la función de densidad de probabilidad. Se probó la metodología en 94 imágenes y se obtuvo una sensibilidad del 88.47 %.

La metodología propuesta por Niemeijer *et al.* (2005) opera sobre el canal verde de las imágenes RGB, a las que se les aplica una corrección de iluminación basada en restar a la imagen del canal verde la respuesta a un filtro mediano de 25×25 de la misma. Elimina las lesiones brillantes convirtiendo a cero todos los píxeles con un valor positivo. Se generaron los candidatos mediante una operación de apertura con un elemento estructural de 9 píxeles rotado 12 veces, se obtuvieron 12 imágenes y se construyó una nueva imagen mapeando el valor de intensidad más alto de las 12 imágenes de cada píxel con lo que se obtiene la vasculatura de la retina; finalmente los candidatos se obtienen al restar de la imagen sin lesiones brillantes la imagen generada anteriormente. Se utilizó un filtro espacial gaussiano (11×11 con $\sigma = 1$) para mejorar el contraste entre fondo y lesiones y se umbralizó para obtener una imagen binaria. Se implementó una operación de crecimiento de regiones y se extrajeron 68 descriptores de cada candidato para clasificar cada pixel con un clasificador k-NN obteniendo una sensibilidad del 100 % y una especificidad del 87 %.

La metodología de Romero *et al.* (2015) opera sobre la proporción del canal verde entre el canal rojo bajo el argumento de que se reduce la iluminación no uniforme, para posteriormente utilizar una técnica de normalización y generar los candidatos a partir de la transformada “Bottom-Hat”. Reduce el número de candidatos utilizando la transformada Hit-Miss y regenera cualquier alteración de los candidatos a microaneursismas utilizando un algoritmo de componentes conectadas. La clasificación se realizó con dos características: 1) circularidad a partir de la proporción de dos eigenvectores y 2) número de veces que la región se encuentra en un valle entre dos picos en el histograma a partir

de una transformada de radón, utilizando dos perceptrones anidados (uno para cada característica). Se probó la metodología con dos bases de datos: DIARETDB1 y ROC; para las que se obtuvieron: especificidad de 93.87 % y 97.47 % respectivamente, sensibilidad de 92.32 % y 88.06 % respectivamente, y precisión de 95.93 % y 92.19 % respectivamente.

La propuesta de Hipwell *et al.* (2000) obtuvo una sensibilidad del 81 % con una especificidad de 93 % para 62 imágenes de prueba (102 de entrenamiento). Su propuesta opera sobre el canal verde, se aplica una corrección de iluminación y posteriormente se eliminan vasos sanguíneos y hemorragias. Se extraen 13 descriptores de intensidad y forma de los candidatos para clasificar como microaneurisma o no microaneurisma en base a un conjunto de reglas establecidas a partir de analizar las 102 imágenes de entrenamiento.

Arenas-Cavalli *et al.* (2015) propusieron una plataforma basada en la web para discriminar imágenes de fondo normales de las que muestran signos de retinopatía diabética. La base de datos que utilizan es privada y consta de 550 imágenes en total. Su propuesta detecta vasos sanguíneos, el disco óptico, lesiones brillantes y microaneurismas y hemorragias. Respecto a la detección de microaneurismas, su método consiste en tres pasos 1) normalizar la imagen, reducir el ruido y mejorar el contraste, 2) eliminar exudados, disco óptico y vasos sanguíneos con los métodos de las etapas anteriores para reducir la cantidad de candidatos y las falsas detecciones, 3) aplicar operaciones morfológicas. No se menciona con que canal de la imagen se trabajó ni las características que seleccionaron ni el clasificador. Reportaron una sensibilidad de 77.7 % y 81.2 % de especificidad.

La metodología de Adal *et al.* (2014) propuso una detección semisupervisada de microaneurismas; trabajaron sobre el canal verde de la imagen, en su etapa de preprocesamiento mejoraron el contraste mediante una descomposición en valores singulares y ecualizando los valores. Mediante la obtención del determinante Hesiano de la imagen extrajeron regiones circulares oscuras. Como características utilizaron un banco de tres filtros gaussianos para obtener 18 descriptores, más 64 descriptores obtenidos mediante el métodos SURF (Speed-Up Robust Features), más otros 3 descriptores extraídos me-

diente la transformada de Radon. Utilizaron cuatro clasificadores: KNN, Random Forest, Naïve Bayes y SVM. Optimizaron los descriptores para cada clasificador y la mayor sensibilidad se obtuvo con la máquina de soporte vectorial, siendo ésta de 81.08 % y la especificidad de 92.31 %.

La metodología de Bhalerao *et al.* (2008) pretende detectar microaneurismas con robustez, para normalizar el contraste y disminuir el ruido utilizan un filtro mediano sobre el canal verde, utilizan el operador Laplaciano de Gaussiano para la detección de regiones, filtrado utilizando una plantilla circular para generar candidatos. No se clasificaron los candidatos, se hizo una búsqueda visual para medir el rendimiento de la segmentación. La sensibilidad repostada es de 82.6 % con una especificidad del 80.2 %.

La propuesta de SujithKumar y Singh (2012) obtuvo resultados aceptables en la determinación del grado de retinopatía diabética a partir del número de microaneurismas detectados; la sensibilidad promedio fue de 94.44 % con un 87.5 % de especificidad. El método consiste en utilizar un filtro mediano de 30×30 sobre el canal verde para obtener una aproximación del fondo de la imagen, para posteriormente restarle el resultado del filtro a la imagen del canal verde, la imagen obtenida se encuentra normalizada. A la imagen resultante se le mejora el contraste utilizando una ecualización de histograma adaptativa seguida de la técnica de umbralización para obtener una imagen binaria. Para obtener las regiones candidatas se emplea el algoritmo de Canny para la detección de bordes y se obtienen descriptores espaciales de cada región. No se especifica el modelo de clasificación; además de detectar los microaneurismas, se predice la etapa de retinopatía diabética en base al número de MA's detectados.

Lazar *et al.* (2010) propusieron un esquema de detección de microaneurismas hasta la etapa de generación de candidatos (segmentación) basándose en un análisis de los perfiles de intensidad de la imagen. El algoritmo fue diseñado para operar con imágenes de máximo 768 píxeles de alto y ancho, por lo que imágenes de mayor resolución son escaladas con una interpolación. Trabajaron sobre el canal verde invertido, de manera que los MA's aparecen como las zonas más brillantes en la imagen, y el perfil de intensidad se asemeja al de una distribución gaussiana en las zonas en las que hay un microaneurisma. Se analizan los perfiles de intensidad en 4 direcciones para obtener perfiles de

intensidad en una dimensión, se umbraliza cada perfil de intensidad y se reconstruye la imagen a partir de éstos. Se aplica un filtrado gaussiano y una umbralización con dos umbrales para obtener los candidatos. No se reporta sensibilidad ni especificidad, a partir de su gráfica se observa que la sensibilidad máxima es alrededor de 62 % con un número promedio de 325 falsos positivos por imagen. Concluyeron que a pesar de que la detección automática de microaneurismas ha sido tema de análisis, estudio e investigación, las metodologías propuestas no son recomendadas en la práctica clínica.

La propuesta de Cree (2008) propone filtrar el canal verde con un filtro de mediana para la corrección de la iluminación, enseguida, la imagen resultante es pasada por un filtro de 3×3 para disminuir el ruido en la imagen. Utilizando la transformada Top-Hat remueven los vasos sanguíneos para posteriormente filtrar la imagen resultante con un filtro gaussiano espacial para detectar objetos redondos pequeños, se umbraliza la imagen y utilizando siete descriptores se clasifican los candidatos mediante el modelo de Naïve-Bayes. La sensibilidad obtenida fue de 85 % con una especificidad de 90 % para 758 imágenes de su propia base de datos.

Sopharak *et al.* (2011) proponen operar sobre el canal verde y mejorarla utilizando un filtro mediano para reducir el ruido, posteriormente ecualizan la imagen con la técnica de CLAHE. Enseguida se disminuye la iluminación no uniforme, se eliminan vasos sanguíneos, exudados y se utilizan operadores morfológicos para generar los candidatos a microaneurismas. Se utilizaron 45 imágenes para conducir el experimento y no se especifica el modelo de clasificación utilizado. La sensibilidad obtenida fue de 81.61 % con un 99.99 % de especificidad, además, la precisión fue del 63.76 % con una exactitud de 99.98 %.

Streeter y Cree (2003) propusieron un método para la detección de microaneurismas que obtuvo un 56 % de sensibilidad con una tasa de falsos positivos de 5.7 por imagen. El preprocesamiento se aplicó sobre el canal verde de la imagen, se dividió ésta entre la respuesta a un filtro de mediana de 56×56 para disminuir la iluminación no uniforme. Se aplicó el operador Top-Hat con un elemento estructural de 25×1 píxeles con ocho orientaciones distintas, la imagen resultada es sometida a dos diferentes filtros acoplados, se binariza la respuesta de ambas y se combinan mediante la operación OR. El resultado

obtenido es la semilla de los candidatos y se aplica el algoritmo de crecimiento de regiones para reconstruir los candidatos a partir de la imagen en el canal verde corregida. Se clasificaron los candidatos en base a 18 descriptores de forma e intensidad mediante el software WEKA, se probaron varios modelos y el que mostró mejores resultados fue el de ADTree (clasificación basada en árboles de decisión).

La metodología de Kande *et al.* (2010) logró obtener una sensibilidad del 100 % para una especificidad del 91 %. El esquema opera sobre la canal verde de la imagen RGB, la cual fue modificada a través de su histograma para que tuviera una distribución similar a la del histograma del canal rojo. Enseguida se le aplicó un estiramiento de histograma a la imagen obtenida anteriormente y se obtuvo la respuesta a un filtro mediano para disminuir variaciones de intensidad. A través de un filtro gaussiano bidimensional se generaron los candidatos a microaneurismas y se umbralizó utilizando una técnica basada en entropía. Finalmente, utilizando la transformada Top-Hat remueven la vasculatura. Utilizando 12 descriptores se clasificaron los candidatos bajo el modelo de máquina de soporte vectorial (SVM). El banco de imágenes se constituyó a partir de 54 imágenes provistas por un hospital y se seleccionaron 35 aleatoriamente de las base de datos públicas DIARETDB1, DIARETDB0 y STARE.

La metodología propuesta por Aravind *et al.* (2013) fue probada con 105 imágenes del *Lotus Eye Care Hospital*. Las técnicas de preprocesamiento se aplicaron sobre el canal verde de la imagen RGB; primeramente se aplicó la técnica de CLAHE para ecualizar el histograma de la imagen y obtener una imagen mejor contrastada. El siguiente procedimiento fue aplicar un filtro de mediana de 3×3 para suavizar la imagen, se mejora el contraste y se umbraliza la imagen. Enseguida se aplican operadores morfológicos para generar los candidatos a MA's y se remueve el disco óptico. Los descriptores extraídos fueron: 1) área, 2) entropía, 3) correlación, 4) energía, 5) contraste, 6) homogeneidad, 7) media y 8) desviación estándar. Se seleccionaron los mejores descriptores para el modelo de clasificación SVM y se obtuvieron 23 VP's, 4 VN's, 1 FP y 2 FN's. La sensibilidad del método logró una sensibilidad del 92 % combinada con un 80 % de especificidad. La exactitud de la metodología fue del 90 %.

Tabla 4: Comparación de metodologías para la detección de microaneurismas.

Autor	Sensitividad	Especificidad	Precisión	Exactitud
Sopharak <i>et al.</i> (2013)	85.68 %	99.99 %	83.34 %	99.99 %
	81.68 %	99.99 %	63.76 %	99.98 %
	87.15 %	96.6 %	65.43 %	96.5 %
Walter <i>et al.</i> (2007)	88.47 %	100 %	No Reportado	No Reportado
Romero <i>et al.</i> (2015)	92.32 % (DIARETDB1), 88.06 % (ROC)	93.87 % (DIARETDB1), 97.47 % (ROC)	No Reportado	No Reportado
Niemeijer <i>et al.</i> (2005)	100 %	87 %	No Reportado	No Reportado
Hipwell <i>et al.</i> (2000)	78 %	91 %	No Reportado	No Reportado
Arenas-Cavalli <i>et al.</i> (2015)	77.7%	81.2 %	No Reportado	No Reportado
Adal <i>et al.</i> (2014)	81.08 %	92.31 %	No Reportado	No Reportado
Bhalerao <i>et al.</i> (2008)	82.6 %	80.2 %	No Reportado	No Reportado
SujithKumar y Singh (2012)	94.44 %	87.5 %	No Reportado	No Reportado
Lazar <i>et al.</i> (2010)	≈ 62%	No Reportado	No Reportado	No Reportado
Cree (2008)	85 %	90 %	No Reportado	No Reportado
Sopharak <i>et al.</i> (2011)	81.61 %	99.99 %	63.76 %	99.98 %
Streeter y Cree (2003)	56 %	No Reportado	No Reportado	No Reportado
Kande <i>et al.</i> (2010)	100 %	91 %	No Reportado	No Reportado
Aravind <i>et al.</i> (2013)	92 %	80 %	No Reportado	90 %

2.3. Fundamentos

Una imagen se refiere a un arreglo bidimensional con diferentes valores de intensidad luminosa (nivel o escala de gris) (Elizondo y Maestre, 2005). Los valores de intensidad son resultado de captar la irradiancia de una fuente sobre un plano. La distribución espacial de la irradiancia es una función continua (Jähne, 2005). Las imágenes digitales son el

resultado de digitalizar una imagen, es decir, convertir la función continua en discreta aplicando un muestreo de las variables en el espacio y la cuantificación de los niveles de gris. Las imágenes digitales típicamente son obtenidas utilizando un sistema de adquisición de imágenes basado en arreglos de sensores.

Los arreglos de sensores (ver Figura 4) son típicamente una malla o rejilla cuadriculada, aunque también existen rejillas hexagonales o triangulares. Cada región (cuadro, hexágono o triángulo) se corresponde a un sensor. El proceso de muestreo depende del arreglo del sensor, donde cada cuadro representa una muestra. En función de la frecuencia de muestreo, estos serán más o menos amplios, abarcando así una mayor o menor área.

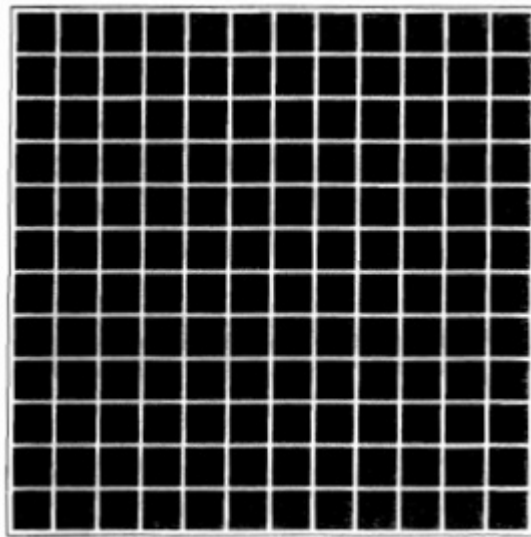


Figura 4: Ejemplo de un arreglo de sensores.

Cada cuadro capta un nivel de intensidad luminosa, y su valor corresponde a la integral de la energía luminosa captada por cada superficie del sensor, asignando así un nivel radiométrico. Dependiendo de la resolución de intensidad del sensor se tiene un número “n” de niveles de gris. Esta resolución depende de la cantidad de bits disponibles para almacenar el valor de intensidad luminosa. La cuantificación se refiere a dividir el rango de la amplitud de intensidad luminosa en diferentes niveles, de manera que a cada rango se le asigna un valor discreto. Si por ejemplo, un sensor capta valores de intensidad en el rango de 0 a 1, y se tienen 2 bits de resolución, habiendo así 4 combinaciones posibles y

por tanto 4 niveles de gris. Por ejemplo: los valores de intensidad luminosa en el rango 0 - 0.25 tendrán una asignación de “00”=0; los valores de intensidad luminosa en el rango 0.26 - 0.50 tendrán una asignación de “01”=1; los valores de intensidad luminosa en el rango 0.51 - 0.75 tendrán una asignación de “10”=2; y finalmente los valores de intensidad luminosa en el rango 0.76 - 1.00 tendrán una asignación de “11”=3.

El proceso de muestreo y cuantificación genera una imagen digital. De acuerdo con Gonzalez y Woods (2011), una imagen digital (ver Figura 5) es una función en dos dimensiones $0 \leq f(x, y) \leq L - 1$, donde x e y son coordenadas en el espacio y la amplitud f en el par coordenado (x, y) se corresponde a la intensidad o nivel de gris en ese punto. La intensidad o valor de gris es un valor discreto en el rango $[0, L - 1]$. El número de niveles de gris L está dado por $L = 2^k$, donde k es el número de bits utilizados para codificar los valores de intensidad luminosa.

El grado de detalle de la imagen, i.e. la resolución, depende de los parámetros con los que se aplica el muestreo y cuantificación a la imagen. Se tiene entonces una resolución de intensidad y una espacial, dependientes de la cantidad de niveles de gris y la finura de la malla del arreglo de sensores respectivamente. A mayor número de niveles de gris y mallas más finas, se tienen imágenes más detalladas, hasta llegar al punto que los cambios entre una región de la malla y de un nivel de gris a otro es imperceptible para el ojo humano. Evidentemente a mayor grado de detalle, la imagen digital se aproxima a la imagen original (función continua). En la figura 6, se observa el resultado de capturar una escena (imagen continua) con un arreglo de sensores, que resulta en una imagen discreta en el espacio y en la amplitud.

	x					
						
y 	$f(0,0)$	$f(0,1)$	$f(0,2)$	$f(0,3)$	$f(0,...)$	$f(0,n-1)$
	$f(1,0)$	$f(1,1)$	$f(1,2)$	$f(1,3)$	$f(1,...)$	$f(1,n-1)$
	$f(2,0)$	$f(2,1)$	$f(2,2)$	$f(2,3)$	$f(2,...)$	$f(2,n-1)$
	$f(3,0)$	$f(3,1)$	$f(3,2)$	$f(3,3)$	$f(3,...)$	$f(3,n-1)$
	$f(4,0)$	$f(4,1)$	$f(4,2)$	$f(4,3)$	$f(4,...)$	$f(4,n-1)$
	$f(5,0)$	$f(5,1)$	$f(5,2)$	$f(5,3)$	$f(5,...)$	$f(5,n-1)$
	$f(...,0)$	$f(...,1)$	$f(...,2)$	$f(...,3)$	$f(...,...)$	$f(...,n-1)$
	$f(m-1,0)$	$f(m-1,1)$	$f(m-1,2)$	$f(m-1,3)$	$f(m-1,...)$	$f(m-1,n-1)$

Figura 5: Imagen digital de $m \times n$, $m=8$ y $n=6$.

Una imagen digital se puede representar como una matriz de $m \times n$. Cada elemento de la matriz, en el dominio de la imagen se conoce como píxel². El píxel es el elemento básico de una imagen digital, la cual es entonces una colección bidimensional finita de pixeles, de los cuales, cada pixel tiene una ubicación particular con su correspondiente valor de intensidad.

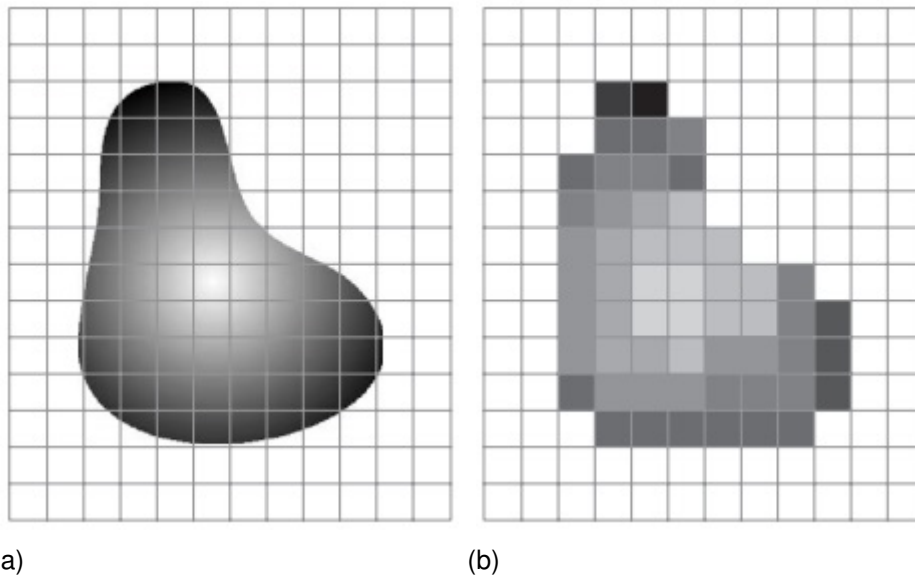


Figura 6: a) Escena real b) Digitalización de la escena. Imagen tomada de Gonzalez y Woods (2011).

Hasta el momento, con lo explicado, se puede hablar de imágenes binarias y en escala

²Del inglés pixel que significa "picture element", i.e., elemento más básico en una imagen

de grises. Binarias cuando se tiene disponible un bit por pixel y solo puede tomar valor de 0 ó 1, blanco o negro; imágenes en escala de grises cuando se tiene más de un bit por pixel, dando pie a cuatro o más niveles de intensidad. Típicamente se utilizan 8 bits por pixel, que se traduce en 256 niveles de gris, con posibles valores en el rango [0, 255].

Existen además imágenes en color, donde para lograr éstas, en cada posición de la rejilla se captura la respuesta a la intensidad en 3 ó más bandas espectrales diferentes. Se utilizan 3 bandas para capturar imágenes en color, típicamente RGB, puesto que se tiene un sensor sensible en la banda espectral del canal rojo, otro sensible al canal verde y uno más al canal azul. Cuando se emplean 4 o más sensores, se da lugar a imágenes multispectrales. Las imágenes a color se corresponden a matrices multidimensionales, $m \times n \times z$, donde m y n será la resolución espacial y el valor de z dependerá del número de canales, tres para el caso de imágenes RGB.

Los sistemas de tratamiento de imágenes tienen el objetivo de tratar las imágenes digitales, destacan tres áreas: visualización, análisis y administración. Cada una de las tres áreas requiere una fase previa de adquisición y mejoramiento de la imagen. El mejoramiento de la imagen se realiza empleando el procesamiento digital de imágenes; las técnicas empleadas difieren de acuerdo al área. Por ejemplo, para la visualización destacan los algoritmos para mejorar la apreciación visual, para el análisis de imágenes, técnicas que destaquen únicamente objetos específicos y para la administración técnicas para aumentar la eficiencia de los medios de almacenamiento, comunicación, transmisión y búsqueda (Gonzalez y Woods, 2011; Deserno, 2011).

La etapa de mejoramiento de la imagen, al ser una fase previa de los tres enfoques que tiene el tratamiento de las imágenes, típicamente se le conoce como preprocesamiento. En la figura 7, se observan las tres principales etapas del tratamiento de imágenes, la formación de imagen, el mejoramiento de la imagen o preprocesamiento y finalmente los tres enfoques: visualización, análisis y administración.

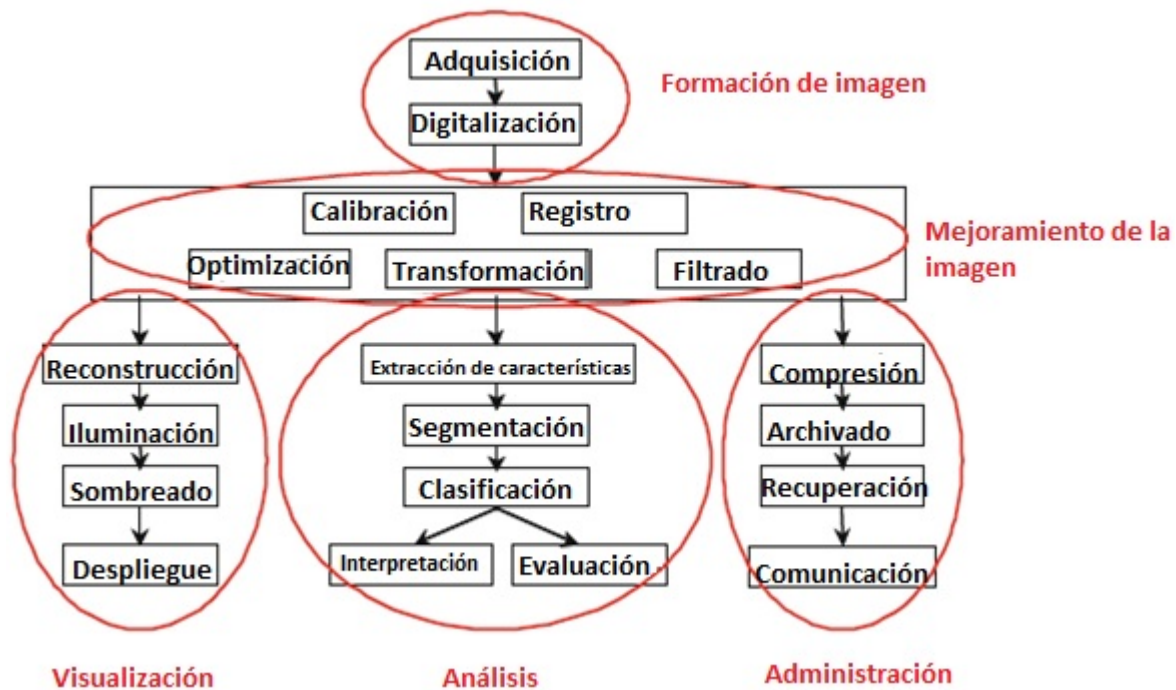


Figura 7: Tratamiento de imágenes.

En el presente trabajo de investigación, se tiene el objetivo de detectar microaneurismas, que son objetos o regiones que se aprecian visualmente en imágenes de fondo de ojo. El enfoque del tratamiento de imágenes de interés es el análisis de imágenes. De acuerdo con Lira (2002), el análisis de imágenes tiene el objetivo de cuantificar las propiedades de objetos o estructuras específicas en las imágenes; Deserno (2011) coincide con Lira (2002), añadiendo que el enfoque de análisis de imágenes comprende el uso de la mayoría de las técnicas de mejoramiento de la imagen y que para aplicar cada una de estas técnicas se requiere de conocimiento a priori de la naturaleza del contenido de la imagen.

El enfoque de tratamiento de imágenes para este trabajo de investigación, es el de análisis de imágenes. El análisis de imágenes se constituye de dos etapas u operadores principales (ver ecuación 1): procesamiento digital de imágenes y reconocimiento de patrones. La etapa de procesamiento digital de imágenes tiene por objetivo el mejoramiento de las imágenes, pone en evidencia y realza las regiones que nos interesan examinar a detalle, en el mejor de los casos las aísla completamente de las regiones que no son de interés. La etapa de reconocimiento de patrones cuantifica las propiedades de las

regiones que se han aislado en la etapa previa, para posteriormente detectar o clasificar las regiones. Para cumplir con el objetivo de analizar imágenes oftálmicas y de detectar microaneurismas, se recurrió al esquema general de la detección de objetos (ver Figura 8), que se basa en el enfoque de clasificación de vectores de características.

$$\text{Análisis de imágenes} = O_{rp}\{O_{pi}(f(x, y))\} \quad (1)$$

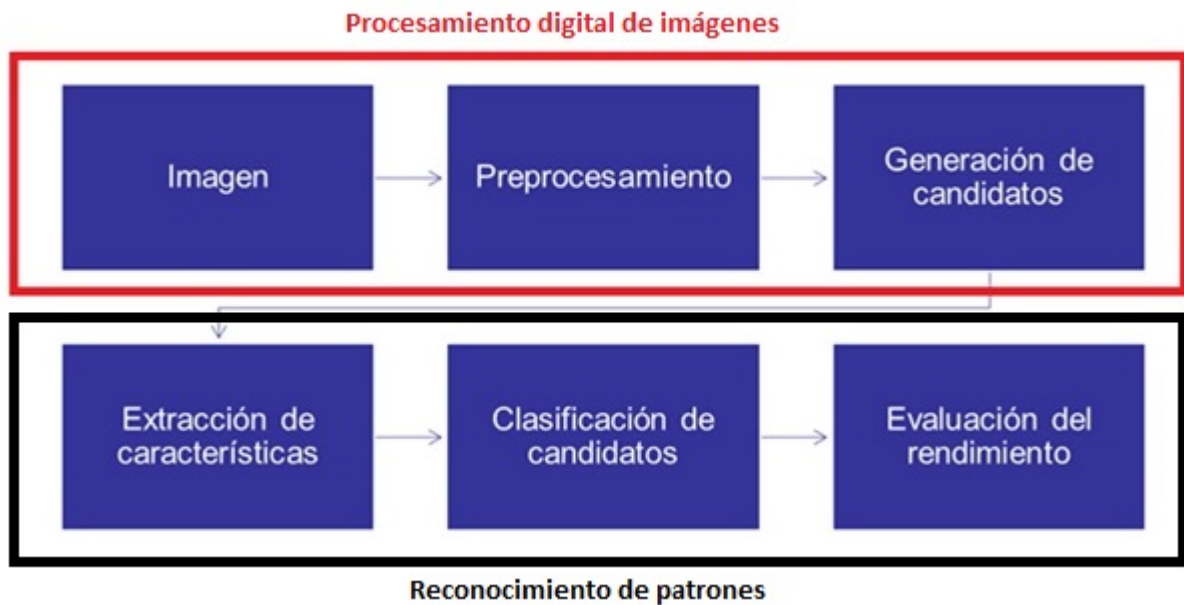


Figura 8: Esquema general de la detección de objetos.

El objetivo de la detección de objetos, es encontrar patrones con características dadas en imágenes. Para esto, en primera instancia se tiene la etapa de formación y adquisición de la imagen, la segunda etapa es el preprocesamiento, i.e. el procesamiento digital de la imagen para su mejoramiento. Esta etapa se subdivide en varias etapas que tienen por objetivo reducir el ruido en la imagen, mejorar el contraste, modificar el histograma, entre otras. La siguiente etapa se refiere a la generación de candidatos, las técnicas utilizadas pertenecen al procesamiento digital de imágenes pero en el esquema general de la detección de objetos se maneja como una etapa diferente para diferenciar las técnicas de mejoramiento y las que aíslan las regiones de interés de las que no se desean analizar. Al proceso de aislar las regiones se le conoce como segmentación, técnica a la que le

procede el umbralizado, que tiene por objetivo generar imágenes binarias indicando las regiones de interés.

Una vez que se tienen las regiones de interés, se procede a la etapa cuatro en el esquema, en la que se cuantifican las propiedades de las regiones de interés, propiedades a las cuales se les conocen como características o descriptores. La quinta etapa consiste en generar un modelo utilizando los descriptores extraídos de cada región. La clasificación requiere tener dos conjuntos de imágenes: entrenamiento y prueba; y consiste en, dada una región y sus descriptores, decidir si es el objeto deseado (en el caso de esta investigación, decidir si es un microaneurisma) o no. El conjunto de entrenamiento es con el que se construye el modelo de clasificación sabiendo a priori qué regiones (candidatos) son microaneurismas y cuáles no. El conjunto de prueba son los datos que se someten al modelo, se analizan sus descriptores y en base a éstos se clasifica cada región como microaneurisma o no microaneurisma. La última etapa del esquema general de la detección de objetos consiste en evaluar el rendimiento del sistema de detección; se establecen métricas de rendimiento (e.g. exactitud, valor predictivo positivo, sensibilidad, curvas ROC, entre otros) que nos permiten comparar los resultados del sistema de detección con otros propuestos en la literatura, validando así el modelo.

2.4. Procesamiento digital de imágenes

En esta sección se describen las sub-etapas (Adquisición de la imagen, mejoramiento y aislamiento de las regiones u objetos de interés) en la etapa del procesamiento digital de imágenes (ver Figura 8).

2.4.1. Adquisición de la imagen

Todo sistema que trata con imágenes digitales, independientemente del objetivo que tenga, ya sea analizarlas, administrarlas o visualizarlas, necesita contar con una etapa de formación de la imagen, etapa en la que se adquiere la imagen. La imagen se puede adquirir de una base de datos de imágenes digitales o formar la imagen por completo. Formar la imagen requiere contar con hardware y software especial; se requiere un arreglo de sensores que permitan capturar una escena para posteriormente con un software especial muestrear y cuantificar la escena para obtener así una imagen digital.

En la figura 9, se muestra el esquema de adquisición de una imagen digital. Una fuente de iluminación irradia energía en forma de luz sobre una escena particular, la escena absorbe parte de esa energía luminosa (partículas sin masa llamados fotones) y la restante es reflejada. Parte de la energía luminosa reflejada es captada por un sistema formador de imágenes, el cual se conforma de dos subsistemas principales: sistema óptico y sistema electrónico.

El sistema óptico consiste en una serie de lentes que emulan los lentes y el funcionamiento del ojo humano; la imagen formada por el sistema óptico es proyectada a un sensor. En cámaras analógicas son películas especiales recubiertas con sustancias químicas que reaccionan a la luz. En cámaras digitales, el sensor forma parte del sistema electrónico; consiste en un arreglo de sensores (rejilla) que capta la energía de los fotones emitidos a cierta frecuencia en cada elemento de la rejilla y convierte esa energía a voltajes. Finalmente, la rejilla discretiza la escena en el espacio, y dependiendo del voltaje captada en cada elemento de la rejilla (i.e. pixel) se asigna un nivel de gris de acuerdo al número de bits disponibles para su representación.

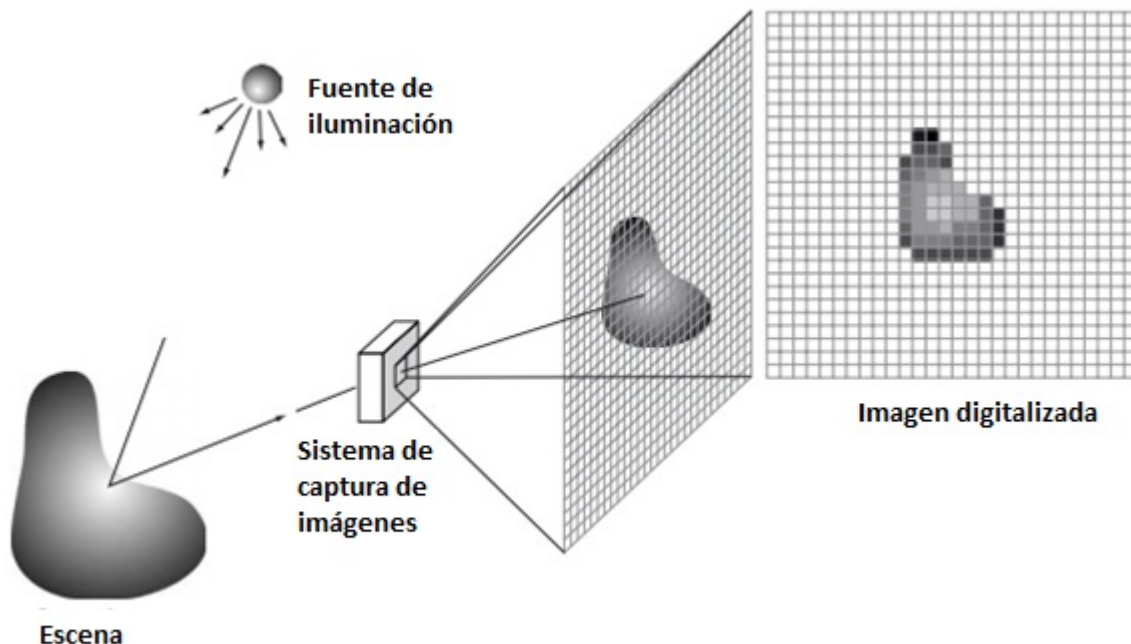


Figura 9: Captura de una escena real. Imagen tomada de Gonzalez y Woods (2011).

2.4.2. Canales y espacios de color

El color es una sensación creada en respuesta a la excitación del sistema visual por el tipo de radiación electromagnética en el rango del espectro visible (Plataniotis y Venetsanopoulos, 2000). Las imágenes a color son representadas por tres componentes de intensidad (Nixon y Aguado, 2008); idealmente son resultado de la combinación de tres diferentes arreglos de sensores, los cuales censan energía luminosa (fotones) en diferentes frecuencias cada uno. En la práctica se cuenta con una sola rejilla, cada elemento de la rejilla capta fotones en diferentes frecuencias y al final se aplican procesos de interpolación para generar los valores de toda la rejilla en todas las frecuencias deseadas.

Las imágenes a color adquiridas por cámaras digitales típicamente son RGB. RGB es un modelo de color³ y se refiere a los términos en inglés RED, GREEN y BLUE; rojo, verde y azul en español, que son las bandas espectrales a las que es sensible el sensor para obtener este tipo de fotografías. El modelo de color RGB (ver Figura 10) es un modelo que se basa en los colores primarios: rojo, verde y azul; y en que a partir de éstos y su respectivo valor de intensidad es posible obtener otros colores resultado de combinar los primarios.

³Modelo de color: se refiere a un modelo matemático para la representación de colores en forma numérica

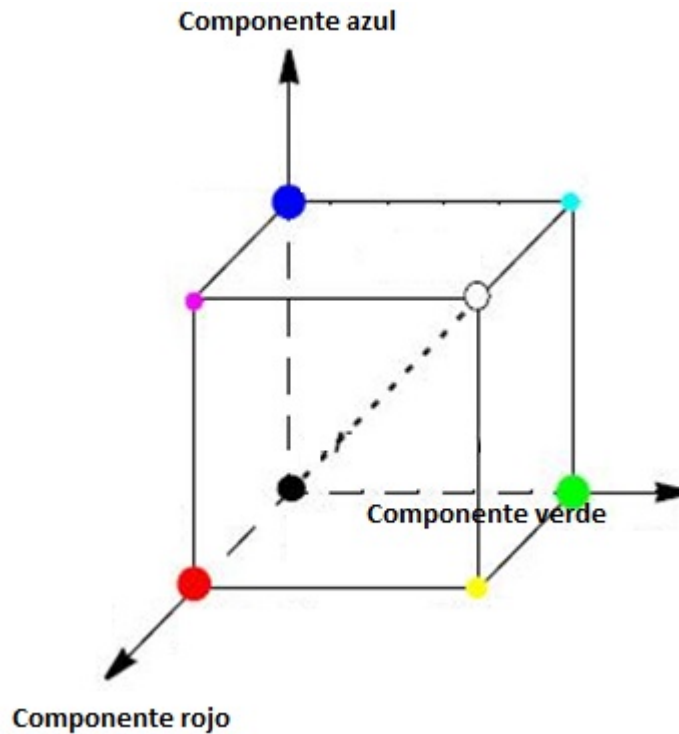


Figura 10: Modelo de color RGB.

Existen diversos modelos de color, de entre los que destacan un grupo de modelos que se basan en la manera en que trabaja el sistema de la visión humana para interpretar el color. Estos modelos se basan en los conceptos de luminosidad⁴, matiz⁵ y saturación⁶, y son utilizados con frecuencia en el área de investigación de visión por computadora con imágenes a color. Se recomienda utilizar estos modelos de color, porque combinan información de intensidad de los tres colores primarios (i.e. rojo, verde y azul) en un solo canal, facilitando la aplicación de técnicas de mejoramiento de la imagen. Trabajar con imágenes en RGB implicaría seleccionar un canal, omitir información del resto de los canales y aplicar las técnicas de mejoramiento sobre cada canal individual (Plataniotis y Venetsanopoulos, 2000).

Los modelos de color que se basan en el sistema de visión y percepción humano (Agoston, 2005) se conocen como la familia HSI (Hue-Saturation-Intensity), los cuales son:

⁴Luminosidad: del inglés *lightness*, se refiere a la percepción de brillo.

⁵Matiz o color: del inglés *hue*, se refiere al color dominante.

⁶Saturación: del inglés *saturation*, se refiere a la cantidad de luz mezclada con cierto color.

1. HSV: Modelo Hue-Saturation-Value (Matiz-Saturación-Valor). Ver Figura 11. En este modelo cada color se forma mediante tres variables, el matiz, la saturación y el valor. El matiz o tinta corresponde a un valor en grados (entre 0 y 360) que define el color predominante. La saturación (valores entre 0 y 1) se refiere a la cantidad de luz (color blanco) mezclada con el color predominante y el valor (valores entre 0 y 1) se refiere a la cantidad de luz de cero energía (color negro) mezclado con la tinta, i.e. intensidad de un color (brillo o mate).

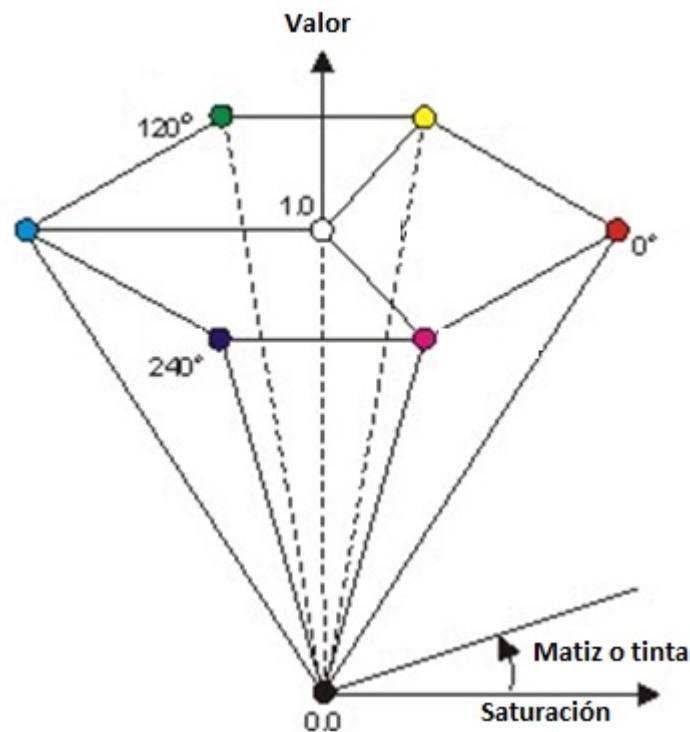


Figura 11: Modelo de color HSV.

2. HSL: Modelo Hue-Saturation-Lightness (Matiz-Saturación-Luminosidad). Ver Figura 12. Es un modelo similar al modelo HSV con la diferencia del parámetro de luminosidad. El concepto de luminosidad se refiere a la percepción de brillo, por lo que este parámetro se refiere a la intensidad del color dominante sobre el resto de los colores (tinta). Los colores se distribuyen en una estructura de cono hexagonal doble, donde el ángulo es el tono, el radio del círculo indica la saturación y la altura la intensidad, obteniendo un blanco total en el extremo superior y un negro total en el extremo inferior.

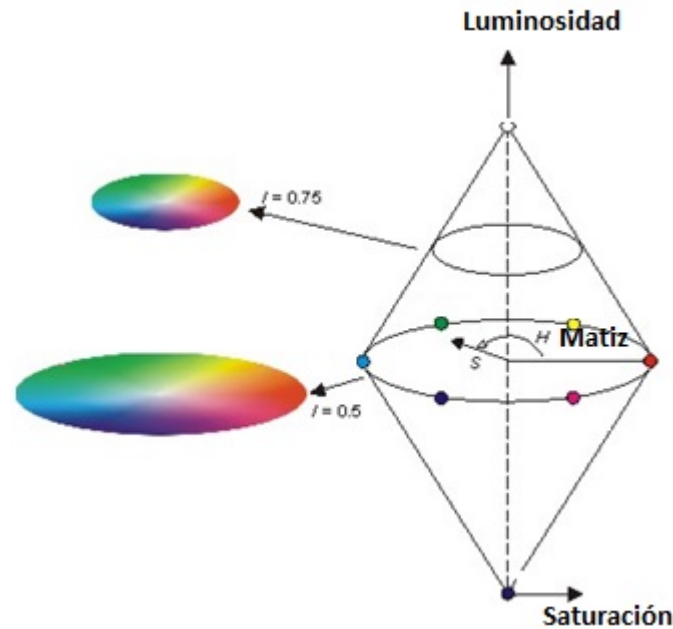


Figura 12: Modelo de color HSL.

2.4.3. Técnicas de reducción de ruido

Una de las técnicas de alta importancia en la fase de mejoramiento de la imagen son las que tienen el objetivo de reducir el ruido. El término de ruido se utiliza para referirse a toda información indeseable en una imagen; de acuerdo a (Lira, 2002) es un patrón espacial ajeno a la escena resultado de un proceso estocástico con una función de probabilidad conocida y asociada al sistema que genera la imagen digital y no debe confundirse con los artefactos, los cuáles se refieren a patrones azarosos ajenos a la escena. El ruido en imágenes, generalmente se manifiesta como píxeles aislados que toman un nivel de gris muy diferente al de sus vecinos (Elizondo y Maestre, 2005).

$$g(x, y) = f(x, y) + n(x, y) \quad (2)$$

La idea que plantea (Gonzalez y Woods, 2011) acerca del ruido, es que consiste en un proceso de degradación de la imagen “Bloque H” (ver Figura 13) al añadirse ruido (ecuación 2). El objetivo de la reducción de ruido es estimar la función de ruido para restaurar la imagen degradada $g(x, y)$ y obtener una estimación $e(x, y)$ de la imagen no

exactamente igual, pero aproximada a la original (escena real) $f(x, y)$.

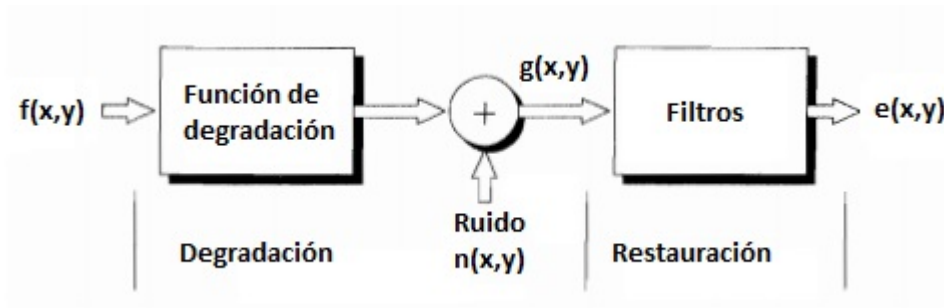


Figura 13: Modelo de ruido en una imagen.

El ruido en las imágenes digitales se debe frecuentemente al proceso de adquisición; influyen las condiciones del medio ambiente, la temperatura del sensor, el nivel de luz de la fuente de iluminación y el proceso de transmisión (interferencias en el canal) (Gonzalez y Woods, 2011). A continuación se detallan los tipos de ruido que degradan frecuentemente las imágenes de acuerdo a los autores Patidar *et al.* (2010); Elizondo y Maestre (2005):

1. Ruido impulsivo o sal y pimienta: El ruido sal y pimienta se caracteriza por la aparición de píxeles brillantes en regiones oscuras y la aparición de píxeles oscuros en zonas brillantes, de manera que el valor del píxel representa un cambio abrupto en la región y se interpreta como un valor no esperado. Las causas de este tipo de ruido van desde errores en el proceso de digitalización hasta problemas en la transmisión; aunque también suelen deberse a fallas en el dispositivo sensor, típicamente por su saturación o la pérdida de la señal. Su función de densidad de probabilidad está dada por la ecuación 3 :

$$p(z) = \begin{cases} P_a & \text{si } z = a \\ P_b & \text{si } z = b \\ 0 & \text{si } \text{otro} \end{cases} \quad (3)$$

2. Ruido Gaussiano: Los píxeles de la imagen cambian su nivel de gris en función de una distribución de probabilidad gaussiana. Este tipo de ruido produce pequeñas

variaciones en la imagen a consecuencia de ruido en el sistema de digitalización (etapa de sensado) y perturbaciones en la transmisión. El canal azul presenta mayor susceptibilidad a ruido de este tipo. El valor final del pixel corresponde al valor ideal de la escena real aumentado por una cantidad de error. Su función de densidad de probabilidad está dada por la ecuación 4:

$$p(z) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(z-\mu)^2}{2\sigma^2}} \quad (4)$$

3. Ruido de disparo o Poisson: Es un tipo de ruido electrónico que ocurre cuando se acumula un cierto número de partículas cargadas con energía, como electrones o fotones, que causan un aumento detectable en el valor del dispositivo sensor, variando el voltaje de salida y por ende los niveles de gris en la imagen capturada respecto a la escena real.
4. Ruido speckle o multiplicativo: La imagen obtenida es el resultado de la multiplicación de dos señales. Es un tipo de ruido que degrada las componentes de alta frecuencia, degrada los detalles finos y la definición de los bordes. Este tipo de ruido hace difícil la identificación de regiones pequeñas y de bajo contraste (Benzarti y Amiri, 2013).
5. Ruido de cuantificación: Este tipo de ruido se debe al proceso de cuantificar una señal analógica, específicamente, es debido a los errores de redondeo cuando no se dispone de suficientes niveles de gris para representar una imagen. La única forma de reducirlo es aumentar la cantidad de bits por pixel disponibles para almacenar el valor radiométrico leído por el sensor.

O’Gorman *et al.* (2008) consideran que no se puede eliminar completamente el ruido en una imagen sin causar alteraciones a la misma, entre más ruido se elimina de la imagen es mayor la reducción de las señales que componen la imagen y puede sufrir distorsiones.

Las técnicas para la reducción de ruido, y en general, las técnicas de mejoramiento de una imagen pueden ser aplicadas en dos dominios: en el dominio espacial y en el

dominio espectral. El procesamiento espacial se refiere a las técnicas que operan directamente sobre los valores de los píxeles de la imagen, sobre la matriz. La ecuación 5 describe el proceso de mejoramiento en el dominio espacial, las técnicas de mejoramiento o transformación T operan sobre la imagen digital $e(x, y)$ para obtener una imagen $S(x, y)$ resultado de la aplicación del operador de transformación.

$$S(x, y) = T(e(x, y)) \quad (5)$$

De acuerdo con Lira (2002), con el objetivo de reducir la complejidad matemática y sobre todo el costo computacional, en lugar de recurrir al dominio de la frecuencia es preferible trabajar en el dominio espacial de la imagen. En el dominio espacial, las técnicas de reducción de ruido utilizadas usualmente son operaciones de grupo. Las operaciones de grupo obtienen un nuevo valor para un pixel dado en base a aplicar una operación sobre el pixel y sus vecinos (ver Figura 14).

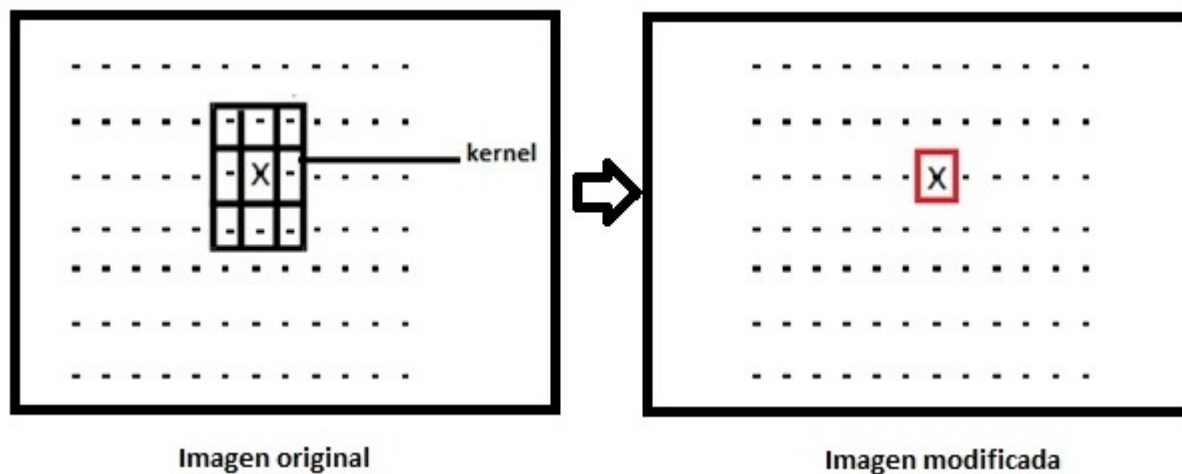


Figura 14: Filtrado espacial.

Se establece un kernel o máscara de convolución de tamaño $m \times n$. Típicamente es una estructura cuadrada donde m es igual a n . El kernel se desplaza a lo largo de la imagen para obtener la respuesta al filtro de cada uno de los píxeles que componen la imagen, y obtener así la imagen modificada, que idealmente debe tener un mejor aspecto visual que la imagen de entrada (original). La respuesta a un filtro espacial está dada por

la ecuación 6, donde x e y son las coordenadas espaciales del pixel, m y n es el tamaño del filtro, $a = m/2 - 1/2$ y $b = n/2 - 1/2$ (Gonzalez y Woods, 2011).

$$e(x, y) = \sum_{s=-a}^a \sum_{t=-b}^b w(s, t) f(x + s, y + t) \quad (6)$$

Es posible establecer la operación que realizan los filtros espaciales. Uno de los filtros utilizados para la reducción de ruido es el filtro de suavizado. El filtro de suavizado de promediado, el cual es una adaptación de un filtro pasa bajo, este filtro difumina la imagen y reduce ruido. La operación que realiza el filtro es un promediado, e.g. si es un filtro de 3 por 3, obtiene la suma de cada píxel que queda delimitado por la máscara de convolución y se divide sobre 9. En la figura 15, se presentan dos ejemplos de máscaras de convolución para un filtro de promedio.

$\frac{1}{9}$	1	1	1
	1	1	1
	1	1	1
a)			

$\frac{1}{16}$	1	2	1
	2	4	2
	1	2	1
b)			

Figura 15: Ejemplo de máscaras para filtro de suavizado.

El filtro gaussiano, al igual que el filtro de promediado es análogo a un filtro pasa bajo que suaviza la imagen y reduce el ruido. Es un filtro de promediado gaussiano. La máscara del filtro gaussiano consiste en aproximar los valores en la máscara de manera que se asemejen a una función gaussiana bidimensional. Conforme la máscara de convolución aumenta de tamaño, el efecto del filtro remueve mayor cantidad de ruido, pero también causa la pérdida de detalle en la imagen.

$\frac{1}{16}$

1	2	1
2	4	2
1	2	1

a)

$\frac{1}{115}$

2	4	5	4	2
4	9	12	9	4
5	12	15	12	5
4	9	12	9	4
2	4	5	4	2

b)

Figura 16: Ejemplo de máscaras para filtro gaussiano.

Otro tipo de filtros que utilizan una máscara de convolución para operar sobre los píxeles de la imagen, son los filtros estadísticos. Uno de los filtros estadísticos que es utilizada con frecuencia para el mejoramiento de las imágenes es el filtro de mediana (dado por la ecuación 7); éste reemplaza el valor del pixel central por el valor de la mediana de los niveles de gris en el vecindario del pixel. El filtro de mediana es utilizado para la reducción del tipo de ruido multiplicativo “speckle” y de sal y pimienta; es de utilidad porque a diferencia de los filtros de promediado preserva los bordes (O’Gorman *et al.*, 2008).

$$e(x, y) = \underset{(s, t) \in S_{x,y}}{\text{mediana}} \{g(s, t)\} \quad (7)$$

Dentro de los filtros estadísticos, comúnmente se encuentran los filtros de máximos y mínimos. El filtro de máximos (dado por la ecuación 8) consiste en asignar al pixel central el valor de nivel de gris más alto del vecindario, ayuda a disminuir el ruido de pimienta. Por el contrario, el filtro de mínimos (dado por la ecuación 9) consiste en asignar el valor de nivel de gris más bajo del vecindario al pixel central, por lo que reduce el ruido de sal.

$$e(x, y) = \underset{(s, t) \in S_{x,y}}{\text{máximo}} \{g(s, t)\} \quad (8)$$

$$e(x, y) = \underset{(s, t) \in S_{x,y}}{\text{mínimo}} \{g(s, t)\} \quad (9)$$

2.4.4. Modificaciones al histograma

El histograma de una imagen es una representación gráfica de la distribución de los valores de intensidad de la imagen (ver Figura 17, i.e. muestra la frecuencia de aparición de cada nivel de gris en la imagen. El histograma tiene un rango $[0, L - 1]$ en donde $L = 2^b$, y b el número de bits con el que se representa cada pixel de la imagen; el histograma está dado por la ecuación 10 donde r_k se refiere al k -ésimo nivel de gris y n_k al número de pixeles en la imagen que tienen la intensidad r_k .

$$h(r_k) = n_k \quad (10)$$

Usualmente se suele normalizar el histograma, obteniendo así un aproximado de la probabilidad $p(r_k)$ de ocurrencia de cada nivel de gris en la imagen (ecuación 11). Las técnicas de transformación basadas en histograma son utilizadas para transformaciones de intensidad, i.e. aumentar o disminuir el brillo en una imagen o de subregiones; es deseable utilizar este tipo de transformaciones para mejorar el contraste en una imagen.

$$p(r_k) = \frac{n_k}{MN} ; \quad \begin{array}{l} M = \text{decolumnas} \\ N = \text{derenglones} \end{array} \quad (11)$$

El contraste se puede definir como el cambio de intensidad entre pixeles y sus vecinos de la imagen, i.e. la cantidad de variabilidad de la intensidad. Mejorar el contraste de una imagen (ver Figura 18) tiene el efecto de esparcir la intensidad sobre todo el rango de los niveles de gris disponibles, mejora la apariencia visual de las imágenes y además da el aspecto de tener una iluminación uniforme en toda la imagen.

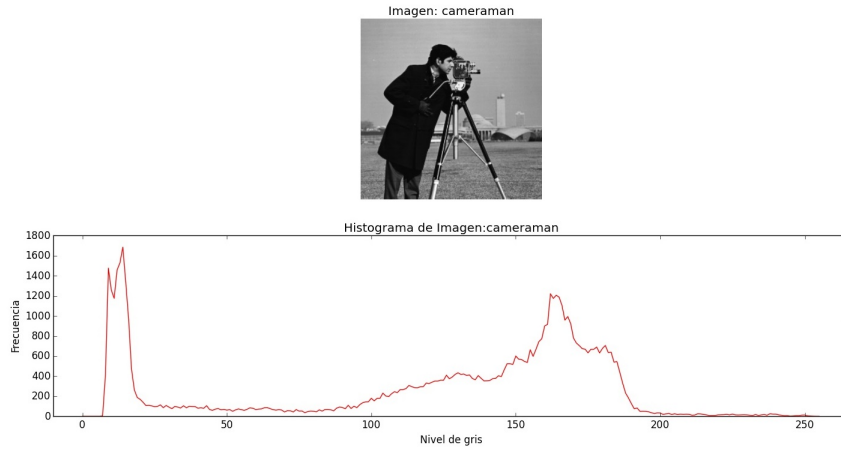


Figura 17: Ejemplo del histograma de una imagen.

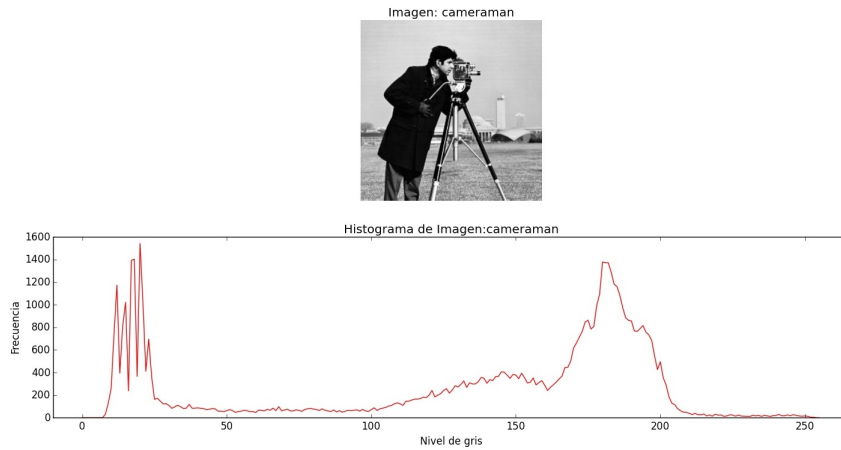


Figura 18: Ejemplo de mejora de contraste (Algoritmo CLAHE).

Las técnicas para mejorar el contraste son basadas en histograma y se conocen como operaciones de punto. Entre estas técnicas, existen operaciones que reducen el rango de niveles de gris, como las funciones logarítmicas; y operaciones que expanden el rango de niveles de gris, como las funciones exponenciales; por otro lado, es posible llevar una imagen a diferentes rangos de niveles de gris utilizando una técnica conocida como normalización del histograma (ecuación 12).

$$N_{x,y} = \frac{N_{max} - N_{min}}{O_{max} - O_{min}} \times (O_{x,y} - O_{min}) + N_{min} \quad \forall x, y \in 1, N \quad (12)$$

La normalización del histograma permite escalar las intensidades de la imagen a un

cierto rango de niveles de gris deseados. Se entiende O por la imagen original y N máximo y mínimo por los valores extremos del rango de intensidad deseado.

Para uniformizar la iluminación de la imagen se utiliza la ecualización del histograma, esta técnica tiene el objetivo de ampliar el rango de niveles de gris en la imagen, y que cada uno de estos tenga una probabilidad de ocurrencia similar. La ecualización del histograma está dada por la ecuación 13.

$$h_{ecu}(i) = round \left[\left(\frac{cdf(k) - min(cdf)}{(M \times N) - min(cdf)} \right) \times L - 1 \right] \quad (13)$$

Donde:

$$\begin{aligned} h(r_k) &= n_k \\ r_k &: k = 0, 1, \dots, L - 1 \\ n_k &: \# \text{ píxeles con nivel de gris } k \\ cdf(k) &= \sum_{k=0}^{L-1} h(r_k) \\ M &: \# \text{ columnas de la imagen} \\ N &: \# \text{ renglones de la imagen} \end{aligned} \quad (14)$$

La operación de ecualización consiste entonces en una especie de *Look – UpTable*, que se construye mediante la aplicación a cada nivel de gris en el histograma de una transformación a un nuevo nivel de gris, para después recorrer cada pixel de la imagen y sustituir su valor original por el nuevo calculado.

La ecualización tiene la desventaja de que redistribuye uniformemente los niveles de gris, es una operación global y aunque reduce la iluminación no uniforme, hace perder detalles en regiones pequeñas. La ecualización del histograma adaptativa soluciona el problema utilizando ventanas (similares a las máscaras de convolución) y aplicando el algoritmo de ecualización del histograma sobre la subregión delimitada por la ventana.

Pizer *et al.* (1987) encontraron que la ecualización de contraste adaptativa tiene el problema de amplificar el ruido y ser un proceso lento, por lo que trabajo en variantes de la ecualización de histograma adaptativa y concluyo que la técnica de ecualización adaptativa del histograma limitada por contraste (CLAHE: Clipped AHE) debería convertirse

en el método predilecto para la mejora de contraste en imágenes médicas.

El CLAHE incluye un factor de recorte que reduce el número de píxeles de los niveles de gris que se encuentren por encima de ese factor y los redistribuye uniformemente sobre todo el histograma. El primer paso consiste en dividir la imagen original en bloques más pequeños, calcular la función de transformación de cada bloque en base al histograma normalizado de la región recortada por el factor de recorte “clip limit” y posteriormente obtener la función de distribución acumulada de cada región. El último paso consiste en realizar una interpolación de los niveles de gris para suavizar la imagen de manera que no se observen visualmente los bloques en los que se dividió la imagen.

2.4.5. Segmentación

La segmentación consiste en subdividir una imagen en regiones u objetos, esto se hace resaltando únicamente regiones con ciertas características sobre el resto de las demás regiones de la imagen y el fondo.

En cierta manera, la segmentación puede definirse como el comienzo de la etapa de detección de objetos, más correctamente una pre-detección, debido a que las imágenes obtenidas de este proceso resaltan las regiones de interés, existiendo casos en los que eliminan completamente los objetos y regiones no deseados.

De acuerdo con Gonzalez y Woods (2011) la segmentación de imágenes es una tarea nada trivial y por el contrario, de las más difíciles en la etapa de procesamiento digital de imágenes, debido a que resulta determinante para la siguiente etapa que es la del reconocimiento de patrones.

La segmentación desde el punto de vista matemático está dada por la ecuación 15, donde R es la imagen completa y cada R_i es una subregión en la imagen. La segmentación se puede realizar de dos diferentes maneras: la primera consiste en encontrar las discontinuidades o cambios notables en el nivel de gris (gradiente), i.e. bordes, la cual se conoce como segmentación basada en bordes; y la segunda consiste en agrupar píxeles con características similares de acuerdo con ciertos criterios que se establecen, e.g. intensidad, parámetros estadísticos, color, textura; la cual se conoce como segmentación

basada en regiones.

$$\bigcup_{i=1}^n R_i = R \quad (15)$$

Las técnicas de segmentación han evolucionado, de manera que se cuentan con otras técnicas que no se pueden clasificar como basadas en regiones o en bordes. Existen técnicas aglomerativas (Clustering) que agrupan píxeles de acuerdo a la información estadística, se suele establecer un número k que indica el número de regiones en las que se desea particionar la imagen. Se inicializan k subregiones aleatoriamente y utilizando la media y la desviación estándar se procede a asignar cada pixel a una subregión. Suelen ser iterativos, de manera que una vez asignado todos los pixeles a una subregión, se calcula el centroide de cada subregión y se asignan nuevamente los píxeles a los nuevos centroides. El algoritmo termina hasta converger, i.e. no hay cambios en las subsecuentes iteraciones.

La técnica de umbralización es a veces incluida como una técnica de segmentación, y otro tipo de segmentación es el de “template matching”, que consiste en tener plantillas de los objetos que se desean encontrar en una imagen y compararla con las regiones de la imagen.

La morfología matemática es una técnicas basada en la teoría de conjuntos para el estudio de la estructura o topología de los objetos, i.e. encontrar patrones espaciales en estructuras bidimensionales, como pueden ser las imágenes. (Lira, 2002)

La morfología matemática cuenta con una serie de operadores que se aplican al objeto deseado utilizando un elemento estructural. El elemento estructural se hace interaccionar con la imagen, y la imagen resultante consiste en el resultado de como dicho objeto encaja o no en la imagen original; de manera que el elemento estructural es el responsable de la alteración de las estructuras en la imagen.

Los operadores morfológicos han sido utilizados para la segmentación de imágenes, en los trabajos Kapur *et al.* (1996); Zana y Klein (2001); Anoraganingrum (1999) se han empleado dichos operadores y probado su efectividad; particularmente porque las opera-

ciones morfológicas simplifican las imágenes conservando características de las estructuras. Los operadores morfológicos son aplicados típicamente a imágenes binarias, y los operadores básicos son los siguientes:

1. Erosión: Es una operación que genera una nueva imagen resaltando únicamente los píxeles de la imagen A en los que el elemento estructural B es un subconjunto de A, de manera que esta operación elimina los bordes. La operación está dada por la ecuación 16 y se observa un ejemplo en la figura 19.

$$A \ominus B = \{x | (B_x) \subseteq A\} \quad (16)$$

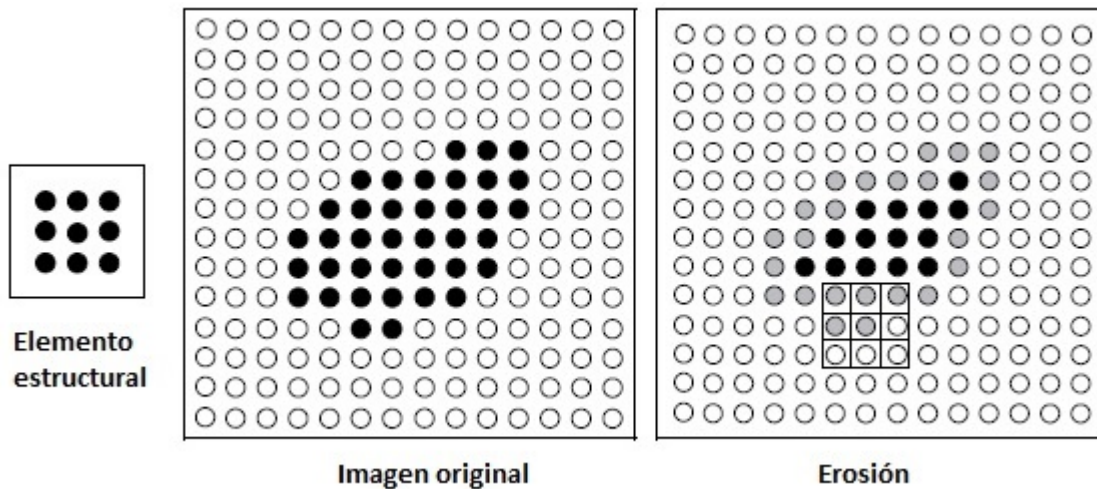


Figura 19: Ejemplo de erosión de una imagen binaria.

2. Dilatación: La operación de dilatación está dada por la ecuación 17, y esta operación consiste en obtener el complemento del elemento estructural para después buscar los píxeles en los que se intersecten A y B, presentando en la imagen de salida los píxeles en los que se intersectan como "1" y "0" en el otro caso. Esta operación agrega píxeles en las fronteras del objeto (ver Figura 20), de manera que al contrario de la erosión, esta operación incrementa o agranda la región o patrón.

$$A \oplus B = \{x | [\hat{B}_x \cap A \neq \emptyset]\} \quad (17)$$

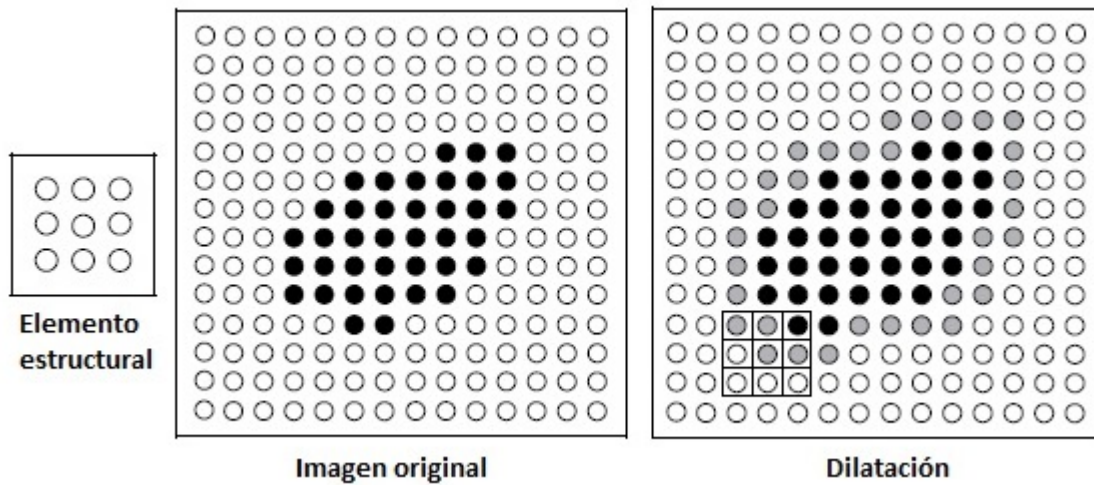


Figura 20: Ejemplo de dilatación de una imagen binaria.

3. **Apertura:** La operación de apertura consiste en aplicar el operador de erosión y a la imagen resultante aplicarle la operación de dilatación. Está dada por la ecuación 18 y ayuda a suavizar los contornos de los objetos eliminando picos más estrechos que el elemento estructural.

$$A \circ B = (A \ominus B) \oplus B \quad (18)$$

4. **Cierre:** Por el contrario a la operación de apertura, el cierre consiste en aplicar el operador de dilatación y a la imagen resultante aplicarle la operación de erosión. Está dada por la ecuación 19 y elimina los valles más estrechos que el elemento estructural. La operación de cierre, al igual que la de apertura se suelen utilizar como filtros para remover ruido del tipo sal y pimienta y para suavizar las formas en las regiones de la imagen.

$$A \cdot B = (A \oplus B) \ominus B \quad (19)$$

5. **Hit-Miss:** La transformada Hit-or-Miss (ecuación 20) consiste en un operador que en la imagen de salida únicamente resalta el pixel central en el que si se coloca el elemento estructural, este es un subconjunto de las regiones en la imagen original, es similar a la operación de erosión. La diferencia radica en que éste es un operador

para encontrar regiones con la estructura que tiene el elemento estructural. Se tiene un elemento estructural que indica donde hay un valor “1” de pixel, un valor “0” y un valor “2” para indicar las zonas en las que no importa el valor del pixel (ver Figura 21). Se puede visualizar entonces como dos operaciones de erosión con dos elementos estructurales distintos B_1 y B_2 , y la intersección de estas genera la imagen resultante.

$$A \circledast B = (A \ominus B_1) \cap (A \ominus B_2) \quad (20)$$

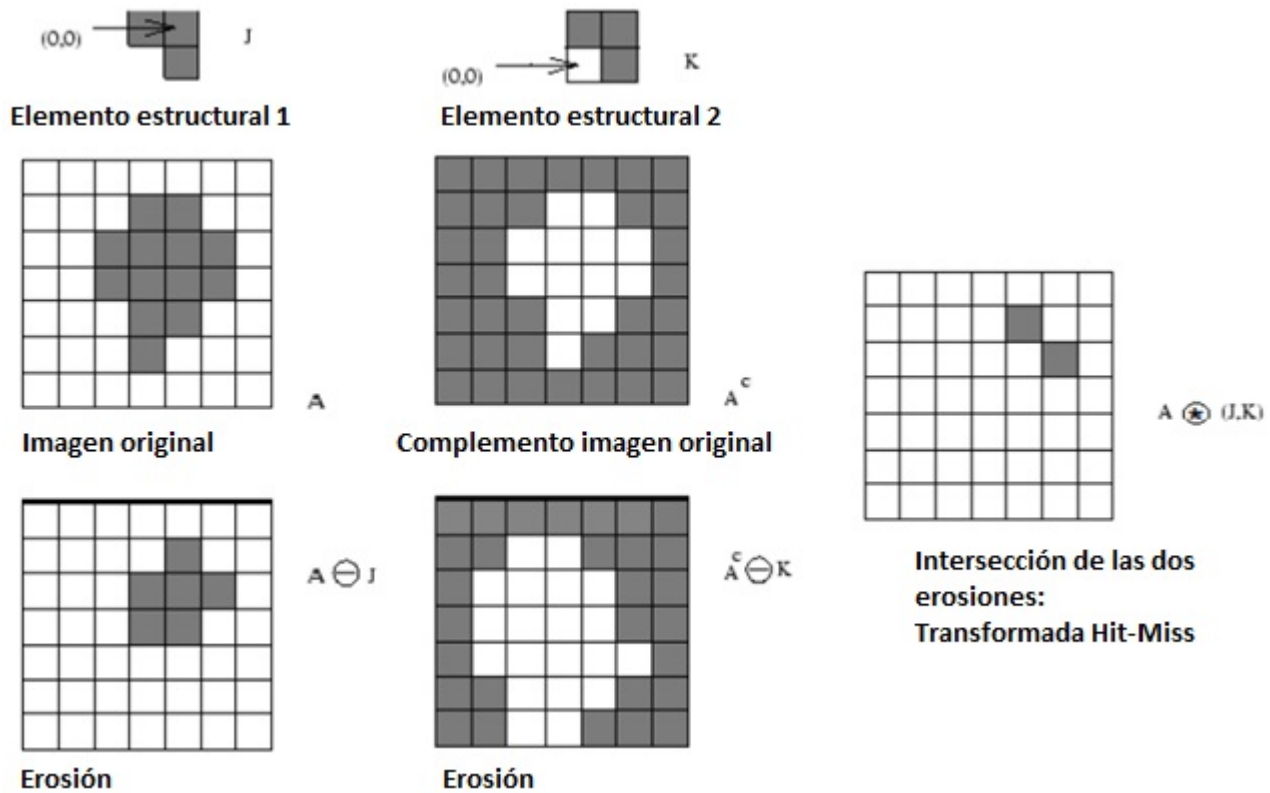


Figura 21: Operador Hit-Miss.

6. Top-Hat: Consisten en operaciones para resaltar objetos con un color contrario al fondo, pudiendo resaltar zonas oscuras o brillantes. El “White Top Hat” está dado por la ecuación 21 y resalta las zonas más brillantes; de manera contraria el “bottom hat” o “black top hat” está dado por la ecuación 22 y resalta las zonas más oscuras.

$$\text{White – Top – Hat} = A - A \circ B \quad (21)$$

$$Bottom - Hat = A \cdot B - A \quad (22)$$

Los operadores morfológicos que se han explicado son exclusivamente para ser utilizados con imágenes binarias; comúnmente se trabaja con imágenes en escala de grises, por lo que se han generalizado estos operadores a este tipo de imágenes y se encuentran en la literatura como “morfolología en niveles de gris”.

La morfolología en niveles de gris es una extensión de los operadores morfológicos básicos, i.e., erosión, dilatación, apertura y cierre. Son utilizados generalmente para detectar bordes y fronteras de regiones, separando así la imagen en subregiones. Las operaciones mantienen la esencia de ampliar o disminuir características en la imagen original; en los operadores binarios se opera espacialmente mientras que en la extensión a niveles de gris opera sobre los perfiles de intensidad, de manera que oscurece o clarea la imagen.

La erosión está dada por la ecuación 23 y la dilatación por otro lado está dada por la ecuación 24; donde A y B se refieren al dominio de f (la imagen) y de el elemento estructural respectivamente.

$$(A \ominus B)(x, y) = \min_{(s,t) \in B} \{A(s + x, t + y) - B(x, y)\} \quad (23)$$

$$(A \oplus B)(x, y) = \max_{(s,t) \in B} \{A(s - x, t - y) + B(x, y)\} \quad (24)$$

Aplicar un operador de dilatación en una imagen en niveles de gris tiene el efecto de reducir los detalles oscuros o incluso eliminarlos, con lo que se obtiene una imagen más clara. Por el contrario, la operación de erosión reduce los detalles brillantes o los elimina, de manera que se obtiene una imagen más oscura que la original. En la figura 22 se muestran ejemplos de estas operaciones.



Figura 22: Ejemplo de operadores morfológicos aplicados a una imagen de niveles de gris.

El concepto de las operaciones de apertura y cierre es exactamente el mismo que para las imágenes binarias. La operación de apertura (ecuación 25) consiste en aplicar el operador de erosión y a la imagen resultante aplicarle la operación de dilatación, elimina los picos más estrechos que el elemento estructural.

$$A \circ B = (A \ominus B) \oplus B \quad (25)$$

La operación de cierre (ecuación 26) consiste en aplicar el operador de dilatación y a la imagen resultante aplicarle la operación de erosión, elimina los valles más estrechos que

el elemento estructural.

$$A \cdot B = (A \oplus B) \ominus B \quad (26)$$

La operación de apertura es utilizada con el objetivo de borrar detalles claros que sean más pequeños que el elemento estructural, pero a diferencia de la operación de erosión no oscurece el resto de la imagen, sino que la mantiene igual. La operación de cierre elimina los detalles oscuros de la imagen al igual que la operación de dilatación, sin embargo, al igual que la apertura, el cierre mantiene el resto de la imagen original, es decir únicamente aclara los detalles oscuros que tienen un tamaño igual o menor al elemento estructural. En la figura 23 se puede observar cómo afecta la operación de apertura y cierre a un perfil de intensidad.

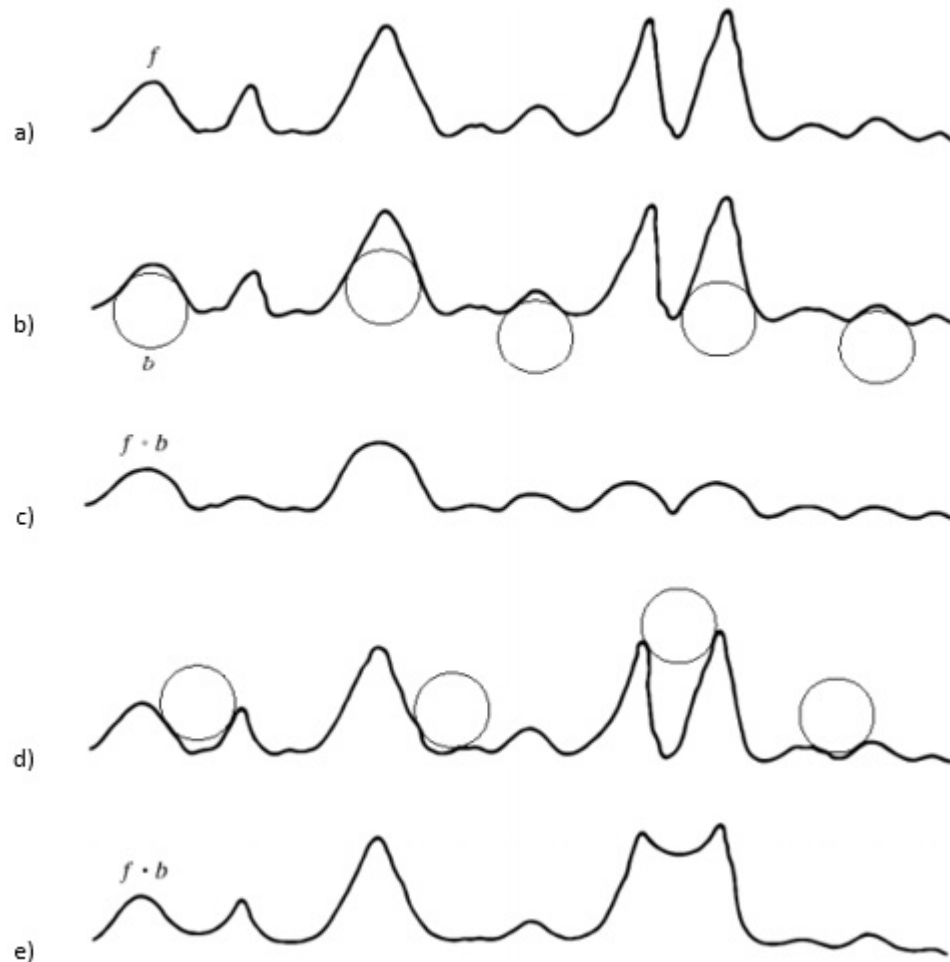


Figura 23: a) Perfil de intensidad b) Apertura c) Resultado de la apertura d) Cierre e) Resultado del cierre.

2.4.5.1. Umbralización

La umbralización o binarización es una técnica que convierte una imagen en escala de grises $f(x,y)$ en una imagen binaria $U(x,y)$, i.e. solo se tienen pixeles con valores igual a uno o cero. La umbralización (ecuación 27) de la imagen en escala de grises consiste en establecer un umbral T , de manera que cada nivel de gris por encima del valor de umbral será igual a 1 y 0 para el caso contrario; o si se desea obtener una imagen binaria invertida, 0 para los niveles de gris por encima del valor de umbral y 1 para el caso contrario.

$$U(x, y) = \begin{cases} 1 & \text{si } f(x, y) > T \\ 0 & \text{si } f(x, y) \leq T \end{cases} \quad (27)$$

Es posible asignar diferente número de umbrales para obtener imágenes en tres o más intensidades. Por ejemplo, si se establecen dos umbrales se obtiene una imagen con 3 diferentes valores de intensidad; con cuatro umbrales se obtienen 5 diferentes valores de intensidad y así sucesivamente. Al igual que la ecualización del contraste, es posible utilizar umbrales adaptativos, de manera que se divide nuevamente la imagen en subregiones y se asigna un umbral diferente a cada región.

2.5. Reconocimiento de patrones

En el contexto del análisis de imágenes, la etapa del reconocimiento de patrones busca el reconocimiento de regiones u objetos que cumplan con ciertas características especificadas para determinar si los objetos de interés pertenecen a una o más clases establecidas.

En el esquema general de detección de objetos, se desea detectar objetos de cierta clase; la detección se refiere a la tarea de determinar si los objetos de la clase deseada se encuentran en la imagen o colección de imágenes analizadas. Las fases de la etapa de reconocimiento de patrones son la extracción de características de cada candidato producto de la segmentación y la clasificación de las características de cada candidato.

2.5.1. Algoritmo de etiquetaje

La etapa de generación de candidatos involucra aplicar técnicas de segmentación y umbralización para resaltar objetos y regiones de interés a la vez que se atenúan objetos y regiones indeseables de acuerdo al objetivo del análisis. Esta etapa se considera una pre-detección, puesto que genera candidatos que son regiones (objetos) de interés más regiones (objetos) similares a éstas en intensidad, textura, y demás características. Para discernir qué objetos son realmente los deseados y cuáles no, se procede a extraer descriptores de cada región, y se conforma un modelo de clasificación que discierne esta interrogante.

Para poder extraer descriptores de las regiones es necesario primeramente identificar cada región y representarla; típicamente se realiza con algoritmos de etiquetaje. El algoritmo de etiquetaje de componentes conectadas agrupa cúmulos de píxeles conectados en regiones (objetos) a partir de una imagen binaria basándose en el concepto de conectividad. Los píxeles que pertenecen a una misma región conectada conforman una región y una etiqueta.

La conectividad se refiere a la proximidad espacial entre píxeles en una imagen, i.e. vecinos. Un pixel tiene ocho vecinos, 4 horizontales y verticales y 4 diagonales. Sea el pixel $p(x, y)$ puede estar conectado por 4 vecinos “4p” (ver ecuación 28) o por 8 vecinos “8p” (ver ecuación 29). El pixel p está conectado con uno de sus vecinos si éstos satisfacen un criterio de similitud, por ejemplo, si su nivel de gris es igual.

$$(x + 1, y), (x - 1, y), (x, y + 1), (x, y - 1) \quad (28)$$

$$(x+1, y), (x-1, y), (x, y+1), (x, y-1), (x+1, y+1), (x+1, y-1), (x-1, y+1), (x-1, y-1) \quad (29)$$

El algoritmo de etiquetaje consiste en recorrer dos veces todos los píxeles de la imagen de izquierda a derecha y de arriba hacia abajo. En el primer recorrido se le asigna

una etiqueta a cada píxel tomando en cuenta a sus vecinos a la izquierda y considerando la conectividad con la que se desea agrupar los píxeles. En el primer recorrido existen tres casos posibles (ver figura 24):

1. Si el píxel no tiene vecinos se le asigna una etiqueta nueva.
2. Si el píxel tiene un vecino se le asigna la etiqueta del vecino.
3. Si el píxel tiene más de un vecino se le asigna una etiqueta de alguno y de indican equivalencias.

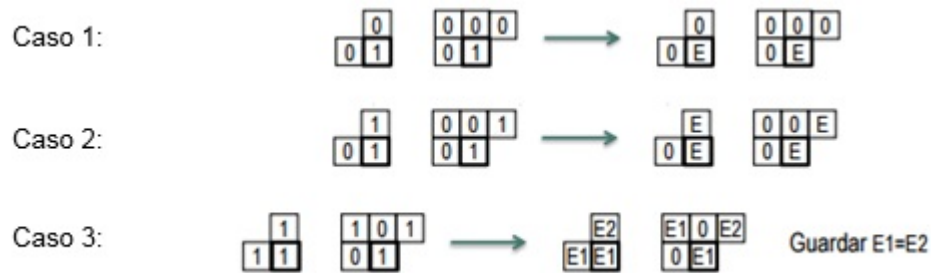


Figura 24: Algoritmo de etiquetaje: Casos posibles.

El segundo recorrido consiste en resolver todas las equivalencias de etiquetas y asignar solo una etiqueta a éstas. En la figura 25 se observa un ejemplo de aplicación del algoritmo de etiquetaje de componentes conectadas.

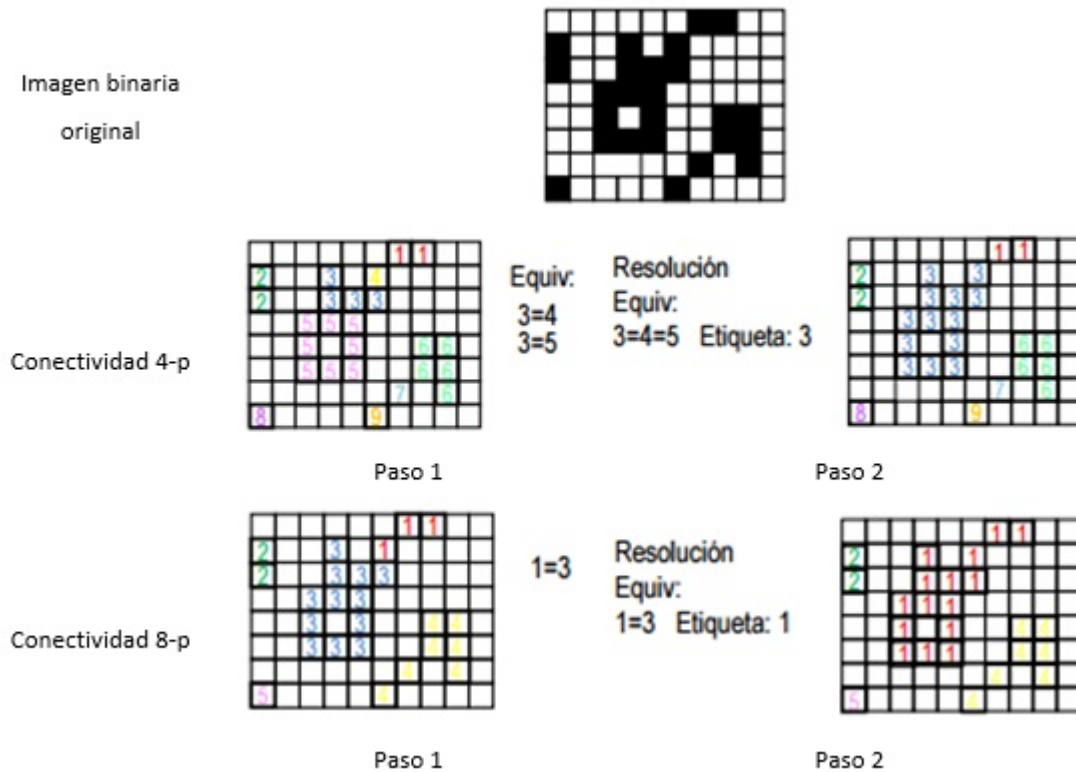


Figura 25: Ejemplo del algoritmo de etiquetaje.

El resultado de aplicar el algoritmo de etiquetaje es la agrupación de píxeles vecinos en regiones, i.e. se obtiene una representación de cada región en la imagen; y a partir de esto se pueden extraer y cuantificar propiedades cada una de las regiones representadas.

2.5.2. Extracción de características

La tarea de la extracción de características consiste en obtener una descripción de los objetos generados como candidatos con el objetivo de obtener un modelo de éste a partir de evaluar las características del objeto. Los descriptores son el resultado de cuantificar propiedades del objeto, tales como la intensidad media, área, perímetro, textura, momentos de inercia, y circularidad por mencionar algunos.

La idea de extraer descriptores de los objetos es la de obtener una representación de cada objeto utilizando un vector de características, el cual se conforma de cada uno de los descriptores cuantificados. El tamaño del vector debe ser suficiente para albergar

las características que sean suficientes para poder discriminar un objeto de otro, o como generalmente se desea, una clase de objetos de otra clase.

El principal problema de la extracción de características es la de seleccionar adecuadamente los descriptores de un objeto, i.e., seleccionar el menor número de descriptores que capturen la esencia del objeto y permitan discernir una clase de otra. Vectores de características con un mayor número de descriptores aumentan la dimensionalidad, con lo que se incrementan los costos de procesamiento, además de que suele disminuir la eficacia de la clasificación correcta debido a que un mayor número de descriptores introducen variaciones y una mayor confusión al modelo de clasificación.

2.5.3. Clasificación

La etapa de clasificación consiste en determinar la clase a la que pertenece un objeto analizando su vector de características. Para poder clasificar los objetos, primeramente es necesario seleccionar un modelo de clasificación y obtener un conjunto de entrenamiento (en caso de clasificación supervisada). Para poder validar el modelo de clasificación se debe conformar un conjunto de prueba y evaluar el desempeño mediante métricas de desempeño, típicamente la exactitud, precisión y sensibilidad.

La clasificación puede ser supervisada o no supervisada. Para la clasificación supervisada, al algoritmo de clasificación se le pone a disposición el conjunto de entrenamiento, el cual contiene distintos vectores de características y su respectiva clase (datos “ground truth⁷”), a partir de esto, el algoritmo puede construir un modelo de clasificación, un plano o hiperplano que separe una clase de otra o una clase de otras en el caso de clasificación multiclase. Para la clasificación no supervisada o “clustering”, el algoritmo no dispone de un conjunto de entrenamiento y agrupa los objetos (generalmente agrupan píxeles y no objetos o regiones) en regiones más grandes de acuerdo a la similitud de sus vectores de características disminuyendo métricas como la varianza entre cada clase o “cluster”.

Bajo el esquema de la detección de objetos, el interés se centra en la clasificación supervisada, ya que es posible generar un conjunto de entrenamiento en el que a cada objeto se le asigne una clase y así con un modelo de clasificación de este tipo sea posible

⁷Término para referirse al resultado correcto que debería dar un clasificador.

asignarle una clase a un nuevo objeto no conocido. Los clasificadores más utilizados son los binarios, de manera que se puede establecer una clase, por ejemplo “Persona” y a los nuevos objetos se les asignará la etiqueta “persona” o “no persona” de acuerdo a que tanta similitud exista entre sus vectores de características con respecto a la de ambas clases. Entre los tipos de clasificadores supervisados, se encuentran los siguientes: vecino más cercano, k-vecinos más cercanos, máquina de soporte vectorial, redes neuronales, adaboost, modelo de mezcla gaussianas, análisis lineal discriminante.

2.5.3.1. Validación

Para la evaluación del rendimiento del sistema de detección de objetos o validación de los modelos de clasificación es necesario construir dos conjuntos de datos, el primero es un conjunto de entrenamiento en el que se sabe a priori la clase de cada objeto y el otro es un conjunto de prueba con el objetivo de evaluar el rendimiento del modelo.

La evaluación del rendimiento consiste en entrenar el modelo de clasificación con el conjunto de entrenamiento y en tratar el conjunto de prueba como casos no conocidos que se someten al modelo de clasificación y éste decide la clase de cada objeto del conjunto. Cada objeto del conjunto de prueba se compara con los datos “ground truth” y en caso de encontrarse correctamente clasificado se cuenta como éxito y de manera contraria se considera un error. Existen diversas métricas para evaluar el rendimiento del clasificador mas la tasa de error es la medida de desempeño más común, la cual es la proporción de los errores entre el número total de predicciones.

Respecto a los conjuntos de entrenamiento y de prueba, la elección de los objetos que quedan en un conjunto u otro es variable y por tanto la estimación del error lo es también, i.e. la tasa de error tiene una alta dependencia de acuerdo a los datos que quedan en el conjunto de entrenamiento y el de prueba. Para la conformación del conjunto de entrenamiento y de prueba se debe dividir el total de datos en dos conjuntos; los métodos utilizados usualmente para la conformación de estos conjuntos son el método “fold out” y la validación cruzada. Generalmente el conjunto de prueba es una tercera parte del conjunto total de datos.

El método “fold out” se refiere a la técnica de separar los datos en un conjunto de

entrenamiento y uno de prueba en base a porcentajes del total de datos. Se establece el porcentaje de los datos que conformarán el conjunto de entrenamiento y se seleccionan aleatoriamente vectores de características hasta cumplir con el porcentaje establecido. Introduce mayores variaciones a la tasa de error pero es más rápido debido a que sólo necesita una iteración.

Por el contrario, la validación cruzada disminuye la variabilidad de la tasa de error a costa de sacrificar rapidez y aumentar el costo computacional, esto es debido a que necesita varias iteraciones. Para la validación cruzada (ver figura 26), el conjunto total de datos se divide en “n” número de segmentos, posteriormente se elige un segmento como datos de prueba y el resto de los “n-1” segmentos conforman el conjunto de entrenamiento. El algoritmo termina hasta que cada segmento ha sido utilizado como conjunto de prueba. La tasa de error es entonces el promedio de los errores de todas las iteraciones; es por esto que es menos susceptible a variabilidad, ya que se garantiza que todos los datos sirven como datos de entrenamiento. El número de segmentos en los que generalmente se particiona el conjunto de datos es en 5, 10 y 20.

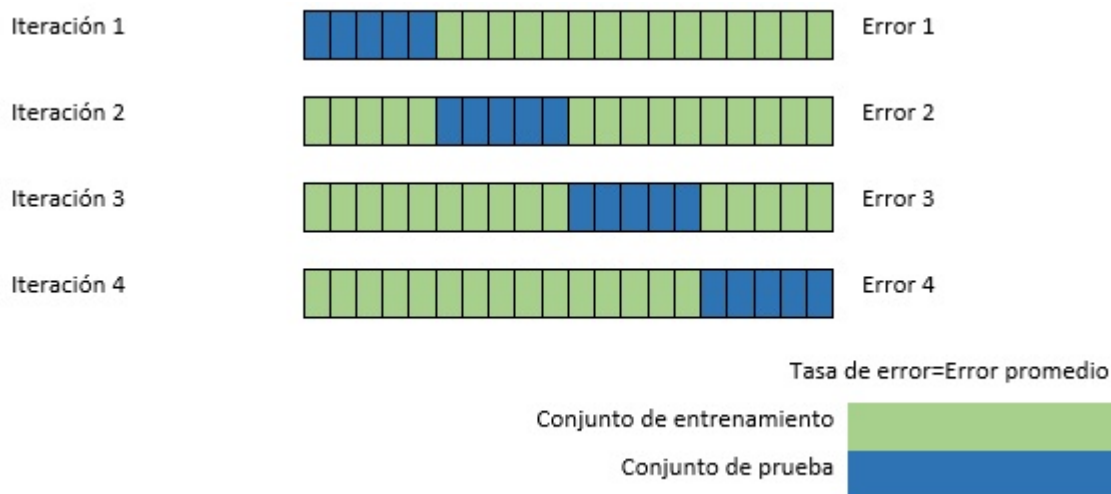


Figura 26: Ejemplo de validación cruzada.

2.5.3.2. Métricas

Las métricas para la evaluación del desempeño de un sistema de detección de objetos, en específico de la etapa de clasificación permiten establecer un punto de com-

paración entre el sistema propuesto con sistemas existentes. Las métricas comúnmente utilizadas se basan en el concepto de matriz de confusión.

La matriz de confusión (ver figura 26) es una herramienta que permite visualizar el nivel de confusión del sistema de detección/clasificación. Relaciona el valor o clase real a la que pertenece un objeto con la predicción del sistema.

		Valor real	
		Positivo	Negativo
Predicción	Positivo	Verdaderos Positivos (VP)	Falsos Positivos (FP)
	Negativo	Falsos Negativos (FN)	Verdaderos Negativos (VN)

Figura 27: Matriz de confusión.

La matriz de confusión y en general los sistemas de detección/clasificación tienen cuatro posibles salidas:

1. Verdaderos Positivos (VP): cuando un objeto positivo se clasifica como un objeto positivo.
2. Verdaderos Negativos (VN): cuando un objeto negativo se clasifica como un objeto negativo.
3. Falsos Positivos (FP): cuando un objeto negativo se clasifica como un objeto positivo (error tipo I).
4. Falsos Negativos (FN): cuando un objeto positivo se clasifica como un objeto negativo (error tipo II).

Los VP's y los VN's conforman las predicciones correctas mientras que los FP's y los FN's conforman el total de las predicciones erróneas. A partir de los cuatro casos se pueden obtener distintas medidas de desempeño, entre las principales se encuentran:

1. Tasa de error: Tasa que indica el número de predicciones erróneamente clasificados por el modelo.

$$Tasa\ de\ error = \frac{FN + FP}{VP + VN + FN + FP} \quad (30)$$

2. Exactitud: Medida que corresponde al porcentaje de predicciones correctas entre el total de predicciones, en sí refleja la proximidad entre el resultado y la clasificación exacta.

$$Exactitud = \frac{\text{predicciones correctas}}{\text{total de predicciones}} = \frac{VP + VN}{VP + FP + VN + FN} \quad (31)$$

3. Precisión: También conocida como valor predictivo positivo (VP+), es una medida de proporción de predicciones positivas correctas con respecto al total de las instancias predichas como positivas; en otras palabras es una medida de calidad de la respuesta positiva del clasificador, i.e., la probabilidad de que el valor realmente sea positivo cuando el diagnóstico es positivo.

$$Precisión = VP+ = \frac{VP}{VP + FP} \quad (32)$$

4. Valor Predictivo Negativo (VP-): Medida de proporción de predicciones negativas correctas. Mide la calidad de la respuesta negativa del modelo de clasificación, i.e. la probabilidad de que la clase del candidato sea realmente negativa cuando la predicción es negativa.

$$VP- = \frac{VN}{VN + FN} \quad (33)$$

5. Sensibilidad (Tasa de verdaderos positivos): Es una medida de proporción que cuantifica las predicciones positivas que han sido correctas respecto a las instancias realmente positivas. Capacidad de detección.

$$Sensibilidad = TPR = \frac{VP}{VP + FN} \quad (34)$$

6. Especificidad (Tasa de verdaderos negativos): Es una medida que indica la fiabilidad de que una predicción negativa sea realmente negativa. Capacidad de discri-

minación.

$$\text{Especificidad} = \text{TNR} = \frac{VN}{FP + VN} \quad (35)$$

7. Tasa de falsos positivos (FPR): Es una medida que cuantifica el error de tipo I.

$$\text{FPR} = \frac{FP}{FP + VN} = 1 - \text{Especificidad} \quad (36)$$

8. Tasa de falsos negativos (FNR): Es una medida que cuantifica el error de tipo II.

$$\text{FNR} = \frac{FN}{VP + FN} = 1 - \text{Sensibilidad} \quad (37)$$

9. Medida F: Media armónica de la precisión y la sensibilidad

$$F_1 = \frac{2VP}{2VP + FP + FN} \quad (38)$$

10. Curvas ROC: Resultado de graficar la tasa de falsos positivos (eje x) y la sensibilidad (eje y).

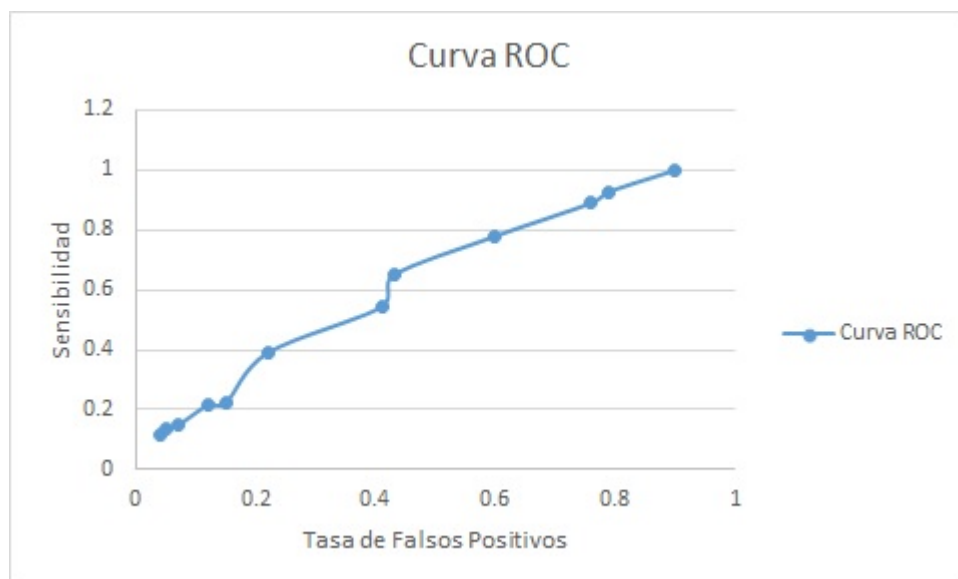


Figura 28: Ejemplo de curva ROC.

Capítulo 3. Algoritmos genéticos

3.1. Introducción

En este capítulo se trata el tema de la técnica de búsqueda de algoritmos genéticos. En primera instancia, a manera de introducción, se define lo que es una técnica de búsqueda y el cómputo evolutivo, para continuar con la descripción de los conceptos y procesos fundamentales que componen a los algoritmos genéticos, para entender las características que éstos heredan de la teoría de la evolución con el objetivo de seguir el principio de Darwin de la “Supervivencia del más apto” (Sivanandam y Deepa, 2008).

3.2. Técnicas de búsqueda

Las técnicas o métodos de búsqueda se refieren a un conjunto de estrategias o técnicas que tienen como objetivo optimizar la solución de un problema, i.e. encontrar la mejor solución de entre un conjunto de soluciones factibles. La mejor solución es aquella que minimiza el costo o maximiza la eficiencia del sistema, e.g. minimización de costo computacional, menor tiempo de procesamiento, etc.

Las técnicas de búsqueda son empleadas en diferentes disciplinas, tanto en el ámbito científico, como en el sector industrial y de desarrollo tecnológico. En específico, destacan las disciplinas de manufactura, aeroespacial, economía, biología y telecomunicaciones, sólo por citar algunas. Un espacio de búsqueda se refiere a las soluciones potenciales del problema.

Las técnicas de búsqueda se encuentran divididas en tres grandes grupos: basadas en cálculos o numéricas, aleatorias y enumerativas. Las técnicas o métodos numéricos encuentran soluciones aproximadas a las óptimas de manera numérica, el problema radica en que muchos de estos métodos no se ajustan a problemas de la vida real, o requieren complejos modelos matemáticos, que al tener espacios de búsqueda pequeños resultan costables, no siendo así cuando son amplios (Bhattacharyya y Maulik, 2013; Bandyopadhyay y Pal, 2007).

Las técnicas enumerativas exploran todas y cada una de las soluciones posibles, lo que resulta en una técnica de búsqueda que garantiza encontrar el óptimo global a un alto

costo computacional y en tiempo. Finalmente, las técnicas de búsqueda aleatorias, que consisten en métodos de búsqueda que exploran aleatoriamente el espacio de búsqueda en lugar de todas y cada una de las posibles soluciones.

En el grupo de las técnicas de búsqueda aleatorias se encuentran las técnicas guiadas, en las cuales a su vez se encuentra un grupo de métodos conocidos como algoritmos o cómputo evolutivo. Estos métodos no exploran un solo punto del espacio de búsqueda, sino que exploran una población de puntos, lo que aumenta las posibilidades de encontrar el óptimo global (Arora, 2015).

3.2.1. Cómputo evolutivo

El cómputo evolutivo es un término acuñado en 1991 (Baeck *et al.*, 2000); en referencia al esfuerzo científico de simular diferentes aspectos de la evolución natural. Se le conoce como una técnica de “Soft-Computing”, puesto que son técnicas tolerantes a incertidumbre e imprecisiones, no son exhaustivas, y gracias a la tolerancia a lo mencionado anteriormente permiten conformar sistemas robustos a bajo costo computacional.

Las técnicas que se clasifican dentro del cómputo evolutivo tienen en común la incorporación de aspectos de la reproducción biológica de los seres vivos, i.e. competencia y selección de individuos, reproducción y variaciones aleatorias; mismas que conducen inevitablemente a la evolución, como resultado de la reproducción de los individuos más competitivos, los cuales en la mayoría de los casos son los mejor adaptados.

Las técnicas del cómputo evolutivo se basan en la teoría de la evolución de Darwin; esta teoría es en esencia un proceso de perfeccionamiento y adaptación de los seres vivos, lo que puede ser traducido a técnicas de búsqueda, como el perfeccionamiento de las soluciones, i.e., optimización. El cómputo evolutivo es entonces una técnica para solucionar problemas que se encuentra inspirada en la naturaleza, concretamente en los mecanismos de la evolución natural. Se rige principalmente bajo los siguientes principios:

1. Supervivencia del más apto.
2. Selección natural y adaptación.

3. Evolución a lo largo del tiempo.

La computación evolutiva tiene cuatro paradigmas: estrategias evolutivas, programación evolutiva, algoritmos genéticos y programación genética; siendo los algoritmos genéticos la técnica más popular.

3.2.2. Algoritmos genéticos

Los algoritmos genéticos fueron desarrollados por Holland en 1975. En su libro “Adaptation in natural and artificial systems”, describió cómo aplicar los fundamentos de la evolución natural para optimizar problemas (Sivanandam y Deepa, 2008). Hoy en día, los algoritmos genéticos se han convertido en una herramienta poderosa frente a los problemas de optimización. Los algoritmos genéticos tienen su base en la genética y la biología, a continuación se describe los fundamentos biológicos de éstos según Sivanandam y Deepa (2008).

1. Célula: unidad anatómica fundamental de los seres vivos, contiene la información genética.
2. Cromosoma: conjunto de genes que se encargan de almacenar la información genética y de codificar las propiedades de las especies. Las posibles propiedades que puede codificar un gen se conocen como alelos, i.e., posibles valores que un gen puede tomar.
3. Genética: se refiere al conjunto de genes, que componen a un individuo. Se le conoce como genotipo. Por otro lado, se encuentra el fenotipo, lo que se refiere al aspecto que tiene el ser vivo al decodificarse el genotipo.
4. Reproducción: se refiere al proceso en el que un individuo engendra y da vida a un nuevo individuo, proceso mediante el cual se copia información genética del padre al hijo.
5. Selección natural: se refiere a la preservación de los seres vivos con características favorables que los hacen más aptos para sobrevivir, continuando así con la preservación de las características favorables y la pérdida de las no favorables.

Los algoritmos genéticos son entonces un método heurístico basado en la supervivencia de las soluciones que mejor resuelven el problema, explorando la aptitud⁸ de un conjunto de soluciones (población de individuos) paralelamente. En la tabla 5 se muestra una comparación entre la terminología para referirse a las representaciones de los individuos en la rama de la biología y de la computación.

Tabla 5: Comparación entre conceptos de evolución natural y la terminología de los algoritmos genéticos.

Evolución natural	Algoritmos genéticos
Cromosoma	Cadena
Gen	Característica o parámetro
Alelo	Valor del parámetro
Locus	Posición de la cadena
Genotipo	Representación codificada de los parámetros
Fenotipo	Estructura decodificada, valores de los parámetros

El esquema general de un algoritmo genético (ver Figura 29) se basa en una población de soluciones, éstas se deben de representar como un cromosoma, típicamente un cromosoma binario. Se genera una población de cromosomas aleatorios, se decodifican y se evalúa qué tan bien o mal cada individuo (resultado de decodificar cada cromosoma) soluciona el problema. En base a la aptitud de cada individuo, se procede a simular la reproducción de los seres vivos utilizando un conjunto de operadores genéticos con el objetivo de crear una nueva población de individuos. Se reemplaza la vieja población por la actual, y se sigue el mismo procedimiento hasta cumplir con las condiciones establecidas para la finalización del algoritmo genético.

3.3. Definiciones y operadores genéticos

3.3.1. Individuos

Un individuo se refiere a una potencial solución al problema que se encuentra dentro del espacio de búsqueda, es decir, se constituye de distintos valores para cada uno de los parámetros de entrada al problema modelado. Los individuos se pueden presentar de dos formas (Figura 30), como cromosomas o como fenotipo.

⁸Entiéndase por función de aptitud, una función cualquiera que sirve como métrica para medir que tan buen desempeño tienen un conjunto de parámetros para la resolución de una tarea o problema.

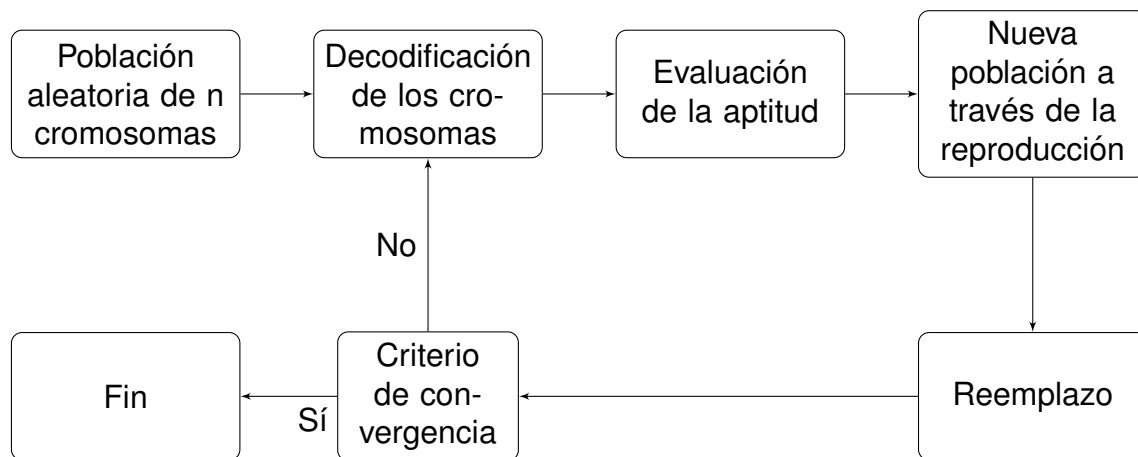


Figura 29: Algoritmo genético simple.

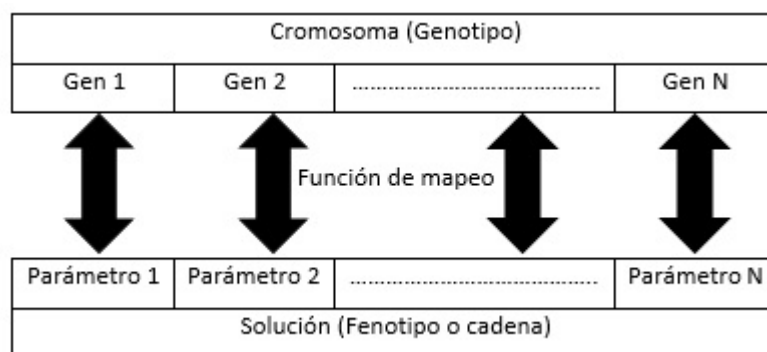


Figura 30: Representación del genotipo y fenotipo (adaptación de Sivanandam y Deepa (2008)).

El cromosoma o genotipo, hace alusión al material genético de la solución, de manera que son datos crudos que sin una función de mapeo hacia una solución válida carece de sentido; típicamente son cadenas de bits. Si se observa desde el punto de vista del modelo, los cromosomas son la forma en que se codifica cada solución potencial, pudiendo emplearse una codificación binaria, octal, hexadecimal, números reales o enteros, entre otras. En la Figura 31, se pueden observar algunos de los métodos de codificación.

Cromosoma binario	1 0 0 0 1 1 1 0 0 1 0 1 1 0 0 1 1
Cromosoma octal	2 4 3 6 7 6 5 4 2 1 0 5 4 0 1 1 2
Cromosoma hexadecimal	F 8 7 3 2 9 A 1 C F D 4 B 6 5 7 A

Figura 31: Métodos de codificación de cromosomas.

La segunda forma de representación es el fenotipo o cadena, que expresa el cromosoma en términos del modelo, es decir, son los valores de los parámetros que entran al modelo, resultado de mapear mediante una función del cromosoma al fenotipo. En la Figura 32, se observa el resultado de establecer una función de mapeo y decodificar los cromosomas de la Figura 31. La función de mapeo establecida es una conversión de binario a decimal, octal a decimal y hexadecimal a decimal, respectivamente, para cada cromosoma.

Fenotipo del cromosoma binario	1 0 0 0 1 1 1 0 0 1 0 1 1 0 0 1 1
Fenotipo del cromosoma octal	2 4 3 6 7 6 5 4 2 1 0 5 4 0 1 1 2
Fenotipo del cromosoma hexadecimal	15 8 7 3 2 9 10 1 12 15 13 4 11 6 5 7 10

Figura 32: Fenotipo de los cromosomas propuestos en la figura 31.

3.3.2. Genes

Los cromosomas se subdividen en genes, de manera que cada cromosoma está compuesto por “n” genes, lo que se traduce en “n” diferentes parámetros codificados dentro de un cromosoma que componen una solución individual al problema modelado. Los genes

son fundamentales para la construcción de los algoritmos genéticos; como ya se mencionó, los cromosomas se componen de un número finito de genes, donde éste número depende del número de parámetros de entrada al problema.

El primer paso para la construcción del algoritmo genético es analizar los parámetros de entrada al problema que se desea modelar, se debe definir el número de parámetros, el tipo de valores y sus respectivos intervalos de operación. Lo anterior, para asegurar que se obtendrán soluciones siempre válidas independientemente del esquema de codificación (representación). El gen es entonces una cadena de datos, en la Figura 33 se puede observar un cromosoma ejemplo, como se puede observar visualmente el cromosoma contiene “n” número de genes, y cada gen está compuesto por una o más posiciones de la cadena (locus o bits). El número de posiciones que se le asignan a un gen aumenta proporcionalmente a medida que el rango de valores que acepta ese parámetro se incrementa.

Cromosoma	1	0	0	0	1	1	1	0	0	1	0	1	1	0	0	1	1
ejemplo	Gen 1				Gen 2				Gen 3		Gen 4		Gen 5		Gen 6		

Figura 33: Estructura de un cromosoma.

En la Figura 34 se observa el fenotipo resultante al utilizar como función de mapeo la conversión de binario a decimal aplicado al cromosoma ejemplo de la Figura 33. El fenotipo (cadena) observado en ésta sugiere que el problema modelado necesita seis parámetros de entrada, y cada uno de éstos tiene intervalos diferentes, razón por la que cada parámetro se compone de un número diferente de posiciones (bits).

Fenotipo	8	57	1	2	0	3
ejemplo	Parámetro 1	Parámetro 2	Parámetro 3	Parámetro 4	Parámetro 5	Parámetro 6

Figura 34: Fenotipo del cromosoma propuesto en la Figura 33.

Retomando la función de mapeo, de acuerdo con Sivanandam y Deepa (2008), ésta es necesaria para convertir el conjunto de soluciones del modelo a una representación que permita al algoritmo genético operar, y para convertir los nuevos individuos obtenidos

como resultado de aplicar los operadores genéticos a valores admisibles del modelo para poder evaluar qué tan aptos son esos nuevos individuos.

3.3.3. Población

El algoritmo genético requiere una población inicial sobre la cual se aplican los operadores genéticos para engendrar nuevos individuos con características adaptadas. La población se refiere a una colección o conjunto de individuos, y una población constituye una generación en el proceso evolutivo del algoritmo genético (cada iteración). De acuerdo con Eiben y Smith (2003), el objetivo de la población es contener posibles soluciones, a las cuales se les evalúa la calidad para resolver el problema modelado. La población es un multiconjunto⁹ de cromosomas (genotipo).

La población, al ser un conjunto de individuos, tiene dos representaciones, “i.e. genotipo y fenotipo”. Al igual que los individuos, el genotipo representa genes y el fenotipo los parámetros. La formulación matemática de ambas poblaciones se define por la ecuación 39 y ecuación 40 respectivamente, donde n es el número de individuos, $len(cromosoma)$ se refiere al tamaño del cromosoma y num_param el número de parámetros que conforman las n soluciones.

$$P_{genotipo} = \begin{bmatrix} bit_geno_{1,1} & bit_geno_{1,2} & . & . & . & bit_geno_{1,len(cromosoma)} \\ bit_geno_{2,1} & bit_geno_{2,2} & . & . & . & bit_geno_{2,len(cromosoma)} \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ bit_geno_{n,1} & bit_geno_{n,2} & . & . & . & bit_geno_{n,len(cromosoma)} \end{bmatrix} \begin{matrix} Cromosoma_1 \\ Cromosoma_2 \\ . \\ . \\ . \\ Cromosoma_n \end{matrix} \quad (39)$$

⁹Del inglés “multiset”: se refiere a que exactamente el mismo individuo puede aparecer en una población dos o más veces.

$$P_{\text{fenotipo}} = \begin{bmatrix} Param_{1,1} & Param_{1,2} & . & . & . & Param_{1,numparam} \\ Param_{2,1} & Param_{2,2} & . & . & . & Param_{2,numparam} \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ Param_{n,1} & Param_{n,2} & . & . & . & Param_{n,numparam} \end{bmatrix} \begin{matrix} Solución_1 \\ Solución_2 \\ . \\ . \\ . \\ Solución_n \end{matrix} \quad (40)$$

La población, entonces, se puede visualizar como una matriz que contiene los individuos como cromosomas o fenotipo. De manera que si se tiene una población de 10 individuos con la configuración del cromosoma de la Figura 33, se obtiene una matriz de 10 x 17, tal y como se puede observar en la Figura 35.

Individuo 1	1	0	0	0	1	1	1	0	0	1	0	1	1	0	0	1	1
Individuo 2	0	1	1	1	0	0	0	1	1	0	0	0	0	1	1	0	0
Individuo 3	1	0	0	0	1	1	1	0	1	1	1	1	1	0	1	1	0
Individuo 4	1	1	0	1	1	1	1	1	0	1	1	1	0	1	1	1	0
Individuo 5	0	0	1	0	1	1	1	1	0	1	1	1	0	1	1	1	1
Individuo 6	1	0	1	0	0	1	1	1	0	1	1	1	1	1	1	0	1
Individuo 7	0	0	0	0	1	1	1	0	1	1	1	0	1	1	1	1	0
Individuo 8	1	0	0	1	0	0	1	1	1	1	1	1	0	1	1	0	1
Individuo 9	1	0	1	0	1	0	1	1	1	0	1	1	0	1	1	1	1
Individuo 10	1	1	0	1	1	0	0	0	1	0	1	1	1	1	0	0	1

Figura 35: Población de 10 individuos.

Las características a tener en cuenta respecto a la población para poder ejecutar el algoritmo genético son:

1. El tamaño de la población: determina la cantidad de individuos que conforman la población de soluciones.

2. Generación de la población inicial: se refiere a la manera en que se crean a los primeros individuos. Existen básicamente dos medios de generación, el aleatorio y el heurístico. El método aleatorio consiste en generar cromosomas aleatoriamente, por ejemplo para la codificación binaria, consiste en inicializar cada posición del cromosoma aleatoriamente a 0 ó 1. De manera contraria, con el método heurístico generalmente se introducen soluciones de las que se conoce tienen un buen desempeño en la resolución del problema modelado. Al tener soluciones aceptables desde la primera generación, la búsqueda tiene la característica de ser guiada y aleatoria, y encuentra buenas soluciones más rápidamente.

Un aspecto a considerar sobre la población, es que se debe tener diversidad en las primeras generaciones, especialmente en la primera generación (población inicial). La diversidad puede definirse como la distancia media entre todos los individuos de la población, a mayor distancia mayor diversidad y viceversa (ver Figura 36).

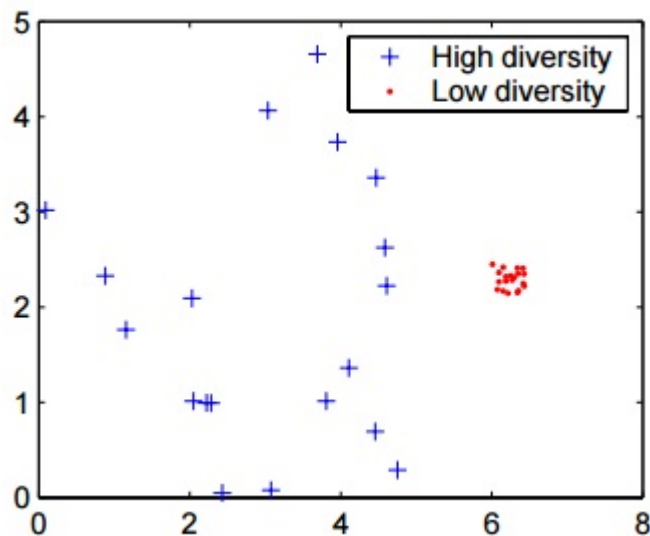


Figura 36: Diversidad en la población. Cruces azules: Alta diversidad. Puntos rojos: Baja diversidad. Adaptación de MathWorks (2005).

El tener diversidad en la población permite explorar en mayor medida el espacio de búsqueda, de manera que es importante tener poblaciones con gran cantidad de individuos como sea posible, se debe aumentar el número de individuos sobre todo cuando se generan de manera heurística. Sivanandam y Deepa (2008) mencionan que una población de alrededor de 100 individuos es frecuente, y que este número aumente en

función del poder de cómputo disponible y las características del modelado del algoritmo genético para con el problema a resolver. De igual manera, Sivanandam y Deepa (2008) advierten el problema del tiempo de convergencia cuando las poblaciones contienen un número grande de individuos, y establece que el tiempo requerido para converger es de orden $(n \log n)$ donde n es el tamaño de la población.

3.3.4. Aptitud

La función de aptitud o función objetivo se refiere a una función de evaluación que se establece para cuantificar qué tan apta es cada solución (individuo fenotipo); asigna una medida de calidad a los individuos.

En cada iteración, se busca que la función de aptitud sea minimizada o maximizada, de acuerdo al criterio con el que calificamos una solución como buena, aceptable o mala. Se maximiza o minimiza la función de aptitud, favoreciendo que las mejores soluciones sean seleccionadas con una mayor probabilidad.

Para Eiben y Smith (2003), la función de aptitud tiene el rol de representar los requerimientos a los que la población se debe de adaptar y define lo que significa un mejor individuo o una mejor población. Conforme la población evoluciona, las soluciones se vuelven cada vez más aptas. Es común escalar la función de aptitud para que se encuentre entre un intervalo de valores (mínimo y máximo), de esta manera, es posible saber qué tan lejos o qué tan cerca se encuentra una solución del óptimo global.

En la Figura 37 se observa un paisaje de aptitud. Mitchell (1999) menciona en su trabajo que: "El concepto de paisaje de aptitud fue originalmente usado por el biólogo Sewell Wright en 1931 para referirse al espacio de todos los posibles genotipos con su respectiva función de aptitud."

El paisaje de aptitud mostrado en la Figura 37 es el resultado de la propuesta del algoritmo genético implementado en este trabajo de investigación. Se tienen poblaciones de 100 individuos, los cuales fueron evolucionados a lo largo de 90 generaciones, información que se puede apreciar en los ejes de la gráfica.

El tercer eje muestra la aptitud de cada solución contenida en cada generación. En la

Figura 37 se puede observar la forma en la que varía la diversidad, alta en las primeras generaciones con tendencia a la baja a medida que el número de generaciones aumenta.

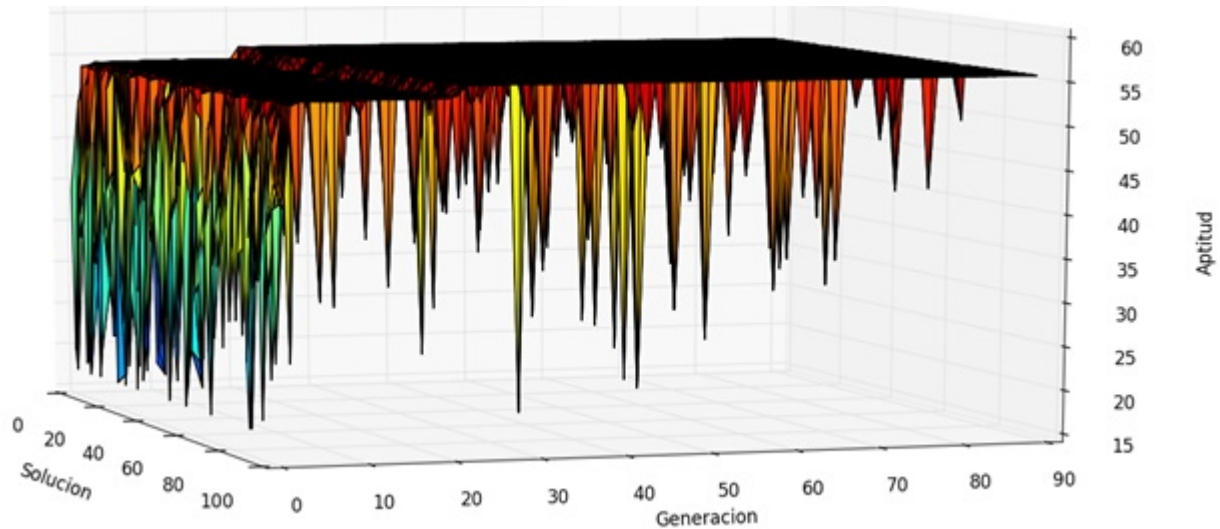


Figura 37: Paisaje de aptitud.

La función de aptitud es relevante para la obtención de buenos resultados del algoritmo genético, es fundamental para la cuantificación de la calidad de las soluciones y por ende tiene un alto peso para el proceso de reproducción. La selección de la función de aptitud no es trivial y es uno de los puntos en el algoritmo genético que debe sintonizarse cuidadosamente para obtener resultados aceptables.

3.3.5. Codificación

La codificación es el proceso de representación de los cromosomas y los genes que los componen. Anteriormente se definió el concepto de individuo y las formas en las que se representan. Un individuo se representa de dos formas distintas: fenotipo (cadena) y genotipo (cromosoma). Existe una función de mapeo que traduce los genotipos en fenotipos y viceversa.

En la Figura 38 se muestra un esquemático sobre la función de mapeo entre el espacio del genotipo hacia el espacio del fenotipo, es decir, del total de combinaciones posibles que se pueden tener de un cromosoma hacia el total de combinaciones que se pueden tener de una solución o cadena. Se debe destacar que la codificación debe ser cuidadosamente seleccionada, así como la función de mapeo, de manera que al decodificar

un cromosoma no se obtenga una solución que no sea válida para la resolución del problema, a esto se le conoce como solución ilegal. De igual manera, debe analizarse el problema a modelar cuidadosamente, para delimitar las partes del espacio de búsqueda que no son factibles, reduciendo así el espacio de soluciones para disminuir el tiempo de convergencia y por ende el tiempo requerido por el algoritmo genético para obtener resultados satisfactorios (Gen y Cheng, 2000).

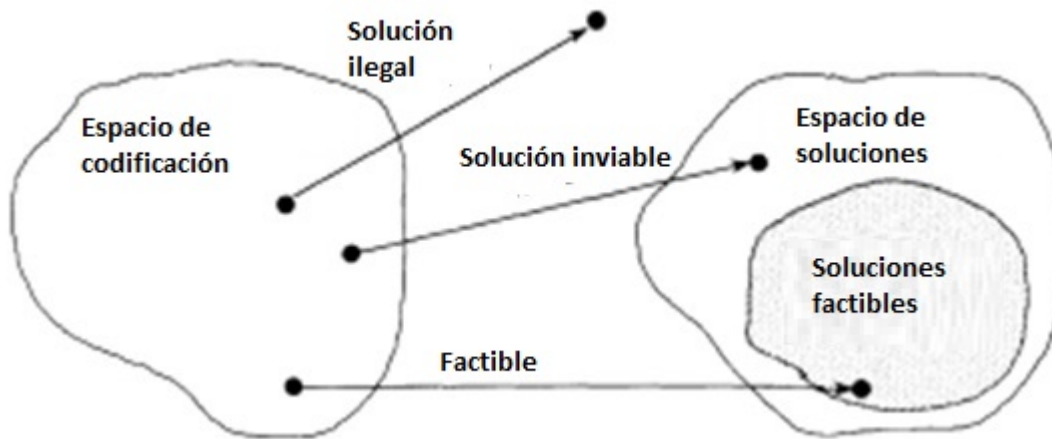


Figura 38: Mapeo del espacio de los cromosomas al espacio de soluciones (adaptación de Gen y Cheng (2000)).

El genotipo se puede representar de distintas formas, algunas de las más comunes son las encontradas en los trabajos de Sivanandam y Deepa (2008); Gen y Cheng (2000); Eiben y Smith (2003):

1. Binaria
2. Octal
3. Hexadecimal
4. Permutación (números enteros)
5. Números reales
6. Valor (datos generales)

Con respecto a la función de mapeo, Gen y Cheng (1997) establecen “Propiedades de las codificaciones” para evaluar los métodos de codificación. El objetivo es construir una

adaptación del algoritmo genético que explore efectivamente el espacio de búsqueda. Las propiedades de las codificaciones son:

1. No redundancia¹⁰: se refiere a que el mapeo entre genotipo y fenotipo deben ser idealmente 1 a 1, es decir, se debe tener el mismo número de combinaciones en ambos espacios, y cada combinación debe corresponderse 1 a 1. Sin embargo, como se advierte en la Figura 39, existen funciones de mapeo en las que un genotipo tiene diversas representaciones en su fenotipo (1 a n) y el caso contrario, en las que diversos genotipos conducen a una única representación en su fenotipo (n a 1).

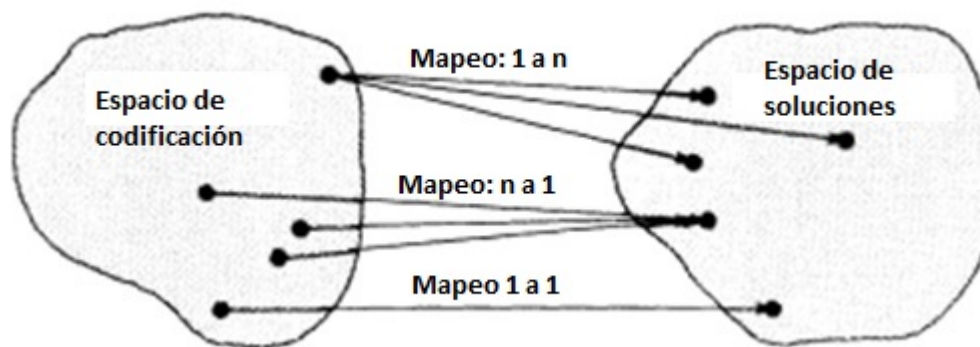


Figura 39: Mapeo de cromosomas a soluciones (adaptación de Gen y Cheng (2000)).

2. Legalidad¹¹: se refiere a que cualquier combinación resultante de un cromosoma se corresponda a una solución válida.
3. Completitud¹²: se refiere a que cada posible solución tenga su representación como cromosoma, para garantizar que todos los puntos contenidos en el espacio de soluciones puedan ser accesibles para la búsqueda del algoritmo genético.
4. Propiedad de Lamarck¹³: se refiere a la capacidad de que las soluciones (padres) hereden a las soluciones (hijos) lo mejor y conduzcan a individuos de mejor calidad, “i.e. individuos más aptos”.
5. Causalidad: se refiere a que pequeños cambios en el espacio del genotipo a causa de una mutación, originen pequeñas variaciones en el espacio del fenotipo.

¹⁰Del inglés “nonredundancy”

¹¹Del inglés “legality”

¹²Del inglés “completeness”

¹³Del inglés “Lamarckian Property”

3.3.6. Reproducción

El proceso de reproducción es la parte fundamental del algoritmo genético, la parte en la que se crean nuevos individuos a partir de aplicar operadores genéticos (selección, cruzamiento, mutación e inversión) sobre una determinada población. La proliferación de la descendencia consiste en un ciclo de cuatro etapas:

1. Selección de los padres
2. Cruzamiento de los padres para crear nuevos individuos “i.e. hijos”
3. Mutación e inversión para simular el proceso evolutivo
4. Reemplazo de todos o una parte de la población (padres) con los nuevos individuos (hijos)

El objetivo del proceso de reproducción es formar una nueva generación de cromosomas cuya aptitud sea mejor que la de la previa generación (Randall, 1995).

3.3.6.1. Selección

La etapa de selección consiste en seleccionar dos individuos de la población y asignarles el papel de padres para posteriormente aplicar el “operador de cruzamiento”. La selección es entonces, la encargada de elegir los individuos de la población, de manera que existen múltiples criterios o esquemas para elegirlos. Esta etapa tiene sus bases en la “Teoría de la evolución de Darwin”, que establece que sólo los individuos más aptos sobreviven y son capaces de reproducirse.

En general, la etapa de selección tiene el objetivo principal de filtrar a los individuos de menor aptitud de la población, de manera que se favorezca la selección de individuos con una alta aptitud. Al seleccionar dos individuos con una aptitud alta en comparación con el resto de la población, se tiene la esperanza de que la descendencia de éstos sea igual o más apta. Una vez seleccionados los dos individuos, son marcados como padres para aplicar el resto de los operadores genéticos que permiten intercambiar información entre ambos.

Aunque existen múltiples métodos para seleccionar los individuos que se reproducirán de una población, entre las propuestas más comunes destacan las descritas por Sivanandam y Deepa (2008), Eiben y Smith (2003), Mitchell (1999), Michalewicz (1994), Goldberg (1989), mismas que se describen a continuación:

1. Selección por ruleta: este método de selección asigna una probabilidad a cada individuo proporcional a su aptitud. La estimación de probabilidad para cada individuo consiste en únicamente dividir su aptitud entre la suma de la aptitud de todos los individuos en la población. El problema radica en que soluciones muy aptas pueden acaparar la mayoría de la ruleta, de manera que son seleccionadas con mayor frecuencia, marginando casi completamente el resto de las soluciones.
2. Selección aleatoria: esta técnica de selección consiste en seleccionar individuos aleatoriamente de la población. No garantiza que se seleccionen los individuos mejor adaptados, sino que favorece la aleatoriedad.
3. Selección por torneo: este método de selección consiste en crear un torneo con un número “n” de individuos, se crea entonces, un subgrupo de la población y los miembros del subgrupo compiten para ser seleccionados como padres. En su versión más simple, el torneo lo gana el individuo con la mejor aptitud, mientras que la otra variante es utilizar la estrategia utilizada en la selección por rango, generar un número aleatorio “R” entre 0 y 1. Se debe especificar un umbral r entre 0 y 1, si $R \geq r$ se utiliza el segundo individuo más apto como padre y si $R < r$ se selecciona el individuo más apto del subgrupo. El tamaño del torneo permite ajustar la presión de selección, Goldberg y Deb (1991) realizaron un estudio y análisis probabilístico para comparar métodos de selección. Encontró que la selección por torneo, cuando $n=2$, es exactamente igual a la selección por clasificación. Además, con sus experimentos concluyó que conforme aumenta el tamaño del torneo, i.e. la presión de selección, aumenta la probabilidad de seleccionar a los individuos más aptos de toda la población, lo que deriva en una convergencia rápida del algoritmo, tal y como se puede observar en la Figura 40.
4. Selección por rango o clasificación (Rank Selection): este método de selección cla-

sifica a todos los individuos asignando un número de 1 hasta N. Las soluciones se enumeran de peor a mejor aptitud. Se seleccionan dos individuos aleatoriamente, así mismo se genera un número aleatorio “R” entre 0 y 1. Se debe especificar un umbral “r” entre 0 y 1, si “ $R \leq r$ ” se utiliza el segundo individuo como padre y si “ $R > r$ ” se selecciona el primero. Otra variante es únicamente seleccionar dos individuos aleatoriamente y asignar como padre al individuo con mayor aptitud. Es una especie de selección por torneo binario. La selección por rango evita la convergencia prematura, de hecho su convergencia es lenta, asegura la diversidad genética y usualmente lleva a encontrar una solución aceptable.

5. Selección de Boltzmann: este método de selección simula el decremento de temperatura de un cuerpo físico. La presión de selección varía con el decremento, por lo que se tiene una presión de selección adaptativa, no es fija, sino que va aumentando a medida que se enfría el cuerpo físico. Esta estrategia permite mantener una alta diversidad genética, y a medida que el número de generaciones aumenta ser más selectivos, i.e. seleccionar individuos más aptos, la diversidad genética de-

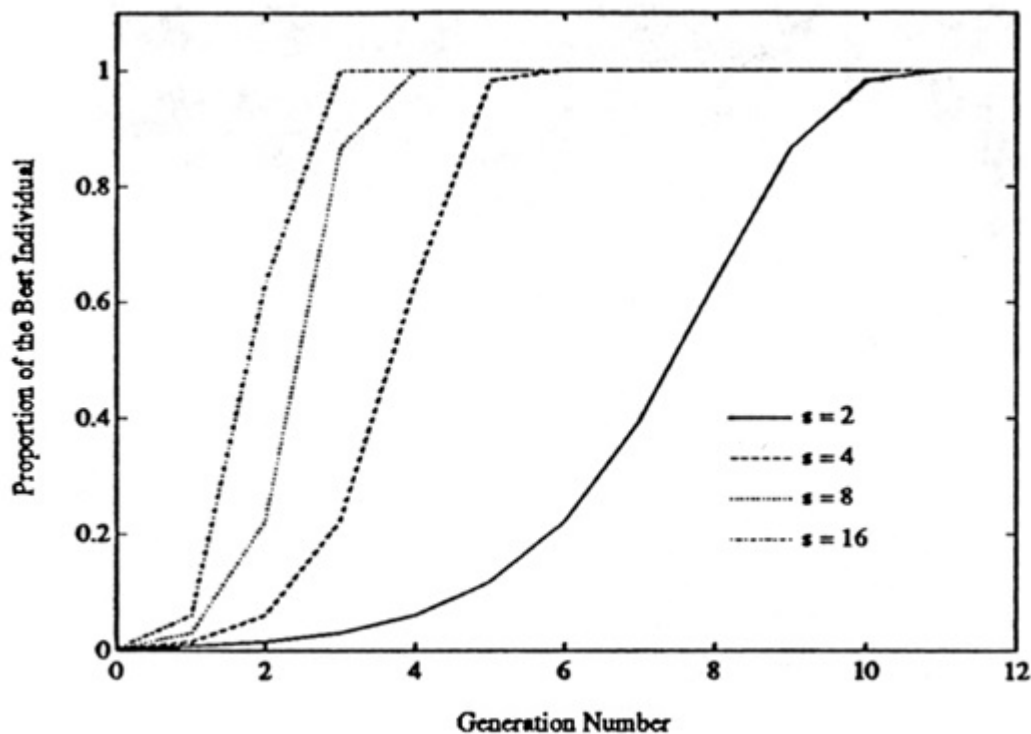


Figura 40: Selección por torneo con diferentes valores de S (tamaño de torneo). Imagen tomada de Goldberg y Deb (1991).

crementa y empieza la convergencia. El algoritmo termina su ejecución cuando la temperatura llega a cero.

6. Selección por muestreo estocástico universal: este método es una variante de la selección por ruleta. Se asigna una porción de la ruleta de acuerdo a la probabilidad de cada individuo. Se declara una variable con valor igual al número de individuos que se desea seleccionar, a partir de lo anterior se divide la ruleta entre el número de individuos deseados con espaciamiento igual y uniforme, y se plasman marcadores en cada posición. Se seleccionaran los individuos en los que haya marcadores. Si por ejemplo, se tiene una ruleta como la mostrada en la Figura 41, y se desea seleccionar 4 individuos, se colocan 4 marcadores espaciados uniformemente. Los cuatro individuos seleccionados son: 1ero. el rojo, 2do. el rojo, 3ero. el azul y 4to. el negro. Este tipo de muestreo permite a las soluciones con baja aptitud tener una mayor oportunidad de ser seleccionadas.



Figura 41: Ejemplo de muestreo estocástico universal.

7. Selección elitista: para este método de selección se toman los “n” individuos más aptos, se copian exactamente iguales a la siguiente generación. El resto de la población ($\text{Población} - n$) es completada con el proceso de reproducción utilizando otro método de selección. Este método de selección garantiza tener al menos uno o más buenos individuos, para que en caso de que el cruzamiento y mutación no

engendren individuos con una buena aptitud, mantener la genética de las buenas soluciones que ya se habían encontrado.

3.3.6.2. Cruzamiento

La etapa de cruzamiento, también llamada de recombinación, es la etapa que simula y emula el proceso de la reproducción natural. Involucra la participación de dos individuos, los cuáles son elegidos en la etapa previa, la de selección. Los dos individuos seleccionados intercambian entre sí segmentos de sus correspondientes cromosomas, de esta manera se crea la descendencia, los hijos, los que son resultado del intercambio genético de sus padres, tal y como en el proceso de la reproducción natural.

La etapa de selección selecciona típicamente los individuos mejor adaptados para solucionar el problema en cuestión, de manera que la etapa de cruzamiento apunta a enriquecer a la población con individuos más aptos, bajo la esperanza de que dos individuos con una buena función de aptitud eventualmente pueden dar origen a un nuevo individuo, cuya aptitud se vea aumentada por vía del operador de cruzamiento, mejorando así la descendencia y las generaciones futuras.

El operador genético de cruzamiento se aplica bajo los siguientes pasos:

1. Seleccionar dos individuos (cromosomas) para el apareamiento
2. Seleccionar aleatoriamente uno o más puntos de cruce en los cromosomas
3. Intercambiar los valores delimitados por los puntos de cruce entre los dos cromosomas

Los métodos de cruzamiento varían de acuerdo al tipo de codificación empleada para representar las posibles soluciones al problema. Usualmente, cuando se utilizan codificaciones binarias, octales, hexadecimales, y de valor, es muy común encontrarse con los métodos de cruzamiento siguientes:

1. Cruzamiento en un punto: se establece un punto de cruce en el cromosoma, el primer hijo será resultado de copiar la información del segmento izquierdo del padre

- 1 y el segmento derecho del padre 2. El segundo hijo, de manera contraria, será una copia del segmento derecho del padre 1 y del segmento izquierdo del padre 2.
2. Cruzamiento en dos puntos: se establecen dos puntos de cruce, de manera que el cromosoma es dividido en tres segmentos. El primer hijo será resultado de copiar los segmentos de los extremos del padre 1 y el segmento central del padre 2. El segundo hijo será una copia de los segmentos extremos del padre 2 y del segmento central del padre 1.
 3. Cruzamiento multipunto: se establece más de un punto de cruce “n”, lo que genera “n+1” segmentos. Los hijos son el resultado de copiar alternadamente segmentos del padre 1 y del padre 2.
 4. Cruzamiento uniforme: se genera un cromosoma vacío de igual tamaño al de los cromosomas. Se llena el cromosoma vacío con valores aleatorios de “0” y “1”. El hijo 1 se genera copiando información del padre 1 si es “0” y del padre 2 si es “1”; de manera contraria, el hijo 2 copia información del padre 1 si es “1” y del padre 2 si es “0”.
 5. Cruzamiento de sustituto reducido (reduced surrogate): para este método de cruzamiento se restringen los puntos de cruce a únicamente las fronteras de los genes, de esta manera permite cruzar genes completos sin alterar sus valores, dejando esta última tarea al operador de mutación.
 6. Cruzamiento multi-padres: es un método de cruzamiento en el que el mecanismo de recombinación requiere más de un individuo padre, e.g. 2,3,4,5,6 o más padres. Para aplicar el operador de cruzamiento utilizando más de 2 padres, existen varias técnicas. Las más comunes son utilizar máscaras aleatorias; y otras utilizan esquemas de votación (Scanning crossover), de manera que para engendrar el hijo, se copia el valor de cada posición en concordancia al valor con mayor frecuencia de aparición en los “n” padres (ver Figura 42) (Eiben *et al.*, 1995; Chambers, 1998; Eiben, 1999).

Cuando se codifican las soluciones empleando números reales, el método mediante el que se cruzan los dos o más individuos suelen ser operaciones aritméticas. Por otro

Padre 1	1	0	0	0	1	1	1	0	0	1	0	1	1	0	0	1	1
Padre 2	0	1	1	0	1	0	1	0	0	0	0	1	0	0	0	0	1
Padre 3	1	1	0	1	1	0	1	0	0	1	0	1	0	1	0	1	1
Hijo	1	1	0	0	1	0	1	0	0	1	0	1	0	0	0	1	1

Figura 42: Cruzamiento "scanning".

lado, cuando se representan las soluciones utilizando números enteros que describen un orden, destacan los métodos de cruce siguientes: cruzamiento de mapeo parcial, propuesto por Golberg y Lingle para la resolución del problema del vendedor viajero¹⁴ (Goldberg y Lingle, 1985); el cruzamiento ordenado (Davis, 1991) y el cruzamiento por ciclo (Oliver *et al.*, 1987).

Independientemente del método empleado para cruzar los individuos seleccionados como padres, el parámetro que juega el papel más determinante es el llamado "Probabilidad de cruzamiento" o "Tasa de cruzamiento". La probabilidad de cruzamiento determina la frecuencia con la que se efectuará el cruzamiento, de manera que existen dos posibilidades: cruzamiento y no cruzamiento.

Para determinar si se aplica el operador de cruzamiento sobre dos las soluciones padres, se toma una decisión pesada. Cuando se determina que no hay cruzamiento, los hijos son copias exactas de los padres; en caso contrario, los hijos son el resultado de aplicar el operador sobre los padres.

De acuerdo a Mitchell (1999), la probabilidad de cruzamiento igual a 0.7 es bastante típica, cabe mencionar, que cuando ésta es igual a 1, siempre se aplica el operador de cruzamiento, de manera contraria, al ser igual a 0, los hijos siempre serán réplicas exactas de sus padres.

Establecer la probabilidad de cruzamiento óptima no es una tarea trivial ni sencilla, y depende de la naturaleza de la función de aptitud (Man *et al.*, 1996). Existen, sin embargo, guías para establecer una probabilidad de cruzamiento adecuada, Man *et al.* (1996) en base a sus investigaciones sugiere utilizar una probabilidad de cruzamiento

¹⁴TSP: Traveling Salesman Problem

igual a 0.9 para poblaciones pequeñas, de alrededor de 30 individuos; mientras que para poblaciones grandes sugiere un valor de 0.6 para poblaciones que contengan 100 o más individuos.

3.3.6.3. Mutación e inversión

La tercera etapa en el proceso de reproducción comprende dos operadores genéticos: mutación e inversión. El operador de mutación se aplica inmediatamente después de la reproducción, y el de inversión después de la mutación, sin embargo, este último operador generalmente no es utilizado en las implementaciones de algoritmos genéticos, e incluso, es pasado por alto en la literatura y no se menciona en el esquema general del algoritmo genético.

El operador genético de mutación altera aleatoriamente los valores en algunas posiciones de los nuevos individuos, “i.e. hijos”. El objetivo de hacer estas alteraciones, principalmente es para prevenir que el algoritmo genético quede atrapado en mínimos locales y juega el rol de recuperarse de la pérdida de material genético (Sivanandam y Deepa, 2008).

Las aportaciones del operador de mutación son:

1. Permite explorar todo el espacio de búsqueda
2. Introducir nuevos individuos en la población
3. Mantener la diversidad de individuos en la población
4. Ayudar a escapar de mínimos o máximos locales

Los métodos de mutación generalmente consisten en cambiar un valor en la cadena, como típicamente se codifican las soluciones utilizando una representación binaria, el volteo (flipping) es uno de los métodos más comunes, que consiste en cambiar 1's por 0's y viceversa. A continuación se describen algunos de los métodos más comunes de mutación:

1. Flipping: consiste en recorrer con un marcador todas las posiciones del cromosoma, tomar una decisión pesada con cierta probabilidad de cruzamiento, si el resultado de la decisión es el de aplicar la mutación, cambia el valor del cromosoma en la posición actual en la que se encuentra el marcador de “0” a “1” o de “1” a “0” según sea el caso.
2. Interchanging: este tipo de mutación consiste en elegir dos posiciones (“q” y “r”) aleatoriamente del cromosoma e intercambiar el valor de la posición “q” por el de la posición “r” y viceversa.
3. Reversing: este tipo de mutación consiste en elegir una posición de manera aleatoria del cromosoma. Entre la posición elegida y la última posición del cromosoma se forma un segmento de menor longitud que el cromosoma, ese segmento es invertido, de manera que la primera posición del segmento pasará a ser la última posición, la segunda posición pasará a ser la penúltima posición del segmento, y así sucesivamente hasta invertir el segmento completamente.
4. Insert: este método de mutación consiste en elegir aleatoriamente dos posiciones “q” y “r” del cromosoma. El paso siguiente es elegir si se mueve la posición “q” o la posición “r”. En caso de elegir la posición “q” y que qr , el valor en esa posición se insertará a la derecha de la posición “r”, teniendo que deslizar los valores de las posiciones contiguas hacia la derecha para hacer espacio. En caso de elegir la posición “r”, su valor se insertará a la derecha de la posición “q”, después de haber aplicado el deslizamiento de los valores hacia la derecha.
5. Scramble: consiste en elegir aleatoriamente dos posiciones del cromosoma; estas dos posiciones forman un segmento central. La mutación del cromosoma consiste en reacomodar de manera aleatoria los valores de ese segmento.

El operador de mutación, al igual que el de cruzamiento, necesita el establecimiento de una probabilidad, en este caso se denomina probabilidad de mutación. El objetivo de este parámetro es decidir qué tan frecuentemente ocurrirá una mutación en el cromosoma. Si no hay mutación, los hijos son el resultado del cruzamiento de los padres (si es que lo hubo), de manera contraria se muta el cromosoma con la probabilidad especificada.

La decisión de la mutación, se toma mediante una decisión pesada. Se debe cuidar de no establecer una probabilidad de 100 %, ya que al hacer esto siempre habrá mutación y entonces el algoritmo genético pasará a ser una búsqueda aleatoria.

Con respecto al valor de la probabilidad de mutación, es un parámetro que debe ser sintonizado; al igual que el valor de la probabilidad de cruzamiento, no resulta un proceso trivial el encontrar este valor. De acuerdo a (Mitchell, 1999), un valor adecuado para una primera ejecución de un algoritmo genético es igual a 0.001. En contraparte, (Sivanandam y Deepa, 2008) afirma que usualmente este valor es igual a $1/L$, siendo L la longitud del cromosoma.

En el estudio de Man *et al.* (1996), se recomienda el valor de 0.001 para poblaciones grandes y 0.01 para poblaciones pequeñas.

El operador genético de inversión, es un operador básico al igual que el operador de selección, cruzamiento y mutación. Este operador consiste en seleccionar un segmento de cromosoma e invertirlo, de tal manera que el último bit del segmento pase a ser el primero, el penúltimo pase a ser el segundo y así sucesivamente.

Para aplicar este operador, se debe cuidar que el esquema mediante el que se codifican las soluciones no sea dependiente de la posición, de manera que una inversión no genere soluciones fuera del espacio de soluciones. Cuando el mapeo de cromosomas a parámetros depende de la posición, hacer esto deriva en que un operador de mutación a gran escala (Whitley, 1994).

De acuerdo a Liepins y Vose (1991); Goldberg y Deb (1991): "El operador genético de inversión fue estudiado por Bagley (1967), Cavicchio (1970) y Frantz (1972); para ser introducido como un operador básico de los algoritmos genéticos por Holland (1975)." Algunos autores suelen pasar este operador por alto, a raíz de que a lo largo de los años se han incorporado nuevas técnicas de cruzamiento y de mutación, y que específicamente en la etapa de mutación existen variantes que hacen precisamente lo que el operador de inversión realiza. Aunque en los algoritmos genéticos no suele hacerse referencia a este último operador, sí que es mencionado y utilizado en otras técnicas de búsqueda mediante cómputo evolutivo como en la evolución diferencial.

3.3.6.4. Reemplazo

La etapa de reemplazo es la última etapa en el ciclo de reproducción. Después de aplicar los operadores de selección, cruzamiento, mutación e inversión, se tienen la población original más los hijos que resultaron de aplicar los operadores antes mencionados. En el caso de que se hayan cruzado todos los individuos, se tiene entonces el doble de la población original; de manera que es necesario filtrar la mitad de esa población para conformar la nueva generación.

En cierta manera, esta etapa puede visualizarse como la última etapa en el ciclo de vida de los seres vivos, en la que los individuos perecen; ya que de igual manera, en la etapa de reemplazo son eliminados ciertos individuos emulando así el deceso natural. La etapa de reemplazo junto con la de selección son determinantes en la convergencia del algoritmo genético.

Generalmente el criterio utilizado para reemplazar a la población, consiste en reproducir todos los individuos, obteniendo un número de hijos igual al número de padres y reemplazando todos los padres por los hijos engendrados. Existen otros criterios para el reemplazo, cada uno influye de manera diferente en la decisión de que individuos permanecen en la población y cuales son reemplazados.

El reemplazo se puede hacer de manera generacional o de manera continua, en el generacional, se reemplaza al final y en el reemplazo continuo se introduce el nuevo individuo en la población en cuanto es engendrado, reemplazando usualmente al individuo menos apto, el cual es seleccionado bajo un esquema similar al de selección por torneo, con la diferencia de que se selecciona al individuo más débil con mayor frecuencia. A continuación se detallan algunas de las técnicas de reemplazo más usuales:

1. Reemplazo aleatorio: en este tipo de reemplazo se aplican los operadores de selección, cruzamiento y mutación tantas veces como el número de hijos deseados. Al tener el número de hijos deseados, se comienza a seleccionar aleatoriamente un individuo de la población actual, se retira y se introduce uno de los hijos engendrados. Este proceso se realiza hasta haber incluido todos los hijos en la población

actual y al terminar se obtiene la población que será una nueva generación.

2. Reemplazo de padres débiles: cuando se seleccionan dos padres y se aplica el operador de reproducción, se tienen los dos padres más los dos hijos, cuatro individuos en total; y únicamente sobreviven para la próxima generación los dos individuos más aptos.
3. Ambos padres: este es el esquema de reemplazo típicamente utilizado. Consiste en aplicar el operador de cruzamiento e inmediatamente sustituir los padres por los hijos, es decir todos los padres serán desechados y la nueva generación estará conformada únicamente por individuos hijos.

3.3.7. Criterio de convergencia

El criterio de convergencia o de paro se refiere a una condición que se debe cumplir para terminar la búsqueda de soluciones aptas, i.e. la ejecución del algoritmo genético. Existen numerosos criterios de convergencia, pudiéndose implementar uno o más en un algoritmo genético. Los criterios de convergencia comúnmente encontrados en la literatura y en implementaciones son los siguientes:

1. Número máximo de generaciones: se establece una variable que especifique el número máximo de generaciones, i.e. qué tantas generaciones permitiremos; se aplicará el proceso de reproducción hasta llegar al número especificado, procediendo a terminar la ejecución del algoritmo genético.
2. Sin cambios en la aptitud: se establece una variable que especifique un número “n” mucho menor que las generaciones esperadas. Al terminar el proceso de reproducción, se buscará el mejor individuo en las últimas “n” generaciones, y si no hay ningún cambio en la aptitud de éste (mejor individuo) se termina la ejecución del algoritmo genético.
3. Estabilidad generacional¹⁵: este criterio de convergencia necesita dos variables. La primera variable indica el número “n” de generaciones sobre la cual se comparará

¹⁵Del inglés: Stall generations

la aptitud promedio de la población. La segunda variable es un número conocido como rango de tolerancia, mientras que los promedios de la aptitud de las “n” últimas generaciones difieran en un número mayor de ese rango de tolerancia el algoritmo continuará ejecutándose, y finalizará una vez que difieran en número menor.

4. Estabilidad en el tiempo¹⁶: este criterio de convergencia, al igual que el de estabilidad generacional, necesita dos variables. La primera variable indica el tiempo “t” en segundos. La segunda variable es el rango de tolerancia, y de igual manera se evalúa que los promedios de la aptitud no difieran en un número mayor o igual de ese rango de tolerancia para finalizar la ejecución. La diferencia radica en que los promedios de la aptitud que se comparan no corresponden a un número “n” de últimas generaciones, sino que corresponden a los promedios de la aptitud de cada generación que se encuentre dentro de un rango (tiempo actual – t , tiempo actual).
5. Umbral para el mejor individuo: para este criterio de convergencia, se establece un parámetro, el cual será un umbral que deberá sobrepasar el mejor individuo de la población para poder terminar la búsqueda.
6. Umbral para el peor individuo: este criterio de convergencia es similar al umbral para el mejor individuo, con la diferencia que quien debe sobrepasar el umbral es el individuo menos apto.
7. Mediana de la aptitud: para este criterio de convergencia, se establece un parámetro umbral. Para que la búsqueda termine, más de la mitad de los individuos de la población deben tener una aptitud mayor a la del umbral.

Existen múltiples criterios de convergencia, los cuales pueden combinarse para evitar ejecutar el algoritmo genético más de lo necesario, para reducir el costo computacional y el tiempo de procesamiento. Aunque algunos criterios de convergencia requieren saber de antemano el comportamiento del problema, otros criterios, como el número de generaciones y de estabilidad permiten explorar soluciones posibles de problemas de los que se desconoce o no se entiende del todo el mecanismo de acción o fenómeno estudiado.

¹⁶Del inglés: Stall time limit

Capítulo 4. Metodología

4.1. Introducción

En el presente capítulo se describe la metodología general de esta investigación, cuyo objetivo general es el reconocimiento de patrones en imágenes oftálmicas, específicamente de las lesiones conocidas como microaneurismas, puesto que son un signo de retinopatía diabética en su etapa no proliferativa, lo que permite tratar en tiempo la enfermedad y evitar complicaciones que deriven en la pérdida de la visión. El objetivo principal es el de mejorar la detección estas regiones de interés (ROI's) a través de una sintonización de parámetros.

En principio, se generó un banco de imágenes a partir de la base de datos pública DIARETDB1 (Kauppi *et al.*, 2007), con las cuales se realizaron diferentes experimentos. Se aplicó procesamiento digital de imágenes a cada una de las imágenes contenidas en el banco con el objetivo de realzar los objetos de interés (microaneurismas) y facilitar la detección de éstos.

En la etapa de procesamiento digital de imágenes, se encontró con el problema de la selección de los parámetros de procesamiento y se implementó un algoritmo genético para encontrar soluciones (parámetros) cuasi-óptimas; finalmente se comparó el desempeño hasta la etapa de generación de candidatos (evaluación de la segmentación) de las soluciones obtenidas contra soluciones heurísticas, y con la solución con los mejores resultados se procedió a extraer características (descriptores) y clasificar los objetos de interés en una de las dos clases posibles: microaneurisma y no microaneurisma; tal y como se observa en el esquema de la figura 43.

Los experimentos fueron realizados en el entorno de trabajo Spyder 2 con el lenguaje de programación Python 2.7.10.0 con las librerías Opencv, matplotlib, scikit-image, numpy; en una computadora Dell E6330 con procesador Intel Core i7-3520M @ 2.90 GHz con 4 Gb de memoria RAM utilizando el SO Windows 7. Posteriormente se configuró el mismo entorno de trabajo (Spyder 2) en la misma computadora pero con el SO Ubuntu 16.04 LTS, con lo que se logró disminuir el tiempo de procesamiento de los experimentos.

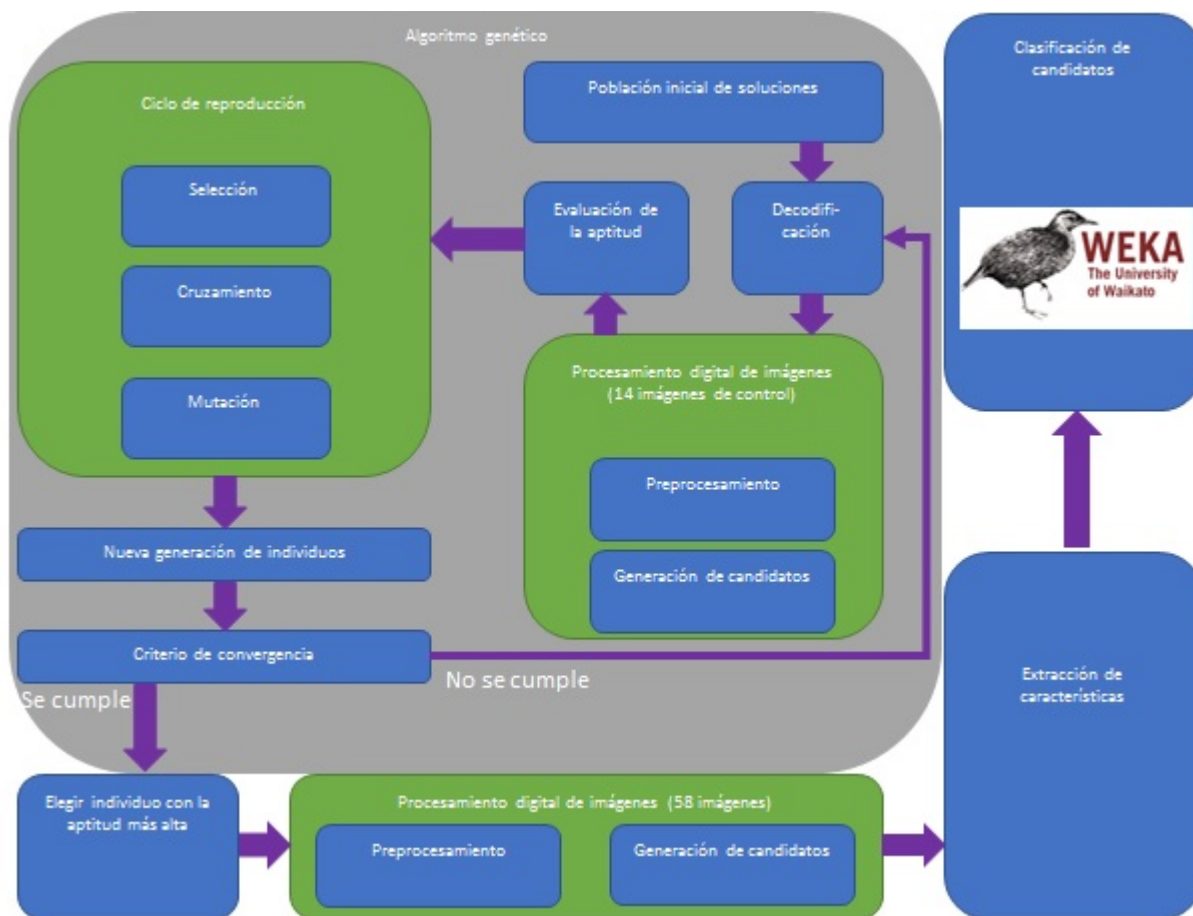


Figura 43: Esquema general del sistema propuesto.

4.2. Banco de imágenes

El banco de datos es un conjunto de datos sobre los que se prueban los procedimientos diseñados, en este caso son imágenes, de ahí el nombre “banco de imágenes”. El banco de imágenes usualmente se encuentra constituido por dos grupos: imágenes de entrenamiento e imágenes de prueba; todas son imágenes y en ocasiones las acompañan archivos “.xml”, “.doc” o incluso imágenes, en las que se brinda información de la localización de las lesiones que se considera como datos “ground truth”, lo que nos sirve para evaluar los procedimientos propuestos en la investigación.

Existen distintas bases de datos públicas referentes a imágenes oftálmicas, HRF¹⁷,

¹⁷HRF: High-Resolution Fundus Image Database

DRIVE¹⁸, MESSIDOR¹⁹, ROC²⁰, STARE²¹, DIARETDB²², entre otras. Cada una de las bases de datos antes mencionadas tiene como objetivo fungir como datos ground truth de diferentes tipos de estructuras en la retina; DRIVE y STARE se centran en evaluar la segmentación de vasos sanguíneos, HRF en evaluar la segmentación de vasos sanguíneos y disco óptico, ROC en evaluar la detección de microaneurismas y hemorragias, y DIARETDB en evaluar la detección de exudados suaves, exudados duros, hemorragias y microaneurismas.

DIARETDB cuenta con tres diferentes versiones: DIARETDB0, DIARETDB1 y DIARETDB1 V2.1. DIARETDB0 contiene 130 imágenes de fondo de ojo con sus respectivas anotaciones "ground truth"; los datos "ground truth" solamente indican si la imagen contiene microaneurismas, hemorragias, exudados suaves, exudados duros y neovascularizaciones, pero no indican el número de detecciones, su ubicación, ni su tamaño y forma. DIARETDB1, es la base de datos con la que optó a trabajar; contiene 89 imágenes y cada imagen tiene 4 imágenes ground truth (ver Figura 44), para exudados duros, hemorragias, microaneurismas y exudados suaves respectivamente. La base de datos DIARETDB1 V2.1 es la segunda versión de la anterior DIARETDB1, se utilizan las mismas 89 imágenes y en lugar de tener 4 distintas imágenes ground truth, se tienen 4 archivos en formato "xml". Cada archivo corresponde a la segmentación realizada por un médico oftalmólogo e indica el tipo de estructura (exudados duros, hemorragias, microaneurismas, exudados suaves, disco óptico, fovea), su forma, tamaño, ubicación (coordenadas espaciales x,y), y en el caso de las lesiones se indica la certeza (baja, media, alta) de la detección. De acuerdo al protocolo planteado por Kauppi *et al.* (2007), se considera certeza baja si el grado de confianza es menor a 50 %, certeza media si es mayor a 50 %, y certeza alta si la confianza es del 100 %.

El banco de imágenes se generó a partir de la base de datos pública DIARETDB1 y DIARETDB1 V2.1 utilizando 58 de las 89 imágenes disponibles.

¹⁸DRIVE: Digital Retinal Images for Vessel Extraction

¹⁹MESSIDOR: Methods to evaluate segmentation and indexing techniques in the field of retinal ophthalmology

²⁰ROC: Retinopathy Online Challenge

²¹STARE: STructured Analysis of the Retina

²²DIARETDB: Standard Diabetic Retinopathy Database

4.3. Procesamiento digital de imágenes oftálmicas

Para realizar el análisis de las imágenes oftálmicas del banco de imágenes, en primera instancia se recurrió a la etapa de procesamiento digital de imágenes, en la que se emplearon distintas técnicas con el objetivo de realzar y poner en evidencia las regiones u objetos a estudiar y analizar (regiones de interés), es decir los microaneurismas. En la figura 45, se observa el esquema que se adoptó para la implementación, los primeros tres bloques corresponden a la etapa de procesamiento digital de imágenes y los últimos tres bloques corresponden a la etapa de reconocimiento de patrones.

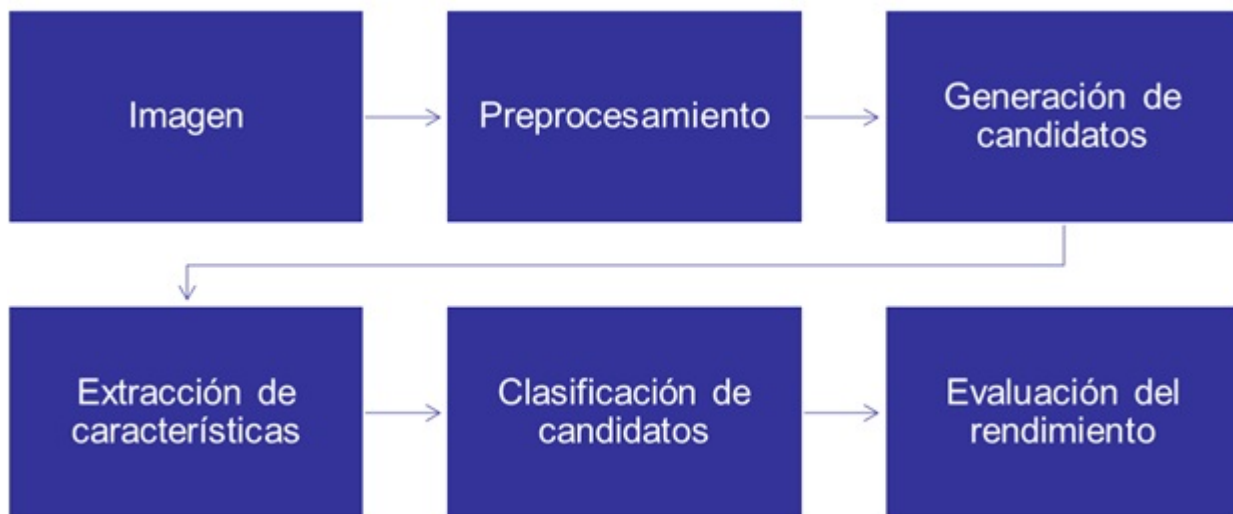


Figura 45: Esquema general de la detección de objetos.

Esta etapa está subdividida en dos conjuntos, donde el primer conjunto de técnicas es el de preprocesamiento, que se refiere a las técnicas que permiten mejorar la calidad de la imagen, comienza con la adquisición de la imagen, continúa con la selección de

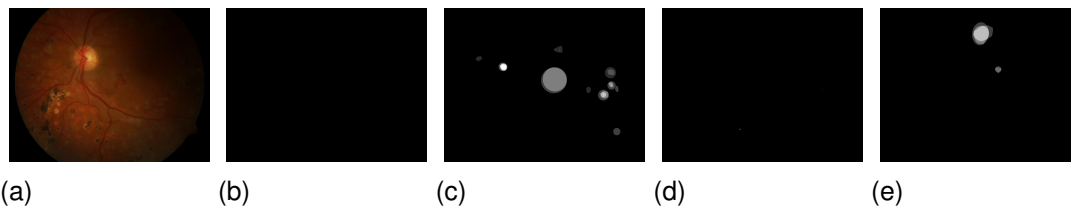


Figura 44: a) Imagen 27 (original), b) Imagen ground truth exudados suaves, c) Imagen ground truth hemorragias, d) Imagen ground truth microaneurismas, e) Imagen ground truth exudados duros.

canales para seguir con técnicas para la reducción de ruido y finaliza con la aplicación de técnicas para mejorar el contraste de la imagen. El segundo conjunto de técnicas es el de la generación de candidatos, cuyo objetivo es el de agrupar los píxeles de la imagen en regiones de acuerdo a la similitud de las características de los píxeles vecinos o de acuerdo a la presencia de estructuras o formas. Es en esta etapa donde se dejan en evidencia las regiones (objetos) de interés a través de imágenes binarias en las que se diferencian las regiones de interés del fondo.

4.3.1. Preprocesamiento

El conjunto de técnicas que componen la etapa de preprocesamiento tiene el objetivo de mejorar la calidad de la imagen, preparando la imagen para facilitar la extracción de las regiones de interés. Las etapas que comprende el preprocesamiento que se analizan en las siguientes subsecciones son las siguientes:

- la adquisición de las imágenes
- la separación de la imagen en sus respectivos canales (bandas espectrales) que la conforman para elegir el canal sobre el que se aplicarán las técnicas de realce
- la reducción de ruido para minimizar la cantidad de información indeseable
- la mejora de contraste para uniformizar el patrón de iluminación

4.3.1.1. Adquisición de imágenes

En este procedimiento se adquiere la imagen a la que se le aplican las técnicas para realzar y para aislar las regiones de interés. La imagen se toma del banco de imágenes que se generó, el cual se encuentra constituido por 58 imágenes de la base de datos DIARETDB1 V2.1. Las 58 imágenes se componen de 1500 x 1152 píxeles, y se tomaron con un campo de visión de 50 grados utilizando una cámara de fondo de ojo ZEISS FF 450 plus y una cámara digital Nikon F5 (Kauppi *et al.*, 2007).

4.3.1.2. Separación de canales

En el capítulo “Análisis de imágenes oftálmicas” se definió el concepto de imagen, y que éstas pueden estar constituidas por una o más bandas espectrales (canales) y que las técnicas para el mejoramiento de la calidad de las imágenes suelen operar sobre un único canal. Las imágenes de la base de datos DIARETDB1 V2.1 son imágenes de tres bandas, tienen un canal en el intervalo de los 0.45 m – 0.51 m, otro en el intervalo de los 0.51 m – 0.60 m, y un último canal en el intervalo de los 0.61 m – 0.69 m, correspondientes al canal azul, verde y rojo respectivamente.

El procedimiento de separación de canales es necesario entonces, para aplicar las técnicas para el mejoramiento de la imagen y la generación de candidatos, el canal elegido puede ser alguno de los que componen a la imagen, canal rojo, verde y azul. De igual manera, se pueden aplicar operadores de transformación a diferentes espacios de color, de manera que se puede trabajar sobre un nuevo canal resultado de la combinación de los tres canales originales.

Para el caso de las imágenes de la retina, trabajos como el de Sopharak *et al.* (2013), Walter *et al.* (2007), Niemeijer *et al.* (2005), por citar algunos, trabajan únicamente en el canal verde. La razón de aplicar técnicas sobre este canal, es debido a la naturaleza del ojo humano, en el cual las distintas estructuras son susceptibles a absorber, dejar penetrar o reflejar ciertas longitudes de onda (ver Figura 46). El canal azul suele ser utilizado para examinar las capas de fibras nerviosas, el canal rojo suele ser utilizado para estudio de las coroides y el canal verde para el estudio de la retina (Manzano, 2004).

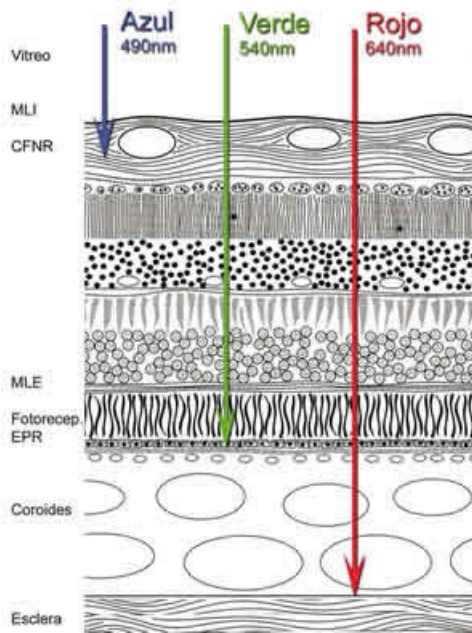


Figura 46: Penetración retiniana de tres longitudes de onda distintas. Imagen tomada de Manzano (2004).

Se descartó operar sobre canales de los espacios HSV, HSI y la propuesta de Romero *et al.* (2015) de aplicar el procesamiento sobre la relación del canal rojo sobre el verde, debido a que se observó que se disminuía la visibilidad a simple vista de algunos microaneurismas. Siguiendo el marco de trabajo de metodologías como la de Walter *et al.* (2007); Sopharak *et al.* (2013); Hipwell *et al.* (2000), se eligió el canal verde.

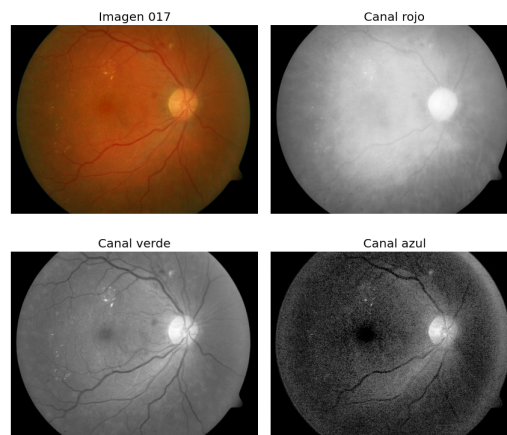


Figura 47: Canales RGB de la imagen 017.

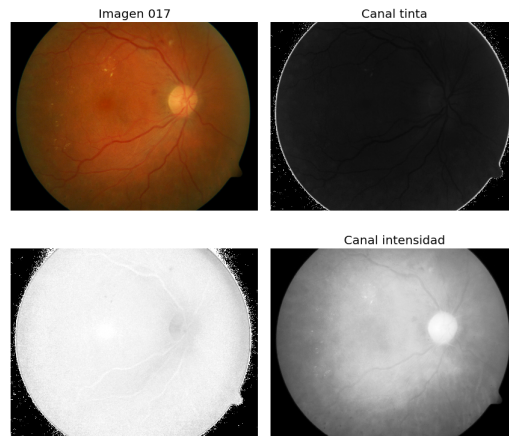


Figura 48: Canales HSV de la imagen 017.

4.3.1.3. Reducción de ruido

El objetivo de la presente etapa se implementó para reducir el ruido en la imagen (canal verde); para la reducción de ruido en imágenes oftálmicas típicamente se emplean filtros gaussianos y filtros de mediana (Romero *et al.*, 2015; Walter *et al.*, 2007; Sopharak *et al.*, 2013; Hipwell *et al.*, 2000). El objetivo es suavizar la imagen para disminuir grandes variaciones que puedan existir entre vecindades de píxeles, como por ejemplo ruido de tipo sal y pimienta y/o speckle. La recomendación es que el tamaño de los filtros no sobrepase el tamaño de la región de interés para evitar disminuir la calidad y el grado de detalle del mismo.



Figura 49: Ejemplo de reducción de ruido.

4.3.1.4. Mejora de contraste

El objetivo de la etapa de mejora de contraste es el de esparcir la intensidad de los píxeles sobre el rango dinámico entero, de manera que se pueda apreciar una variabilidad de la intensidad, i.e., se mejora el contraste. La mejora del contraste contribuye a disminuir la iluminación no uniforme.

Para la mejora de contraste se empleó la ecualización adaptativa del histograma limitada por contraste (CLAHE), una técnica de ecualización adaptativa comúnmente utilizada en el procesamiento de imágenes médicas, tales como ultrasonografías, angiografías o radiografías.



Figura 50: Ejemplo de CLAHE.

4.3.2. Generación de regiones candidatas a MA

La etapa de generación de candidatos comienza una vez que se han aplicado las técnicas de corrección en la imagen original, es decir, una vez que ha sido realizada. Para lograr dejar las regiones de interés en evidencia, primeramente se realiza una segmentación y enseguida se le aplica el proceso de umbralización a la misma para obtener imágenes binarias, donde las regiones de interés serán las correspondientes a los “unos lógicos”, o “1’s”.

4.3.2.1. Segmentación

Para la obtención de las regiones de interés, se procedió a segmentar las imágenes utilizando la adaptación de los operadores morfológicos para niveles de gris. Estos operadores han sido previamente utilizados en la identificación de lesiones en ultrasonografías,

posteriormente en angiografías, y en imágenes de fondo de ojo. Es de las técnicas más comunes en el estado del arte y numerosas metodologías han sido propuestas utilizando estos operadores como mecanismo de segmentación.

La segmentación se basa en la operación de cierre utilizando un elemento estructural circular de radio sintonizable. Al resultado del cierre se le resta la imagen obtenida en la fase previa, es decir, se obtiene la transformada bottom-hat. La operación de cierre elimina todos los valles que son más estrechos que el elemento estructural, de manera que se obtiene una imagen en donde las componentes oscuras han sido eliminadas, a esta imagen se le resta la imagen de la etapa previa, de manera que queda una imagen donde se resaltan las regiones más oscuras de esa imagen, generando así los candidatos a microaneurismas, que son de las regiones más oscuras en una imagen de fondo de ojo.

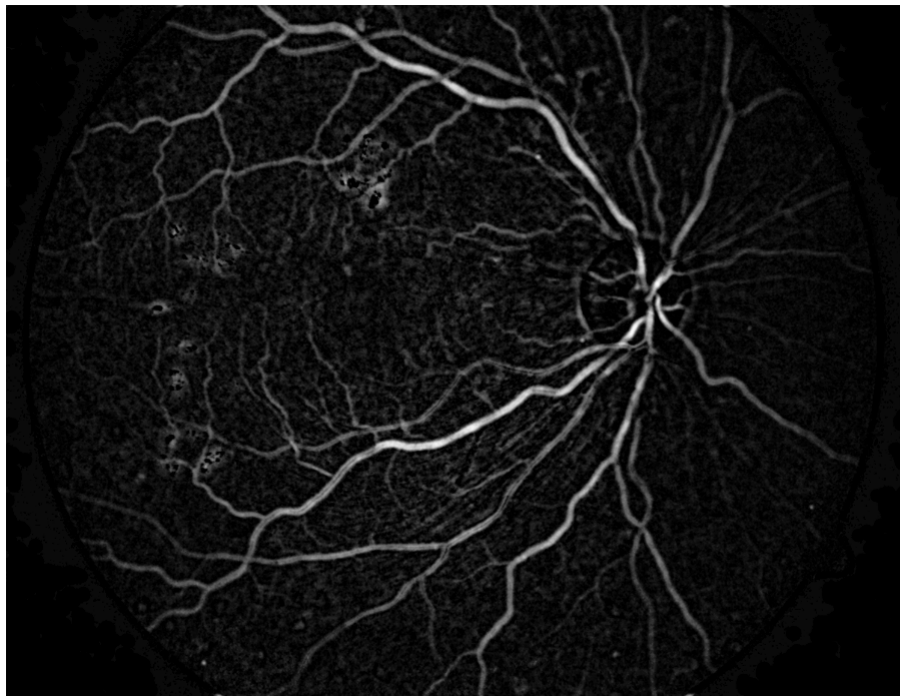


Figura 51: Segmentación mediante la transformada de bottom-hat.

4.3.2.2. Umbralización

En la etapa anterior se generan los candidatos, sin embargo sigue siendo una imagen en niveles de gris. El próximo paso es generar una imagen binaria donde destaquen los

candidatos del resto de la imagen; para esto se aplica una operación de umbralización.

Debido a la variabilidad de una imagen a otra (luminosidad), sin aplicar técnicas de normalización, y aun aplicándolas, no resulta factible establecer un umbral fijo. Se propuso emplear una técnica de umbralización adaptativa (método de OTSU), que en base a técnicas estadísticas determina el mejor umbral para separar la imagen en los dos segmentos deseados.

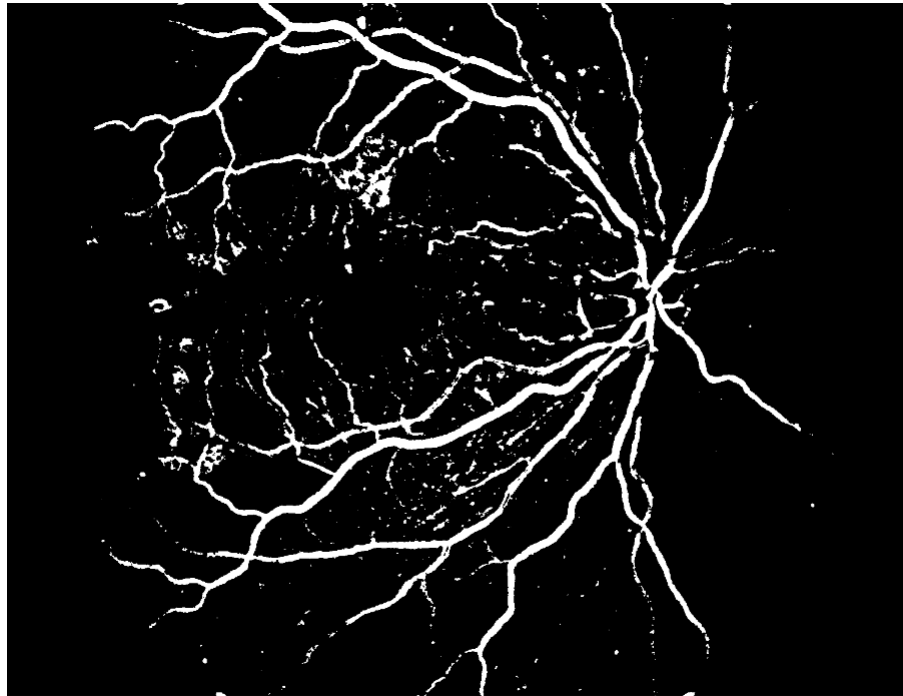


Figura 52: Umbralización de la imagen resultado de la transformada de bottom-hat.

4.4. Búsqueda de soluciones cuasi-óptimas

Una vez que se integraron las diferentes técnicas de preprocesamiento y el método de generación de regiones candidatas, se procedió a procesar cada una de las imágenes contenidas en el banco con parámetros reportados en la literatura. En este punto del trabajo de investigación surge la problemática de encontrar los parámetros de procesamiento que en su conjunto cumplan el objetivo de realzar las regiones de interés y maximizar la cantidad de microaneurismas presentes en los candidatos obtenidos, resultado de la segmentación y umbralizado. Se observó que mientras que para algunas imágenes ciertos parámetros generaban dentro de los candidatos la mayoría de los mi-

croaneurismas de acuerdo a las anotaciones (datos ground truth), para otras sucedía todo lo contrario, y en el peor de los casos se perdían todos los microaneurismas.

Se comenzó a buscar los parámetros óptimos intuitivamente, para lo cual se propuso una solución y se evaluó. El método de evaluación propuesto ha sido utilizado en trabajos previos, en los que evalúan la segmentación de imágenes de ultrasonido para la detección de cáncer de mama (Adán, 2012). Para evaluar la segmentación se tomaron 10 imágenes del banco de datos (58 imágenes de la base de datos DIARETDB1) y a partir de sus datos "ground truth" se generaron 10 plantillas para ser utilizadas como imágenes para el control de la segmentación. El número total de MA's existentes en las 10 imágenes es 29.

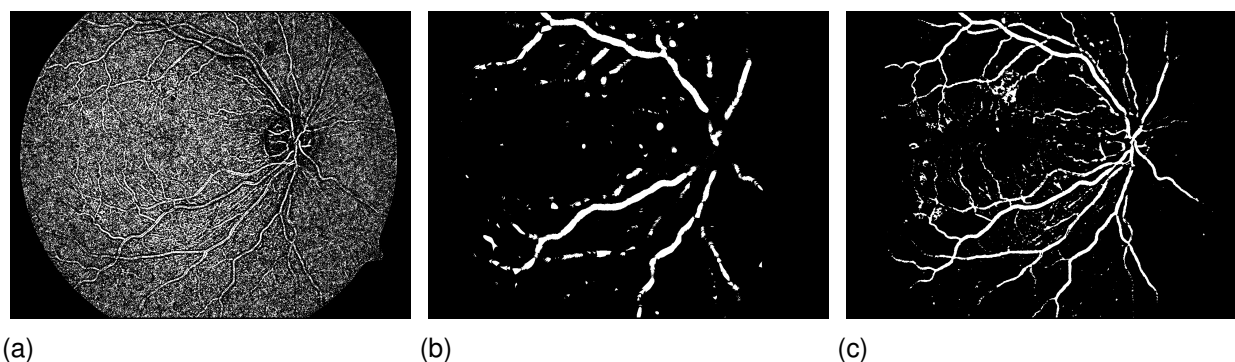


Figura 53: Importancia de la optimización de parámetros: a) Imagen 17 sin preprocesamiento, b) Imagen 17 procesada con parámetros sub-óptimos, c) Imagen 17 procesada con parámetros cuasi-óptimos.

El método de evaluación consiste en encerrar dentro de una ventana las regiones que son un microaneurisma en la imagen ground truth, y en sobreponer cada una de las ventanas en la imagen procesada. Enseguida se comparan las regiones que quedan encerradas en la misma ventana en ambas imágenes, que son el objeto real y la detección (Figura 54 a)), las métricas de evaluación (Figura 54 b)) nos indican el porcentaje de área detectado correctamente (verdadero positivo), el porcentaje de área no detectado (falso negativo), el porcentaje de área incorrectamente detectado (falso positivo) y el porcentaje de área correctamente no detectado (verdadero negativo).

A partir de esta primera evaluación los resultados obtenidos no fueron favorables debido a la gran cantidad de microaneurismas no detectados; de manera que se optó por la implementación de un algoritmo genético bajo la premisa de que puede encontrar soluciones no óptimas pero sí muy cerca de éstas. Aunque existen diferentes técnicas

de búsqueda, se optó por la utilización de un algoritmo genético porque a diferencia de otras técnicas, debido a sus operadores y su aleatoriedad, son capaces de explorar rápidamente el paisaje de búsqueda y son menos susceptibles a quedar atrapados en óptimos locales. Pueden explorar múltiples zonas de óptimos locales al mismo tiempo, de manera que son capaces de encontrar soluciones aceptables, aun cuando no se tiene información del problema ni de cómo resolverlo (Sivanandam y Deepa, 2008).

4.4.1. Propuesta de algoritmo genético

Se implementó un algoritmo genético a partir del esquema del algoritmo genético simple, el cual se muestra en la figura 55. La adaptación del algoritmo genético que se implementó, se programó en el entorno de programación “Python XY 2.7.10.0”, y se encuentra compuesto por una función para cada etapa del AG.

En las siguientes subsecciones se describirá la propuesta para cada bloque en mayor detalle y por el momento se explicará el esquema general. Se tiene el problema de evaluar qué tan exacta es la generación de candidatos (segmentación) respecto a las imágenes procesadas y los datos ground truth. Anteriormente se describieron las técnicas utilizadas para el procesamiento digital de imágenes, mismas que se encuentran implementadas en una misma función, y cada valor posible de los parámetros de procesamiento constituye una solución; el objetivo es evaluar el desempeño de cada solución para encontrar la que mejor se adapta a nuestro banco de imágenes, de manera que aumente la capacidad

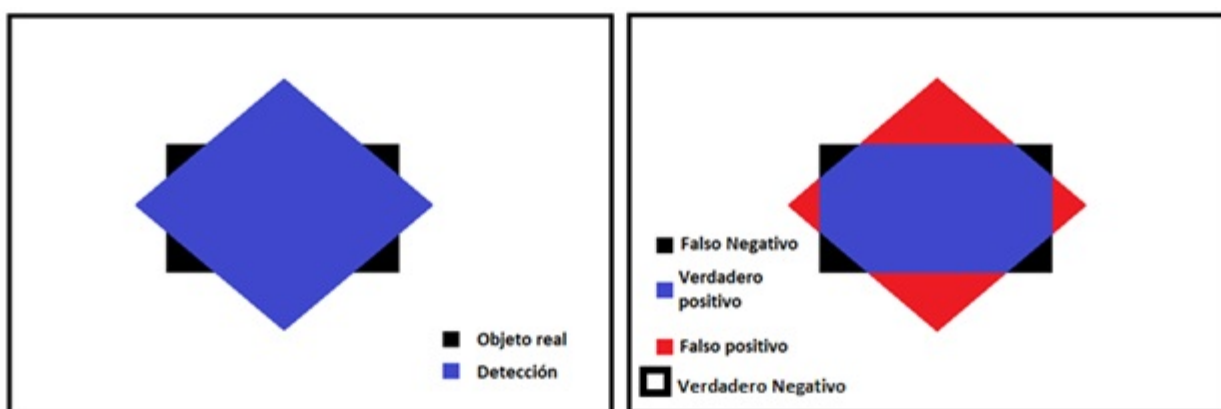


Figura 54: Evaluación de la generación de candidatos: a) Objeto a detectar (cuadro negro) y detección (azul), b) Métricas de evaluación de la segmentación.

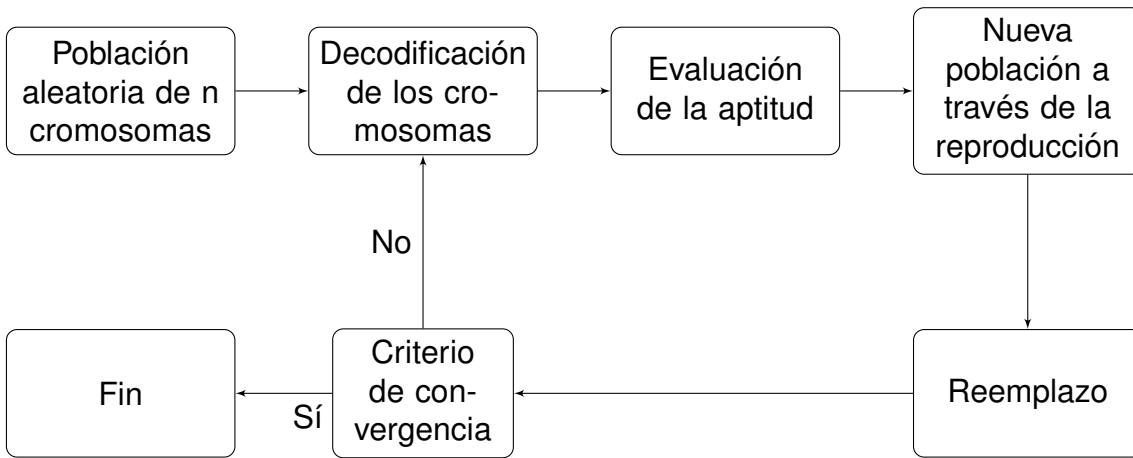


Figura 55: Esquema general del algoritmo genético simple.

de detección de microaneurismas. El primer paso es generar una población aleatoria de “n” soluciones, en adelante, individuos, los cuales representan una solución (parámetros para el procesamiento digital de imágenes); el segundo paso es ingresar “n” número de individuos manualmente a la población, a manera de guiar el algoritmo hacia zonas del espacio de búsqueda de las que heurísticamente creemos se encuentran cerca de óptimos locales o incluso del óptimo global. A esta primera población se le denomina la *generación 1*, la cual evoluciona hacia una nueva generación por el resultado de aplicar los operadores genéticos de la reproducción (Selección, cruzamiento, mutación) una vez que sea determinado qué tan apto es cada individuo, y de reemplazar la generación anterior por los nuevos individuos generados, y así sucesivamente para crear la generación 2, 3, ..., n, como se muestra en la ecuación 41.

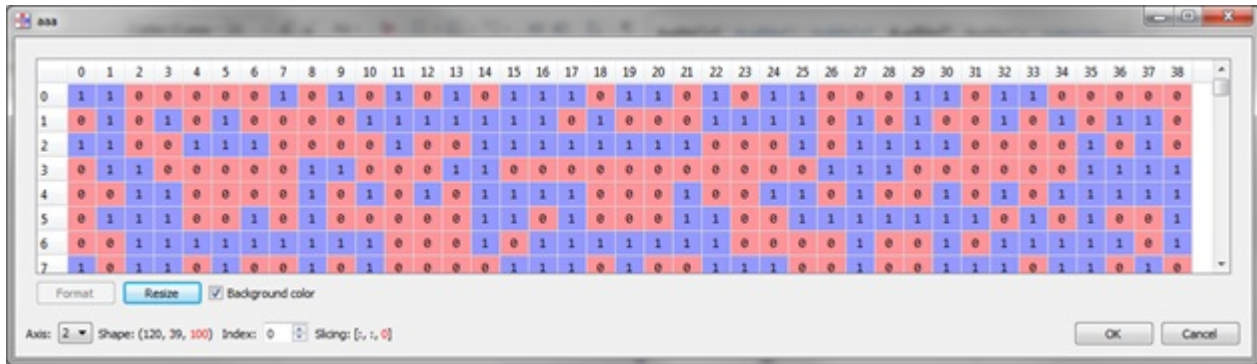
$$gen(num_gen) = \begin{cases} población = random(genotipo, n) & si \quad num_gen = 1 \\ población = OG \{gen(num_gen - 1)\} & si \quad num_gen > 1 \end{cases} \quad (41)$$

Donde: *OG* = Operadores Genéticos

En la figura 56 b) se muestra la matriz multidimensional del fenotipo de cada individuo, cada renglón representa un individuo diferente dentro de la población. Las columnas representan un parámetro diferente, la columna 0 es el parámetro de entrada para el filtro gaussiano, la columna 1 el parámetro de entrada para el filtro mediano, la columna 2 y 3

para el algoritmo CLAHE, y finalmente la columna 4 el tamaño del elemento estructural; por otra parte, las columnas 5 y 6 corresponden a los valores de la aptitud de cada individuo, con la columna 5 para el esquema en el que se filtra el ruido con el filtro gaussiano y la columna 6 para el esquema que involucra el filtro mediano.

(a)



(b)

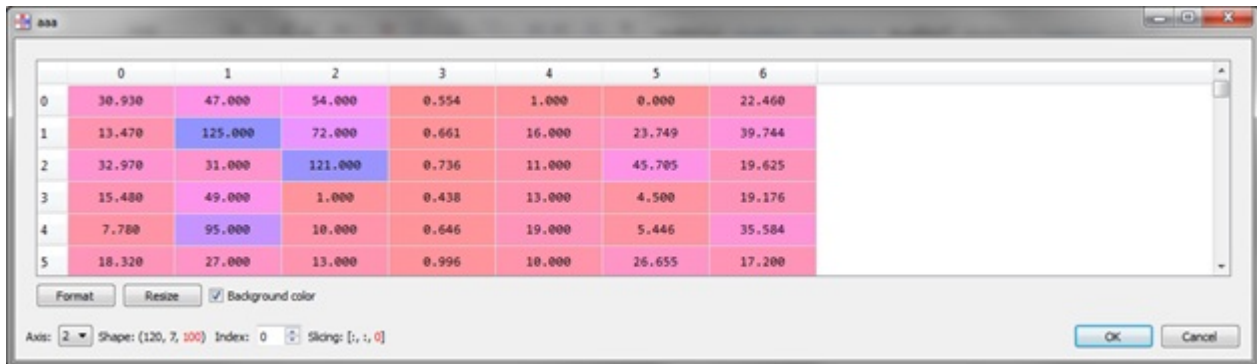


Figura 56: a) Matriz multidimensional genotipo, b) Matriz multidimensional fenotipo.

4.4.1.1. Codificación (Representación)

La etapa de codificación o representación es una de las etapas de mayor importancia, en la cual establecemos el método de codificar los parámetros del problema a optimizar (parámetros de PDI) y su intervalo de valores, ya sean números reales, enteros, naturales; valores de tipo lógico o caracteres.

En la figura 57 se muestra la estructura del cromosoma propuesto. Se optó por utilizar una codificación de tipo binaria, de manera que cada posición en el cromosoma puede tomar dos valores diferentes (1's ó 0's lógicos). Los parámetros del filtro gaussiano y el

“clip limit” toman valores reales y son continuos, para lo que se estableció un rango con límite inferior y límite superior de acuerdo a los valores de la literatura más un margen para permitir explorar una mayor cantidad de soluciones.

Filtro Gaussiano												Filtro mediano						Ventana_CLAHE							Clip_Limit									Radio E. Estructural				
1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31	32	33	34	35	36	37	38	39

Figura 57: Cromosoma propuesto.

Para el filtro gaussiano se consideró un límite inferior igual a 0, un límite superior igual a 40.95 y 12 bits, de manera que se tiene la función de mapeo descrita en la ecuación 42:

$$fenotipo_{f_{gaussiano}} = \frac{40.95 \times [bin_a_dec(gen_filtro_gaussiano)]}{2^{12} - 1} \quad (42)$$

Donde *gen_filtro_gaussiano* es una cadena binaria, es decir el gen correspondiente al filtro gaussiano, y para obtener el valor del parámetro (fenotipo) se convierte la cadena a decimal, se multiplica por 40.95 y se divide entre el número de combinaciones que se pueden obtener con 12 bits, de manera que el valor de entrada al filtro gaussiano se encuentra entre 0 y 40.95 con aumentos de 0.01.

Para el parámetro “clip limit” se utilizan 9 bits, un límite inferior igual a 0 y un límite superior igual a 1, de manera que se tiene la función de mapeo descrita en la ecuación 43:

$$fenotipo_{clip_limit} = \frac{1 \times [bin_a_dec(gen_clip_limit)]}{2^9 - 1} \quad (43)$$

Donde *gen_clip_limit* es la cadena binaria correspondiente al parámetro “clip limit”, y para obtener su fenotipo se convierte la cadena a decimal, se multiplica por 1 y se divide entre el número de combinaciones que se pueden obtener con 9 bits, de manera que el valor que puede tomar esté parámetro se encuentra entre 0 y 1 con aumentos de 0.0019. Los otros tres parámetros, del filtro mediano, de la ventana del algoritmo CLAHE y del radio del elemento estructural para el operador morfológico toman valores pertenecientes

a los números naturales. El parámetro para el tamaño de la ventana del filtro mediano toma únicamente valores impares, de manera que se construyó un diccionario o “Look Up Table” en el que se definen cada una de las 64 posibles combinaciones que se pueden obtener con una cadena de 6 bits (ver Figura 58).

```
dict_f_med = {"000000" : 1, "000001" : 3, "000010" : 5, "000011":7,
              "000100" : 9, "000101" : 11, "000110" : 13, "000111":15,
              "001000" : 17, "001001" : 19, "001010" : 21, "001011":23,
              "001100" : 25, "001101" : 27, "001110" : 29, "001111":31,
              "010000" : 33, "010001" : 35, "010010" : 37, "010011":39,
              "010100" : 41, "010101" : 43, "010110" : 45, "010111":47,
              "011000" : 49, "011001" : 51, "011010" : 53, "011011":55,
              "011100" : 57, "011101" : 59, "011110" : 61, "011111":63,
              "100000" : 65, "100001" : 67, "100010" : 69, "100011":71,
              "100100" : 73, "100101" : 75, "100110" : 77, "100111":79,
              "101000" : 81, "101001" : 83, "101010" : 85, "101011":87,
              "101100" : 89, "101101" : 91, "101110" : 93, "101111":95,
              "110000" : 97, "110001" : 99, "110010" : 101, "110011":103,
              "110100" : 105, "110101" : 107, "110110" : 109, "110111":111,
              "111000" : 113, "111001" : 115, "111010" : 117, "111011":119,
              "111100" : 121, "111101" : 123, "111110" : 125, "111111":127}
```

Figura 58: Diccionario para el tamaño de la ventana del filtro mediano.

El radio del elemento estructural se codifica en una cadena de 5 bits, y al igual que con el parámetro del filtro mediano, se estableció un diccionario para mapear cada una de las posibles combinaciones a un valor de radio (ver Figura 59); el intervalo se estableció de 1 a 19 y algunos valores aparecen más de una vez para favorecer la probabilidad de obtener esa característica (alelo). Para la ventana del algoritmo CLAHE, su función de mapeo consiste en transformar la cadena de 7 bits a decimal y sumarle 1, de manera que los valores posibles se encuentran en el rango de 1 a 128.

```
dict_el_est = {"00000" : 1, "00001" : 2, "00010" : 3, "00011":4,
               "00100" : 5, "00101" : 6, "00110" : 7, "00111":8,
               "01000" : 9, "01001" : 10, "01010" : 11, "01011":11,
               "01100" : 12, "01101" : 12, "01110" : 13, "01111":13,
               "10000" : 14, "10001" : 14, "10010" : 15, "10011":15,
               "10100" : 16, "10101" : 16, "10110" : 16, "10111":17,
               "11000" : 17, "11001" : 17, "11010" : 17, "11011":18,
               "11100" : 18, "11101" : 19, "11110" : 19, "11111":19}
```

Figura 59: Diccionario para el elemento estructural.

4.4.1.2. Selección

En esta etapa del algoritmo genético se determinan los individuos a los que se les permitirá reproducirse para crear una nueva generación. El método de selección desempeña un papel crucial en la aleatoriedad del algoritmo y por ende en la capacidad de exploración del espacio de búsqueda; elegir un método que tiende a favorecer los individuos mejores adaptados puede conducir a la convergencia prematura debido a un estancamiento en óptimos locales. De igual manera no se deben favorecer individuos con baja aptitud puesto que nunca se llegaría a encontrar una solución aceptable.

El método de selección implementado en el algoritmo desarrollado es el método selección por torneo, cuyo funcionamiento se basa en generar un subgrupo del total de la población de manera aleatoria, y elegir el individuo más apto para ser padre. La ventaja de este método de selección es que el número de individuos que entran al torneo (número de individuos en cada subgrupo) permite ajustar qué tanto se favorece a los individuos más aptos respecto a los menos aptos. En el algoritmo genético desarrollado, se incorpora un factor de elitismo, el cual nos permite conservar a los mejores individuos de la generación actual e incluirlos directamente en la siguiente generación sin aplicarles ningún operador genético. De esta manera, la nueva población se constituye de acuerdo a la ecuación 44. El incorporar elitismo en la propuesta garantiza que al menos se tendrá un cierto número de soluciones aceptables en la siguiente generación.

$$Nueva_población = \begin{cases} 90 \% \text{ hijos} + 10 \% \text{ elitismo} & \text{si } num_individuos = \text{par} \\ 87.1 \% \text{ hijos} + 12.8 \% \text{ elitismo} & \text{si } num_individuos = \text{impar} \end{cases} \quad (44)$$

Para obtener los hijos, se eligen los padres sin reemplazo, de manera que una vez que un individuo ha sido elegido para ser padre, no podrá ser nuevamente padre. Los individuos que ya fueron seleccionados como padres se almacenan en un vector de restricciones que les impiden volver a ser elegidos como padres. Además, en el algoritmo se incorporó un factor que determina si el padre elegido de cada subgrupo es la mejor solución o la segunda mejor solución, con lo que se espera favorecer a soluciones no tan

bien adaptadas; este procedimiento se realiza generando un número aleatorio entre 0 y 1, se fija el valor del factor (presión de selección) y si el número aleatorio es menor que el factor, el padre será el individuo más apto, mientras que si es mayor el padre será el segundo individuo más apto. Cuando el factor es igual a 1, el individuo más apto siempre será el padre y cuando es menor a 0.5 privilegia al segundo más apto.

4.4.1.3. Cruzamiento

La etapa de cruzamiento es en la que se genera la nueva población, a partir de cruzar los individuos seleccionados como padres para generar nuevos individuos (hijos), de esta manera de conforma la descendencia. Los tipos de cruzamiento que se incorporaron en el algoritmo genético desarrollado, son en un punto y en dos puntos.

Para el cruzamiento en un punto (ver Figura 60), se selecciona aleatoriamente de la cadena binaria una de las 39 posiciones, en la posición seleccionada se establece una división de manera que la cadena queda dividida en dos segmentos, y los hijos se crean a partir de compartir información genética de ambos padres; para el primer hijo se toma el primer segmento del padre número uno y el segundo segmento de padre número dos, mientras que para el segundo hijo, se toma el segundo segmento del padre número uno y el primer segmento del padre número dos.

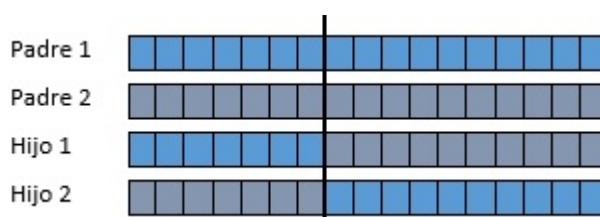


Figura 60: Cruzamiento en un punto.

A diferencia del cruzamiento en un punto, el de dos puntos (ver Figura 61) selecciona dos diferentes posiciones en la cadena de manera aleatoria, las dos posiciones dividen en tres segmentos a ambos padres, y los hijos son el resultado de cruzar los segmentos de los extremos del padre número uno y el segmento central del padre número dos para el hijo uno, mientras que para el hijo número dos es lo contrario, de manera que es el resultado de seleccionar el segmento central del padre número uno y los segmentos extremos del padre número dos.

El procedimiento propuesto es el siguiente, la etapa de selección genera el padre número uno y el padre número dos, se selecciona aleatoriamente un valor de tipo lógico “Sí” o “No” con probabilidad igual a 0.7 para el “Sí” y el complemento de 0.3 para el “No”. A la probabilidad de 0.7 se le conoce como la probabilidad de cruzamiento, la cual es otro de los factores en el algoritmo genético que nos permite tener mayor o menor aleatoriedad y por ende diversidad genética, explorando así y ampliando el paisaje de búsqueda. La probabilidad de cruzamiento fue seleccionada de 0.7 de acuerdo a la recomendación de Mitchell (1999). Cuando el valor seleccionado es “No”, los padres no son cruzados, es decir, los hijos son copias exactas de sus respectivos padres, en caso de que el valor seleccionado sea “Sí”, se selecciona aleatoriamente uno de los dos cruzamientos, en un punto o en dos puntos, con igual probabilidad, y el cruzamiento seleccionado será la forma en la que se generarán los hijos; en este punto termina el algoritmo de cruzamiento y los hijos entran al siguiente proceso, el de mutación.

4.4.1.4. Mutación

El operador genético conocido como mutación, es un operador que se incorpora con el objetivo de mantener la diversidad genética; básicamente consiste en cambiar valores en las posiciones de la cadena (cromosoma) con cierta probabilidad denominada probabilidad de mutación. Se toma una decisión pesada para seleccionar un valor de tipo lógico “Sí” o “No”, donde el “Sí” tiene una probabilidad de 0.0256 y 0.9743 para el “No”. La probabilidad de mutación es entonces la probabilidad de obtener un “Sí” para el valor lógico. Los valores se eligieron de acuerdo a las recomendaciones en la literatura sobre algoritmos genéticos (Mitchell, 1999; Sivanandam y Deepa, 2008) en donde se señala que un valor de probabilidad de mutación recomendado es igual a 0.1 o bien el resultado de dividir la máxima probabilidad (1) entre el número de elementos que componen

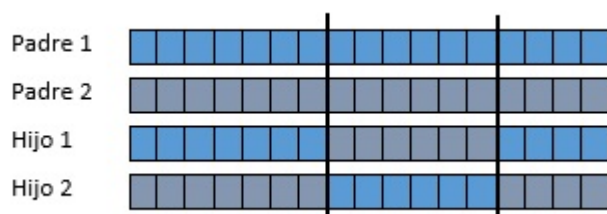


Figura 61: Cruzamiento en dos puntos.

la cadena (cromosoma) resultando en las probabilidades previamente mencionadas para el “Sí” y el “No”. Se elige ese respectivo valor para la probabilidad de mutación bajo el argumento de que es un valor adecuado para una primera ejecución del algoritmo, y que al igual que la probabilidad de cruzamiento, ésta debe ser ajustada de acuerdo al comportamiento observado, si es que hay poca diversidad genética, si tarda mucho en converger o si su convergencia es prematura. La mutación evita que el algoritmo quede atrapado en óptimos (mínimos o máximos) locales; incrementar esta probabilidad aumenta la diversidad y disminuirla ralentiza la evolución de los individuos. Al utilizar una codificación binaria, mutar se traduce en cambiar un “uno” por un “cero” y viceversa. Los métodos de mutación que fueron incorporados al algoritmo son: “flipping”, “interchanging” y “reversing”. El método “flipping” se puede visualizar en la figura 62, consiste en recorrer cada posición de la cadena (cromosoma) y tomar una decisión pesada para determinar si ese valor es mutado o no, de manera que se crea una cadena (cromosoma) de mutación, si el resultado de la decisión es sí, el valor en la posición en la que se encuentra es modificado, un “uno” es cambiado por un “cero” y viceversa. Para la implementación, la probabilidad de mutación se ha establecido de acuerdo a lo recomendado por Mitchell (1999), por lo que se le ha asignado un valor igual a $1/39$, de manera que en promedio para cada individuo se suscitará mutación de tipo “flipping” en una de las 39 posiciones de la cadena.

Hijo antes de la mutación	1 0 0 0 1 1 1 0 0 1 0 1 1 0 0 1 1
Cadena de mutación (0:no, 1:sí)	1 1 1 0 0 0 1 0 0 0 1 0 1 0 0 1 0
Hijo después de la mutación	0 1 1 0 1 1 0 0 0 1 1 1 0 0 0 0 1

Figura 62: Mutación de tipo “flipping”.

Para el método “interchanging” , se seleccionan dos posiciones distintas aleatoriamente de la cadena (cromosoma), y la información contenida en ambas posiciones es intercambiada. En la figura 63 se puede observar un ejemplo gráfico de este tipo de mutación. Para este tipo de mutación, se estableció una probabilidad de mutación igual a $10/39$, de manera que habrá mutación una vez por cada 3.9 veces que este método de mutación sea seleccionado.

El método “reversing” consiste en seleccionar aleatoriamente una posición de la cadena (cromosoma) e invertir la sub-cadena de la posición seleccionada hacia el final de la cadena (cromosoma). En la figura 64 se muestra una ejemplificación de este método de mutación. Al igual que para el método de “interchanging”, la probabilidad de mutación se estableció igual a 10/39, de manera que el 25 % de las veces que este método de mutación sea seleccionado, se presentarán alteraciones en el hijo engendrado en la fase previa (cruzamiento).

Finalmente, únicamente se aplica un método de mutación a cada hijo; el método de mutación se elige mediante una decisión pesada (equiprobable) en la que cada uno de los tres métodos tiene igual probabilidad de ser seleccionada.

4.4.1.5. Reemplazo

La etapa de reemplazo es la última en el ciclo de reproducción y tiene el objetivo de crear nuevas generaciones de individuos. En la adaptación que se desarrolló del algoritmo genético simple, se empleó una combinación del esquema de reemplazo de ambos padres hasta integrar un 87.1 % o 90 % de la población de acuerdo a la ecuación 44, y los porcentajes restantes integran a los mejores individuos de la generación. Se implementó la etapa de reemplazo de ésta manera para que en caso de que los nuevos

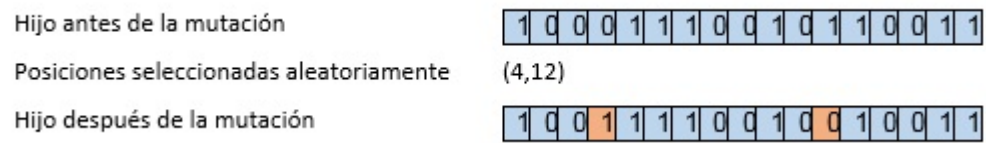


Figura 63: Mutación de tipo “interchanging”.

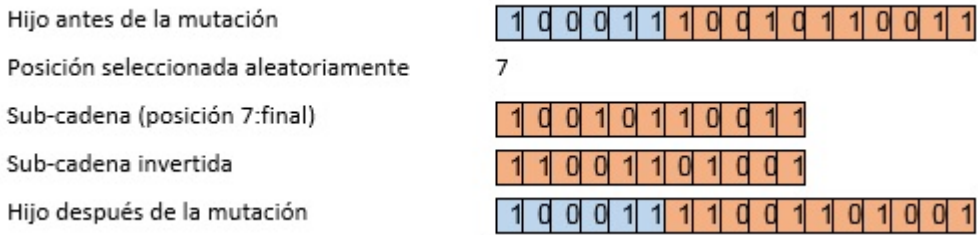


Figura 64: Mutación de tipo “reversing”.

individuos resultado de aplicar los operadores genéticos (hijos) resulten ser menos aptos que los padres, al menos asegurar conservar un porcentaje de individuos aceptables, permitiendo que sus genes estén disponibles para ser usados nuevamente en otro ciclo de reproducción para engendrar una nueva generación.

4.4.1.6. Convergencia

El criterio de convergencia que se empleó en la adaptación del algoritmo genético simple, fue de acuerdo al establecimiento de un número máximo de generaciones. Para las primeras aproximaciones, el número máximo de generaciones fue establecido a 90. En estas primeras aproximaciones se tenía el desconocimiento de qué tan bien trabajaría el algoritmo genético para este problema en particular, además de acuerdo a los resultados obtenidos se fueron ajustando algunos de los parámetros de la adaptación del AG simple, como las probabilidades de mutación, de cruzamiento, los porcentajes de hijos engendrados y elitismo, el número de generaciones, y sobre todo de la función objetivo.

El número de generaciones se fue disminuyendo como consecuencia de observar en las primeras ejecuciones del algoritmo, que este convergía alrededor de la generación 25, debido a esto se optó por seguir experimentando con diferentes parámetros, para lo que se favoreció la aleatoriedad del algoritmo y se redujo el número máximo de generaciones a 70 y posteriormente a 50 generaciones.

4.4.2. Función objetivo

La función objetivo o función de aptitud, es la función que se desea maximizar o minimizar con el objetivo de alcanzar el máximo o el mínimo global según el planteamiento del problema. En el caso de esta implementación, se busca maximizar la función, la cual inicialmente se propuso como una evaluación de la segmentación de acuerdo a las métricas observadas en la figura 54.

Se debe tener en cuenta, que para poder evaluar la segmentación, inicialmente, para los experimentos en los que se analizaron distintas soluciones de forma heurística y la primera ejecución del algoritmo genético (experimento 1 y 2), se propusieron 10 imágenes de control que contenían 29 lesiones en total. Sin embargo, para la evaluación de la

segmentación resultante de cada individuo (implementación en adaptación del algoritmo genético simple) se utilizaron las anotaciones de la base de datos DIARETDB1 V2.1. Primeramente se concentró la información de los 4 oftalmólogos en un solo documento, conservando únicamente las anotaciones referentes a microaneurismas, para cada imagen. A partir de los archivos generados, se procedió a la creación de 14 plantillas, es decir, 14 nuevas imágenes de control que en total contienen 29 lesiones. Las 14 imágenes se escogieron cuidadosamente, de manera que para cada lesión se tuviera una certeza media o alta de acuerdo al oftalmólogo que la identificó, descartando así lesiones con certeza baja. Lo anterior, se realizó pensando en que una lesión con certeza baja puede no ser realmente una lesión, y que con dicha restricción se disminuye la confusión en la selección de parámetros de procesamiento digital de imágenes y se guía así la evolución de los individuos hacia parámetros que resalten las lesiones de las que se tiene una mayor certeza.

Se condujeron múltiples experimentos en los que el principal parámetro del algoritmo genético que se modificó fue la función objetivo, para encontrar la que mejor adapta a las soluciones hacia una mejor segmentación, de entre las que se utilizaron se encuentran: la sensibilidad, la precisión, la exactitud, y una ponderación de verdaderos positivos y verdaderos negativos; y se utilizó cada una de ellas para ambos esquemas (reducción de ruido con filtro gaussiano y reducción de ruido con filtro mediano).

4.4.3. Clasificación

La clasificación de regiones candidatas se realizó en el software Weka 3.6.13. Después de ejecutar múltiples veces el algoritmo genético propuesto, con diferentes parámetros (probabilidad de mutación, probabilidad de cruzamiento, porcentaje de elitismo, función objetivo, número de generaciones) se seleccionó el esquema en el que después de la generación de candidatos persistió el mayor número de microaneurismas reales con respecto a los datos “ground truth”.

Con ese conjunto de parámetros se procedió a procesar cada una de las imágenes del banco de datos, se procedió a extraer características de cada candidato, manualmente se marcó cada uno de los candidatos como “microaneurisma” y “no microaneurisma”, y

se procedió a crear un archivo “.arff” desde Python 2.7.10.0 para poder ser importado en Weka 3.6.13.

4.4.4. Extracción de características (descriptores)

En la sección anterior se especificó que el algoritmo genético fue utilizado múltiples veces con diferentes parámetros. Se evaluaron las posibles soluciones utilizando como métrica el número de lesiones “pre-detectadas”. La solución que se eligió generó como candidatos a microaneurismas 207 de las 246 lesiones reales.

Se procesaron todas las imágenes con los respectivos parámetros de procesamiento de esa solución, y a continuación se utilizó un algoritmo de etiquetaje para identificar el número de regiones candidatas.

De cada región se procedió a extraer 27 descriptores. Los descriptores fueron elegidos de acuerdo a los propuestos por Sopharak *et al.* (2013); Niemeijer *et al.* (2005); Walter *et al.* (2007). A continuación se enlista cada uno de éstos:

1. Área
2. Perímetro
3. Relación de aspecto
4. Circularidad
5. Diámetro circular equivalente
6. Intensidad media de canal rojo
7. Intensidad media de canal verde
8. Intensidad media de canal azul
9. Intensidad media de canal hue (tinte o matiz)
10. Intensidad media de canal saturación
11. Intensidad media de canal intensidad

12. Intensidad media imagen mejorada hasta la etapa de mejora de contraste
13. Media de la respuesta a un filtro gaussiano con $\sigma = 1$
14. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 1$
15. Media de la respuesta a un filtro gaussiano con $\sigma = 2$
16. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 2$
17. Media de la respuesta a un filtro gaussiano con $\sigma = 4$
18. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 4$
19. Media de la respuesta a un filtro gaussiano con $\sigma = 8$
20. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 8$
21. Diferencia de la media del filtro con $\sigma = 1$ y $\sigma = 2$
22. Diferencia de la media del filtro con $\sigma = 2$ y $\sigma = 4$
23. Diferencia de la media del filtro con $\sigma = 4$ y $\sigma = 8$
24. Número de píxeles de la región candidata
25. Valor de intensidad máximo de la región
26. Valor de intensidad mínimo de la región
27. Rango dinámico de la región

4.4.5. Clasificación de candidatos

La etapa de clasificación consiste en determinar en base al vector de características de cada candidato, si éste es o no un microaneurisma. Esta etapa se realizó en la plataforma de Weka 3.6.13, el primer paso fue importar el archivo “.arff” generado en la fase previa e inmediatamente después se procedió a utilizar diferentes modelos de clasificadores (22) con el objetivo de obtener una comparación del desempeño de cada uno de éstos para el problema de la detección de microaneurismas.

Para la evaluación de los modelos de clasificación, fue necesario dividir los datos en dos conjuntos, uno de entrenamiento y uno de prueba. No fue posible utilizar la validación cruzada debido al desbalance existente entre el número de regiones candidatas que no eran MA's y las que sí eran MA's. Utilizar la validación cruzada genera un sobreentrenamiento para asignar la etiqueta de "no microaneurisma" a las regiones candidatas, debido a que únicamente 207 regiones eran microaneurismas y 91,353 regiones no eran microaneurismas.

El conjunto de entrenamiento se conformó con 107 regiones que si eran microaneurismas y 107 que no lo eran. El conjunto de pruebas se conformó con 91,346 regiones, de las cuáles únicamente 100 correspondían a microaneurismas reales.

Capítulo 5. Experimentación y resultados

5.1. Introducción

En el presente capítulo se proponen y describen los experimentos para la evaluación de la metodología de procesamiento digital de imágenes propuesta, los parámetros utilizados, los mecanismos empleados para la optimización de éstos, la sintonización de estos mecanismos, y la cuantificación del número de candidatos generados respecto a los datos ground truth para determinar la solución que permite identificar el mayor número de lesiones.

5.2. Experimento 1

5.2.1. Descripción

En la metodología se describieron las etapas y técnicas para el procesamiento de las imágenes; cada una de las técnicas requiere de parámetros de entrada, e.g. tamaño de los filtros, tamaño del elemento estructural, tamaño de las ventanas. El primer experimento consistió en utilizar parámetros similares a los encontrados en el estado del arte (ver Capítulo 2).

Para el primer experimento se propuso evaluar los parámetros de procesamiento con un conjunto de 10 imágenes (ver Anexo A) que se seleccionaron del banco de imágenes, y que de acuerdo a los datos ground truth se generó una plantilla (imagen binaria) donde únicamente destacan las lesiones que son microaneurismas. El objetivo de generar las plantillas (imágenes de control) es el de tener un control de la segmentación (generación de candidatos), i.e. poder determinar el número de microaneurismas reales que generan los parámetros de procesamiento y qué tanta área de la lesión generan éstos. El total del número de lesiones presentes en las imágenes de control es de 29 microaneurismas.

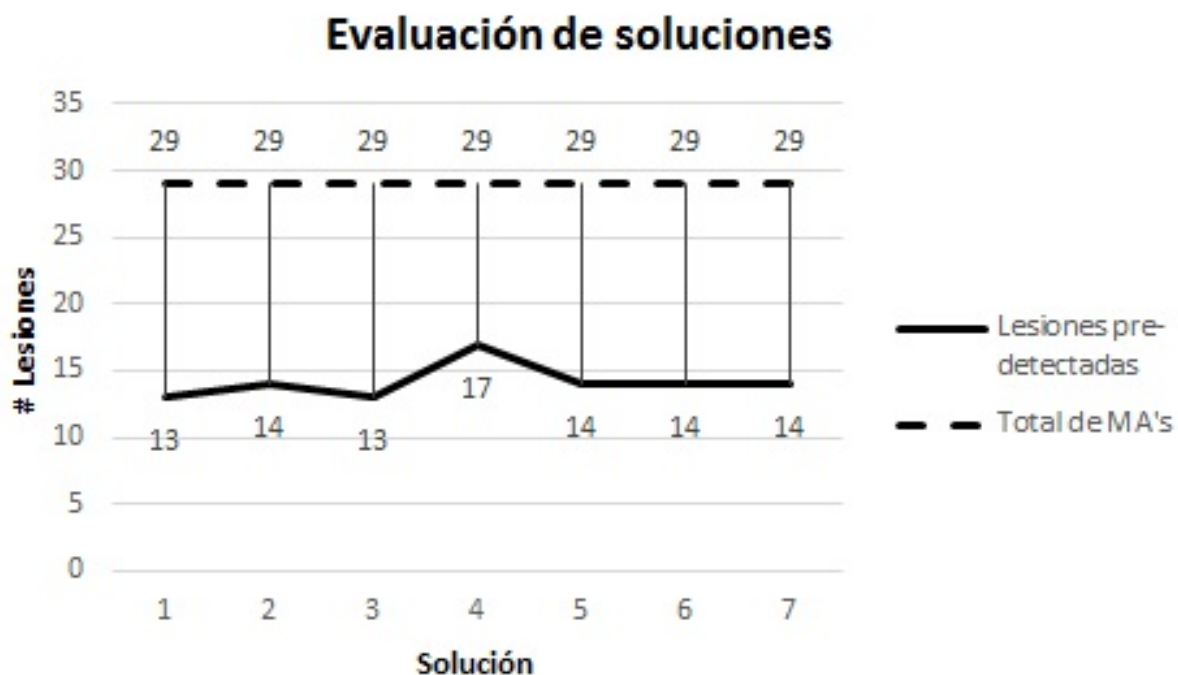
El primer experimento consistió en proponer un parámetro de procesamiento y en hacer pequeñas variaciones para observar su respuesta, además de realizar una comparación de éstos con la generación de candidatos cuando no se incorpora preprocesamiento. Las soluciones evaluadas se muestran en la tabla 6.

Tabla 6: Soluciones propuestas.

Solución	Filtro gaussiano	Filtro mediano	Ventana CLAHE	Clip Limit	Elemento estructural
1		No preprocesamiento			4
2	1.3	NA	20	0.1	10
3	2	NA	20	0.5	10
4	1	NA	20	0.2	10
5	1.3	NA	20	0.03	11
6	1	NA	20	0.05	12
7	0.8	NA	20	0.05	15

5.2.2. Resultados

Con las siete soluciones propuestas se procedió a procesar las 10 imágenes seleccionadas para el presente experimento con el objetivo de compararlas con las plantillas construidas previamente. El objetivo de la comparación es el de determinar el número de microaneurismas reales que se encuentran en los candidatos a microaneurismas generados por cada solución y para cada imagen. Una vez que se procesaron las imágenes se cuantificaron el número de microaneurismas segmentados (ver Figura 65).

**Figura 65:** Número de microneurismas detectados por cada solución.

La mejor solución de las que se propusieron para el presente experimento logró generar 17 candidatos que efectivamente son microaneurismas. Además del conteo de microaneurismas, se procedió a evaluar las seis soluciones que emplean preprocesamiento para determinar, además del número de microaneurismas "pre-detectados", el porcentaje de área de microaneurisma segmentado con respecto al área de la lesión original; para esto se establecieron tres métricas: 1) el porcentaje de píxeles falsos positivos, 2) el porcentaje de píxeles falsos negativos, y 3) el porcentaje de píxeles verdaderos positivos; con las que se evaluó cada imagen resultado de procesamiento con su respectiva plantilla. En la tabla 7 se presentan los resultados y se puede observar que la solución que se desempeño mejor fue la número 4 con 17 microaneurismas "pre-detectados", sin embargo, en promedio, menos de la mitad del área total de cada microaneurisma es identificado.

Tabla 7: Porcentaje de área de microaneurismas reales detectada.

Solución	Lesiones detectadas	Lesiones no detectadas	M1 - FP		M2 - FN		M3 - VP	
			Media	Dev_Std	Media	Dev_Std	Media	Dev_Std
2	14	15	1.91	5.76	29.87	26.56	33.73	29.04
3	13	16	0.23	0.62	26.39	28.62	25.09	27.48
4	17	12	1.47	2.74	38.16	30.60	43.16	31.31
5	14	15	0.27	0.82	31.94	30.49	26.39	26.38
6	14	15	1.35	3.27	32.54	29.56	31.61	28.50
7	14	15	1.35	3.32	31.80	29.69	32.35	29.47

5.3. Experimento 2

5.3.1. Descripción

A partir del primer experimento es posible percatarse de la necesidad de procesar las imágenes con parámetros adecuados; en el presente trabajo de investigación se plantea la hipótesis de que la generación de candidatos es un problema de optimización que puede resolverse con la técnica de cómputo evolutivo: algoritmos genéticos.

Para determinar si la hipótesis se ha formulado correctamente es necesario evaluar

las soluciones a las que converge la implementación del algoritmo genético propuesta. El presente experimento se realizó utilizando el algoritmo genético explicado previamente en la metodología. Para esta primera ejecución del algoritmo, se estableció un tamaño de población igual a 100, el punto de convergencia se fijó a un número máximo de generaciones igual a 90, la función de aptitud por la métrica de $aptitud = 0.3 * ld + 0.4 * vpt + 0.3 * vnt$, donde ld es un valor que toma un valor de 0 ó 100 dependiendo de si el número de píxeles verdaderos positivos de cada microaneurisma en la imagen procesada comparada con su plantilla es igual o mayor a 3, vpt es el porcentaje de verdaderos positivos y vnt el porcentaje de verdaderos negativos, de manera que la función de aptitud se encuentra entre 0 y 100; los demás parámetros del algoritmo genético se mantuvieron de acuerdo a los valores estipulados en la metodología.

5.3.2. Resultados

En la figura 66a se observa la evolución del individuo más apto (línea continua), la media de la aptitud de los individuos (línea de guión) y la evolución del individuo con menor aptitud (línea continua con marcadores de estrella); en la figura 66b se observa el paisaje de aptitud de las soluciones exploradas por el algoritmo genético. A partir de las gráficas se puede observar una convergencia temprana del algoritmo, i.e., individuos con aptitud similar al de mayor aptitud aparecieron alrededor de la generación número 20, y en adelante, los incrementos de aptitud no resultaron significativos. Respecto al individuo de menor aptitud, se observa un fenómeno deseado, i.e., su aptitud se encuentra oscilando hacia una mayor y una menor aptitud a lo largo de todas las generaciones; este comportamiento es signo de aleatoriedad y permite introducir eventualmente variaciones para explorar otras soluciones que puede que sean más aptas que el individuo más apto actual.

Para evaluar las soluciones obtenidas con el algoritmo genético, se seleccionó aleatoriamente una solución de las primeras cinco generaciones y una solución de las últimas 5 generaciones; en la tabla 8 se muestran éstas. Con ambas soluciones se procesaron las 10 imágenes seleccionadas (ver Anexo A) para compararlas con su respectiva plantilla y cuantificar el número de microaneurismas "pre-detectados".

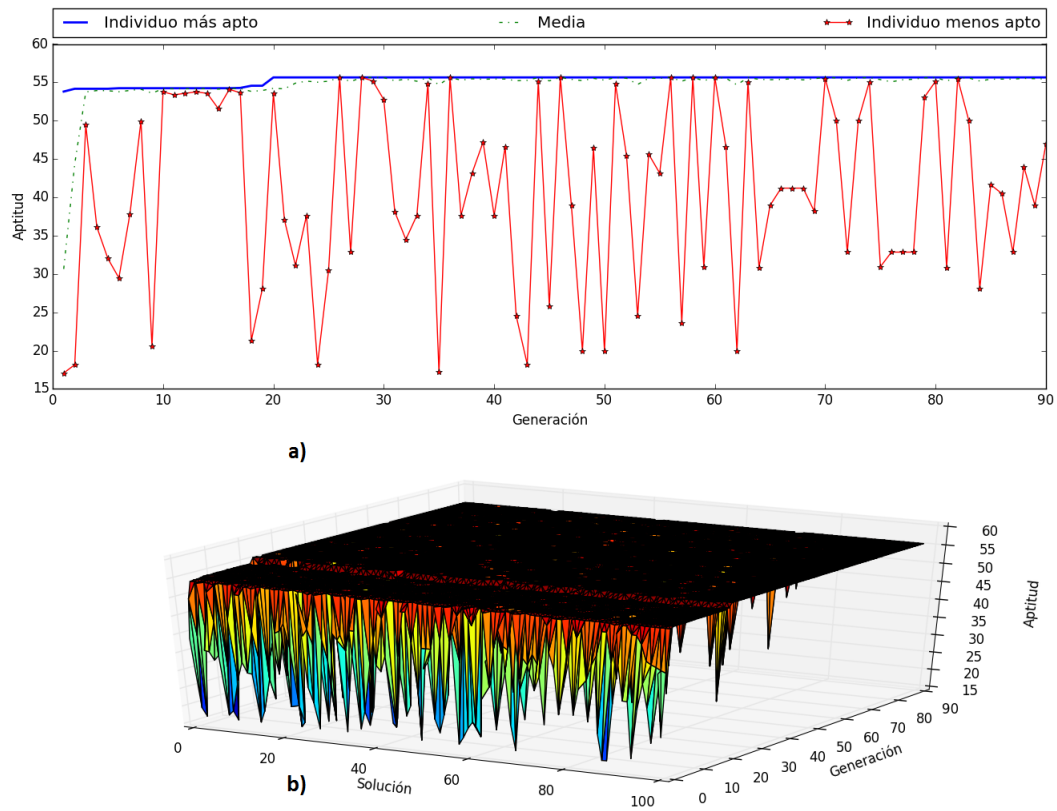


Figura 66: a) Evolución de la población, b) Paisaje de aptitud.

La solución elegida aleatoriamente de entre los individuos de las últimas cinco generaciones mejoró en $\approx 35\%$ el número de lesiones detectadas con respecto a la mejor solución encontrada en el experimento 1, logrando "pre-detectar" 23 microaneurismas de los 29 existentes (ver Tabla 9). La mejora de las soluciones con respecto al número de lesiones detectadas permite validar el uso de algoritmos genéticos como técnica de optimización de parámetros de preprocesamiento.

Tabla 8: Soluciones propuestas.

Solución	Filtro gaussiano	Filtro mediano	Ventana CLAHE	Clip Limit	Elemento estructural
8	10.97	NA	30	0.38	15
9	1	NA	75	0.593	11

Tabla 9: Número de microaneurismas reales pre-detectados.

Solución	Lesiones detectadas	Lesiones no detectadas
8	1	28
9	23	6

5.4. Experimento 3

5.4.1. Descripción

Para el experimento 3 se crearon 14 nuevas plantillas (ver Anexo A) para servir como control de la segmentación en base a 14 imágenes en una nueva implementación de un algoritmo genético. Para el presente experimento se probaron 4 funciones de aptitud utilizando dos esquemas diferentes para cada una, de manera que se ejecutó 8 veces el algoritmo genético, 4 implementaciones para optimizar utilizando el esquema que involucra el filtro mediano como técnica de reducción de ruido, y 4 implementaciones para optimizar utilizando el esquema que involucra el filtro gaussiano. Para la evaluación de cada ejecución se seleccionaron tres diferentes soluciones, una de entre los individuos más aptos (últimas 20 generaciones), otra de los individuos entre la generación 20 y 30, y la última de las primeras 20 generaciones.

El número de individuos de la población se estableció a 130 con un criterio de convergencia de número máximo de generaciones igual a 50. Para esta implementación se redujo el número de generaciones para disminuir el tiempo de ejecución del algoritmo; para compensar el bajo número de generaciones se favoreció la aleatoriedad del algoritmo aumentando la probabilidad de cruzamiento y de mutación al doble de los valores fijados en la implementación anterior y se ajustó la presión de selección para favorecer a los individuos menos aptos. El mecanismo de selección y el porcentaje de elitismo no sufrió alteraciones.

En cuanto a la función de aptitud, se propusieron las siguientes:

- Experimento 3a: Optimización utilizando un filtro mediano con función de aptitud: $aptitud = 0.3 \times Id + 0.7 \times exactitud$. Solución 10-12

- Experimento 3c: Optimización utilizando un filtro mediano con función de aptitud:
 $aptitud = 0.3 \times ld + 0.7 \times precisión$. Solución 13-15
- Experimento 3e: Optimización utilizando un filtro mediano con función de aptitud:
 $aptitud = 0.3 \times ld + 0.7 \times sensibilidad$. Solución 16-18
- Experimento 3g: Optimización utilizando un filtro mediano con función de aptitud:
 $aptitud = 0.3 \times ld + 0.4 \times vpt + 0.3 \times vnt$. Solución 19-21
- Experimento 3b: Optimización utilizando un filtro gaussiano con función de aptitud:
 $aptitud = 0.3 \times ld + 0.7 \times exactitud$. Solución 22-24
- Experimento 3d: Optimización utilizando un filtro gaussiano con función de aptitud:
 $aptitud = 0.3 \times ld + 0.7 \times precisión$. Solución 25-27
- Experimento 3f: Optimización utilizando un filtro gaussiano con función de aptitud:
 $aptitud = 0.3 \times ld + 0.7 \times sensibilidad$. Solución 28-30
- Experimento 3h: Optimización utilizando un filtro gaussiano con función de aptitud:
 $aptitud = 0.3 \times ld + 0.4 \times vpt + 0.3 \times vnt$. Solución 31-33

La evaluación de las soluciones encontradas por cada una de las 8 ejecuciones consistió en procesar cada una de las 58 imágenes del banco de datos para estas soluciones y las obtenidas en el experimento 1 y 2, y comparar el número de MA's reales predetecados por cada solución.

5.4.2. Resultados

En el Anexo A se pueden observar los paisajes de aptitud y las gráficas de la evolución de la población para cada ejecución del algoritmo, de las que se puede inferir la necesidad de disminuir el porcentaje de elitismo, la convergencia del algoritmo hacia una solución es prematura, y aunque el algoritmo sigue generando individuos aleatorios aún en las últimas generaciones, se puede observar un estancamiento en la aptitud más alta. Con cada una de las soluciones se procedió a procesar cada una de las imágenes (58) para cuantificar el número de microaneurismas reales persistentes hasta la etapa de segmentación.

En la tabla 10 se compara el número de candidatos que generó cada una de las soluciones evaluadas, así como los verdaderos positivos, falsos positivos, falsos negativos y la sensibilidad de cada una de éstas para las 58 imágenes del banco de datos.

Tabla 10: Evaluación de las soluciones encontradas por el Algoritmo Genético.

Solución	# de candida- tos	VP's	FP's	FN's	Sensibilidad
1	2430609	17	2430592	229	6.91
2	29673	160	29513	86	65.04
3	18091	137	17954	109	55.69
4	58968	165	58803	81	67.06
5	27748	151	27597	95	61.38
6	48274	153	48121	93	62.2
7	73482	150	73332	96	60.98
8	22678	19	22659	227	7.72
9	91560	207	91349	39	84.14
10	235725	1	235724	245	0.41
11	526053	0	526053	246	0
12	13245	18	13227	228	7.32
13	346900	1	346899	245	0.41
14	105855	6	105849	240	2.44
15	10255	20	10235	226	8.13
16	764817	3	764814	243	1.22
17	30550	32	30518	214	13.01
18	15081	30	15051	216	12.2
19	959173	8	959165	238	3.25
20	88163	187	87976	59	76.02
21	12157	22	12135	224	8.94
22	13261	133	13128	113	54.07
23	347708	164	347544	82	66.67
24	53425	158	53267	88	64.23

Tabla 10: Continuación...

Solución	# de candida- tos	VP's	FP's	FN's	Sensibilidad
25	77389	27	77362	219	10.98
26	117026	175	116851	71	71.14
27	478752	1	478751	245	0.41
28	398576	47	398529	199	19.11
29	15337	96	15241	150	39.02
30	40044	189	39855	57	76.83
31	313830	168	313662	78	68.29
32	23204	71	23133	175	28.86
33	515433	127	515306	119	51.63

La mejor solución fue encontrada por el algoritmo genético del experimento 2 y logró predelectar 207 MA's. En el experimento planteado en esta sección se obtuvieron 24 soluciones a los parámetros de procesamiento, ninguna superó la sensibilidad lograda por la mejor solución del experimento 2, aunque hubo soluciones que generaron un número de candidatos cercano (189 la más cercana con una reducción de candidatos del 43.73 % respecto a la solución 9). La optimización sobre el filtro mediano únicamente generó una solución aceptable (187 MA's), mientras que la implementación sobre el filtro gaussiano, aunque no superó los 207 MA's, encontró 4 soluciones que se pueden considerar aceptables (189, 175, 168, 164).

5.5. Experimento 4: Selección de características

5.5.1. Descripción

El presente experimento consistió en utilizar el software WEKA para la evaluación y selección de características. El objetivo de la presente etapa es reducir la dimensionalidad de los vectores para disminuir los tiempos de procesamiento. Para reducir la dimensionalidad se desechan descriptores en los que no existe un margen de separación suficiente para las instancias que son positivas y negativas.

Las regiones candidatas fueron obtenidas a partir de procesar las 58 imágenes del banco de datos con los parámetros de la solución 9 (ver Tabla 8) y de aplicar el algoritmo de etiquetaje. Después de extraer las características de cada región se procedió a marcar manualmente los vectores de características como "sí" o "no" dependiendo de si la región era o no un microaneurisma.

La selección de características se realizó con tres diferentes métodos implementados en WEKA: BestFisrt, búsqueda exhaustiva y algoritmo genético. En el experimento 5 se comparan los resultados de clasificación de los tres conjuntos de descriptores encontrados, más el resultado obtenido con los 27 descriptores.

5.5.2. Resultados

El conjunto de descriptores encontrados para cada método se enumeran a continuación:

Método BestFirst

1. Relación de aspecto
2. Circularidad
3. Intensidad media de canal rojo
4. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 1$
5. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 4$
6. Diferencia de la media del filtro con $\sigma = 2$ y $\sigma = 4$
7. Valor de intensidad mínimo de la región
8. Rango dinámico de la región

Método de búsqueda exhaustiva

1. Relación de aspecto

2. Circularidad
3. Intensidad media de canal rojo
4. Media de la respuesta a un filtro gaussiano con $\sigma = 1$
5. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 1$
6. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 4$
7. Diferencia de la media del filtro con $\sigma = 2$ y $\sigma = 4$
8. Rango dinámico de la región

Método del algoritmo genético

1. Relación de aspecto
2. Circularidad
3. Diámetro circular equivalente
4. Intensidad media de canal rojo
5. Intensidad media de canal verde
6. Intensidad media de canal hue (tinte o matiz)
7. Media de la respuesta a un filtro gaussiano con $\sigma = 1$
8. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 1$
9. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 4$
10. Desviación estándar de la respuesta a un filtro gaussiano con $\sigma = 8$
11. Diferencia de la media del filtro con $\sigma = 1$ y $\sigma = 2$
12. Diferencia de la media del filtro con $\sigma = 2$ y $\sigma = 4$
13. Rango dinámico de la región

5.6. Experimento 5: Clasificación de candidatos

5.6.1. Descripción

El presente experimento se realizó en el software WEKA y consistió en entrenar con el conjunto de entrenamiento (214 regiones: 107 MA's y 107 no MA's) distintos modelos de clasificación (22 modelos) y en evaluarlos con el conjunto de pruebas (91,346 regiones: 100 MA's y 91,246 no MA's).

Cada evaluación del modelo de clasificación se realizó cuatro veces: 1) con los 27 descriptores, 2) con los 9 descriptores encontrados por el método BestFirst, 3) con los 8 descriptores encontrados por el método de búsqueda exhaustiva y 4) con los 13 descriptores encontrados por el método del algoritmo genético.

5.6.2. Resultados

Las métricas con mayor relevancia en cuanto a la detección de microaneurismas, son la sensibilidad y la especificidad. La sensibilidad indica qué tan apto es el modelo para detectar los microaneurismas y la especificidad indica qué tan apto es el modelo para discriminar los falsos positivos. Entre mayor sea la sensibilidad menor será el número de microaneurismas no detectados (se reduce el error de tipo II), cuando la sensibilidad es del 100 % significa que el modelo ha sido capaz de encontrar todas las lesiones sin descartar ninguna, i.e. el número de falsos negativos es cero. Por otro lado, la especificidad se traduce en la capacidad de discriminación del modelo, cuando ésta tiende a 100 % significa que el número de falsos positivos (error de tipo I) ha sido reducido a cero; la importancia de mantener esta métrica por encima de un umbral permite mantener el número de falsos positivos en un rango tolerable.

Para la detección de microaneurismas en programas a escala nacional, la Asociación Británica de Diabetes ha propuesto un desempeño mínimo referente a sensibilidad del 80 % y especificidad mínima del 95 % (Horton *et al.*, 2016).

En las tablas 11, 12, 13 y 14 se muestran los resultados obtenidos para cada modelo de clasificación utilizando diferentes conjuntos de descriptores de acuerdo a los métodos de selección.

De entre los modelos de clasificación utilizados en la comparación, la mayoría logró superar el umbral de la sensibilidad del 80 %, sin embargo no lograron ni igualar ni superar el umbral de especificidad propuesto por la Asociación Británica de Diabetes.

El modelo de clasificación que obtuvo el mayor porcentaje de especificidad fue el de "Näive Bayes Multinomial Updateable" utilizando 28 descriptores con 77.95 %, sin embargo su porcentaje de sensibilidad estuvo por debajo de la métrica deseada, siendo igual al 72 %. En la figura 67 se observa la curva ROC de este modelo.

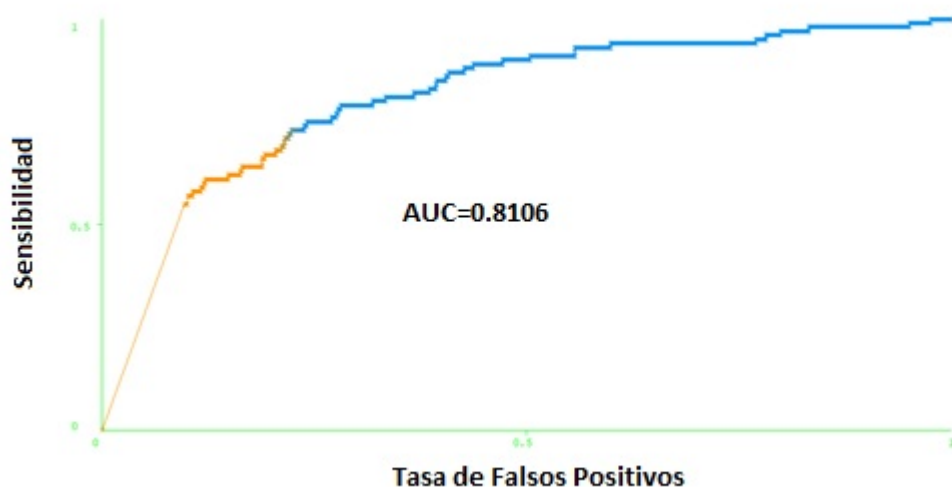


Figura 67: Curva ROC para el modelo "Näive Bayes Multinomial Updateable" con 28 descriptores.

Por otro lado, el método de FilteredClassifier (8 y 9 descriptores) logró obtener el 100 % de sensibilidad a un costo de una disminución significativa en porcentaje de especificidad (33.39 %). Los modelos OneR, J48 y Decision Stump, utilizando los 28 descriptores lograron obtener un desempeño equilibrado, cada uno de éstos superó el umbral de sensibilidad deseado. OneR: 85 % y 73.32 %, J48: 92 % y 67.26 %, Decision Stump: 87 % y 69.26 %.

Aunque no es posible realizar una comparación justa debido a que cada una de las metodologías utiliza diferentes bases de datos, en la tabla 15 se comparan los resultados de clasificación más significativos de la metodología propuesta con algunos del estado del arte.

Tabla 11: Evaluación de 22 modelos de clasificación con 28 descriptores.

Modelo	VP	FP	FN	VN	Tasa de error	Exactitud	Precisión	Sensibilidad	Especificidad	VP-	FPR	FNR	F1
Adaboost	87	33143	13	58103	36.30	63.70	0.26	87.00	63.68	99.98	36.32	13.00	0.52
Bagging	85	28112	15	63154	30.78	69.22	0.30	85.00	69.20	99.98	30.80	15.00	0.60
Filtered classifier	99	88800	1	2446	97.21	2.79	0.11	99.00	2.68	99.96	97.32	1.00	0.22
Iterative Classifier optimizer	95	37770	5	53476	41.35	58.65	0.25	95.00	58.61	99.99	41.39	5.00	0.50
LogitBoost	97	51230	3	40016	56.09	43.91	0.19	97.00	43.86	99.99	56.14	3.00	0.38
Random Subspace	86	25250	14	65996	27.66	72.34	0.34	86.00	72.33	99.98	27.67	14.00	0.68
Kstar	78	30229	22	61017	33.12	66.88	0.26	78.00	66.87	99.96	33.13	22.00	0.51
LWL	87	28041	13	63205	30.71	69.29	0.31	87.00	69.27	99.98	30.73	13.00	0.62
Perceptrón multica	74	34502	26	56744	37.80	62.20	0.21	74.00	62.19	99.95	37.81	26.00	0.43
BayesNet	100	80162	0	11084	87.76	12.24	0.12	100.00	12.15	100.00	87.85	0.00	0.25
Naïve Bayes	72	20118	28	71128	22.05	77.95	0.36	72.00	77.95	99.96	22.05	28.00	0.71
Updateable													
SGD	85	36314	15	54932	39.77	60.23	0.23	85.00	60.20	99.97	39.80	15.00	0.47
SimpleLogistic	80	26493	20	64303	29.17	70.83	0.30	80.00	70.82	99.97	29.18	20.00	0.60
SMO	83	37765	17	53481	41.36	58.64	0.22	83.00	58.61	99.97	41.39	17.00	0.44
Voted Perceptron	78	35084	22	56162	38.43	61.57	0.22	78.00	61.55	99.96	38.45	22.00	0.44
oneR	85	24337	15	66909	26.66	73.34	0.35	85.00	73.33	99.98	26.67	15.00	0.69
PART	94	43931	6	47315	48.10	51.90	0.21	94.00	51.85	99.99	48.15	6.00	0.43
Decision Stump	87	28041	13	63205	30.71	69.29	0.31	87.00	69.27	99.98	30.73	13.00	0.62
J48	92	29869	8	61377	32.71	67.29	0.31	92.00	67.27	99.99	32.73	8.00	0.61
LMT	80	26943	20	64303	29.52	70.48	0.30	80.00	70.47	99.97	29.53	20.00	0.59
RandomTree	94	45368	6	45878	49.67	50.33	0.21	94.00	50.28	99.99	49.72	6.00	0.41
REPTree	84	25613	16	65633	28.06	71.94	0.33	84.00	71.93	99.98	28.07	16.00	0.65

Tabla 12: Evaluación de 22 modelos de clasificación con 9 descriptores.

Modelo	VP	FP	FN	VN	Tasa de error	Exactitud	Precisión	Sensibilidad	Especificidad	VP-	FPR	FNR	F1
Adaboost	90	65940	10	25306	72.20	27.80	0.14	90.00	27.73	99.96	72.27	10.00	0.27
Bagging	96	80569	4	10677	88.21	11.79	0.12	96.00	11.70	99.96	88.30	4.00	0.24
Filtered classifier	100	88199	0	3047	96.55	3.45	0.11	100.00	3.34	100.00	96.66	0.00	0.23
Iterative Classifier optimizer	90	65682	10	25564	71.92	28.08	0.14	90.00	28.02	99.96	71.98	10.00	0.27
LogitBoost	89	69878	11	21368	76.51	23.49	0.13	89.00	23.42	99.95	76.58	11.00	0.25
Random Subspace	92	77774	8	13472	85.15	14.85	0.12	92.00	14.76	99.94	85.24	8.00	0.24
Kstar	77	52164	23	39082	57.13	42.87	0.15	77.00	42.83	99.94	57.17	23.00	0.29
LWL	74	60894	26	30352	66.69	33.31	0.12	74.00	33.26	99.91	66.74	26.00	0.24
Perceptrón multica	85	52873	15	38373	57.90	42.10	0.16	85.00	42.05	99.96	57.95	15.00	0.32
BayesNet	98	82840	2	8406	90.69	9.31	0.12	98.00	9.21	99.98	90.79	2.00	0.24
Naïve Bayes													
Multinomial Updateable	82	30831	18	60415	33.77	66.23	0.27	82.00	66.21	99.97	33.79	18.00	0.53
SGD	81	28856	19	62390	31.61	68.39	0.28	81.00	68.38	99.97	31.62	19.00	0.56
SimpleLogistic	80	27138	20	64108	29.73	70.27	0.29	80.00	70.26	99.97	29.74	20.00	0.59
SMO	81	29070	19	62176	31.84	68.16	0.28	81.00	68.14	99.97	31.86	19.00	0.55
Voted Perceptron	75	34837	25	56409	38.16	61.84	0.21	75.00	61.82	99.96	38.18	25.00	0.43
oneR	85	24337	15	66909	26.66	73.34	0.35	85.00	73.33	99.98	26.67	15.00	0.69
PART	97	46832	3	44414	51.27	48.73	0.21	97.00	48.68	99.99	51.32	3.00	0.41
Decision Stump	87	28041	13	63205	30.71	69.29	0.31	87.00	69.27	99.98	30.73	13.00	0.62
J48	87	32664	13	58582	35.77	64.23	0.27	87.00	64.20	99.98	35.80	13.00	0.53
LMT	80	27138	20	64108	29.73	70.27	0.29	80.00	70.26	99.97	29.74	20.00	0.59
RandomTree	93	38099	7	53147	41.72	58.28	0.24	93.00	58.25	99.99	41.75	7.00	0.49
REPTree	85	24251	15	66995	26.56	73.44	0.35	85.00	73.42	99.98	26.58	15.00	0.70

Tabla 13: Evaluación de 22 modelos de clasificación con 8 descriptores.

Modelo	VP	FP	FN	VN	Tasa de error	Exactitud	Precisión	Sensibilidad	Especificidad	VP-	FPR	FNR	F1
Adaboost	90	65940	10	25306	72.20	27.80	0.14	90.00	27.73	99.96	72.27	10.00	0.27
Bagging	96	80569	4	10677	88.21	11.79	0.12	96.00	11.70	99.96	88.30	4.00	0.24
Filtered classifier	100	88199	0	3047	96.55	3.45	0.11	100.00	3.34	100.00	96.66	0.00	0.23
Iterative Classifier optimizer	90	65682	10	25564	71.92	28.08	0.14	90.00	28.02	99.96	71.98	10.00	0.27
LogitBoost	89	69878	11	21368	76.51	23.49	0.13	89.00	23.42	99.95	76.58	11.00	0.25
Random Subspace	92	77774	8	13472	85.15	14.85	0.12	92.00	14.76	99.94	85.24	8.00	0.24
Kstar	77	52164	23	39082	57.13	42.87	0.15	77.00	42.83	99.94	57.17	23.00	0.29
LWL	74	60894	26	30352	66.69	33.31	0.12	74.00	33.26	99.91	66.74	26.00	0.24
Perceptrón multicapa	85	52873	15	38373	57.90	42.10	0.16	85.00	42.05	99.96	57.95	15.00	0.32
BayesNet	98	82840	2	8406	90.69	9.31	0.12	98.00	9.21	99.98	90.79	2.00	0.24
Naïve Bayes													
Multinomial Updateable	82	30831	18	60415	33.77	66.23	0.27	82.00	66.21	99.97	33.79	18.00	0.53
SGD	81	28856	19	62390	31.61	68.39	0.28	81.00	68.38	99.97	31.62	19.00	0.56
SimpleLogistic	80	27138	20	64108	29.73	70.27	0.29	80.00	70.26	99.97	29.74	20.00	0.59
SMO	81	29070	19	62176	31.84	68.16	0.28	81.00	68.14	99.97	31.86	19.00	0.55
Voted Perceptron	75	34837	25	56409	38.16	61.84	0.21	75.00	61.82	99.96	38.18	25.00	0.43
oneR	85	24337	15	66909	26.66	73.34	0.35	85.00	73.33	99.98	26.67	15.00	0.69
PART	97	46832	3	44414	51.27	48.73	0.21	97.00	48.68	99.99	51.32	3.00	0.41
Decision Stump	87	28041	13	63205	30.71	69.29	0.31	87.00	69.27	99.98	30.73	13.00	0.62
J48	87	32664	13	58582	35.77	64.23	0.27	87.00	64.20	99.98	35.80	13.00	0.53
LMT	80	27138	20	64108	29.73	70.27	0.29	80.00	70.26	99.97	29.74	20.00	0.59
RandomTree	93	38099	7	53147	41.72	58.28	0.24	93.00	58.25	99.99	41.75	7.00	0.49
REPTree	85	24251	15	66995	26.56	73.44	0.35	85.00	73.42	99.98	26.58	15.00	0.70

Tabla 14: Evaluación de 22 modelos de clasificación con 13 descriptores.

Modelo	VP	FP	FN	VN	Tasa de error	Exactitud	Precisión	Sensibilidad	Especificidad	VP-	FPR	FNR	F1
Adaboost	87	36145	13	55101	39.58	60.42	0.24	87.00	60.39	99.98	39.61	13.00	0.48
Bagging	84	27535	16	63711	30.16	69.84	0.30	84.00	69.82	99.97	30.18	16.00	0.61
Filtered classifier	97	88383	3	2863	96.76	3.24	0.11	97.00	3.14	99.90	96.86	3.00	0.22
Iterative Classifier optimizer	90	40088	10	51158	43.90	56.10	0.22	90.00	56.07	99.98	43.93	10.00	0.45
LogitBoost	90	40088	10	51158	43.90	56.10	0.22	90.00	56.07	99.98	43.93	10.00	0.45
Random Subspace	95	75560	5	15686	82.72	17.28	0.13	95.00	17.19	99.97	82.81	5.00	0.25
Kstar	76	37302	24	53944	40.86	59.14	0.20	76.00	59.12	99.96	40.88	24.00	0.41
LWL	87	28041	13	63205	30.71	69.29	0.31	87.00	69.27	99.98	30.73	13.00	0.62
Perceptrón multica	79	33653	21	57593	36.86	63.14	0.23	79.00	63.12	99.96	36.88	21.00	0.47
BayesNet	99	85969	1	5277	94.11	5.89	0.12	99.00	5.78	99.98	94.22	1.00	0.23
Naïve Bayes													
Multinomial Updateable	76	23823	24	67423	26.11	73.89	0.32	76.00	73.89	99.96	26.11	24.00	0.63
SGD	80	32912	20	58334	36.05	63.95	0.24	80.00	63.93	99.97	36.07	20.00	0.48
SimpleLogistic	80	26943	20	64303	29.52	70.48	0.30	80.00	70.47	99.97	29.53	20.00	0.59
SMO	81	32435	19	58811	35.53	64.47	0.25	81.00	64.45	99.97	35.55	19.00	0.50
Voted Perceptron	99	81392	1	9854	89.10	10.90	0.12	99.00	10.80	99.99	89.20	1.00	0.24
oneR	85	24337	15	66909	26.66	73.34	0.35	85.00	73.33	99.98	26.67	15.00	0.69
PART	92	39073	8	52173	42.78	57.22	0.23	92.00	57.18	99.98	42.82	8.00	0.47
Decision Stump	87	28041	13	63205	30.71	69.29	0.31	87.00	69.27	99.98	30.73	13.00	0.62
J48	77	57561	23	33685	63.04	36.96	0.13	77.00	36.92	99.93	63.08	23.00	0.27
LMT	80	26943	20	64303	29.52	70.48	0.30	80.00	70.47	99.97	29.53	20.00	0.59
RandomTree	90	78316	10	12930	85.75	14.25	0.11	90.00	14.17	99.92	85.83	10.00	0.23
REPTree	84	25613	16	65633	28.06	71.94	0.33	84.00	71.93	99.98	28.07	16.00	0.65

En la mayoría de los modelos de clasificación se obtuvo una sensibilidad aceptable, sin embargo, en términos de especificidad los modelos no tuvieron un desempeño aceptable. Cabe destacar que no se implementaron técnicas de reducción de candidatos, por lo que el bajo porcentaje de especificidad era de esperarse. Cabe mencionar que los 39 microaneurismas que se perdieron en la etapa de segmentación no fueron contabilizados para la etapa de clasificación en WEKA. Si se desea evaluar el esquema completo, la sensibilidad se reduce en $\approx 39\%$.

El principal objetivo del presente trabajo de investigación es demostrar la viabilidad del algoritmo genético como técnica de optimización de parámetros para el procesamiento de imágenes oftálmicas para mejorar la identificación de microaneurismas. Los resultados mostrados en el experimento 1, 2 y 3 validan la hipótesis planteada. La etapa de clasificación utilizando WEKA es una etapa que se decidió agregar para analizar los puntos débiles de la metodología propuesta, y en un futuro, al implementar una metodología de detección de microaneurismas tener una guía de los descriptores y los modelos de clasificación que pueden ser viables.

Tabla 15: Comparación de la metodología propuesta con metodologías en el estado del arte.

Autor	Sensitividad	Especificidad	Exactitud
Walter <i>et al.</i> (2007)	88.47 %	100 %	No Reportado
Romero <i>et al.</i> (2015)	92.32 % (DIARETDB1), 88.06 % (ROC)	93.87 % (DIARETDB1), 97.47 % (ROC)	No Reportado
Niemeijer <i>et al.</i> (2005)	100 %	87 %	No Reportado
Hipwell <i>et al.</i> (2000)	78 %	91 %	No Reportado
Arenas-Cavalli <i>et al.</i> (2015)	77.7%	81.2 %	No Reportado
Adal <i>et al.</i> (2014)	81.08 %	92.31 %	No Reportado
Bhalerao <i>et al.</i> (2008)	82.6 %	80.2 %	No Reportado
SujithKumar y Singh (2012)	94.44 %	87.5 %	No Reportado
Lazar <i>et al.</i> (2010)	≈ 62%	No Reportado	No Reportado
Cree (2008)	85 %	90 %	No Reportado
Sopharak <i>et al.</i> (2011)	81.61 %	99.99 %	99.98 %
Streeter y Cree (2003)	56 %	No Reportado	No Reportado
Kande <i>et al.</i> (2010)	100 %	91 %	No Reportado
Aravind <i>et al.</i> (2013)	92 %	80 %	90 %
Näive Bayes Multi-nomial Updateable: 28 descriptores	72.00 %	77.95 %	77.94 %
Filtered Classifier: 8 y 9 descriptores	100 %	33.39 %	3.44 %
OneR: 28 descriptores	85.00 %	73.32 %	73.34 %
J48: 28 descriptores	92.00 %	67.26 %	67.29 %

5.7. Conclusiones

En el presente capítulo se presentaron los resultados de los experimentos propuestos. Sobre los experimentos 1, 2 y 3 se puede concluir que la adopción de la técnica de búsqueda de algoritmos genéticos conduce a una mejora significativa de la segmentación. El esquema que utiliza el filtro gaussiano como técnica de reducción de ruido obtuvo la

mayor sensibilidad, logrando detectar 207 de los 246 microaneurismas totales.

En base a los resultados, la optimización mediante el algoritmo genético logró aumentar hasta en un 51 % el número de microaneurismas detectados con respecto a las soluciones heurísticas. Aunque el algoritmo genético mostró su efectividad, es recomendable seguir probando diferentes funciones de aptitud, así como aumentar las probabilidades de cruzamiento y mutación, a la vez que se disminuya el porcentaje de elitismo. Por cuestiones de tiempo de procesamiento no fue posible realizar más experimentos. Para trabajos futuros se propone implementar una etapa en el AG que guarde las soluciones encontradas y su respectiva aptitud, para que al momento de generar la misma solución, evite procesar nuevamente y copie directamente el valor de la aptitud.

Por otro lado, emplear una función de aptitud multiobjetivo es una opción a explorar que podría conducir a mejores soluciones. Por ejemplo podría buscarse maximizar el valor predictivo negativo, y disminuir la tasa de falsos positivos y negativos.

Con respecto a la clasificación de los candidatos, se obtuvieron buenas tasas de sensibilidad, sin embargo, la especificidad de éstos no es comparable a otras metodologías en el estado del arte. Es altamente recomendable incluir una etapa de reducción de candidatos.

Finalmente, con respecto a los conjuntos de entrenamiento y de prueba, es recomendable ampliar el banco de imágenes y preferentemente contar con instancias positivas y negativas a MA's balanceadas, lo anterior para poder emplear métodos, como la de validación cruzada en el que se garantice que cada instancia ha sido utilizada en ambos conjuntos (entrenamiento y de prueba), obteniendo así métricas más estables y con menor variabilidad.

Capítulo 6. Conclusiones y trabajo a futuro

6.1. Introducción

El presente capítulo tiene el objetivo de exponer las conclusiones derivadas del trabajo de investigación, presentar las aportaciones de este trabajo de tesis para la resolución del problema propuesto y al estado del arte, y brindar recomendaciones y propuestas para trabajar a futuro sobre esta línea de investigación.

6.2. Conclusiones

La alta prevalencia de la diabetes en los últimos años conlleva a una alta demanda de personal médico con especialidad en enfermedades oftálmicas debido a que las enfermedades en el ojo son complicaciones frecuentes en personas diabéticas.

En este trabajo de investigación se propuso diseñar una metodología que permita identificar anormalidades en imágenes de fondo de ojo. Existen distintos tipos de anormalidades y el trabajo se centró en detectar específicamente las lesiones llamadas microaneurismas. Los microaneurismas evidencian la retinopatía diabética en fase temprana, siendo el objetivo detectar el problema antes de que se agrave.

De acuerdo al esquema general de detección de objetos se procedió a diseñar algoritmos de procesamiento digital de imágenes para mejorar la calidad de la imagen y en la etapa de segmentación aislar los microaneurismas de otras estructuras presentes en las imágenes de fondo de ojo; en esta etapa se encontró con el problema de seleccionar los parámetros adecuados de las técnicas de procesamiento digital de imágenes.

Para solucionar este problema se implementó un método de optimización basado en el cómputo evolutivo, específicamente el método de algoritmos genéticos. Se generaron 14 plantillas a partir de 14 imágenes de la base de datos DIARETDB1 con ayuda de los datos “ground truth”, y se implementó un algoritmo genético para optimizar la segmentación de esas 14 imágenes tomando como métrica la similitud con la plantilla generada previamente. El resultado fue una mejora en la segmentación de las imágenes, i.e. una alta cantidad del número de regiones que eran microaneurismas se mostraron en la imagen final y la cantidad de regiones que no eran microaneurismas disminuyó. Con esto

se logró demostrar que los algoritmos genéticos encuentran soluciones aceptables para los parámetros de procesamiento digital de imágenes y encaminan a una mejora en la segmentación de las imágenes.

Posteriormente, con el objetivo de determinar si las soluciones aceptables para 14 imágenes de la base de datos son válidas para el resto de las imágenes de ésta, se procedió a procesar 58 imágenes de la base de datos DIARETDB1 con la mejor solución encontrada por los algoritmos genéticos. Los resultados obtenidos fueron satisfactorios, validando a los algoritmos genéticos como técnica de optimización para el procesamiento de imágenes oftálmicas. La mejor solución encontrada aumentó en $\approx 35\%$ el número de microaneurismas "pre-detectados" en las 10 plantillas con respecto a la mejor solución encontrada heurísticamente.

Con cada una de las soluciones se procesó cada una de las imágenes del banco de datos, la mejor solución encontrada por los algoritmos genéticos generó hasta un 51.09 % más de candidatos que son MA's con respecto a las 6 soluciones heurísticas evaluadas. De igual manera, la mejor solución encontró 207 de los 246 microaneurismas, obteniendo así una sensibilidad del 84.14 %.

Continuando con el esquema general de detección de objetos, se procedió a extraer descriptores de cada candidato a microaneurisma obtenido en la etapa de segmentación con el objetivo de generar vectores de características. Los vectores de características se integraron de distintos descriptores comúnmente utilizados en el reconocimiento de microaneurismas. La etapa de clasificación de candidatos consistió en clasificar cada región candidata a microaneurisma como microaneurisma o no microaneurisma utilizando la clasificación por vectores de características y se implementó para evaluar el comportamiento de distintos modelos de clasificación utilizando la clasificación por vectores de características. Los modelos de clasificación utilizados fueron métodos implementados en el software Weka con sus parámetros por defecto, de entre los que se utilizaron, los siguientes obtuvieron mejores métricas: el de Näive Bayes Multinomial Updateable, OneR, J48 y Decision Stump. El clasificador con mejores porcentajes de detección de microaneurismas fue el de Filtered Classifier (8 y 9 descriptores) con un 100 % de MA's detectados, sin embargo la especificidad de éste se encontró muy por debajo del estándar

con un 33.39 %.

En la metodología propuesta no se implementó una etapa de reducción de candidatos, por lo que cada imagen del banco de datos, al ser procesada generó una alta cantidad de regiones candidatas, de manera que los bajos porcentajes para la métrica de especificidad eran de esperarse. A pesar de no implementar una etapa de reducción de candidatos, algunos modelos obtuvieron una especificidad que logró discriminar una alta cantidad de regiones candidatas. El modelo Näive Bayes Multinomial Updateable logró una especificidad del 77.95 % con una sensibilidad del 72 %, logrando discriminar hasta 77.86 % regiones candidatas, por lo que se espera que en un trabajo futuro que contenga la etapa de reducción de candidatos se obtenga una especificidad muy cercana a la propuesta en el estándar.

Durante el desarrollo del trabajo de investigación se observó que la identificación de anormalidades oftalmológicas es un problema complejo y un reto para los oftalmólogos; de igual manera se observó que la detección automática de microaneurismas es un problema no resuelto aún, de alta complejidad para las áreas de aprendizaje de máquina, tele-oftalmología y visión por computadora.

Sobre las técnicas de procesamiento digital de imágenes (etapa de preprocesamiento), a pesar de no ser técnicas desarrolladas específicamente para el objetivo deseado, sino técnicas generalmente implementadas en sistemas de visión por computadora, se mostraron robustas al combinarse con técnicas de optimización en cuanto a la generación de candidatos se refiere.

Sobre la etapa de la optimización de los parámetros, se encontró una convergencia prematura en las primeras ejecuciones del algoritmo. Se sintonizó y se optó por favorecer una convergencia rápida debido a los altos tiempos de ejecución. Acerca de la función de aptitud, la que presentó la mejor solución fue la de ponderar el porcentaje de región de microaneurisma detectado verdaderamente, más el porcentaje de región verdadera negativa detectada como no microaneurisma. El problema con la función de aptitud fue que pequeñas variaciones en aptitud presentaban enormes diferencias en la generación asertiva de candidatos, por lo que en trabajos futuros se recomienda investigar y pro-

poner nuevas funciones de aptitud. La optimización sobre el filtro gaussiano encontró mejores soluciones con respecto a las encontradas por el algoritmo cuando se optimizó sobre el filtro mediano. La solución que generó el mayor número de candidatos logró una sensibilidad del 84.14%.

La mayoría de las metodologías encontradas en el estado del arte utilizan diferentes bases de datos, de manera que no es posible realizar una comparación justa de la metodología propuesta con éstos. El estado del arte de la detección automática de microaneurismas presenta metodologías con resultados aceptables, el problema es que se suelen utilizar bases de datos de imágenes homogéneas, tomadas bajo las mismas condiciones, e.g. iluminación, cámara, campo de visión por lo que al operar con imágenes heterogéneas adquiridas mediante diferentes cámaras y condiciones de iluminación; disminuye las tasas de detección de microaneurismas. De igual manera, en el presente trabajo se operó sobre imágenes homogéneas y falta validar la metodología con imágenes heterogéneas.

Es imprescindible contar con sistemas de detección de retinopatía diabética en el primer nivel de atención a la salud, con el objetivo de identificar la complicación en fases tempranas, detener su evolución y evitar la pérdida de la visión. Se considera que con la metodología presentada puede formularse un protocolo de evaluación y captura de imágenes de fondo de ojo. Esto implica establecer una cámara, capacitar personal técnico para obtener imágenes de la retina lo más homogéneas posibles; y evaluar pacientes en el primer nivel de atención a gran escala. Se requiere además generar bases de datos de las imágenes y generar plantillas con ayuda de oftalmólogos que sirvan como datos “ground truth” para que con la metodología propuesta se obtengan parámetros óptimos y los modelos de clasificación con mejores tasas de detección. Esto sería un primer paso hacia un sistema de tele-oftalmología operativo para la detección automática de microaneurismas.

6.3. Aportaciones

Las principales aportaciones de este trabajo de investigación para el desarrollo de futuros temas de tesis y los grupos de trabajo de visión son:

- Desarrollo de una metodología robusta para la detección de lesiones oftálmicas utilizando imágenes homogéneas.
- Se demostró que es posible mejorar los resultados en la etapa de preprocesamiento de las imágenes de fondo de ojo con técnicas de optimización basadas en poblaciones de soluciones. Se considera que la metodología resulta un esquema modular, en el que es posible cambiar las técnicas de procesamiento digital de imágenes, la técnica de optimización, los modelos de clasificación.
- Se utilizaron los datos “ground truth” de los cuatro oftalmólogos proporcionados por DIARETDB1 para la creación de 14 plantillas a partir de 14 imágenes de la base de datos DIARETDB1 con únicamente los microaneurismas, y 58 plantillas en las que cada una de éstas muestra regiones ovaladas delimitando los microaneurismas para facilitar la evaluación visual y contabilización de los éstos. Las plantillas serán proporcionadas al CICESE para futuros trabajos de investigación sobre esta línea.
- Un algoritmo genético sintonizable y funcional para la optimización de parámetros de preprocesamiento, con función de aptitud e imágenes a optimizar modificables.
- Una comparación de distintos modelos de clasificación utilizando WEKA.
- Un programa semi-automatizado para contabilizar microaneurismas hasta la etapa de segmentación, y para generar conjuntos de entrenamiento y prueba.

6.4. Trabajo a futuro

Para el desarrollo de futuros temas de investigación sobre esta línea, es recomendable considerar y enfocarse en lo siguiente:

- Ampliar el banco de imágenes.
- Desarrollar una base de datos e interfaz gráfica para delimitar manualmente las lesiones, el tipo de lesión y generar imágenes “ground truth” y plantillas, que sirvan para optimizar la segmentación y para evaluar de manera visual y con menor dificultad las etapas de segmentación y de detección.

- Proponer funciones de aptitud diferentes y determinar cuál presenta mejores resultados en relación a la exactitud de la segmentación y el tiempo de ejecución del algoritmo.
- Investigar, ampliar y/o probar nuevas técnicas en la etapa de procesamiento digital de imágenes. Es altamente recomendable implementar una etapa de reducción de candidatos.
- Trabajar conjuntamente con una clínica de oftalmología para probar la metodología propuesta en un ambiente médico real.
- Siguiendo la metodología propuesta, robustecer el análisis de imágenes y validar el esquema para imágenes heterogéneas.

Literatura citada

- AAO: American Academy of Ophthalmology (2015). Eye health statistics.
- ADA: American Diabetes Association (2015). Standards of medical care in diabetes—2016: Summary of revisions. *Diabetes Care*, **39**(Supplement 1): S4–S5.
- Adal, K. M., Sidibé, D., Ali, S., Chaum, E., Karnowski, T. P., y Mériaudeau, F. (2014). Automated detection of microaneurysms using scale-adapted blob analysis and semi-supervised learning. *Computer methods and programs in biomedicine*, **114**(1): 1–10.
- Adán, E. A. (2012). *Metodología computacional para la segmentación de ultrasonografías de mama basada en la red neuronal pulso acoplada*. Tesis de maestría, CINVESTAV.
- Agoston, M. (2005). *Computer Graphics and Geometric Modelling*. Computer Graphics and Geometric Modeling. Springer.
- Amin, J., Sharif, M., y Yasmin, M. (2016). A review on recent developments for detection of diabetic retinopathy. *Scientifica*, p. 20.
- Anoraganingrum, D. (1999). Cell segmentation with median filter and mathematical morphology operation. En: *Proceedings 10th International Conference on Image Analysis and Processing*. pp. 1043–1046.
- Aravind, C., Ponnibala, M., y Vijayachitra, S. (2013). Automatic detection of microaneurysms and classification of diabetic retinopathy images using svm technique. En: *IJCA Proceedings on international conference on innovations in intelligent instrumentation, optimization and electrical sciences ICIIIOES (11)*. pp. 18–22.
- Arenas-Cavalli, J. T., Ríos, S. A., Pola, M., y Donoso, R. (2015). A web-based platform for automated diabetic retinopathy screening. *Procedia Computer Science*, **60**: 557 – 563. Knowledge-Based and Intelligent Information amp; Engineering Systems 19th Annual Conference, KES-2015, Singapore, September 2015 Proceedings.
- Arora, R. (2015). *Optimization: Algorithms and Applications*. CRC Press.
- Arroyo, M. (2015). *Estadificación de retinopatía diabética en el servicio de oftalmología del hospital general de Querétaro durante septiembre-diciembre del 2014*. Tesis de doctorado, Universidad Autónoma de Querétaro.
- Bae, J. P., Kim, K. G., Kang, H. C., Jeong, C. B., Park, K. H., y Hwang, J.-M. (2011). A study on hemorrhage detection using hybrid method in fundus images. *Journal of digital imaging*, **24**(3): 394–404.
- Baeck, T., Fogel, D., y Michalewicz, Z. (2000). *Evolutionary Computation 1: Basic Algorithms and Operators*. Basic algorithms and operators. Taylor & Francis.
- Bandyopadhyay, S. y Pal, S. K. (2007). *Classification and learning using genetic algorithms : applications in bioinformatics and web intelligence*. Natural computing series. Springer. Berlin, New York.
- Benzarti, F. y Amiri, H. (2013). Speckle noise reduction in medical ultrasound images. *arXiv preprint arXiv:1305.1344*.

- Bhalerao, A., Patanaik, A., Anand, S., y Saravanan, P. (2008). Robust detection of microaneurysms for sight threatening retinopathy screening. En: *Computer Vision, Graphics & Image Processing, 2008. ICVGIP'08. Sixth Indian Conference on*. IEEE, pp. 520–527.
- Bhattacharyya, S. y Maulik, U. (2013). *Soft Computing for Image and Multimedia Data Processing*. Springer Science+Business Media. Springer Publishing Company, Incorporated.
- Chambers, L. (1998). *Practical Handbook of Genetic Algorithms: Complex Coding Systems*. Número v.3 en: *Practical Handbook of Genetic Algorithms*. CRC-Press.
- Chia, D. y Yap, E. (2004). Comparison of the effectiveness of detecting diabetic eye disease: diabetic retinal photography versus ophthalmic consultation. *Singapore medical journal*, **45**: 276–279.
- Coleman, A. L., Staff, A., Emptage, N. P., Mizuiri, D., Kealey, S., Lum, F. C., Regillo, C. D., y Hyman, L. (2016). American academy of ophthalmology retina/vitreous panel. preferred practice pattern ® guidelines. diabetic retinopathy.
- Cree, M. J. (2008). The waikato microaneurysm detector. *the University of Waikato, Tech. Rep.*
- Davis, L. (1991). *Handbook of genetic algorithms*. VNR computer library. Van Nostrand Reinhold.
- Deserno, T. (2011). *Biomedical Image Processing*. Biological and Medical Physics, Biomedical Engineering. Springer Berlin Heidelberg.
- Dr. Ariel Prado-Serrano, Dra Marilu Anahí Guido-Jiménez, D. J. T. C.-B. (2009). Prevalencia de retinopatía diabética en población mexicana. pp. 261–277.
- Eiben, A. (1999). Experimental results on the effects of multi-parent recombination: an overview. *Practical Handbook of Genetic Algorithms Complex Coding Systems*, **3**.
- Eiben, A. E. y Smith, J. E. (2003). *Introduction to Evolutionary Computing*. SpringerVerlag.
- Eiben, A. E., van Kemenade, C. H. M., y Kok, J. N. (1995). *Orgy in the computer: Multi-parent reproduction in genetic algorithms*, pp. 934–945. Springer Berlin Heidelberg, Berlin, Heidelberg.
- Elizondo, J. J. E. y Maestre, L. E. P. (2005). *Fundamentos para el procesamiento de imágenes*. UABC.
- FID: Federación Internacional de la Diabetes (2011). *Atlas de la diabetes de la FID*. Federación Internacional de la Diabetes.
- FID: Federación Internacional de la Diabetes (2013). *Atlas de la diabetes de la FID*. Federación Internacional de la Diabetes.
- Fleming, A. D., Philip, S., Goatman, K. A., Williams, G. J., Olson, J. A., y Sharp, P. F. (2007). Automated detection of exudates for diabetic retinopathy screening. *Physics in medicine and biology*, **52**(24): 7385–96.

- Gen, M. y Cheng, R. (1997). *Genetic Algorithms and Engineering Design*. A Wiley Interscience publication. Wiley.
- Gen, M. y Cheng, R. (2000). *Genetic Algorithms and Engineering Optimization*. A Wiley-Interscience publication. Wiley.
- Goldberg, D. E. (1989). *Genetic Algorithms in Search, Optimization and Machine Learning*. Addison-Wesley Longman Publishing Co., Inc., primera edición. Boston, MA, USA.
- Goldberg, D. E. y Deb, K. (1991). A comparative analysis of selection schemes used in genetic algorithms. En: *Foundations of Genetic Algorithms*. Morgan Kaufmann, pp. 69–93.
- Goldberg, D. E. y Lingle, Jr., R. (1985). Alleles, loci and the traveling salesman problem. En: *Proceedings of the 1st International Conference on Genetic Algorithms*, Hillsdale, NJ, USA. L. Erlbaum Associates Inc., pp. 154–159.
- Gonzalez, R. y Woods, R. (2011). *Digital Image Processing*. Pearson Education.
- Gutiérrez, J., Rivera-Dommarco, J., Villalpando Hernández, S., Franco, A., y Cuevas-Nasu, L. (2012). Encuesta nacional de salud y nutrición: Resultados nacionales 2012.
- Hipwell, J. H., Strachan, F., Olson, J. A., McHardy, K. C., Sharp, P. F., y Forrester, J. V. (2000). Automated detection of microaneurysms in digital red-free photographs: a diabetic retinopathy screening tool. *Diabetic Medicine*, **17**(8): 588–594.
- Horton, M. B., Silva, P. S., Cavallerano, J. D., y Aiello, L. P. (2016). Clinical components of telemedicine programs for diabetic retinopathy. *Current Diabetes Reports*, **16**(12): 129.
- ICO: International Council of Ophthalmology (2012). Number of ophthalmologists in practice and training worldwide.
- IMSS: Instituto Mexicano del Seguro Social (2014). *Tratamiento de la Diabetes Mellitus tipo 2 en el primer nivel de atención*. Instituto Mexicano del Seguro Social.
- Jähne, B. (2005). *Digital Image Processing*. Número v. 1 en: Digital Image Processing. Springer Berlin Heidelberg.
- Kande, G. B., Savithri, T. S., y Subbaiah, P. V. (2010). Automatic detection of microaneurysms and hemorrhages in digital fundus images. *Journal of digital imaging*, **23**(4): 430–437.
- Kapur, T., Grimson, W. L., III, W. M. W., y Kikinis, R. (1996). Segmentation of brain tissue from magnetic resonance images. *Medical Image Analysis*, **1**(2): 109 – 127.
- Kauppi, T., Kalesnykiene, V., Kamarainen, J.-K., Lensu, L., Sorri, I., Raninen, A., Voutilainen, R., Uusitalo, H., Kälviäinen, H., y Pietilä, J. (2007). The diaretdb1 diabetic retinopathy database and evaluation protocol. En: *BMVC*. pp. 1–10.
- Lazar, I., Hajdu, A., y Quareshi, R. J. (2010). Retinal microaneurysm detection based on intensity profile analysis. En: *8th International Conference on Applied Informatics*.

- Liepins, G. E. y Vose, M. D. (1991). Deceptiveness and genetic algorithm dynamics. En: G. J. E. Rawlins (ed.), *Foundations of Genetic Algorithms*. Morgan Kaufmann Publishers.
- Lima Gómez, V. (2006). Retinopatía diabética simplificada: la escala clínica internacional. *Revista del Hospital Juárez de México*, **73**(4): 170–174.
- Lira, J. (2002). *Introducción al tratamiento digital de imágenes*. Fondo de Cultura Económica.
- Man, K. F., Tang, K. S., y Kwong, S. (1996). Genetic algorithms: concepts and applications [in engineering design]. *Industrial Electronics, IEEE Transactions on*, **43**(5): 519–534.
- Manzano, G. (2004). Fotografía de fondo de ojo con filtros.
- MathWorks (2005). *Genetic Algorithm and Direct Search Toolbox: For Use with MATLAB : Computation, Visualization, Programming : User's Guide, Version 2*. The MathWorks.
- Michalewicz, Z. (1994). *Genetic Algorithms + Data Structures = Evolution Programs (2Nd, Extended Ed.)*. Springer-Verlag New York, Inc. New York, NY, USA.
- Mitchell, M. (1999). *An introduction to genetic algorithms*. First MIT Press, quinta edición.
- Niemeijer, M., van Ginneken, B., Staal, J., Suttorp-Schulten, M. S. A., y Abramoff, M. D. (2005). Automatic detection of red lesions in digital color fundus photographs. *IEEE Transactions on Medical Imaging*, **24**(5): 584–592.
- Niemeijer, M., van Ginneken, B., Russell, S. R., Suttorp-Schulten, M. S., y Abramoff, M. D. (2007). Automated detection and differentiation of drusen, exudates, and cotton-wool spots in digital color fundus photographs for early diagnosis of diabetic retinopathy iovs-06-0996 accepted version.
- Nixon, M. S. y Aguado, A. S. (2008). *Feature extraction and image processing*. Academic Press, segunda edición.
- O’Gorman, L., Sammon, M., y Seul, M. (2008). *Practical Algorithms for Image Analysis with CD-ROM*. Cambridge University Press.
- Oliver, I. M., Smith, D. J., y Holland, J. R. C. (1987). A study of permutation crossover operators on the traveling salesman problem. En: *Proceedings of the Second International Conference on Genetic Algorithms on Genetic Algorithms and Their Application*, Hillsdale, NJ, USA. L. Erlbaum Associates Inc., pp. 224–230.
- Patidar, P., Gupta, M., Srivastava, S., y Nagawat, A. K. (2010). Image de-noising by various filters for different noise. *International Journal of Computer Applications (0975–8887) Volume*.
- Pizer, S. M., Amburn, E. P., Austin, J. D., Cromartie, R., Geselowitz, A., Greer, T., ter Haar Romeny, B., Zimmerman, J. B., y Zuiderveld, K. (1987). Adaptive histogram equalization and its variations. *Computer vision, graphics, and image processing*, **39**(3): 355–368.

- Plataniotis, K. y Venetsanopoulos, A. (2000). *Color Image Processing and Applications*. Digital Signal Processing. Springer Berlin Heidelberg.
- Randall, M. (1995). The future and applications of genetic algorithms. En: *Proceedings of the Electronic Directions to the Year 2000 Conference, IEEE Computer*. Society Press, pp. 471–475.
- Riordan-Eva, P. (2012). *Vaughan y Asbury: oftalmología general (18a. ed.)*.
- Romero, R. R., Carballido, J. M., Capistrán, J. H., y Valencia, L. U. (2015). A method to assist in the diagnosis of early diabetic retinopathy: Image processing applied to detection of microaneurysms in fundus images. *Computerized Medical Imaging and Graphics*, **44**(68): 41–53.
- Sender Palacios, M. J. (2007). Exploración oftalmológica básica del paciente diabético en la atención primaria. *Claves de diabetología*, **3**: 24.
- Sivanandam, S. N. y Deepa, S. N. (2008). *Introduction to genetic algorithms*. Springer Science+Business Media. Springer Berlin Heidelberg New York.
- Sopharak, A., Uyyanonvara, B., Barman, S., y Williamson, T. (2011). Automatic microaneurysm detection from non-dilated diabetic retinopathy retinal images. En: *Proceedings of the World Congress on Engineering*. Vol. 2, pp. 6–8.
- Sopharak, A., Uyyanonvara, B., y Barman, S. (2013). Simple hybrid method for fine microaneurysm detection from non-dilated diabetic retinopathy retinal images. *Computerized medical imaging and graphics : the official journal of the Computerized Medical Imaging Society*, **37**(5-6): 394–402.
- Sorrentino, F. S., Allkabes, M., Salsini, G., Bonifazzi, C., y Perri, P. (2016). The importance of glial cells in the homeostasis of the retinal microenvironment and their pivotal role in the course of diabetic retinopathy. *Life Sciences*, **162**: 54 – 59.
- SSA: Secretaría de Salud (2013). *Diagnóstico y Tratamiento de Diabetes Mellitus en el Adulto Mayor Vulnerable*. CENETEC.
- SSA: Secretaría de Salud (2014). *Detección de retinopatía diabética en primer nivel de atención*. Número IMSS-735-14. CENETEC.
- Streeter, L. y Cree, M. J. (2003). Microaneurysm detection in colour fundus images. *Image Vision Comput. New Zealand*, pp. 280–284.
- SujithKumar, S. y Singh, V. (2012). Automatic detection of diabetic retinopathy in non-dilated rgb retinal fundus images. *International Journal of Computer Applications*, **47**(19).
- Tran, K., Mendel, T. A., Holbrook, K. L., y Yates, P. A. (2012). Construction of an inexpensive, hand-held fundus camera through modification of a consumer “point-and-shoot” camerapoint-and-shoot hand-held fundus camera. *Investigative ophthalmology & visual science*, **53**(12): 7600–7607.

- Walter, T., Massin, P., Erginay, A., Ordonez, R., Jeulin, C., y Klein, J.-C. (2007). Automatic detection of microaneurysms in color fundus images. *Medical image analysis*, **11**(6): 555–66.
- Whitley, D. (1994). A genetic algorithm tutorial. *Statistics and Computing*, **4**: 65–85.
- WHO: World Health Organization (2016). *Global report on diabetes*. World Health Organization. pp. 3–8.
- Yogesana, K., Kumar, S., Goldschmidt, L., y Cuadros, J. (2008). *Teleophthalmology*. Springer.
- Zana, F. y Klein, J. C. (2001). Segmentation of vessel-like patterns using mathematical morphology and curvature evaluation. *IEEE Transactions on Image Processing*, **10**(7): 1010–1019.

Anexos

A.1. Banco de imágenes

Tabla A.1: Imágenes utilizadas para los experimentos con su respectivo número de microaneurismas de acuerdo a los datos ground truth.

Imagen	# de MA's	Nombre	# de MA's
image001	25	image057	0
image002	14	image058	2
image003	39	image059	3
image004	27	image060	0
image017	8	image061	3
image023	7	image062	0
image024	6	image063	4
image027	2	image068	2
image028	0	image069	1
image029	3	image070	3
image030	3	image071	2
image031	4	image072	0
image032	4	image073	1
image033	0	image074	1
image034	1	image075	1
image035	4	image076	1
image036	2	image077	2
image037	3	image078	4
image038	9	image079	3
image039	3	image080	1
image040	2	image081	1
image045	2	image082	3
image046	1	image083	3
image047	0	image084	0

Tabla A.1: Continuación...

Imagen	# de MA's	Nombre	# de MA's
image048	2	image085	5
image049	0	image086	6
image050	2	image087	7
image051	2	image088	2
image056	5	image089	5
Microaneurismas totales			246

A.2. Imágenes de control

Tabla A.2: Imágenes de control de la segmentación del experimento 1.

Plantilla propuesta	Nombre	# de MA's
	image017	8

Tabla A.2: Continuación...

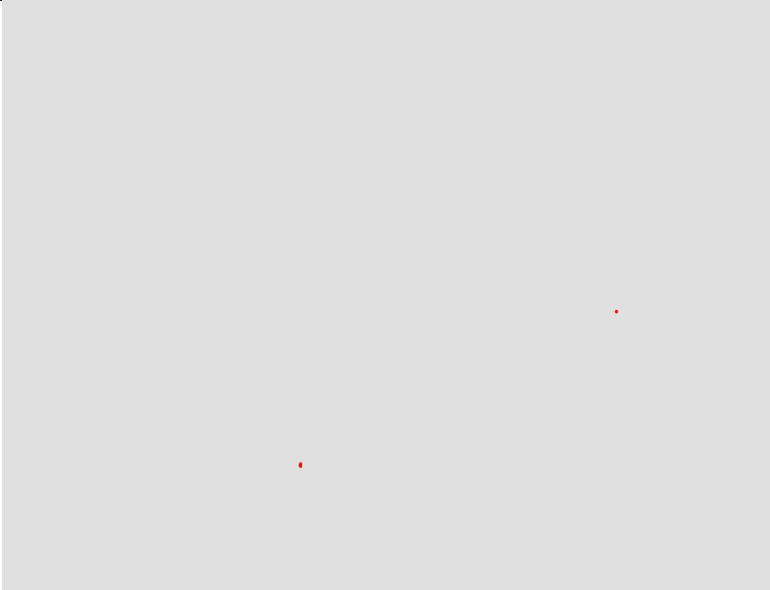
Plantilla propuesta	Nombre	# de MA's
	image027	2
	image031	4

Tabla A.2: Continuación...

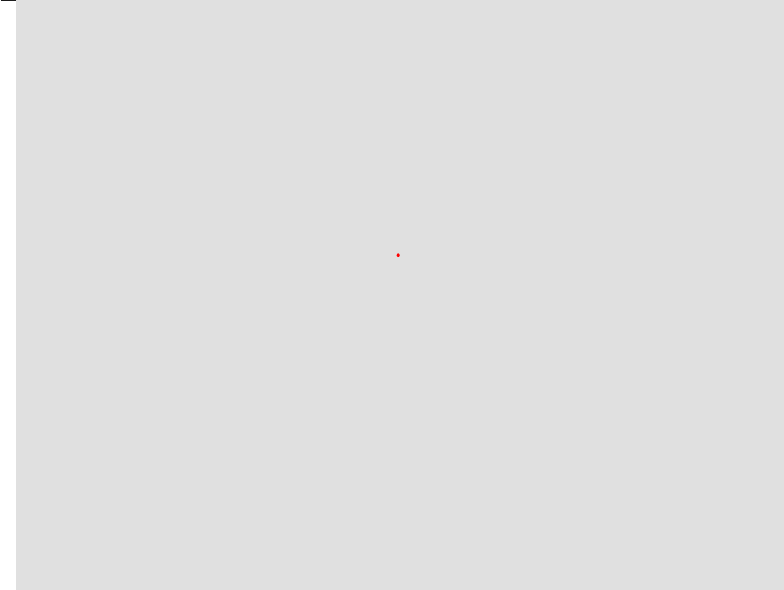
Plantilla propuesta	Nombre	# de MA's
	image034	1
	image035	4

Tabla A.2: Continuación...

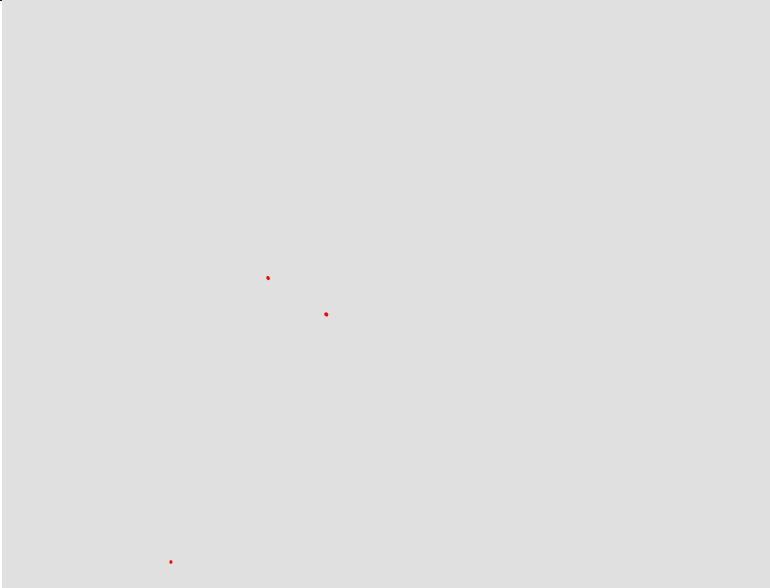
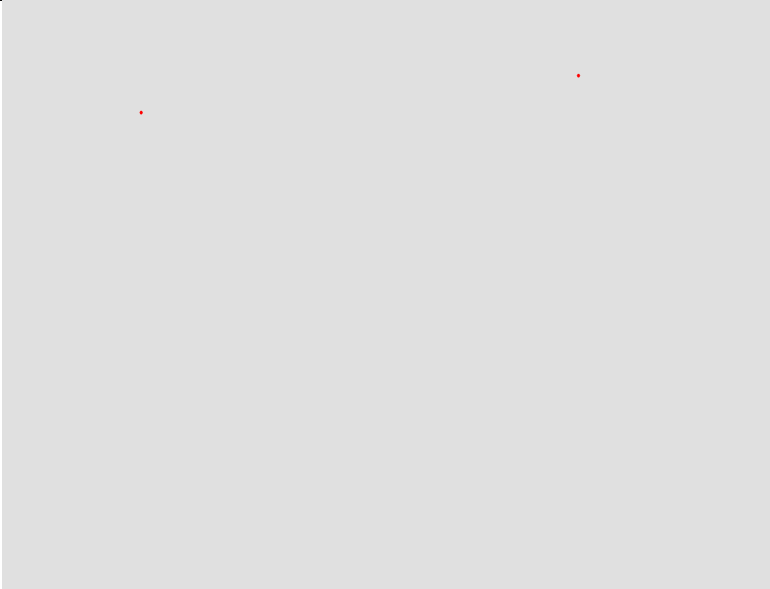
Plantilla propuesta	Nombre	# de MA's
	image039	3
	image040	2

Tabla A.2: Continuación...

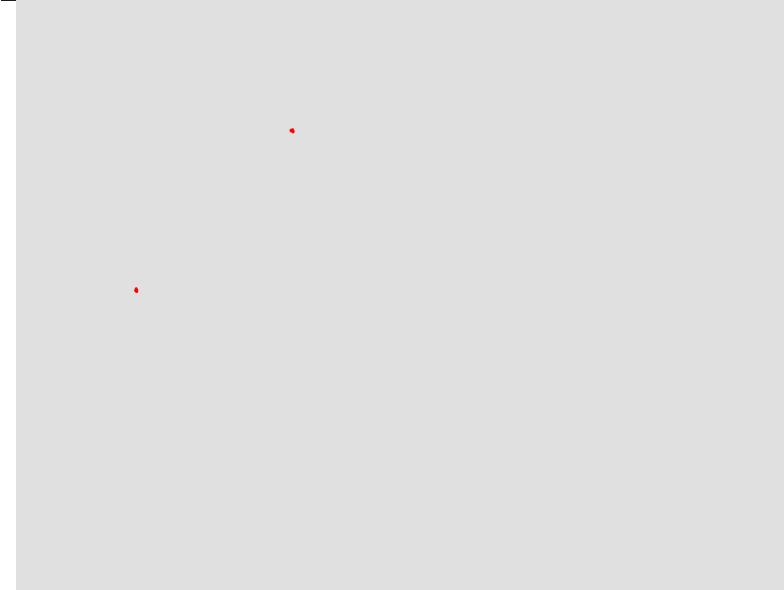
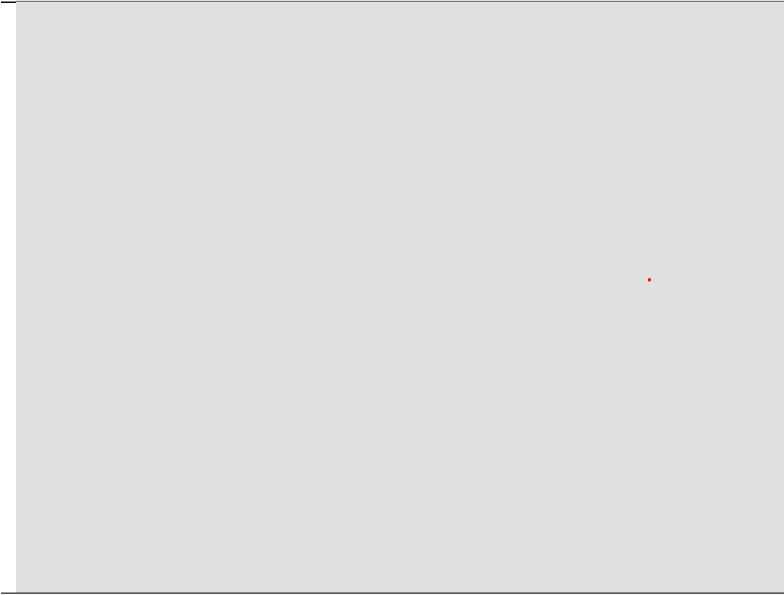
Plantilla propuesta	Nombre	# de MA's
	image045	2
	image046	1

Tabla A.2: Continuación...

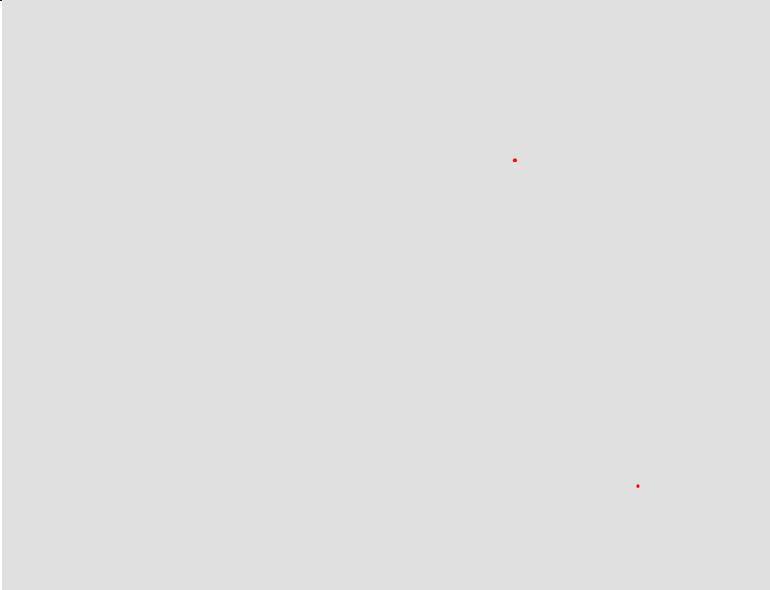
Plantilla propuesta	Nombre	# de MA's
	image048	2
Total de microaneurismas		29

Tabla A.3: Imágenes de control de la segmentación del experimento 2.

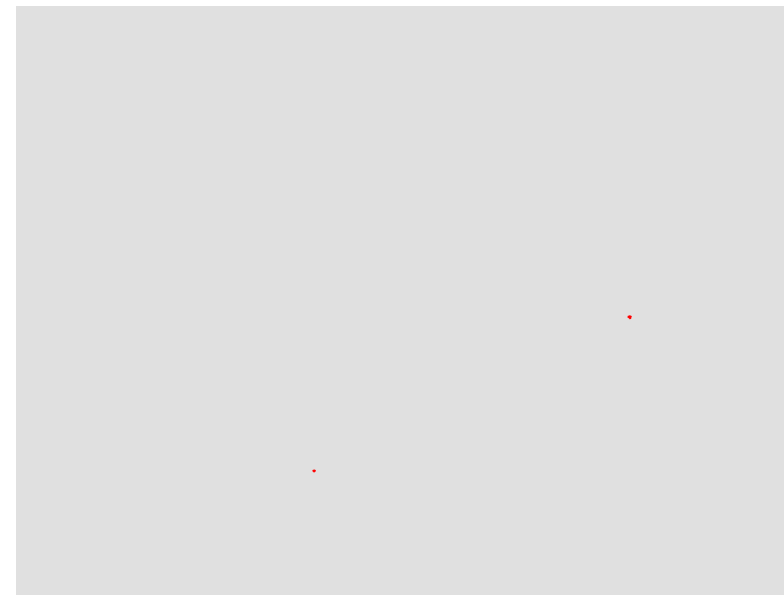
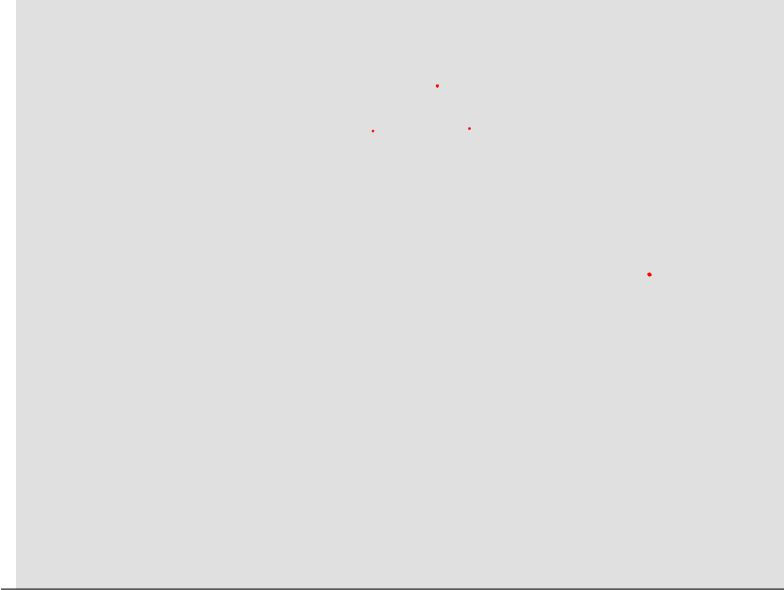
Plantilla propuesta	Nombre	# de MA's
	image027	2
	image031	4

Tabla A.3: Continuación...

Plantilla propuesta	Nombre	# de MA's
	image035	4
	image039	3

Tabla A.3: Continuación...

Plantilla propuesta	Nombre	# de MA's
	image040	2
	image045	2

Tabla A.3: Continuación...

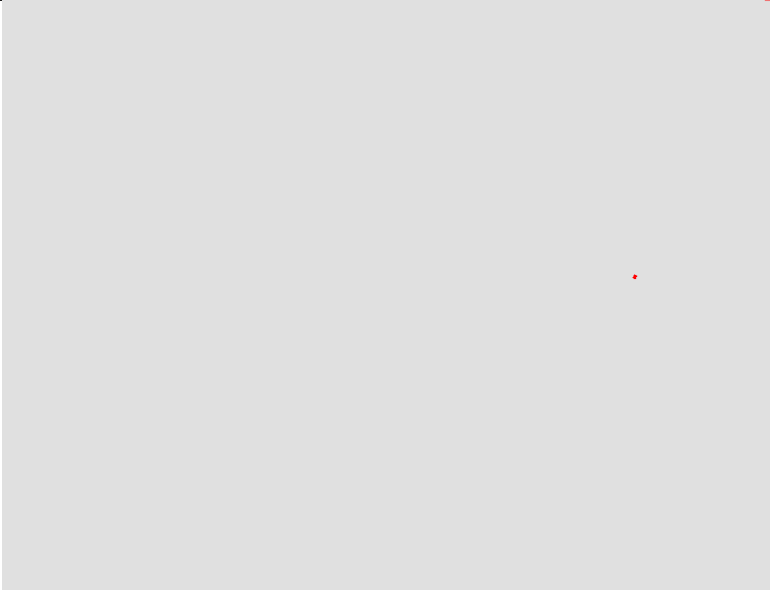
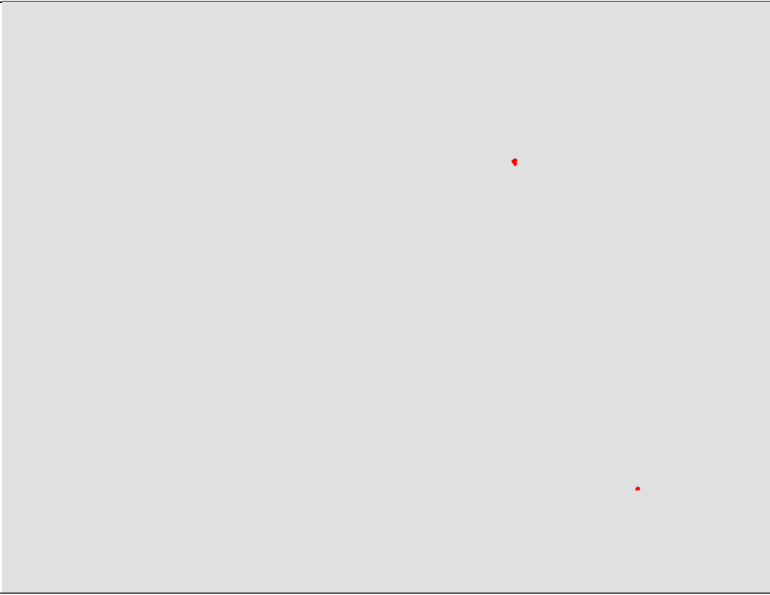
Plantilla propuesta	Nombre	# de MA's
	image046	1
	image048	2

Tabla A.3: Continuación...

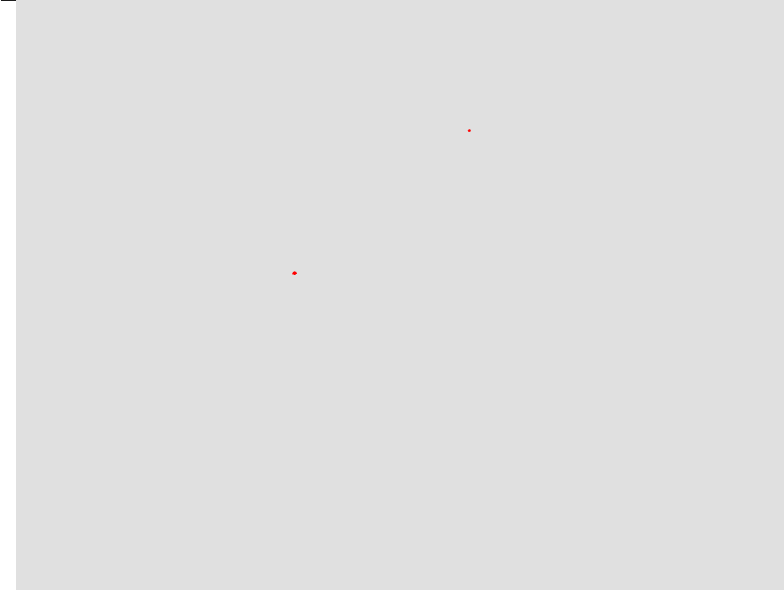
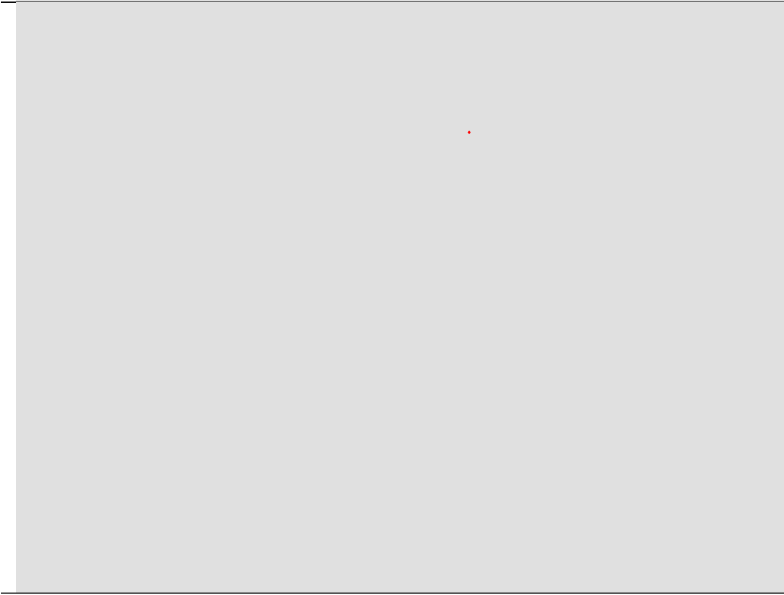
Plantilla propuesta	Nombre	# de MA's
	image058	2
	image074	1

Tabla A.3: Continuación...

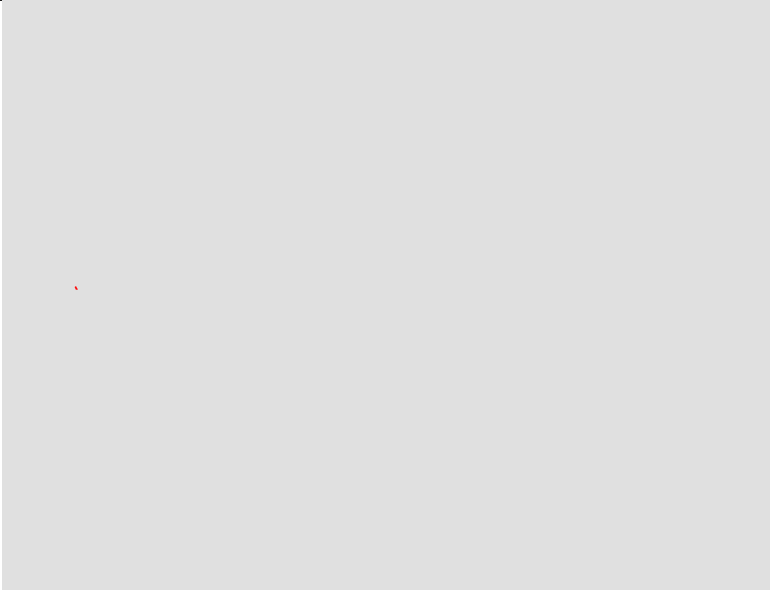
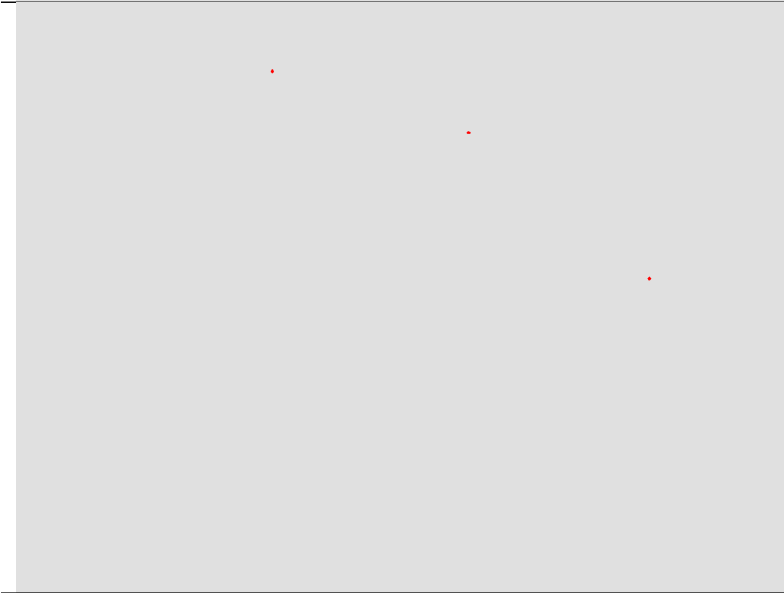
Plantilla propuesta	Nombre	# de MA's
	image075	1
	image076	1

Tabla A.3: Continuación...

Plantilla propuesta	Nombre	# de MA's
	image081	1
	image082	3
Total de microaneurismas		29

A.3. Experimento 3: Gráficas

A.3.1. Experimento 3a

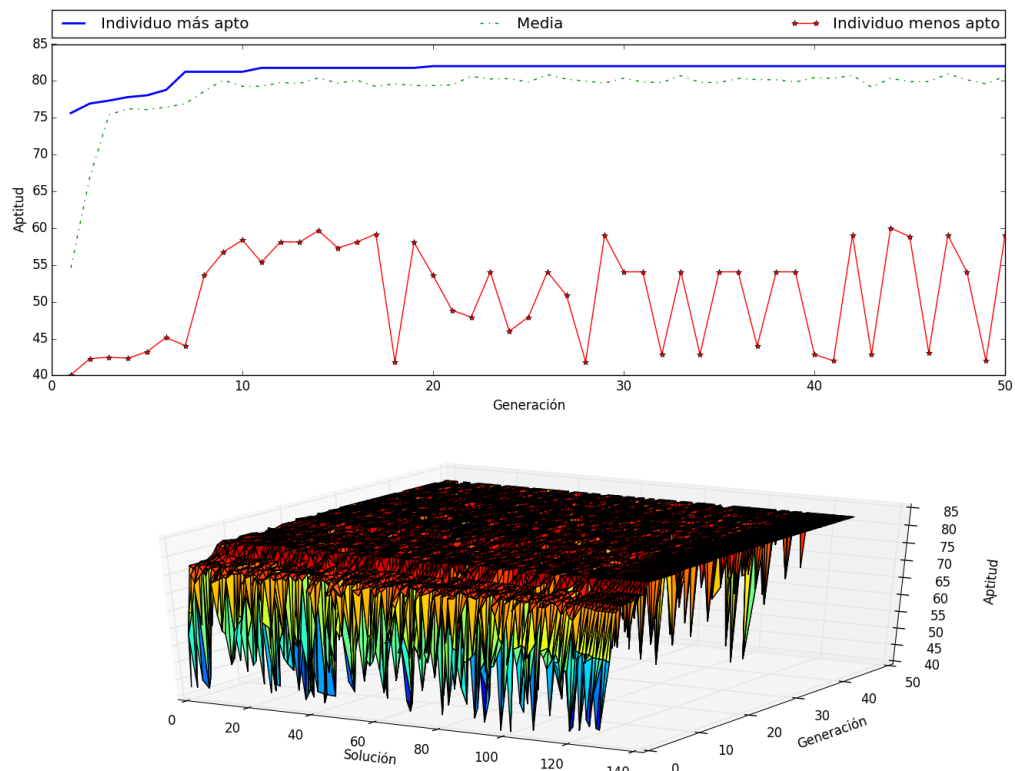


Figura A.1: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.

A.3.2. Experimento 3c

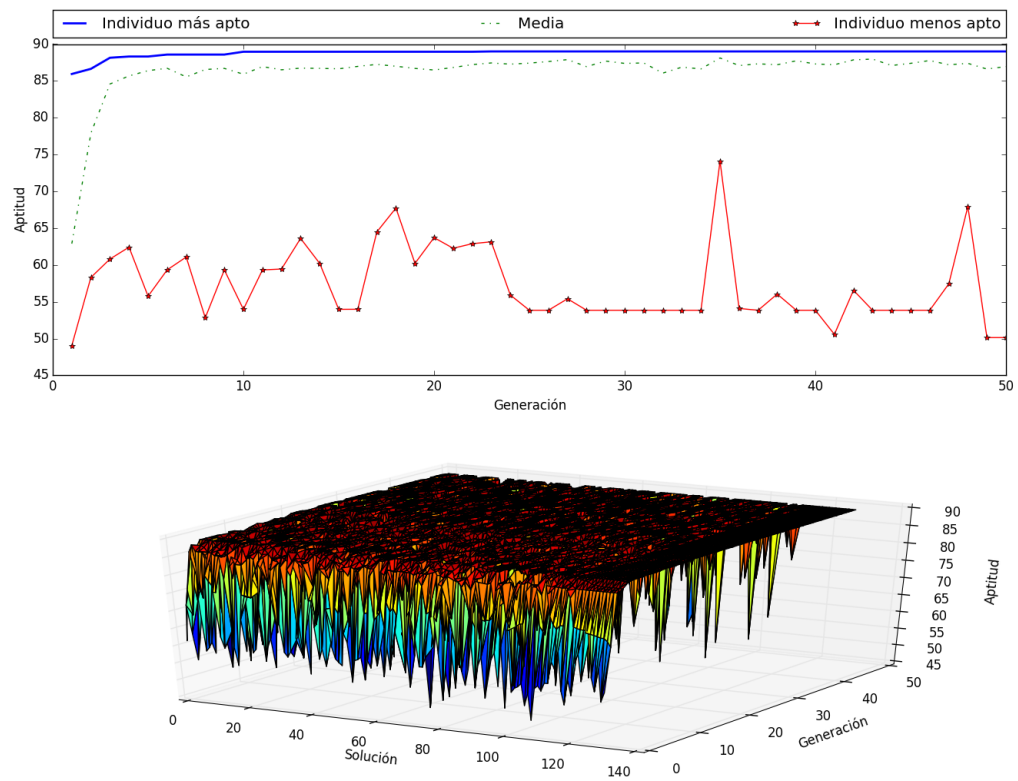


Figura A.2: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.

A.3.3. Experimento 3e

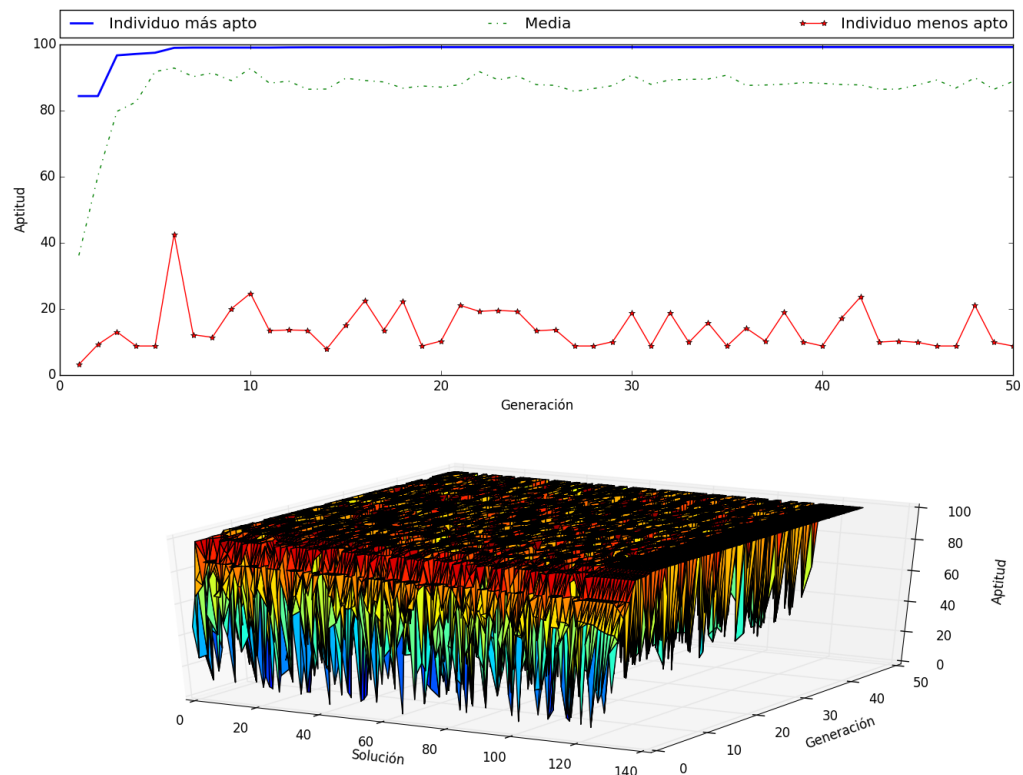


Figura A.3: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.

A.3.4. Experimento 3g

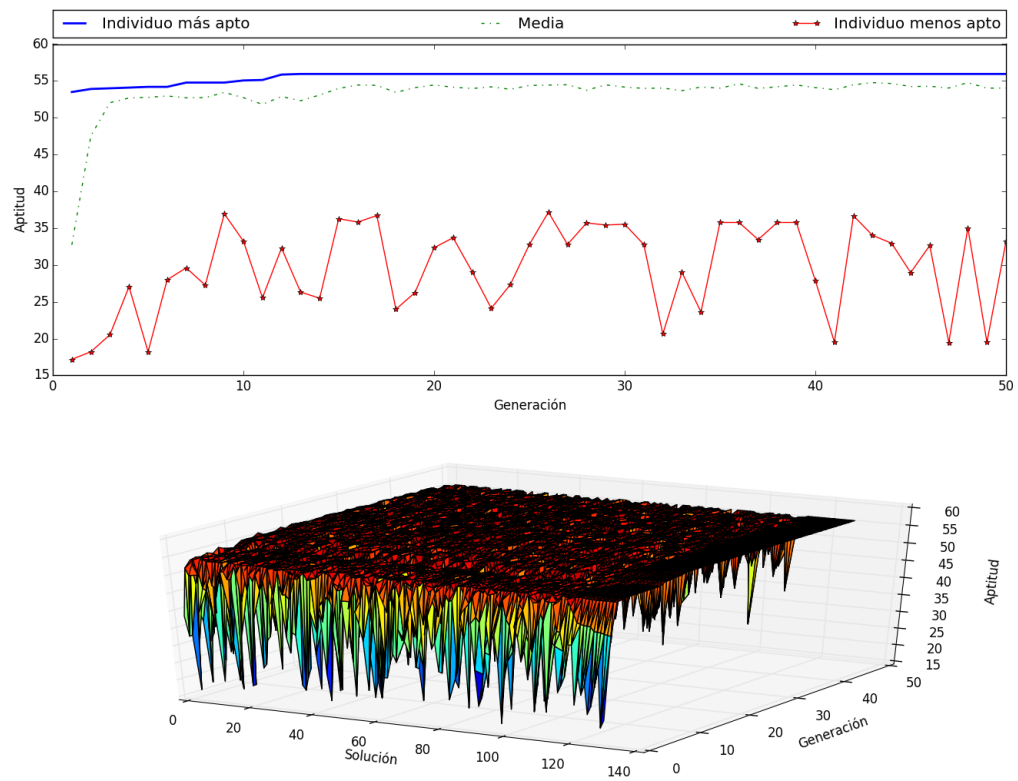


Figura A.4: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.

A.3.5. Experimento 3b

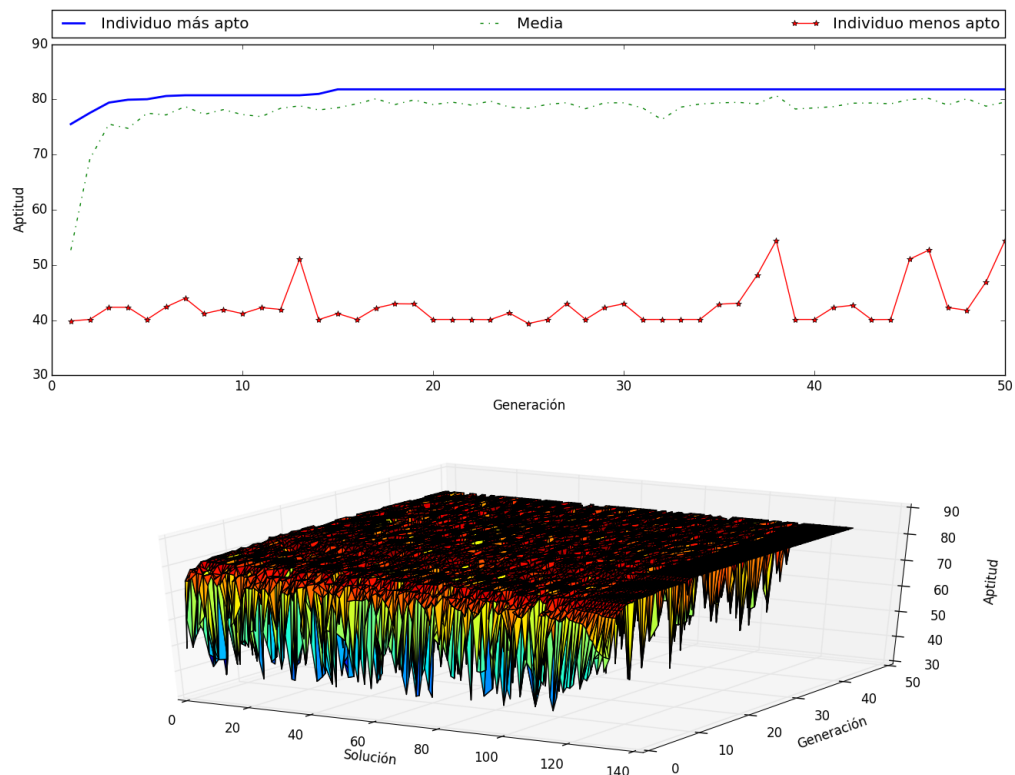


Figura A.5: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.

A.3.6. Experimento 3d

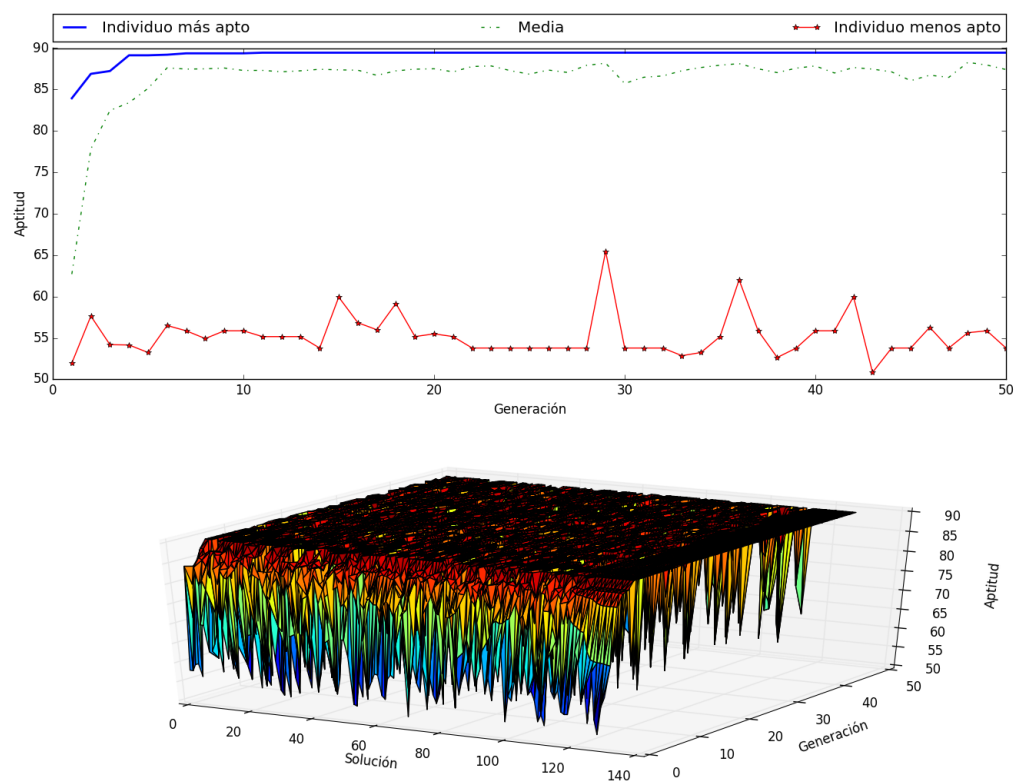


Figura A.6: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.

A.3.7. Experimento 3f

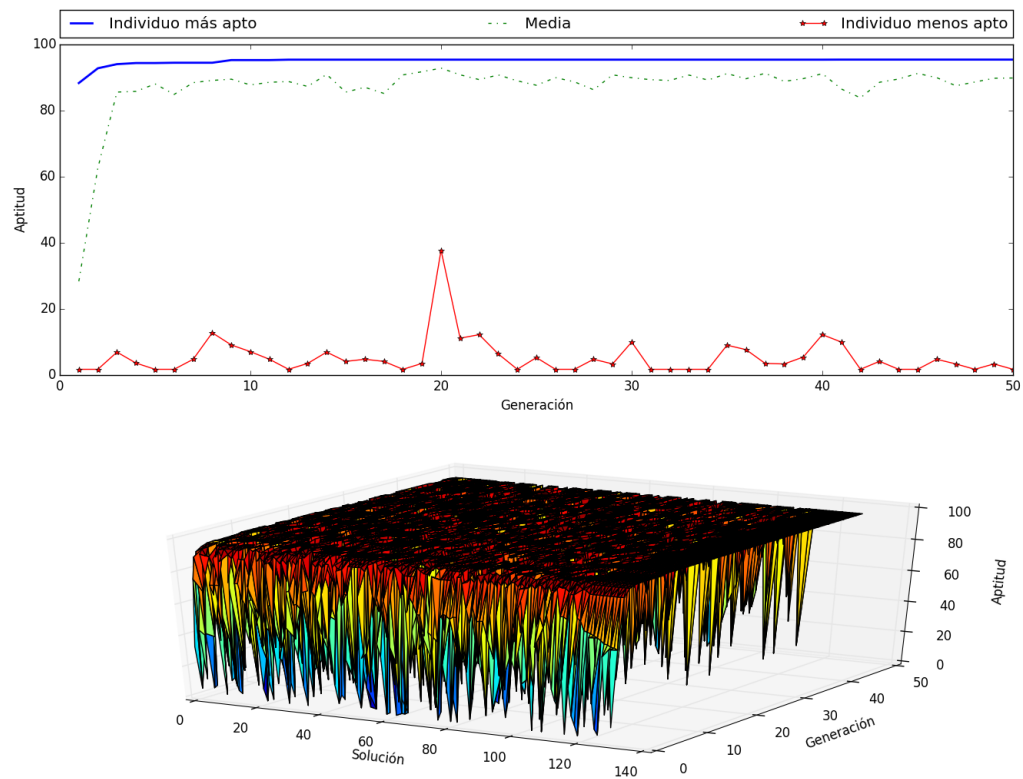


Figura A.7: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.

A.3.8. Experimento 3h

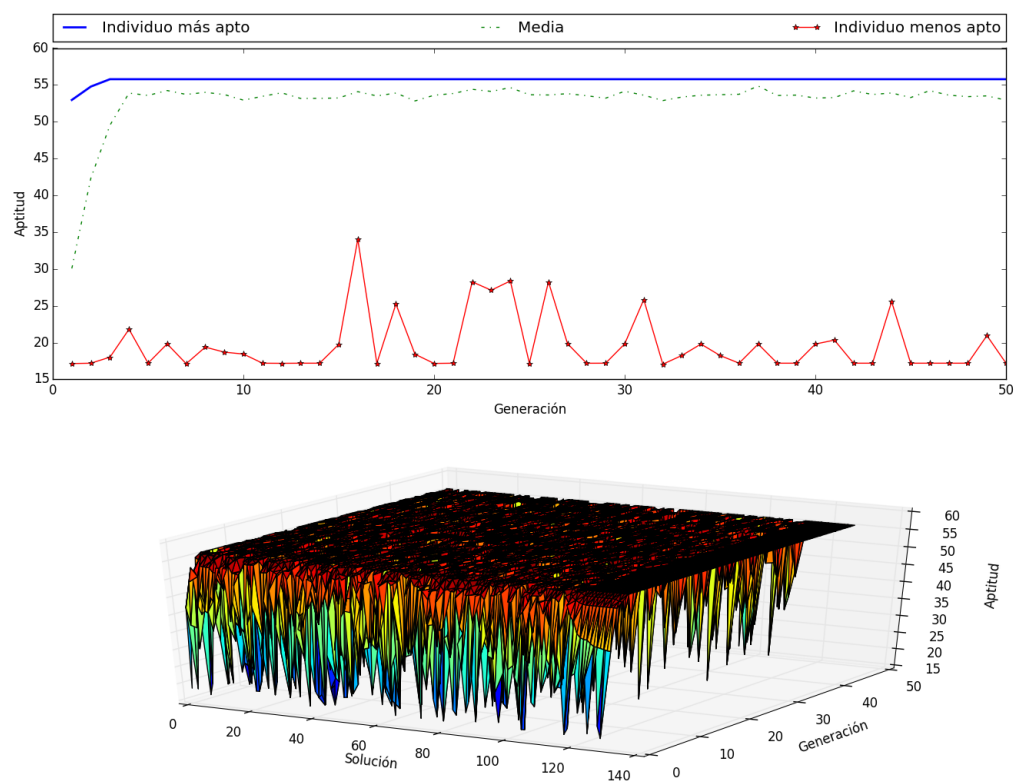


Figura A.8: Imagen superior: Evolución de la población. Imagen inferior: Paisaje de aptitud.