



الجمهورية العربية السورية

جامعة دمشق

كلية الهندسة الميكانيكية والكهربائية

قسم الهندسة الطبية

التنبؤ والكشف المبكر عن مرض الزهايمر باستخدام تعلم الآلة

رسالة مقدمة لنيل درجة الماجستير في الهندسة الطبية

إعداد:

م. جورج رياشي

إشراف:

أ.د.م. رشا مسعود

2025-2024



الجمهورية العربية السورية

جامعة دمشق

كلية الهندسة الميكانيكية والكهربائية

قسم الهندسة الطبية

قرار لجنة الحكم

قدم هذا البحث لمتطلبات الحصول على درجة الماجستير في الهندسة الطبية في قسم الهندسة الطبية - كلية الهندسة الميكانيكية والكهربائية - جامعة دمشق.

لجنة الحكم:

الدكتور	الاختصاص	التوقيع
د. رشا مسعود	تحكم طبي حيوي	
د. أيمن صابوني	الأجهزة المخبرية	
د. وائل الإمام	إحصاء حيوي	

تاريخ المناقشة: في تمام الساعة (11 صباحاً) يوم 14 / 1 / 2024 في قاعة السيمينار
في قسم الهندسة الطبية.

الشكر والتقدير

الشكر والحمد لله على فضله في إنارة العقل، وتجاوز العقبات، وإتاحة إنجاز هذا البحث.

كل الشكر والتقدير لجامعة دمشق العريقة على كل ما قدمته وتقدمه في سبيل التطور والتقدم، وكذلك الشكر والعرفان لكلية الهندسة الميكانيكية والكهربائية، وأخص بالشكر قسم الهندسة الطبية من أعضاء الهيئة التدريسية من أساتذة دكاترة ومهندسين لما قدموه من علم ومعرفة.

الشكر الأكبر والامتنان للمشرفة على هذا البحث الأستاذة الدكتورة رشا مسعود الداعم الأكبر، والأم المعطاء، والتي أغنت البحث بعلمها وتوجيهاتها ومتابعتها لإتمام هذا البحث والارتقاء به للمستوى المطلوب.

الشكر لمبادرة التصوير العصبي لمرض الزهايمر ADNI على السماح للوصول إلى قاعدة بياناتها التي تم العمل عليها في هذا البحث.

الشكر والتقدير للجنة التحكيم الكريمة الدكتور أيمن صابوني والدكتور وائل الإمام على جهودهم في تقييم البحث وتقديم الملاحظات القيمة.

والشكر لكل شخص ساعدنا وكان عوناً لنا في هذا البحث، ولم يبخل علينا بأي معلومة يعرفها.

الإهداء

إلى من عجزت كلماتي عن وصفها فهي أعظم هبات الله، إلى من دعاؤها سر
نجاحي.....

أمي (رينه، ثناء)

إلى من تعلمت منه الاستمرار حتى أصل إلى ما أريد، إلى السند الحقيقي والدعم
الدائم.....

أبي (خليل، إبراهيم)

إلى يدي اليمنى من أشارك معهم ذكريات الطفولة وأحلام الشباب، إلى سندي وقوتي وعوني في
هذه الحياة.....

(بطرس، سارة، بولس)

إلى من تذوقت معهم أجمل اللحظات، إلى من تشاركت معهم الكثير فأصبحوا جزءاً من
عائلتي.....

(إيمي، سارة)

إلى الأصدقاء الذين تعرفت عليهم خلال هذا البحث فكانوا نعم الرفاق وخير السند والدعم.....
(بيان، ثريا، صالح، عمر، بطرس، طلاب الدراسات العليا)

فهرس المحتويات

المحتوى	الصفحة
الفصل الأول: الإطار العام للبحث	
1-1 مقدمة	2
2-1 مشكلة البحث	2
3-1 تساؤلات البحث	3
4-1 فرضيات البحث	3
5-1 مسلمات البحث	3
6-1 هدف البحث	3
7-1 أهمية البحث	4
8-1 مبررات البحث	4
9-1 محدوديات البحث	4
10-1 مخطط البحث	5
11-1 منهج البحث	6
12-1 أقسام الرسالة	6
13-1 خاتمة	7
الفصل الثاني: القسم النظري	
1-2 مقدمة	9
2-2 الزهايمر	9
3-2 التغيرات الدماغية لمرض الزهايمر	11
4-2 تعلم الآلة	13
1-4-2 آلة متجه الدعم	14
2-4-2 بايز البسيط	15
3-4-2 الجار الأقرب	16
4-4-2 شجرة القرار	18

21	 5-4-2 الغاية العشوائية
23	 6-4-2 التعزيز الكيفي
23	 7-4-2 تعزيز التدرج
24	 5-2 مقاييس الأداء
27	 6-2 ضبط المحددات الفائقة
28	 7-2 تقليل السمات
29	 8-2 خاتمة

الفصل الثالث: أدبيات البحث

31	 1-3 مقدمة
31	 2-3 الدراسات المرجعية
31	 1-2-3 استخدام طرق تعلم الآلة من أجل التنبؤ المبكر بالزهايمر والكشف عنه
35	 2-2-3 استخدام السجلات الصحية الالكترونية للكشف عن الزهايمر غير المشخص ..
37	 3-2-3 استخدام التعلم العميق للكشف عن مرض الزهايمر والتنبؤ المبكر به
40	 3-3 خاتمة

الفصل الرابع: منهج البحث وإجراءاته

42	 1-4 مقدمة
42	 2-4 المنهجية
44	 3-4 قاعدة البيانات
44	 4-4 معالجة وتهيئة البيانات
44	 1-4-4 القيم المفقودة
46	 2-4-4 القيم الشاذة
47	 3-4-4 موازنة العينات
48	 4-4-4 الترميز
49	 5-4-4 تقسيم البيانات
50	 5-4 بناء النموذج
50	 1-5-4 آلة متجه الدعم

51	 2-5-4 بايز البسيط
52	 3-5-4 الجار الأقرب
53	 4-5-4 شجرة القرار
54	 5-5-4 الغابة العشوائية
55	 6-5-4 التعزيز الكيفي
56	 7-5-4 تعزيز التدرج
56	 6-4 تدريب النماذج وتقييمها
57	 7-4 تقليل السمات
57	 1-7-4 اختيار K الأفضل
58	 2-7-4 الشبكة المرنة
58	 3-7-4 استبعاد السمات العودية
59	 8-4 خاتمة

الفصل الخامس: النتائج والمناقشة

61	 1-5 مقدمة
61	 2-5 المحددات الفائقة
61	 1-2-5 آلة متجه الدعم
62	 2-2-5 بايز البسيط
63	 3-2-5 الجار الأقرب
63	 4-2-5 شجرة القرار
64	 5-2-5 الغابة العشوائية
65	 6-2-5 التعزيز الكيفي
66	 7-2-5 تعزيز التدرج
67	 3-5 أداء النماذج مع كل السمات
68	 1-3-5 آلة متجه الدعم
71	 2-3-5 بايز البسيط
74	 3-3-5 الجار الأقرب

77	 4-3-5 شجرة القرار
80	 5-3-5 الغابة العشوائية
83	 6-3-5 التعزيز الكيفي
86	 7-3-5 تعزيز التدرج
89	 4-5 أداء النماذج بعد تقليل السمات
89	 1-4-5 مع سمات طريقة اختيار K الأفضل
91	 1-1-4-5 الغابة العشوائية
94	 2-1-4-5 تعزيز التدرج
96	 2-4-5 مع سمات طريقة الشبكة المرنة
98	 1-2-4-5 الغابة العشوائية
100	 2-2-4-5 تعزيز التدرج
103	 3-4-5 مع سمات طريقة استبعاد السمات العودية
104	 1-3-4-5 الغابة العشوائية
107	 2-3-4-5 تعزيز التدرج
109	 5-5 مقارنة بين أداء جميع النماذج قبل وبعد تقليل السمات
111	 6-5 مقارنة مع الدراسات المرجعية
112	 7-5 خاتمة

الفصل السادس: الخاتمة

114	 1-6 مقدمة
114	 2-6 الاستنتاجات
114	 3-6 مميزات البحث
115	 4-6 التوصيات والمقترحات
115	 5-6 خاتمة

فهرس الأشكال

الشكل	الصفحة
الشكل (1-1) استمرارية مرض الزهايمر	2
الشكل (2-1) المخطط العام للبحث	5
الشكل (1-2) مراحل تقدم مرض الزهايمر	10
الشكل (2-2) مقارنة بين a الدماغ السليم و b الدماغ عند مرض الزهايمر	12
الشكل (3-2) آلية عمل خوارزمية SVM	15
الشكل (4-2) آلية عمل خوارزمية KNN	17
الشكل (5-2) آلية عمل خوارزمية DT	20
الشكل (6-2) آلية عمل خوارزمية RF	22
الشكل (7-2) آلية عمل خوارزمية AB	23
الشكل (8-2) آلية عمل خوارزمية GB	24
الشكل (9-2) مصفوفة الارتباك متعددة الأصناف	25
الشكل (10-2) طريقة البحث الشبكي	27
الشكل (11-2) آلية اختيار الخصائص العامة	28
الشكل (1-3) دقة النماذج دراسة Shastri	32
الشكل (2-3) دقة الشبكة العصبونية في دراسة ElZawawi	35
الشكل (3-3) مخطط عمل دراسة Shahbaz	38
الشكل (1-4) منهجية البحث المتبعة	43
الشكل (2-4) القيم المفقودة	45
الشكل (3-4) القيم المتطرفة في بعض المتغيرات	47
الشكل (4-4) عدد العينات في كل صنف بعد الموازنة	48
الشكل (5-4) Label Encoding للخروج DX	49
الشكل (1-5) القيم الأفضل للمحددات الفائقة لنموذج SVM الأفضل	61
الشكل (2-5) القيم الأفضل للمحددات الفائقة لنموذج GNB الأفضل	62

63	الشكل (3-5) القيم الأفضل للمحددات الفائقة لنموذج KNN الأفضل
64	الشكل (4-5) القيم الأفضل للمحددات الفائقة لنموذج DT الأفضل
65	الشكل (5-5) القيم الأفضل للمحددات الفائقة لنموذج RF الأفضل
66	الشكل (6-5) القيم الأفضل للمحددات الفائقة لنموذج AB الأفضل
67	الشكل (7-5) القيم الأفضل للمحددات الفائقة لنموذج GB الأفضل
68	الشكل (8-5) مصفوفة الارتباك للنموذج SVM
69	الشكل (9-5) تقرير التصنيف للنموذج SVM
70	الشكل (10-5) مقاييس الأداء للنموذج SVM
71	الشكل (11-5) مصفوفة الارتباك للنموذج GNB
72	الشكل (12-5) تقرير التصنيف للنموذج GNB
73	الشكل (13-5) مقاييس الأداء للنموذج GNB
74	الشكل (14-5) مصفوفة الارتباك للنموذج KNN
75	الشكل (15-5) تقرير التصنيف للنموذج KNN
76	الشكل (16-5) مقاييس الأداء للنموذج KNN
77	الشكل (17-5) مصفوفة الارتباك للنموذج DT
78	الشكل (18-5) تقرير التصنيف للنموذج DT
79	الشكل (19-5) مقاييس الأداء للنموذج DT
80	الشكل (20-5) مصفوفة الارتباك للنموذج RF
81	الشكل (21-5) تقرير التصنيف للنموذج RF
82	الشكل (22-5) مقاييس الأداء للنموذج RF
83	الشكل (23-5) مصفوفة الارتباك للنموذج AB
84	الشكل (24-5) تقرير التصنيف للنموذج AB
85	الشكل (25-5) مقاييس الأداء للنموذج AB
86	الشكل (26-5) مصفوفة الارتباك للنموذج GB
87	الشكل (27-5) تقرير التصنيف للنموذج GB
88	الشكل (28-5) مقاييس الأداء للنموذج GB

92	 الشكل (29-5) مصفوفة الارتباك للنموذج RF مع سمات SKB
93	 الشكل (30-5) تقرير التصنيف للنموذج RF مع سمات SKB
94	 الشكل (31-5) مصفوفة الارتباك للنموذج GB مع سمات SKB
95	 الشكل (32-5) تقرير التصنيف للنموذج GB مع سمات SKB
98	 الشكل (33-5) مصفوفة الارتباك للنموذج RF مع سمات EN
99	 الشكل (34-5) تقرير التصنيف للنموذج RF مع سمات EN
101	 الشكل (35-5) مصفوفة الارتباك للنموذج GB مع سمات EN
102	 الشكل (36-5) تقرير التصنيف للنموذج GB مع سمات EN
105	 الشكل (37-5) مصفوفة الارتباك للنموذج RF مع سمات RFE
106	 الشكل (38-5) تقرير التصنيف للنموذج RF مع سمات RFE
107	 الشكل (39-5) مصفوفة الارتباك للنموذج GB مع سمات RFE
108	 الشكل (40-5) تقرير التصنيف للنموذج GB مع سمات RFE
109	 الشكل (41-5) قيم الدقة للنماذج قبل وبعد تقليل السمات
110	 الشكل (42-5) قيم الحساسية للنماذج قبل وبعد تقليل السمات
110	 الشكل (43-5) قيم الإحكام للنماذج قبل وبعد تقليل السمات
111	 الشكل (44-5) قيم درجة F1 للنماذج قبل وبعد تقليل السمات

فهرس الجداول

الجدول	الصفحة
الجدول (1-3) دقة وحساسية الخوارزميات دون التحقق المتقاطع في دراسة Dhakal	33
الجدول (2-3) السمات المختارة في دراسة ElZawawi	34
الجدول (3-3) دقة وحساسية كل من هذه الخوارزميات في دراسة Stephan	35
الجدول (4-3) السمات العشرة الأعلى ارتباطاً بالإصابة بالزهايمر في دراسة Stephan ...	36
الجدول (5-3) السمات العشرة الأعلى ارتباطاً بالإصابة بالزهايمر في دراسة Park	37
الجدول (6-3) دقة الخوارزميات المستخدمة في دراسة Shahbaz	39
الجدول (1-4) السمات Features والخرج Label	49
الجدول (2-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج SVM	51
الجدول (3-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج NB	51
الجدول (4-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج KNN	52
الجدول (5-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج DT	53
الجدول (6-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج RF	54
الجدول (7-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج AB	55
الجدول (8-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج GB	56
الجدول (9-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج SKB	57
الجدول (10-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج EN	58
الجدول (11-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج RFE	58
الجدول (1-5) حالات Positive و Negative للنموذج SVM	69
الجدول (2-5) حالات Positive و Negative للنموذج GNB	72
الجدول (3-5) حالات Positive و Negative للنموذج KNN	75
الجدول (4-5) حالات Positive و Negative للنموذج DT	78
الجدول (5-5) حالات Positive و Negative للنموذج RF	81
الجدول (6-5) حالات Positive و Negative للنموذج AB	84

87	 الجدول(5-7) حالات Positive و Negative للنموذج GB
89	 الجدول(5-8) مقارنة بين أداء النماذج المدربة باستخدام 21 سمة
90	 الجدول(5-9) أفضل 9 سمات بحسب SKB
91	 الجدول(5-10) مقارنة بين أداء النماذج المدربة باستخدام سمات SKB
92	 الجدول(5-11) حالات Positive و Negative للنموذج RF مع سمات SKB
95	 الجدول(5-12) حالات Positive و Negative للنموذج GB مع سمات SKB
97	 الجدول(5-13) أفضل 9 سمات بحسب EN
97	 الجدول(5-14) مقارنة بين أداء النماذج المدربة باستخدام سمات EN
99	 الجدول(5-15) حالات Positive و Negative للنموذج RF مع سمات EN
101	 الجدول(5-16) حالات Positive و Negative للنموذج GB مع سمات EN
103	 الجدول(5-17) أفضل 9 سمات بحسب RFE
104	 الجدول(5-18) مقارنة بين أداء النماذج المدربة باستخدام سمات RFE
105	 الجدول(5-19) حالات Positive و Negative للنموذج RF مع سمات RFE
108	 الجدول(5-20) حالات Positive و Negative للنموذج GB مع سمات RFE
111	 الجدول(5-21) المقارنة هذا البحث مع الدراسات السابقة

فهرس المعادلات

المعادلة	الصفحة
المعادلة (1-2) قانون بايز	16
المعادلة (2-2) ناتج بايز البسيط	16
المعادلة (3-2) المسافة الإقليدية	18
المعادلة (4-2) مسافة مانهاتن	18
المعادلة (5-2) طريقة Gini	21
المعادلة (6-2) طريقة Entropy	21
المعادلة (7-2) ناتج الغابة العشوائية	22
المعادلة (8-2) الدقة	26
المعادلة (9-2) الحساسية	26
المعادلة (10-2) الإحكام	26
المعادلة (11-2) درجة F1	26
المعادلة (12-2) الحساسية الكلية	26
المعادلة (13-2) الإحكام الكلية	26
المعادلة (14-2) درجة F1 الكلية	26
المعادلة (1-4) الحدود الدنيا	46
المعادلة (2-4) الحدود العليا	46

قائمة المصطلحات

المصطلح باللغة العربية	المصطلح باللغة الإنكليزية
الدقة	Accuracy
مرض الزهايمر	Alzheimer Disease
بروتين بيتا أميلويد	Beta-Amyloid
لويحات بيتا أميلويد	Beta-Amyloid Plaques
القشرة المخية	Cerebral Cortex
الأصناف	Classes
تقرير التصنيف	Classification Report
مصفوفة الارتباك	Confusion Matrix
السمات	Features
درجة F1	F1-Score
الحصين	Hippocampus
المحددات الفائقة	Hyperparameters
تابع الخطأ	Loss Function
تعلم الآلة	Machine Learning
ضعف الإدراك المعتدل	Mild Cognitive Impairment
عصبي تنكسي	Neurodegeneration
الخلايا العصبية	Neurons
الاختبارات العصبية النفسية	Neuropsychological Tests
الإدراك الطبيعي	Normal Cognitive
مقاييس الأداء	Performance Measures
الإحكام	Precision
الحساسية	Sensitivity
المشابك العصبية	Synapses
بروتين تاو	Tau
تشابكات تاو	Tau Tangles
البطينات	Ventricles

قائمة الاختصارات

AB	Adaboost
AD	Alzheimer Disease
ADNI	Alzheimer Disease Neuroimaging Initiative
DT	Decision Tree
EN	Elastic Net
GB	Gradient Boosting
KNN	K-Nearest Neighbors
MCI	Mild Cognitive Impairment
MRI	Magnetic Resonance Imaging
NB	Naive Bayes
NL	Normal Cognitive
PET	Positron Emission Tomography
RF	Random Forest
RFE	Recursive Feature Elimination
SKB	Select K Best
SVM	Support Vector Machin

الملخص

يعتبر مرض الزهايمر مرض عصبي تنكسي يزداد سوءاً مع مرور الوقت، ولا تظهر أعراضه السريرية في وقت مبكر وإنما في مراحله المتأخرة، كما أنه لا يوجد علاج مؤكد لهذا المرض وإنما تخفيف لأعراضه فقط، وانطلاقاً من أهمية الكشف المبكر عن الزهايمر في محاولة إيجاد العلاج المناسب، وللسماح للمرضى في اختيار المشاركة في التجارب السريرية، والتخطيط لحياتهم مستقبلاً، عمل هذا البحث على بناء نموذج للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر، حيث تم استخدام سبعة نماذج تعلم آلة، وهي: آلة متجه الدعم، بايز البسيط، الجار الأقرب، شجرة القرار، الغابة العشوائية، التعزيز الكيفي، تعزيز التدرج، وذلك للكشف عن ثلاث حالات: مرض الزهايمر (AD)، وضعف الإدراك المعتدل (MCI)، والإدراك الطبيعي (NL).

تم استخدام قاعدة بيانات (ADNI) ومعالجتها وتجهيزها حيث تم معالجة كل من القيم المفقودة والشاذة، تحويل السمات الاسمية إلى سمات رقمية، وموازنة عدد العينات في الأصناف، وبالإضافة لذلك ضبط المحددات الفائقة لكل النماذج باستخدام البحث الشبكي والتحقق المتقاطع ثم تدريبها واختبارها وتقييم أدائها لتحديد النموذج الأفضل، وكما تم تقليل السمات باستخدام ثلاثة طرق وإعادة تدريب النماذج السابقة للوصول إلى أفضل تسع سمات مع أفضل نموذج.

تفوق نموذج تعزيز التدرج ونموذج الغابة العشوائية على باقي النماذج سواءً قبل تقليل السمات أو بعدها وكانا الأفضل أداء وكفاءة، حيث بالنسبة لنموذج RF كانت قيمة الدقة والحساسية 98.81% ولنموذج GB كانت قيمة الدقة والحساسية 99.05%، ويشير هذا إلى نتائج واعدة لاستخدام هذه النماذج لدعم عملية اتخاذ القرارات الطبية.

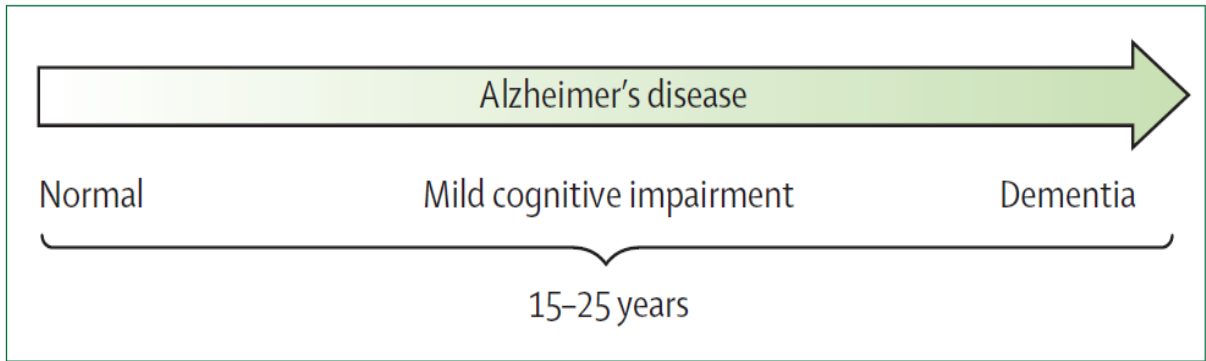
الكلمات المفتاحية: مرض الزهايمر، ضعف الإدراك المعتدل، الكشف المبكر عن الزهايمر، تعلم الآلة، المحددات الفائقة، آلة متجه الدعم، بايز البسيط، الجار الأقرب، شجرة القرار، الغابة العشوائية، التعزيز الكيفي، تعزيز التدرج.

الفصل الأول:

الإطار العام للبحث

1-1 مقدمة:

شهدت تقنيات تعلم الآلة تطوراً كبيراً في السنوات الأخيرة، وأصبحت قادرة على تحليل وفهم البيانات الضخمة بطرق فعالة ودقيقة، ومع تزايد عدد المصابين بمرض الزهايمر، يعتبر الكشف المبكر عن هذا المرض أمراً حيوياً لتحسين فرص العلاج والإدارة الفعالة للمرضى، لذلك تم العمل في هذه الدراسة على تسخير هذه التقنيات المختلفة لتحليل البيانات المتاحة واختيار الأفضل منها للوصول إلى النماذج المناسبة للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر ذات كفاءة ودقة عالية حيث تم بناء سبعة نماذج باستخدام تعلم الآلة للكشف عن ثلاث حالات: مرض الزهايمر (Alzheimer Disease (AD)، وضعف الإدراك المعتدل (Mild Cognitive Impairment (MCI)، والإدراك الطبيعي (Normal Cognitive (NL)، ويوضح الشكل (1-1) استمرارية مرض الزهايمر.



الشكل (1-1) استمرارية مرض الزهايمر. [11]

1-2 مشكلة البحث:

مرض الزهايمر من الأمراض العصبية التنكسية التقدمية أي يزداد سوءاً مع مرور الوقت، وكما وإنه لا يوجد علاج ناجح ومؤكد لهذا المرض، وبالإضافة لظهور أعراضه السريرية الواضحة في وقت متأخر من مراحله حيث التزايد الكبير لتراكم اللويحات والتشابكات وتدمير عدد كبير من الخلايا العصبية، مما يعني صعوبة التحكم في المرض، وضعف المساعدة الطبية التي يمكن تقديمه للمريض عند التشخيص المتأخر للزهايمر، وتدهور حالته الصحية بشكل سريع، لذلك لابد من العمل على الكشف المبكر عن المرض أو التنبؤ بخطر حدوثه.

1-3 تساؤلات البحث:

من التساؤلات التي تم طرحها:

- 1- هل يمكن الاستفادة من تعلم الآلة في بناء نموذج للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر؟
- 2- ما هي أفضل خوارزميات تعلم الآلة في بناء نموذج للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر؟
- 3- ما هي أهم السمات لبناء نموذج للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر؟
- 4- ما هي أفضل طريقة من طرق تعلم الآلة لتقليل سمات قاعدة البيانات مع المحافظة على أداء نموذج التنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر؟

1-4 فرضيات البحث:

لابد من وجود عدة عوامل وسمات مرتبطة بشكل وثيق مع خطر الإصابة بمرض الزهايمر، وتساعد في بناء نموذج للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر، مما يتيح للأفراد فرصة أكبر للمشاركة في تجارب العلاج السريرية، أو حتى الإدارة الجيدة للمرض مما يسمح للمريض أن يعتمد على نفسه لأطول فترة ممكنة.

1-5 مسلمات البحث:

من الأمور التي يمكن اعتبارها مسلمات لهذا البحث: الكشف عن الأشخاص المعرضين لخطر الإصابة بالزهايمر أو المصابين في وقت مبكر أمر مهم جداً وأكد عليه جميع الباحثين والأطباء كمفتاح للعلاج الناجح، لأنه يجعلهم المستفيد الأكبر من التجارب والاختبارات السريرية للعلاج.

1-6 هدف البحث:

نتيجة ازدياد عدد الأفراد الكبار في السن، وازدياد عدد المصابين بمرض الزهايمر، وتأخر ظهور الأعراض السريرية لهذا المرض، وصعوبة التحكم به في مراحله المتأخرة، سيتم العمل في هذا البحث على بناء نموذج باستخدام خوارزميات تعلم الآلة، وذلك للتنبؤ المبكر عن مرض الزهايمر وذلك من

خلال الكشف عن الأفراد المعرضين لخطر الإصابة بهذا المرض، وذلك باستخدام البيانات الطبية المتاحة، واستخلاص أهم السمات منها، وربطها مع مرض الزهايمر.

1-7 أهمية البحث:

يعد التنبؤ والكشف المبكر عن الإصابة بمرض الزهايمر والأفراد المعرضين لخطر الإصابة به أمر بالغ الأهمية لإدارة المرض بشكل صحيح وسليم، وكذلك يسهل عملية مراقبة المرض وفهم آلية تطوره بشكل جيد، مما يسمح باتخاذ الإجراءات المناسبة من أجل توفير العلاج الأفضل والوقاية من المضاعفات والعمل على الحفاظ على قدرة الفرد للقيام بالوظائف اليومية والاكتفاء والاعتماد على الذات لأطول فترة ممكنة، كما أنه يعطي الأفراد فرصة اختيار المشاركة في التجارب والدراسات السريرية التي تعمل على التأكد من مفعول وتأثير العلاجات الجديدة لهذا المرض والتي تتطلب مراحل مبكر من الزهايمر وحتى فترات ما قبل الإصابة، وأيضاً يتيح للفرد وعائلته وضع الخطط المستقبلية والتعامل مع الشؤون القانونية والمالية والعمل على تدابير المعيشة والحفاظ على السلامة.

1-8 مبررات البحث:

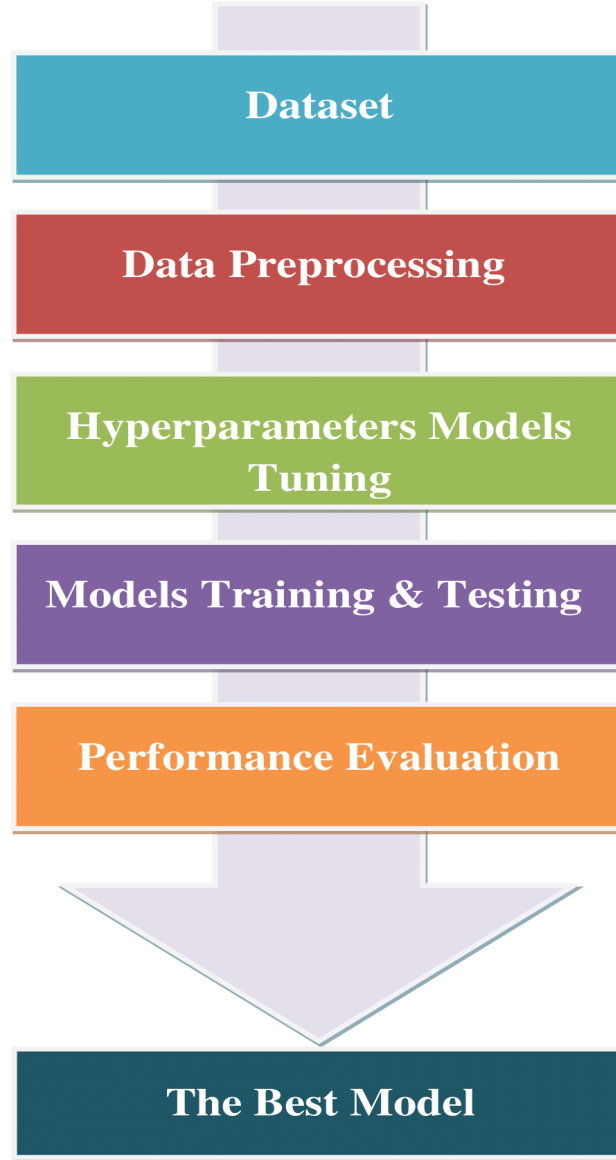
من المبررات ازدياد عدد الأفراد الكبار في السن وعدد المصابين بمرض الزهايمر، وتأخر ظهور الأعراض السريرية، وتأكيد الباحثين على أهمية التنبؤ والتشخيص المبكر للإصابة بمرض الزهايمر وحاجتهم لأفراد مصابين ببداية المرض لتحقيق أفضل النتائج إيجاد العلاج، وذلك كون المرحلة ما قبل الإصابة أو المرحلة المبكرة أو مرحلة ما قبل الأعراض قد تكون المفتاح للعلاج الناجح، فالتجارب السريرية الحالية تؤكد أن الاستراتيجيات الوقائية والعلاجية تعتمد بشكل كبير على المرحلة المبكرة لمرض الزهايمر أو ما قبل الإصابة لذلك إيجاد أداة تكشف الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر يعتبر أمراً مهماً.

1-9 محددات البحث:

سيتم التعامل مع البيانات الطبية والسمات المتوفرة بحسب قاعدة البيانات التي يمكن تأمينها، ولن يتطرق البحث إلى مراحل ما بعد الإصابة بمرض الزهايمر وتصنيفها.

10-1 مخطط البحث:

يوضح الشكل (1-2) المخطط العام الذي تم العمل وفقه في هذا البحث.



الشكل (1-2) المخطط العام للبحث.

حيث بعد دراسة وتحليل المراجع العلمية والدراسات المرجعية التي تناولت مرض الزهايمر والتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر تم العمل على تأمين قاعدة البيانات المناسبة لهذا البحث حيث استخدم قاعدة بيانات.

ثم باستخدام لغة Python اللغة البرمجة الأكثر شهرة والعمل ضمن بيئة Google Colab التي تتوفر فيها جميع المكتبات المطلوبة لهذا البحث تم تحليل ومعالجة وتهيئة قاعدة البيانات وضبط المحددات الفائقة (Hyperparameters) لبناء نماذج تعلم الآلة حيث عمل هذا البحث على بناء سبعة نماذج،

وهي: Decision ،K-Nearest Neighbors ،Naive Bayes ،Support Vector Machine ،Tree ،Random Forest ،Adaboost ،Gradient Boosting ،بعد ذلك تم تدريب النماذج واختبارها وتقييم أدائها، ومقارنة النتائج للوصول إلى النموذج الأفضل. كما تم العمل في هذه الدراسة كخطوة إضافية على استخدام عدة طرق لتقليل عدد السمات وإعادة تدريب النماذج السابقة واختبارها وتقييم أدائها، ومقارنة النتائج للوصول إلى أفضل أقل عدد من السمات مع أفضل نموذج.

11-1 منهج البحث:

المنهج المتبع في هذا البحث هو المنهج المتبع في البحوث العلمية الهندسية، حيث يعتبر هذا البحث دراسة تحليلية وتجريبية وبرمجية، فبعد دراسة وتحليل المعلومات في المراجع العلمية والدراسات التي تناولت مرض الزهايمر، والتنبؤ بالأفراد المعرضين لخطر الإصابة بمرض الزهايمر، أو الكشف المبكر عنهم، والمتغيرات الحيوية والفيزيولوجية المرتبطة به، تم جمع وتأمين قاعدة البيانات التي ستخدم هذا البحث، ثم دراسة وتحليل هذه البيانات التي تم جمعها، واستخلاص أهم السمات منها، بعد ذلك بناء نماذج التنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر بتجريب عدة خوارزميات تعلم الآلة وتقييم أدائها لتحديد النموذج الأفضل لهذا البحث، بالإضافة تجريب عدة طرق لتقليل السمات وإعادة تدريب النماذج وتقييم أدائها لتحديد النموذج الأفضل، وفي النهاية تحليل جميع النتائج، ووضع التوصيات المستقبلية، وكتابة الرسالة.

12-1 أقسام الرسالة:

تقسم هذه الرسالة إلى خمسة فصول رئيسية، وهي:

- 1- الفصل الأول: الإطار العام وفيها فكرة عامة عن البحث، وتناولت مشكلة البحث وهدفه وأهميته ومبرراته، بالإضافة إلى منهج البحث والمخطط العام لخطوات العمل.
- 2- الفصل الثاني: القسم النظري ويستعرض معلومات عن مرض الزهايمر، وكذلك تعريف بتعلم الآلة، والأسس النظرية للخوارزميات المستخدمة في بناء نموذج البحث، وطرق تقييم الأداء، بالإضافة لشرح عن طرق تقليل السمات المستخدمة.

3- الفصل الثالث: أدبيات البحث ويحتوي على شرح للدراسات المرجعية وما قام به الباحثين الآخرين بهذه الدراسات ذات الصلة بهذا البحث.

4- الفصل الرابع: منهج البحث وإجراءاته ويحتوي شرح مفصل لجميع الخطوات المتبعة في بناء النموذج المطلوب من قاعدة البيانات ومعالجتها إلى ضبط المحددات الفائقة للنماذج وبنائها وتقييم أدائها بالإضافة إلى الطريقة المتبعة لتقليل الخصائص.

5- الفصل الخامس: النتائج والمناقشة ويستعرض ويناقش نتائج ضبط المحددات الفائقة، وعرض ومناقشة نتائج تقييم النماذج من مصفوفات الارتباك (Confusion Matrix) وتقارير التصنيف (Classification Reports) ومقاييس الأداء (Performance Measures)، وكما يعرض ويناقش نتائج طرق تقليل السمات وتقييم النماذج المعاد تدريبها حسب سمات كل طريقة، وفي النهاية مقارنة هذه النتائج مع بعضها البعض والمقارنة النموذج الأفضل لهذا البحث مع النموذج الأفضل للدراسات المرجعية من حيث الدقة (Accuracy)، الحساسية (Sensitivity)، الإحكام (Precision)، ودرجة F1 (F1-Score).

6- الفصل السادس: الخاتمة وتحتوي على ملخص للبحث، واستنتاجات البحث وميزاته، والتوصيات والآفاق المستقبلية.

13-1 خاتمة:

عمل هذا الفصل على إعطاء فكرة عامة عن البحث حيث تم العمل على تسخير تقنيات تعلم الآلة من أجل بناء نموذج للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر، وحيث أعطت النماذج نتائج واعدة.

الفصل الثاني:

القسم النظري

1-2 مقدمة:

يقدم هذا الفصل القسم النظري لهذا البحث حيث يستعرض بعض المعلومات عن مرض الزهايمر بالإضافة للتعريف عن تعلم الآلة والأسس النظرية للنماذج التي تم استخدامها في هذه الدراسة وكيفية إيجاد مقاييس تقييم أدائها وكفاءتها.

2-2 الزهايمر:

يعد مرض الزهايمر مرض عصبي تنكسي أي يزداد سوءاً مع مرور الوقت يصيب الجهاز العصبي عند كبار السن وتختلف سرعة تقدمه والقدرات التي تتأثر به من شخص إلى آخر ومع مرور الوقت يتضرر عدد أكبر من الخلايا العصبية وتتأثر مناطق أكثر من الدماغ.[2]

يعد Alöis Alzheimer أول من لاحظ هذا المرض عام 1907 عند فحصه مجهرياً دماغ مريضته البالغة من العمر 51 سنة التي عانت من فقدان الذاكرة حاد، وانخفاض كبير في الإدراك والانتباه، واضطرابات في الشخصية حيث لاحظ تراكم بروتيني بيتا أميلويد (Beta-Amyloid) وتاو (Tau) واعتلال في الأوعية الدموية ووصف الحالة بأنها مرض خطير في المخ وأطلق Emil Kraepelin على هذه الحالة الطبية مرض الزهايمر لأول مرة في كتابه Handbook of Psychiatry حيث كتب يمكن أن يحدث فقدان تدريجي للوظائف الإدراكية بسبب اضطراب دماغي مثل مرض الزهايمر أو عوامل أخرى مثل التسمم والالتهابات والخلل في الجهازين الرئوي والدوري.[3]

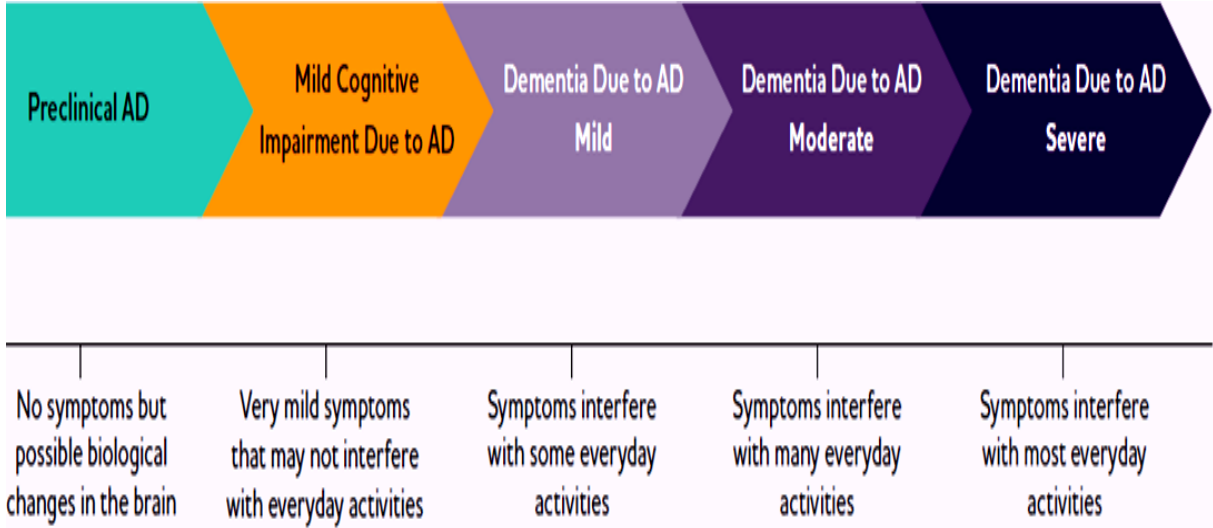
تعد التغيرات الدماغية الناجمة عن مرض الزهايمر العامل الأكثر شيوعاً للإصابة بالخرف وهو يمثل ما يقدر إلى 80% من حالات الخرف.[2] لا تظهر أعراض مرض الزهايمر في مراحله المبكرة وإنما في مراحله المتأخرة، وكما يعتقد أن التغيرات الدماغية التي تسبب مشاكل الذاكرة واللغة والتفكير تبدأ قبل 20 سنة أو أكثر من ظهور الأعراض.[4]

يصنف مرض الزهايمر كمرض خطير وقاتل حيث تشير الدراسات إلى أن الأشخاص الذين تبلغ أعمارهم 65 وأكثر يعيشون في المتوسط لمدة تتراوح بين (4-8) سنوات بعد تشخيص إصابتهم بمرض الزهايمر مع العلم أنه قد يعيش البعض القليل لما يصل إلى أكثر من 10 سنوات.[2]

يزداد عدد المصابين بمرض الزهايمر بسبب ازدياد متوسط عمر الإنسان وعدد الأشخاص الكبار في السن حيث بحسب منظمة الصحة العالمية سيتضاعف عدد الأشخاص الذين تجاوزوا 60 سنة من 1

مليار في عام 2020 إلى 2.1 مليار في عام 2050، ومن المتوقع أن يصل عدد الأشخاص الذين تجاوزوا 85 سنة في عام 2050 حوالي 426 مليون وبالتالي يتوقع زيادة عدد الأشخاص المصابين بمرض الزهايمر من 55 مليون في عام 2019 إلى 139 مليون في عام 2050.[5] لا يوجد علاج مؤكد لمرض الزهايمر حتى الآن وتقتصر الأدوية المعطاة للمرضى على علاج أعراض مرض الزهايمر وتخفيف منها وذلك لتسهيل حياة المريض، مع العلم على وجود عدة أدوية قيد التطوير لمحاولة معالجة الزهايمر.[6]

يمر كل مرضى الزهايمر بمرحلة ضعف الإدراك المعتدل حيث تعتبر مرحلة ما قبل المرض وهي المرحلة المهمة في الكشف المبكر عن مرض الزهايمر أو الأشخاص المعرضين للإصابة به، حيث يتحول 15% من الأشخاص المصابين بضعف الإدراك المعتدل بعد حوالي سنتين من الإصابة، (-40% 30%) منهم بعد حوالي 5 سنوات إلى مرضى زهايمر، ونسبة قليلة قد لا تتحول إلى مرضى زهايمر، ويعاني مرضى ضعف الإدراك المعتدل من مشاكل في التفكير واللغة والذاكرة لكن بدرجة خفيفة وتكون غير ملحوظة عند الآخرين ولا تؤثر على الإدراك العام وقدرتهم على القيام بالأنشطة اليومية، ويوضح الشكل (2-1) مراحل تقدم مرض الزهايمر.[2]



الشكل (2-1) مراحل تقدم مرض الزهايمر.[2]

ويستخدم في تشخيص الزهايمر الاختبارات العصبية النفسية (Neuropsychological Tests) التي تختبر الذاكرة والإدراك واللغة والقيام بالأنشطة اليومية، بالإضافة إلى التصوير بالرنين المغناطيسي (MRI) للخلايا العصبية ولقياس الضمور الحاصل، وكذلك التصوير بالإصدار البوزيتروني (PET).[7]

2-3 التغيرات الدماغية لمرض الزهايمر:

يحتوي دماغ الشخص البالغ السليم على مليارات الخلايا العصبية Neurons، ولكل منها امتدادات طويلة ومتفرعة، وتمكن هذه الامتدادات الخلايا العصبية الفردية من تكوين اتصالات مع الخلايا العصبية الأخرى، وتسمى هذه الوصلات بالمشابك العصبية Synapses حيث تتدفق عندها المعلومات في دفعات صغيرة من المواد الكيميائية التي تفرزها خلية عصبية وتستقبلها خلية عصبية أخرى. ويحتوي الدماغ على تريليونات المشابك العصبية التي تسمح للإشارات بالتنقل بسرعة عبر الدماغ، وتخلق هذه الإشارات الأساس الخلوي للذكريات والأفكار والأحاسيس والعواطف والحركات والمهارات. [2]

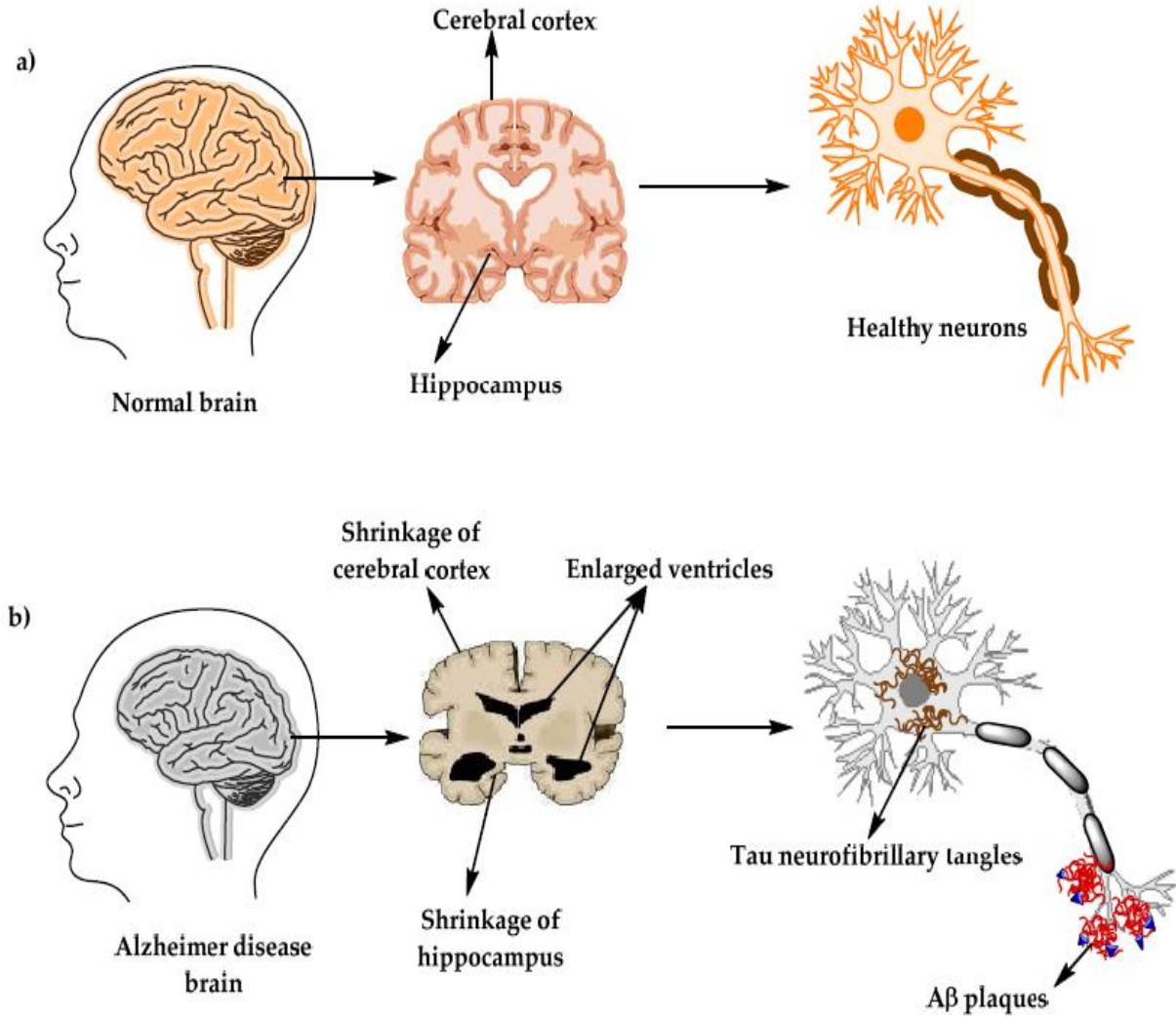
أبرز تغييرين من التغيرات الدماغية المرتبطة بمرض الزهايمر: تراكم بقايا البروتين بيتا أميلويد في كتل تسمى لويحات بيتا أميلويد (Beta-Amyloid Plaques) خارج الخلايا العصبية، وتراكم بشكل غير طبيعي من بروتين تاو يسمى تشابكات تاو (Tau Tangles) داخل الخلايا العصبية، وتتبع هذه التغيرات تلف الخلايا العصبية وتدميرها، وهذا ما يسمى بالتكس العصبي، وهو يعد سمة أساسية لمرض الزهايمر. [2]

يلعب بروتين بيتا أميلويد وبروتين تاو أدواراً مختلفة في مرض الزهايمر حيث تتسبب لويحات بيتا أميلويد في إتلاف الخلايا العصبية عن طريق التدخل في الاتصال بين الخلايا العصبية عند المشابك، وتمنع تشابكات تاو داخل الخلايا العصبية نقل العناصر الغذائية والجزيئات الأخرى الضرورية للأداء الطبيعي وبقاء الخلايا العصبية، وعلى الرغم من أن التسلسل الكامل للأحداث غير واضح، فقد يبدأ بروتين بيتا أميلويد في التراكم قبل بروتين تاو، ويرتبط تراكم بيتا أميلويد بزيادة بروتينات لاحقاً في تاو. [4]

إن وجود تراكم بروتينات بيتا أميلويد وتاو ينشط خلايا الجهاز المناعي في الدماغ تسمى الخلايا الدبقية الصغيرة، حيث تحاول هذه الخلايا إزالة تراكم تلك البروتينات وكذلك الحطام المنتشر من الخلايا الميتة والمحتضرة، وقد يبدأ الالتهاب المزمن عندما لا تستطيع الخلايا الدبقية الصغيرة القيام بدورها، وكما يحدث الضمور بسبب فقدان الخلايا ويتأثر أداء الدماغ الطبيعي بشكل أكبر بسبب انخفاض قدرة الدماغ على استقلاب الجلوكوز الذي يعتبر مصدر طاقته الرئيسي. [2]

تشمل التغيرات الدماغية الأخرى المرتبطة بمرض الزهايمر والتهاب وضمور حجم المخ، وكذلك انكماش حجم القشرة المخية (Cerebral Cortex)، وانكماش حجم حصين (Hippocampus)، وتمدد حجم

البطينات (Ventricles)، ويوضح الشكل (2-2) مقارنة بين الدماغ السليم والدماغ عند مرض الزهايمر. [7]



الشكل (2-2) مقارنة بين a الدماغ السليم و b الدماغ عند مرض الزهايمر. [7]

بالإضافة لما سبق من تغيرات يصبح هناك أيضاً عدم قدرة على تخزين ذكريات جديدة وكما يحدث فقدان الذاكرة القديمة تدريجياً، وضعف النشاط اليومي حتى الوصول إلى عدم القدرة على القيام بالمهام اليومية، وأعراض أخرى مثل فقدان القدرة على الكلام (ضعف اللغة)، وفقدان القدرة على الحركة (اضطراب في المهارات الحركية)، وفقدان القدرة على الإدراك والانتباه والتركيز، وكل هذه الأعراض تستمر بالتزايد سوءاً كلما استمر مرض الزهايمر في التقدم إلى مراحله المتقدمة. [7]

2-4 تعلم الآلة Machine Learning:

هو أحد فروع الذكاء الاصطناعي الذي يركز على تطوير الأنظمة التي تمكن أجهزة الحاسوب من التعلم والتكيف وتحسين أدائها من خلال التجربة، يعتمد تعلم الآلة بشكل أساسي على البيانات حيث يتم استخراج المعلومات الهامة والتعلم منها مما يعطي القدرة لأجهزة الحاسوب على اتخاذ القرارات والتنبؤ بنتائج مستقبلية، ويوجد عدة تعريفات لتعلم الآلة حيث يعرف أيضاً بأنه مجال الدراسة الذي يمنح أجهزة الحاسوب القدرة على التعلم دون برمجتها صراحة [8].

يتضمن تعلم الآلة بشكل أساسي تغذية النظام بالبيانات والتي يمكن أن تكون منظمة مثل الجداول أو غير منظمة مثل الصور أو النصوص، ثم تبدأ عملية التعلم عن طريق تحديد الأنماط والعلاقات بين البيانات باستخدام خوارزميات مختلفة، بعد ذلك إنشاء نموذج بناء على الأنماط المكتسبة والتي يمكن استخدامها بعد ذلك للتنبؤ أو التصنيف على بيانات جديدة، وأخيراً يمكن تحسين النموذج وتطويره بمرور الوقت مع تلقي المزيد من البيانات والتغذية الراجعة على تنبؤاته، وتصنف خوارزميات تعلم الآلة بشكل عام إلى ثلاثة أنواع، وهي:

1- التعلم الموجه (Supervised Learning): وهو أحد الأنواع الأساسية لتعلم الآلة الذي تتوفر فيه المدخلات والمخرجات في مجموعة البيانات المستخدمة للتدريب، ويمكن اعتبار هذه البيانات بمثابة معلم يوجه الخوارزمية ويشرف على عملية التعلم من خلال الإجابات (المخرجات) الصحيحة المعروفة مسبقاً، والتي يجب أن تسعى الخوارزمية للوصول إليها في مخرجاتها، بالتالي فإن الخوارزمية في التعلم الموجه تقوم بإجراء التنبؤات بشكل متكرر على بيانات التدريب، وتصحيح الأوزان والمحددات بناءً على الأخطاء الناتجة عن هذه التنبؤات، وتستمر عملية التعلم حتى تصل الخوارزمية إلى مستوى مقبول من الأداء، وهذا ما يسمح لها بالتعميم بشكل جيد على بيانات جديدة لم تستخدم في التدريب [8].

2- التعلم غير الموجه (Unsupervised Learning): في هذا النوع من تعلم الآلة لا تتوفر المخرجات الصحيحة أو المعلومات المعروفة مسبقاً، بدلاً من ذلك تحاول الخوارزمية اكتشاف الأنماط والهيكليلات الموجودة في البيانات من خلال تحليل المدخلات فقط، وعلى عكس التعلم الموجه فلا يوجد معلم يوجه عملية التعلم وإنما الخوارزمية تقوم بتحديد أوجه التشابه والاختلاف بين المدخلات مما يسمح بتجميع العناصر المتشابهة معاً في مجموعات أو أصناف معينة تعتمد عليها الخوارزمية في اتخاذ القرارات بالنسبة للبيانات الجديدة [8].

3- التعلم المعزز (Reinforcement learning): يجمع هذا النوع من تعلم الآلة بين خصائص التعلم الموجه والتعلم غير الموجه حيث تتلقى الخوارزمية مكافآت أو عقوبات بناءً على أداؤها، فيتم مكافأتها عندما تحقق نتائج صحيحة ومعاقبتها عند ارتكاب أخطاء، ولكن دون توجيهها بشكل مباشر حول كيفية تصحيح تلك الأخطاء، ويتوجب على الخوارزمية استكشاف بيئتها وتجربة احتمالات مختلفة حتى تتوصل إلى كيفية الحصول على الإجابة الصحيحة، أي أن التعلم المعزز يعتمد على مفهوم التجربة والخطأ مما يسمح للخوارزمية بتطوير استراتيجيات فعالة مع مرور الوقت لتحقيق الأهداف المرجوة.[8]

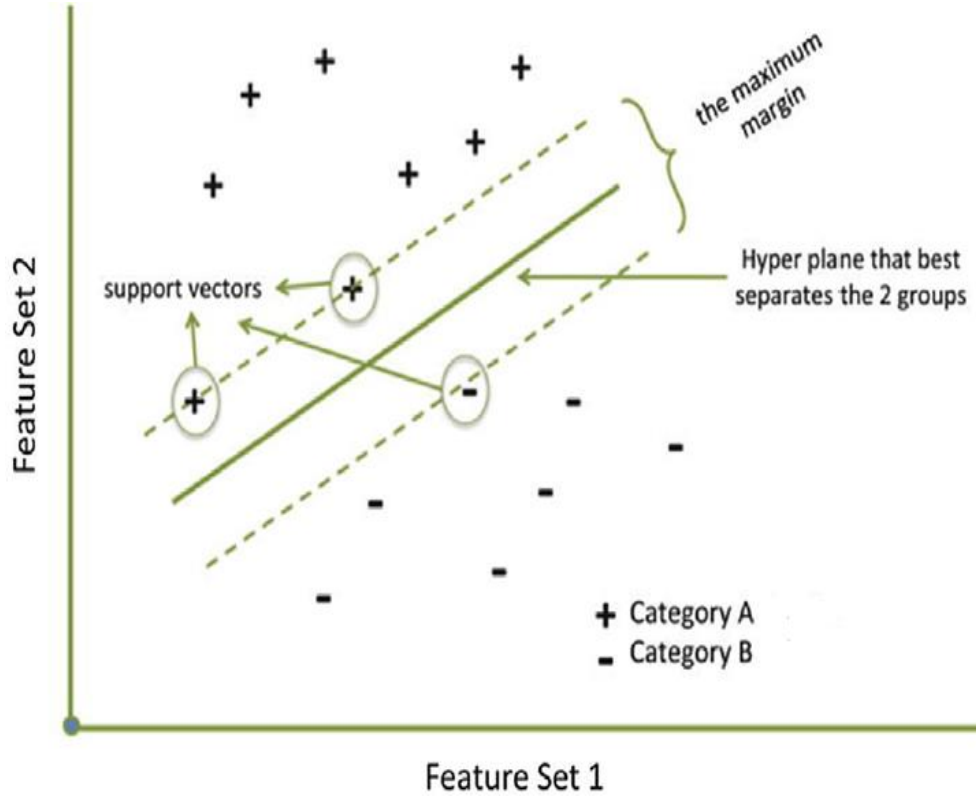
ومن خوارزميات تعلم الآلة الموجه:

2-4-1 آلة متجه الدعم (Support Vector Machine):

تعد آلة متجه الدعم (SVM) هي من إحدى خوارزميات تعلم الآلة الموجه وتتخلص عملية تدريبها في تحديد حدود القرار التي تعد الحد الفاصل بين الأصناف المختلفة في البيانات وتدعى هذه الحدود المستوي الفائق (Hyperplane) والذي يجب أن يزيد المسافة لأقصى حد بين نقاط متجهات الدعم (Support Vectors) لأصناف الخرج المختلفة، حيث هذه المتجهات هي نقاط من عينات البيانات التي تقع بالقرب من المستوي الفائق وعلى حدود الهامش (Margin) وهي المسافة بين نقاط متجهات الدعم والمستوي الفائق) وتعد نقاط متجهات الدعم النقاط الأكثر أهمية في تحديد موقع هذا المستوي حيث أي تغير في موقع نقاط متجهات الدعم يؤثر على موقع المستوي الفائق، ويوضح الشكل (2-3) آلية عمل خوارزمية SVM.[9]

إن أفضل المستوي الفائق هو الذي يحقق أكبر هامش وذلك لإتاحة إمكانية تعميم خوارزمية SVM، وكما تتعلق درجة تعقيد هذه الخوارزمية بعدد السمات (Features) المستخدمة ففي حالة البيانات ثنائية الأبعاد 2D يكون المستوي الفائق عبارة عن خط أي 1D وفي حالة البيانات ثلاثية الأبعاد 3D يكون المستوي الفائق عبارة عن مستوي أي 2D وهكذا.[9]

غالباً ما تكون خوارزمية SVM من النوع الخطي (Linear SVM) بغض النظر عن درجة التعقيد ويكون المستوي الفائق مستقيم وليس منحنى كما في خوارزمية SVM اللاخطية.[9]



الشكل (2-3) آلية عمل خوارزمية SVM. [9]

2-4-2 بايز البسيط Naive Bayes:

يعد بايز البسيط (NB) أحد خوارزميات تعلم الآلة الموجه والتي تعتمد آلية عملها على حساب الاحتمالات وفقاً لنظرية Bayes، ووصفت بكلمة (Naive) كونها تفترض استقلالية جميع السمات في مجموعة البيانات عن بعضها البعض، مما يعني أن كل سمة تؤثر على قرار الخوارزمية بشكل مستقل عن السمات الأخرى، فتغير قيم سمة معينة لا يؤثر على سمة أخرى وهو افتراض قد يكون صعب التحقيق في العديد من الحالات، ولكن مع ذلك تؤدي هذه الخوارزمية أداء وكفاءة جيدة في مجموعة متنوعة من التطبيقات. [10]

تستند خوارزمية NB إلى نظرية بايز التي تعتبر طريقة رياضية لحساب احتمالية حدوث هدف معين في التجارب المستقبلية بناءً على حدوثه في التجارب السابقة، ووفقاً لمنطق بايز يمكن قياس حالة عدم اليقين فقط من خلال تحديد احتمالات الأحداث المختلفة، ولذلك تعتبر نظرية بايز (خوارزمية NB) مقياس لعدم اليقين (Quantifying Uncertainty) حيث تسعى لزيادة درجة احتمال صحة الفرضية من خلال وضع معلومات أساسية وأدلة إضافية. [11]

وتوضح المعادلة (2-1) قانون بايز الذي تعمل وفقه خوارزمية NB: [10]

$$P(\omega_j | \mathbf{x}_i) = \frac{P(\mathbf{x}_i | \omega_j) \cdot P(\omega_j)}{P(\mathbf{x}_i)} \quad \dots\dots(1-2)$$

حيث:

$\mathbf{x}_i$: هو متجه السمات للعينة i حيث $i: \{1, 2, \dots, n\}$

ω_j : هو الصنف j حيث $j: \{1, 2, \dots, m\}$

$P(\omega_j | \mathbf{x}_i)$: احتمال وقوع الحدث ω_j شرط وقوع الحدث $\mathbf{x}_i$ (أي احتمال أن تنتمي العينة i إلى الصنف ω_j بناءً على السمات $\mathbf{x}_i$)، ويدعى الاحتمال اللاحق (Posterior Probability) والذي يعبر عن احتمال حدوث حدث معين بعد أخذ المعلومات الجديدة (السمات) في الاعتبار.

$P(\mathbf{x}_i | \omega_j)$: احتمال وقوع الحدث $\mathbf{x}_i$ شرط وقوع الحدث ω_j (أي احتمال وجود السمات $\mathbf{x}_i$ شرط أن تكون العينة i تنتمي إلى الصنف ω_j)، وتدعى الاحتمالية (Likelihood) والتي تعبر عن مدى احتمال ملاحظة السمات المعطاة إذا كانت العينة تنتمي إلى الصنف المعني.

$P(\omega_j)$: وهو احتمال الصنف ω_j ، ويدعى الاحتمال السابق (Prior Probability) والذي يعبر عن تقدير الاحتمال لحدوث حدث معين قبل أن تأخذ في الاعتبار المعلومات الجديدة، أي يعتمد هذا الاحتمال على المعرفة السابقة حول توزيع الأصناف.

$P(\mathbf{x}_i)$: وهو احتمال السمات $\mathbf{x}_i$ ، وتدعى الأدلة (Evidence) والتي تعبر عن مدى احتمال ملاحظة مجموعة السمات المعطاة بض النظر عن أي صنف.

وبالنهاية يتم اختيار الفرضية (الصنف الذي تنتمي إليه العينة) ذات قيمة الاحتمال الأكبر، كما موضح

في المعادلة (2-2): [10]

$$\text{predicted class label} \leftarrow \arg \max_{j=1 \dots, m} P(\omega_j | \mathbf{x}_i) \quad \dots\dots(2-2)$$

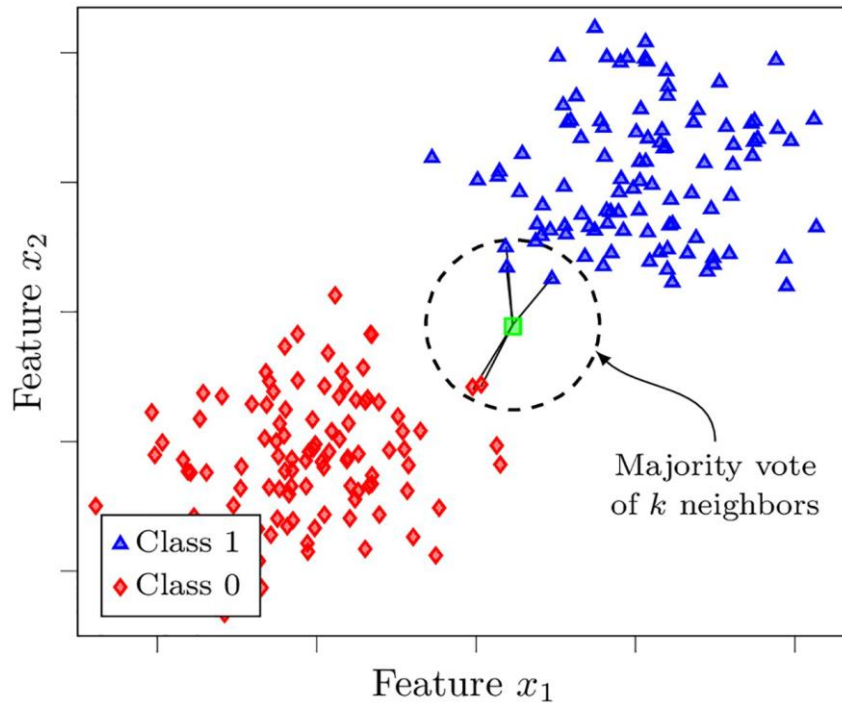
3-4-2 الجار الأقرب K-Nearest Neighbors:

يعد الجار الأقرب (KNN) أحد خوارزميات تعلم الآلة الموجه تعتمد آلية عملها على مثل "الطيور على أشكالها تقع" أي أن هذه الخوارزمية تعتمد في تصنيف عينة غير معروفة الصنف بناءً على تصنيف العينات الجيران لها حيث هذه الجيران هي عينات من مجموعة التدريب معروفة الصنف أي تصنيف كل عينة على نحو العينات المحيطة بها وبالتالي عندما تعمل هذه الخوارزمية على تصنيف

عينة جديدة غير معروفة الصنف يتم التنبؤ بها من خلال النظر إلى تصنيف أقرب عينات جيران لها. [12]

تعتمد خوارزمية KNN على تحديد الجار الأقرب بحساب جميع المسافات بين العينة غير معروفة الصنف وعينات مجموعة التدريب معروفة الصنف وحيث يتوافق صنف هذه العينة المجهولة مع عينة مجموعة التدريب معروفة الصنف التي تحقق المسافة ذات القيمة الأصغر وبذلك تصنف هذه الخوارزمية العينة المجهولة حسب صنف الجارة الأقرب. [12]

إن الاعتماد على عينة واحدة معروفة الصنف يعطي ضعف في قاعدة تصنيف هذه الخوارزمية من الممكن أن تكون جيدة ذلك في حال العينة المجهولة محاطة بعينات معروفة لها نفس التصنيف ولكن في حال كانت محاطة بعينات معروفة لها تصنيفات مختلفة أي تنتمي صنفين أو أكثر يؤدي إلى ضعف في التصنيف لذلك ومن أجل زيادة مستوى الدقة يتم الاعتماد على عدد K جار أقرب ولذلك سميت هذه الخوارزمية K-Nearest Neighbors وبالتالي تصبح قاعدة تصنيف هذه الخوارزمية تأخذ K عينة من العينات المحيطة معروفة التصنيف وذات المسافات الأصغر بعين الاعتبار لتصنيف العينة المجهولة، وبالتالي تصنف خوارزمية KNN العينة المجهولة للصنف الذي يحتوي على معظم أقرب جيرانها البالغ عددهم K (أي تتم عملية تصويت Voting ويتم اختيار الصنف الذي كانت تنتمي إليه أغلب عينات الجيران K الأقرب)، [12] ويوضح الشكل (2-4) آلية عمل خوارزمية KNN.



الشكل (2-4) آلية عمل خوارزمية KNN. [13]

إن التصنيف على أساس المسافة والتصويت بين عدد الجيران الأقرباء قد لا يكون دقيقاً بشكل كبير في بعض الحالات فقد يكون بعض الجيران الأقرباء تقع في نطاق واسع أي بعيدة عن العينة المجهولة المراد تصنيفها لذلك لرفع الأداء يتم تثقيف صوت كل عينة من عدد الجيران الأقرباء حسب مسافتها أي زيادة أهمية العينات الأقرب أثناء التصنيف.[12]

تعتبر خوارزمية KNN متعلم كسول لأنها لا تولد مصنف من بيانات مجموعة التدريب بل تخزنها وتعود إليها في كل مرة تحتاج إجراء التصنيف، وهذا يجعل الطريقة أسهل ولكنها مكلفة حسابياً وذلك للحاجة إلى حساب المسافات بين العينات المعروفة وغير المعروفة في كل مرة.[12]

يوجد طريقتين رياضيتين تستخدمهما خوارزمية KNN في حساب المسافات:[14]

الطريقة الأولى حساب المسافة الإقليدية Euclidean Distance الموضحة في المعادلة (3-2)

التالية:[14]

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad \dots\dots(3-2)$$

الطريقة الأولى حساب المسافة مانهاتن (Manhattan Distance) الموضحة في المعادلة (4-2)

التالية:[14]

$$d(x, y) = \sum_{i=1}^n |x_i - y_i| \quad \dots\dots(4-2)$$

حيث x العينة الأولى و y العينة الثانية.

4-4-2 شجرة القرار Decision Tree:

تعد شجرة القرار (DT) إحدى خوارزميات تعلم الآلة الموجه المكون من مجموعة الأسئلة منظمة بشكل هرمي على شكل شجرة وهذا الهيكل الانسيابي يسهل فهم عملية اتخاذ القرار، وتتألف شجرة القرار من المكونات الأساسية التالية:

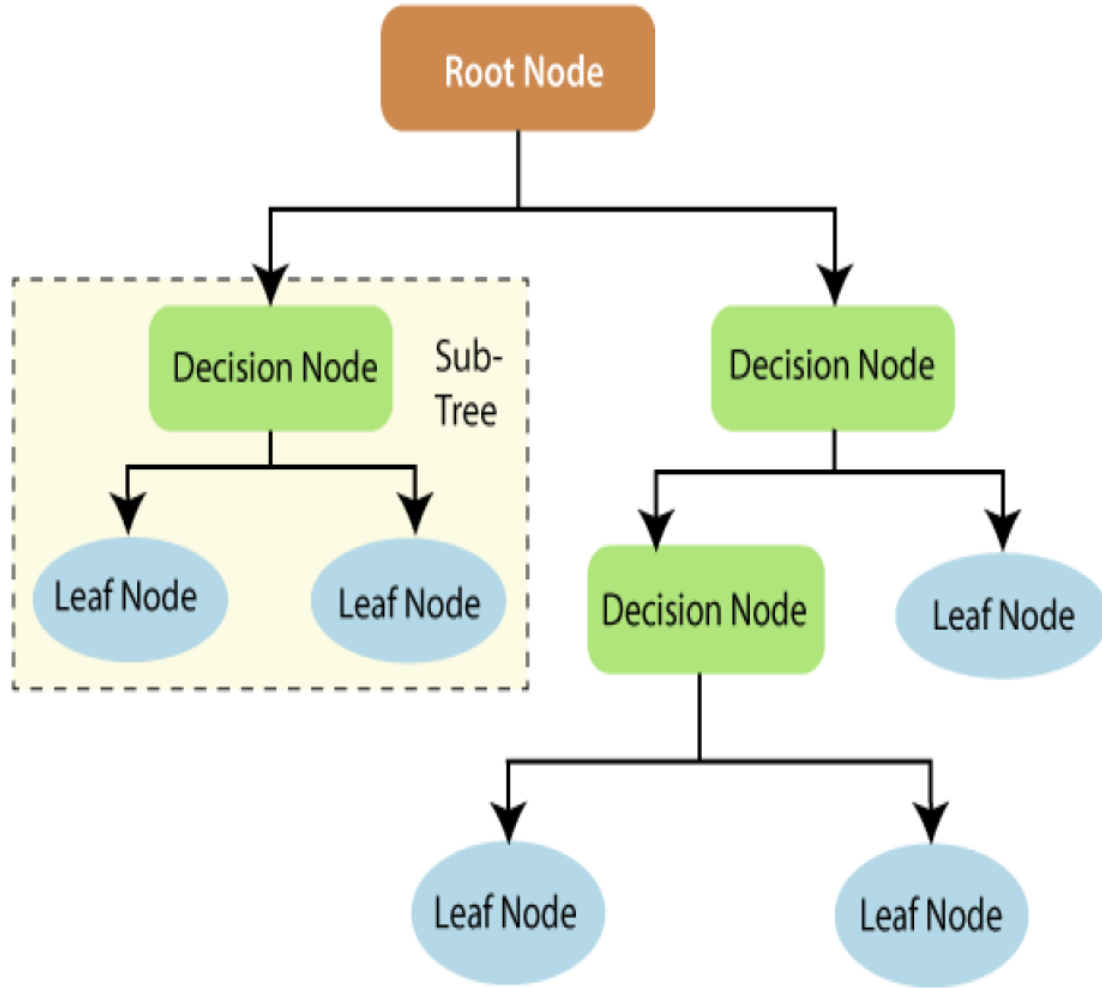
1- عقدة الجذر (Root Node): تمثل جذر شجرة القرار فهي نقطة البداية التي تحتوي على مجموعة البيانات كاملة والسؤال (القرار الأولي) الذي يجب اتخاذ قرار بشأنه، وتعتبر العقدة الأولى التي يتم عندها بداية الانقسامات بناءً على أول لعملية اتخاذ قرار. [15]

2- الفروع (Branches): تمثل المسارات بين العقد حيث ينبثق أول فروع من العقدة الجذرية وتمتد إلى العقد الأخرى، ويمثل كل فرع الخيار الممكن اتخاذه بناءً على القرار في العقدة الجذرية، ويمكن أن تتفرع الفروع إلى المزيد من العقد كدلالة على الخيارات المختلفة والنتائج الممكنة لكل خيار حيث يؤدي كل فرع إلى عقدة جديدة. [15]

3- عقد القرار (Decision Nodes): تمثل العقد الفرعية المنبثقة من العقدة الجذرية أي الأسئلة والقرارات الفرعية الناتجة عنها، وتحتوي كل عقدة قرار على مجموعة فرعية من البيانات وذلك بناءً على اختبار مجموعة سمات فرعية معينة، ويمكن أن يكون هناك عدة عقد داخلية على حسب عدد الخيارات المتاحة، ومن الممكن أن تتفرع هذه العقد إلى عقد جديدة. [15]

4- العقد النهائية (Leaf Nodes): وهي النقاط النهائية في شجرة القرار التي لا يوجد بعدها تفرعات أو انقسامات، حيث تمثل هذه العقد النتائج والقرارات النهائية لكل مسار من الخيارات المتخذة. [15]

تبدأ عملية اتخاذ القرار (تصنيف العينة الجديدة مثلاً) من العقدة الجذرية، وبناءً على السؤال المطروح في هذه العقدة يتم اختيار أحد الفروع للانتقال إلى العقدة الداخلية التالية، ويتم الانتقال عبر الفروع من عقدة داخلية إلى عقدة داخلية أخرى بالاعتماد على الأسئلة المطروحة في كل عقدة، وهكذا تستمر العملية حتى الوصول إلى العقدة الورقية النهائية التي تمثل النتيجة النهائية (الصنف الذي تنتمي إليه العينة)، ويوضح الشكل (2-5) آلية عمل خوارزمية DT. [15]



الشكل (2-5) آلية عمل خوارزمية DT. [15]

تتألف كل عقدة من ثلاث مكونات أساسية: [16]

1- $C(T)$: مجموعات فرعية من الأصناف (Classes Subset) التي يمكن الوصول إليها من العقدة t .

2- $F(T)$: مجموعة فرعية من السمات (Features Subsets) للعقدة t .

3- $D(T)$: قاعدة اتخاذ القرار (Decision Rule) للعقدة t .

يتم الانقسام في العقد بالاعتماد على معيار معين يعمل على تقييم جودة الانقسام، وله نوعان الأكثر شيوعاً Gini و Entropy تعملان على قياس درجة عدم النقاء (Impurity) في مجموعة البيانات لتحديد أفضل النقاط التي يجب تقسيم البيانات عندها وعندما تكون العينات من نفس الصنف يكون عدم النقاء مساوياً للصفر أي المجموعة نقية وبهذه الحالة لا يعود هناك داعي للمزيد من الانقسامات لأن دقة التصنيف لهذا الصنف تصبح مرتفعة وبالتالي لتقليل حدوث الأخطاء لابد من الوصول إلى مجموعات

نقية قدرة المستطاع. [17]

وتوضح المعادلة (5-2) طريقة Gini: [18]

$$Gini(S) = 1 - \sum_{j=1}^m P_j^2 \quad \dots\dots(5-2)$$

وتوضح المعادلة (6-2) طريقة Entropy: [18]

$$Entropy(S) = - \sum_{j=1}^m P_j \log P_j \quad \dots\dots(6-2)$$

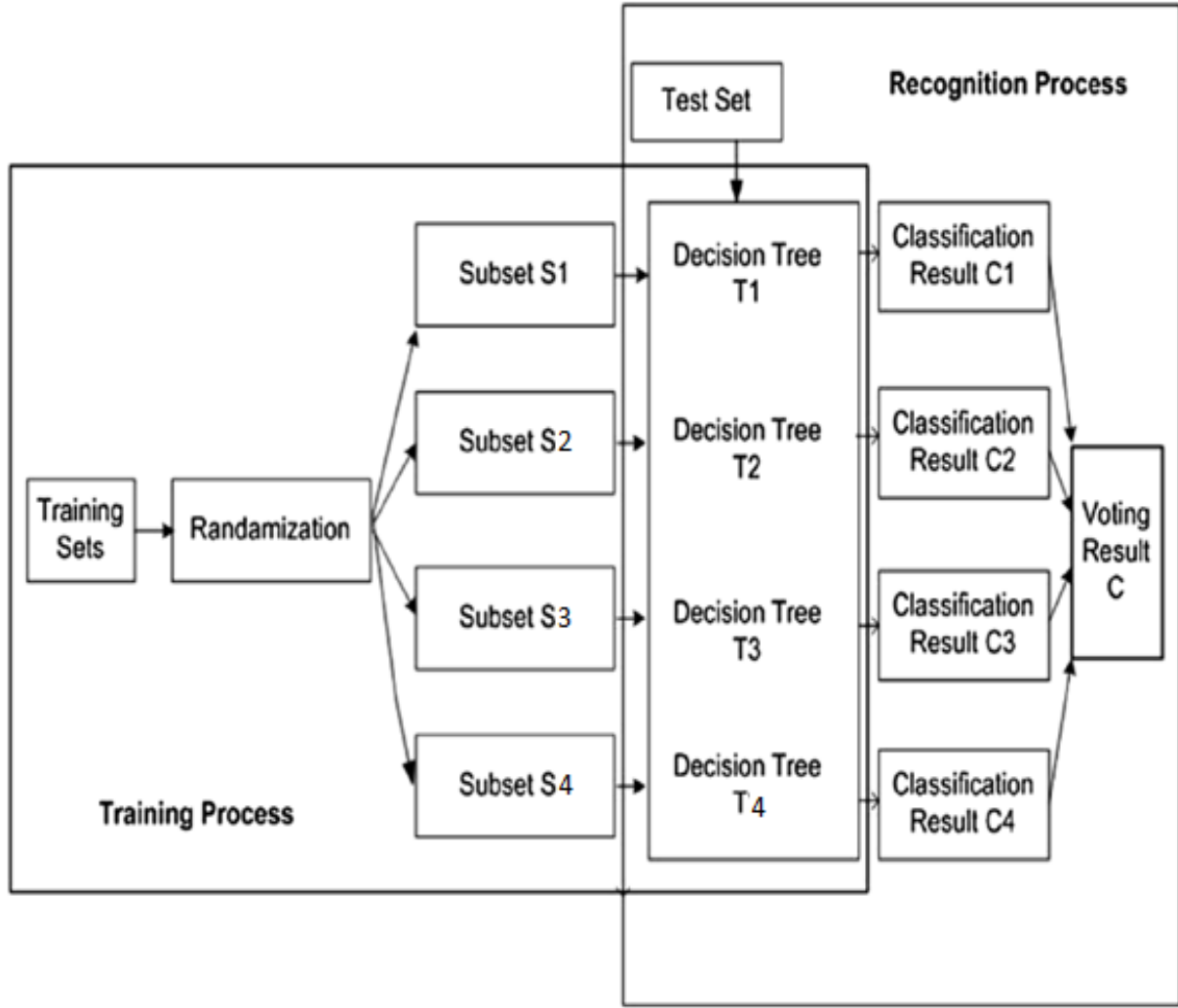
حيث P_j هي تردد الصنف j في مجموعة البيانات S في العقدة.

5-4-2 الغابة العشوائية (Random Forest):

أول اقتراح خوارزمية RF كان من قبل Breiman في جامعة كاليفورنيا عام 2001 وهي عبارة عن مجموعة مصنفات أساسية وهذه المصنفات هي أشجار قرار DT والتي تكون مستقلة عن بعضها البعض بشكل كامل فعند اختيار عينة جديدة يتم تحديد صنفها بناءً على تصويت كل من تلك المصنفات التي تعطي صوتها (قرارها) بشكل مستقل حيث الصنف الذي يحصل على أكثر أصوات يكون هو صنف تلك العينة أي أن النتيجة النهائية لخوارزمية RF تعتمد على التصويت (Voting).

إن تسمية Forest جاءت كون هذه الخوارزمية تتألف من عدد كبير من مصنفات شجرة القرار، وبالنسبة إلى Random نتيجة للعملية العشوائية أثناء عملية البناء عند اختيار مجموعات التدريب الفرعية (Training Subsets) من مجموعة عينات التدريب الكاملة واختيار مجموعات السمات الفرعية (Features Subsets) وذلك لكل مصنف DT مكون لهذه الخوارزمية وهذا ما يحقق الاستقلالية للمصنفات ويحسن من دقة التصنيف ويساعد على التعميم (Generalization) بشكل أفضل، ويوضح

الشكل (6-2) آلية عمل خوارزمية RF. [19]



الشكل (2-6) آلية عمل خوارزمية RF. [19]

كل ما ذكر في الفقرة السابقة ينطبق على أشجار القرار المكونة لخوارزمية RF، وكون بناء DT واحدة سريع والعمل المتوازي والاستقلالية المحققان في إنشاء DT المكونة لخوارزمية RF يعطي سرعة كبيرة في البناء ووقت أقل للتدريب والتصنيف بمعدل أعلى من المصنفات ولها قدرة على تجنب فرط التدريب (Overfitting)، وهذا ما يعطي خوارزمية شعبيتها وزيادة استخدامها يوماً بعد يوم، [19] وتوضح

المعادلة (2-7) عمل خوارزمية RF: [20]

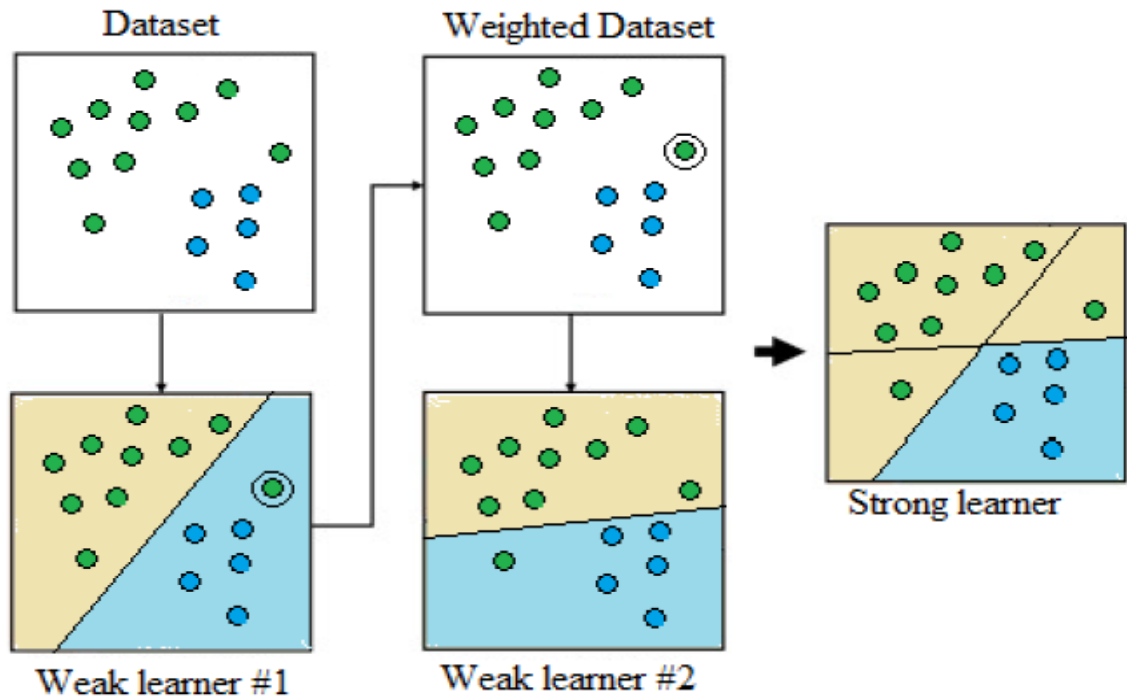
$$H(x) = \arg \max_Y \sum_{i=1}^k I(h_i(x) = Y) \quad \dots\dots(7-2)$$

حيث:

Y: خرج النموذج، k: عدد DT، h_i : النموذج DT رقم i.

6-4-2 التعزيز الكيفي (Adaboost):

يعتبر التعزيز الكيفي (AB) من أحد خوارزميات تعلم الآلة الموجه التي اقترحت في التسعينات من قبل Freund و Schapire، والتي تعتمد في عملها على إنشاء مجموعة من المتعلمين الضعفاء (Weak Learners) ودمجها في النهاية بالمتعلم القوي (Strong Learner) أي النموذج النهائي، حيث يتم تدريب المتعلمين الضعفاء بشكل تسلسلي على أخطاء بعضها البعض حتى نصل إلى المتعلم القوي أي يتم تدريب المتعلم الأول على مجموعة البيانات الأصلية وتقييم أدائه ومنح الحالات التي صنفت بشكل خاطئ أوزان أكبر وتمرر للمتعم الثاني الذي يتدرب مع التركيز على هذه الحالات وهكذا، كما يحدد معدل التعلم (learning rate) درجة تأثير كل متعلم ضعيف بالقرار في النموذج النهائي القوي، إذا بشكل أساسي تركّز خوارزمية AB على تصحيح الأخطاء من خلال تعديل الأوزان، ويوضح الشكل (7-2) آلية عمل خوارزمية AB. [21]

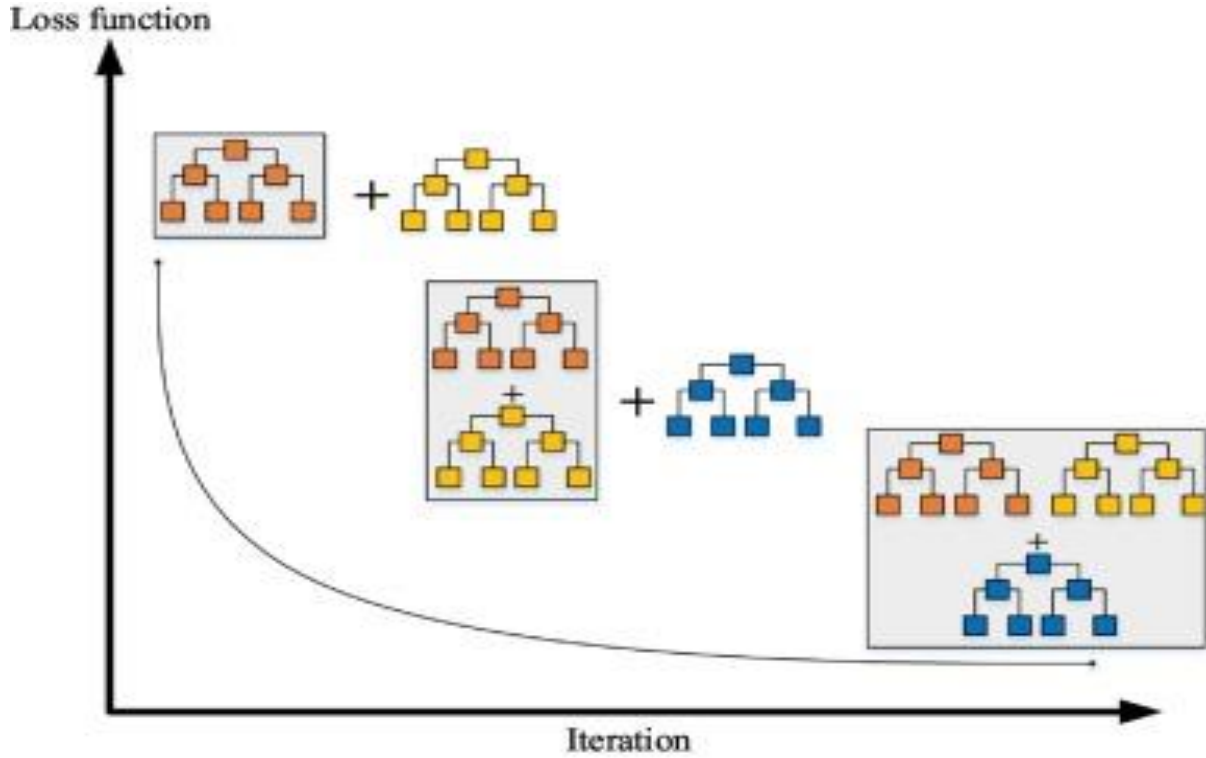


الشكل (7-2) آلية عمل خوارزمية AB. [22]

7-4-2 تعزيز التدرج (Gradient Boosting):

يعد (GB) إحدى خوارزميات تعلم الآلة الموجه وكان أول اقتراح لهذه الخوارزمية من قبل Friedman عام 2001، وتعمل على بناء النموذج النهائي القوي من خلال دمج مجموعة من المتعلمين

الضعفاء Weak Learners والتي يتم تدريبها بشكل تسلسلي بحيث كل متعلم جديد يعمل على تحسين الأخطاء التي ارتكبها المتعلم الذي قبله وذلك باستخدام تابع الخسارة (Loss Function) في تحديد المناطق التي تحتاج إلى تصحيح، كما يحدد معدل التعلم درجة تأثير كل متعلم ضعيف بالنموذج النهائي القوي، إذا بشكل أساسي تركّز خوارزمية GB على تقليل تابع الخسارة تدريجياً من خلال تحسين المتعلم، ويوضح الشكل (2-8) آلية عمل خوارزمية GB.[23]



الشكل (2-8) آلية عمل خوارزمية GB.[24]

5-2 مقاييس الأداء :

تعد مقاييس الأداء مفيدة للغاية في تقييم ومقارنة نماذج التعلم الآلة المختلفة واختبار قدرة المصنف حيث تستخدم لمقارنة أداء نموذجين مختلفين أو تحليل سلوك نفس النموذج عن طريق ضبط محدداته الفائقة المختلفة، وتستند هذه المقاييس إلى مصفوفة الارتباك لأنها تحتوي على جميع المعلومات ذات الصلة حول أداء النماذج وقاعدة التصنيف، ومن هذه المقاييس الدقة، الحساسية، الإحكام، ودرجة F1.[25]

مصفوفة الارتباك عبارة عن جدول متقاطع يسجل عدد مرات ظهور الحالات بين الأصناف الحقيقية بحسب مخرجات البيانات المعطاة للخوارزمية والأصناف التي قامت الخوارزمية بالتنبؤ به، ويقع على

القطر الرئيسي الحالات التي يتوافق فيها الصنف المتنبئ به مع الحقيقي، ويبين الشكل (2-9) مصفوفة الارتباك متعددة الأصناف في حالة الصنف b هو موضع الاهتمام. [25]

		Predicted Class			
		C_1	C_2	...	C_N
Actual Class	C_1	$C_{1,1}$	FP	...	$C_{1,N}$
	C_2	FN	TP	...	FN
	...	...	...	...	...
	C_N	$C_{N,1}$	FP	...	$C_{N,N}$

الشكل (2-9) مصفوفة الارتباك متعددة الأصناف. [26]

حيث أن:

TP هي True Positive تشير إلى الحالات التي تنبأت الخوارزمية تصنيفها بشكل صحيح على أنها تنتمي إلى الصنف b ، مما يعني أن التصنيف الذي تنبأت به الخوارزمية يتطابق مع التصنيف الحقيقي. FP هي False Positive تشير إلى الحالات التي تنبأت الخوارزمية تصنيفها بشكل خاطئ على أنها تنتمي إلى الصنف b بينما في الواقع تنتمي إلى صنف آخر، مما يعني أن التصنيف الذي تنبأت به الخوارزمية غير متطابق مع التصنيف الحقيقي.

TN هي True Negative تشير إلى الحالات التي تنبأت الخوارزمية تصنيفها بشكل صحيح على أنها لا تنتمي إلى الصنف b ، مما يعني أن الخوارزمية قد تنبأت بشكل صحيح أن هذه الحالات تنتمي إلى الأصناف الأخرى بشكل مطابق للتصنيف الحقيقي لهذه الحالات.

FN هي False Negative تشير إلى الحالات التي تنبأت الخوارزمية تصنيفها بشكل خاطئ على أنها لا تنتمي إلى الصنف b ، بينما هي في الواقع تنتمي إليه بحسب التصنيف الحقيقي لهذه الحالات.

عندما يتم الانتقال إلى صنف آخر نقوم بإعادة حساب القيم Positive و Negative مرة أخرى حيث تتغير تسمية مربعات مصفوفة الارتباك متعددة الأصناف حسب الصنف المستهدف، وهذا يعني إعادة

تقييم كل من TP، FP، TN، و FN بناءً على الصنف الجديد الذي يتم تحليله. [25]

ويتم حساب الدقة وفق المعادلة (2-8): [25]

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad \dots\dots(8-2)$$

ويتم حساب الحساسية وفق المعادلة (9-2): [25]

$$Recall_k = \frac{TP_k}{TP_k + FN_k} \quad \dots\dots(9-2)$$

ويتم حساب الإحكام وفق المعادلة (10-2): [25]

$$Precision_k = \frac{TP_k}{TP_k + FP_k} \quad \dots\dots(10-2)$$

ويتم حساب درجة F1 وفق المعادلة (11-2): [25]

$$F1-Score = \left(\frac{2}{precision^{-1} + recall^{-1}} \right) = 2 \cdot \left(\frac{precision \cdot recall}{precision + recall} \right) \quad \dots(11-2)$$

وتحسب هذه المقاييس للنموذج ككل وفي حالة العينات المتوازنة في جميع الأصناف والمحققة في هذا البحث وفق المعادلات التالية (12-2) (13-2) (14-2): [25]

$$MacroAveragePrecision = \frac{\sum_{k=1}^K Precision_k}{K} \quad \dots\dots(12-2)$$

$$MacroAverageRecall = \frac{\sum_{k=1}^K Recall_k}{K} \quad \dots\dots(13-2)$$

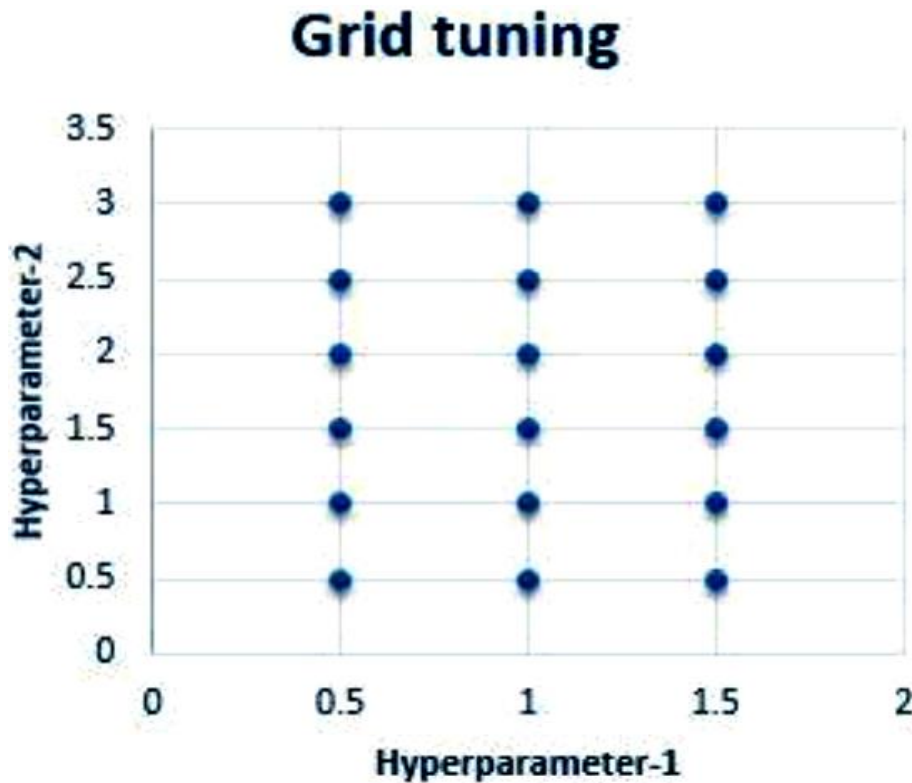
$$Macro F1-Score = 2 * \left(\frac{MacroAveragePrecision * MacroAverageRecall}{MacroAveragePrecision^{-1} + MacroAverageRecall^{-1}} \right) \quad \dots(14-2)$$

حيث k هي الصنف وK هي عدد الأصناف.

2-6 ضبط المحددات الفائقة:

يعد ضبط المحددات الفائقة (Hyperparameters Tuning) خطوة أساسية في العثور على أفضل تركيبة من القيم للمحددات الفائقة لنماذج تعلم الآلة، وهي من الخطوات المهمة التي يمكن أن تؤثر بشكل مباشر على نتيجة التدريب، علماً أن الضبط الدقيق يزيد من الدقة التنبؤية للنموذج [27]. وعلى عكس المحددات (parameters) التي يحددها النموذج أثناء التدريب يجب إسناد قيم للمحددات الفائقة قبل بدء التدريب، وتحديدتها يدوياً يستغرق قدراً كبيراً من الوقت خاصةً عند الحاجة لضبط عدد كبير من هذه المحددات، لذلك استخدام الطرق الآلية يوفر الكثير من الوقت ويبحث في مجموعة قيم أكبر للوصول إلى أفضل تركيبة من القيم للمحددات الفائقة [27].

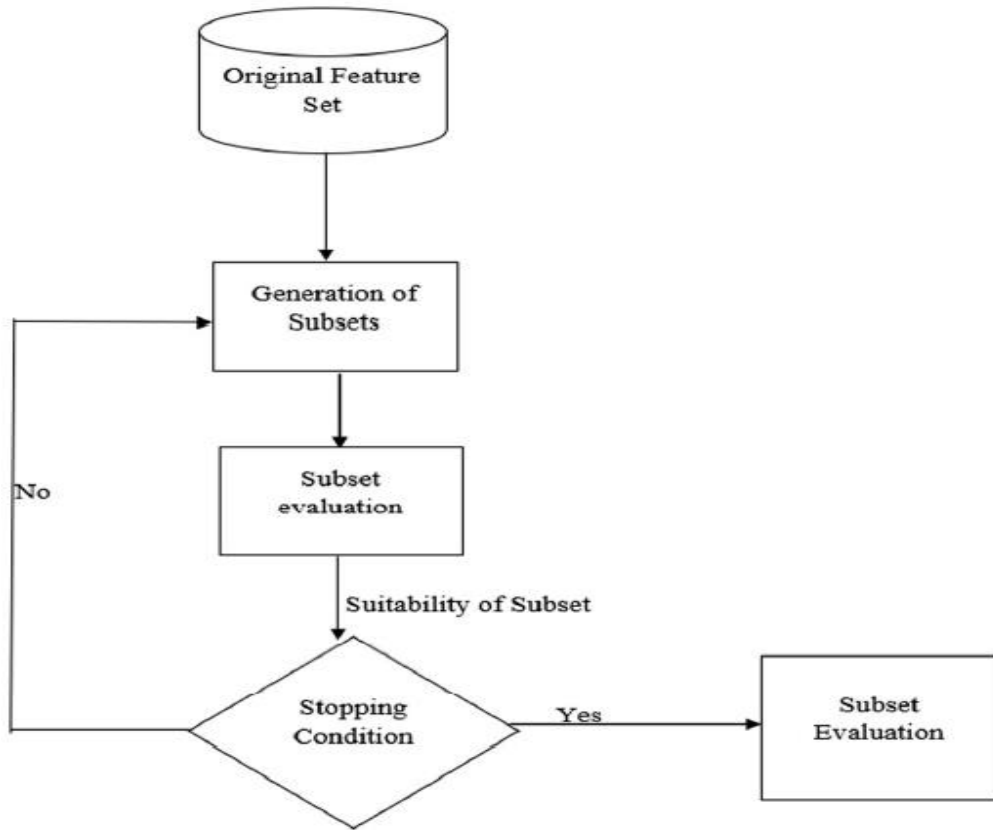
تم استخدام البحث الشبكي Grid Search مع التحقق المتقاطع Cross Validation في هذا البحث، وهو خوارزمية استكشافية مستقلة شاملة تأخذ في الاعتبار كل مجموعة ممكنة من القيم المعطاة لاختيار المحددات الفائقة منها، حيث تختبر وتقيم أداء كل مجموعة فريدة من المحددات الفائقة في مساحة البحث لتحديد المجموعة التي تعطي أفضل أداء، بالإضافة لكونها طريقة بسيطة ومباشرة يمكن تنفيذها بسهولة [28] ويوضح الشكل (2-10) طريقة البحث الشبكي.



الشكل (2-10) طريقة البحث الشبكي [27]

7-2 تقليل السمات:

أو ما يعرف باختبار السمات (Feature Selection) وهي إحدى جوانب هندسة السمات (Feature Engineering) التي تهدف إلى اختيار أفضل مجموعة فرعية من السمات ذات الصلة من مجموعة سمات البيانات الأصلية باستخدام معايير تقييم محددة، وذلك لتقليل أبعاد المشكلة التي يتم معالجتها مما يسهل عملية التصنيف حيث تم تقليل القياس في مجموعات البيانات ذات الأبعاد الكبيرة مما يعطي أداء أفضل في التعلم، ويوضح الشكل (2-11) آلية Feature Selection العامة. [29]



الشكل (2-11) آلية اختيار الخصائص العامة. [28]

بشكل عام، تم تطوير طرق مختلفة من Feature Selection للحصول على مجموعات فرعية مثالية، وهي ثلاث طرق رئيسية:

1- طريقة المرشح (Filter Approach): يتضمن مقياساً مستقلاً لتقييم المجموعات الفرعية من السمات دون إشراك إحدى خوارزميات تعلم الآلة، وهو فعال وسريع حسابياً ولكن من الممكن أن تفوت ميزات ليست مفيدة في حد ذاتها ولكنها يمكن أن تكون مفيدة جداً عند دمجها مع ميزات أخرى. [30]

وتعد طريقة اختيار K الأفضل (Select K Best (SKB)) تابعة لهذا النهج، حيث تتطلب تحديد الطريقة الإحصائية المناسبة (مثل Chi-Square، T-test ، F-test) لتكون المقياس لتقييم المجموعات الفرعية من السمات وكذلك تحديد عدد السمات المطلوب الوصول إليه.

2- طريقة التغليف (Wrapper Approach): يستخدم إحدى خوارزمية تعلم الآلة لتقييم المجموعة الفرعية من السمات حيث يتطلب تدريب الخوارزمية عدة مرات، ويختار مجموعة السمات الفرعية المناسبة لخوارزمية التعلم الآلة بالشكل الأفضل لذلك يكون أدائه أفضل عادةً. [30]

وتعد طريقة استبعاد السمات العودية (Recursive Feature Elimination (RFE)) تابعة لهذا النهج وحيث تتطلب تحديد خوارزمية تعلم الآلة (غالباً الغابة العشوائية أو الانحدار اللوجستي) لتكون المقياس لتقييم المجموعات الفرعية من السمات وكذلك تحديد عدد السمات المطلوب الوصول إليه.

3-: طريقة المضمنة (Embedded Approach): وهو دمج بين طريقة المرشح وطريقة التغليف حيث يختار مجموعة السمات الفرعية المناسبة بشكل مباشر أثناء التدريب. [30]

وتعد طريقة الشبكة المرنة (Elastic Net (EN)) تابعة لهذا النهج وتتبع خوارزميتي تعلم الآلة Lasso وRidge في السمات عن طريق تحديد عامل للعقوبة أثناء التدريب.

8-2 الخاتمة:

تم العمل في هذا الفصل على تغطية الجانب النظري لهذا البحث حيث استعرض معلومات عن مرض الزهايمر وتعلم الآلة والأسس النظرية للنماذج التي تم استخدامها في هذه الدراسة وكيفية إيجاد مقاييس تقييم أدائها وكفاءتها بالإضافة لتوضيح تقليل السمات.

الفصل الثالث:

أدبيات البحث

3-1 مقدمة:

جذب انتشار الأمراض العصبية التنكسية مثل مرض الزهايمر اهتمام الباحثين في جامعات العالم، والذين يعملون على تطوير منهجيات ذات أداء عالي للتشخيص والعلاج والوقاية واكتشاف الأدوية المستهدفة، وذلك انطلاقاً من أهمية الكشف المبكر عن الزهايمر الذي من الممكن أن يكون مفتاح للعلاج الفعال، أو على الأقل التخطيط المستقبلي لعلاج مراحل هذا المرض، ورفع سوية الحياة والرعاية عند المرضى، ويستعرض هذا الفصل الدراسات المرجعية.

3-2 الدراسات المرجعية:

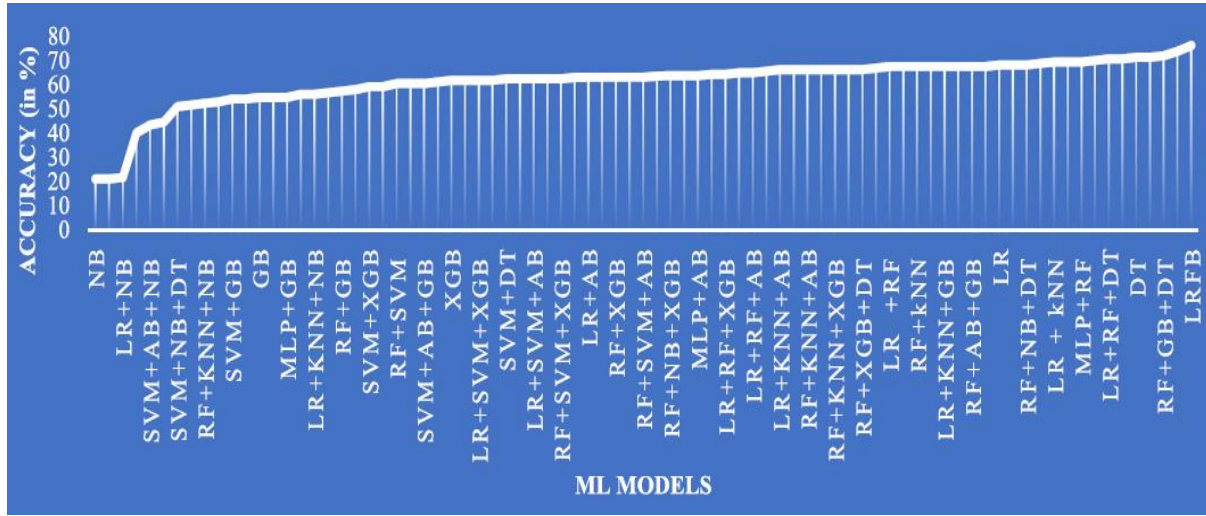
3-2-1 استخدام طرق تعلم الآلة من أجل التنبؤ المبكر بمرض الزهايمر والكشف عنه:

في دراسة تهدف للتنبؤ والتشخيص المبكر بمرض الزهايمر عمل Shastry وزملاؤه عام 2023 على تطوير نموذج يجمع عدة خوارزميات تعلم الآلة لتشخيص مرض الزهايمر وتحديد الأشخاص المعرضين لخطر الإصابة به بهدف تحسين نوعية الحياة لهؤلاء المرضى، وذلك كون الزهايمر مرض تنكسي خطير من الصعب تشخيصه بشكل مبكر، تم استخدام البيانات من قاعدة بيانات مبادرة التصوير العصبي لمرض الزهايمر (ADNI) للتصنيف والتنبؤ بثلاث حالات: الإدراك الطبيعي (CN)، تدهور الإدراك المعتدل المتأخر (LMCI)، الزهايمر (AD). [31]

تضمنت آلية العمل الموضوعية مرحلتين: وتتضمن المرحلة الأولى (معالجة البيانات): اختيار السمات، ترميزها، تقسيمها للتدريب والاختبار، وتقييمها، بينما تتضمن المرحلة الثانية (التصنيف): تدريب خوارزميات تعلم الآلة، اختبارها، تقييم أدائها، واختيار أفضلها لتصنيف الحالات الثلاثة. [31]

اقتصرت الدراسة على 628 سجلاً من سجلات قاعدة بيانات ADNI بنسبة 80% للتدريب (502 سجلاً)، و20% للاختبار (126 سجلاً)، وتم استخدام k-fold cross validation حيث $K=10$ لتدريب مجموعة كبيرة من خوارزميات تعلم الآلة، وهي: Decision Tree، K-Nearest، Naïve، Multi-Layer Perceptron، Support Vector Machine، XGBoost، Neighbour، Gradient Boost، Random Forest، Logistic Regression، AdaBoost، Bayes، واستخدام هذه الخوارزميات بشكل مفرد بالإضافة إلى الدمج فيما بينها، وكان أفضل نموذج تم التوصل إليه هو LRFB (Logistic Random Forest Boosting) والناتج من دمج الخوارزميات الثلاث

LR، RF، وGB، وذلك بدقة 76.19% وحساسية 73.89% ، ويبين الشكل(3-1) دقة النماذج التي تم تدريبها في هذه الدراسة.[31]



الشكل(3-1) دقة النماذج دراسة Shastry.[31]

وكانت هذه الدراسة جيدة حاولت المقارنة بين أغلب نماذج تعلم الآلة ولكن لم تقم بمعالجة القيم المتطرفة ولم تذكر عدد العينات في كل صنف بشكل صريح وإذا ما تم الموازنة فيما بينها. وفي دراسة أخرى قام Dhakal وزملاؤه عام 2023 على تصميم وتطوير نموذج للتنبؤ بالخرف الناتج عن الزهايمر عن طريق توظيف أساليب تعلم الآلة للوصول إلى أعلى دقة وأفضل أداء، وذلك انطلاقاً من كون الخرف يشكل مصدر تحدي رئيسي فيما يخص صحة كبار السن، وصعوبة رفع سوية الرعاية الصحية المقدمة لهم.[32]

استخدمت البيانات التي تم جمعها بواسطة مشروع سلسلة الوصول المفتوح لدراسات التصوير (OASIS) والذي أتاحه مركز أبحاث مرض الزهايمر بجامعة واشنطن، وذلك للتصنيف والتنبؤ بثلاث حالات: الأشخاص السليمين (Nondemented)، المصابين بالخرف الزهايمر (Demented)، الأشخاص المصابين بضعف الإدراك المعتدل والذين سيتحولون لمصابين بالخرف (Converted)، وكانت أعداد العينات ذكوراً وإناثاً في كل صف 72، 64، و14 بنفس الترتيب على التوالي، وحيث كانت أعمارهم تتراوح بين 60 و96 عاماً.[32]

تضمنت آلية العمل الموضوعية 4 مراحل: المرحلة الأولى: وتتضمن جمع البيانات، المرحلة الثانية: تهيئة البيانات من خلال معالجة القيم المفقودة (طريقة المتوسط)، حذف القيم الشاذة، التقييس، التحويل، واختيار السمات حيث تم استخدام خوارزميتين لتقليل السمات وهما LASSO و Chi-square، والمرحلة

الثالثة: تقسيم البيانات بنسبة 70% لمجموعة التدريب و30% لمجموعة الاختبار، والمرحلة الرابعة: تدريب النماذج (باستخدام السمات OASIS كاملة، سمات LASSO، وسمات Chi كلاً على حد)، وتقييم أداء هذه النماذج بالاعتماد على الدقة، واستخدام cross validation حيث $K=10$ ، وبالنهائية الوصول إلى النتيجة المطلوبة. [32]

تم تنفيذ تسعة من تقنيات التعلم الآلة الخاضعة للإشراف، وهي: Decision Tree، Extra Tree، K-Nearest Neighbour، Support Vector Machine، Naïve Bayes، AdaBoost، Logistic Regression، Random Forest، Gradient Boost. [32]

الجدول (3-1) دقة وحساسية الخوارزميات المدربة دون التحقق المتقاطع في دراسة Dhakal [31]

	All Features		Chi Features		Lasso Features	
	Accuracy	Recall	Accuracy	Recall	Accuracy	Recall
AdaBoost	88.70	87.93	88.77	87.93	88.70	87.93
Decision Tree	89.51	89.60	91.93	89.65	89.51	89.65
Extra Tree	83.87	87.93	91.12	89.65	76.61	73.77
Gradient Boost	88.70	87.93	93.54	87.93	89.51	87.93
K-Nearest Neighbour	95.16	91.37	94.35	87.93	95.16	100
Logistic Regression	94.35	87.93	94.35	87.93	94.35	87.93
Naïve Bayes	94.35	87.93	94.35	87.93	94.35	87.93
Random Forest	94.35	87.93	94.35	87.93	94.35	87.93
Support Vector Machine	96.77	87.93	94.35	87.93	94.35	87.93

وأظهرت أغلب الخوارزميات دون التحقق المتقاطع (Cross Validation) دقة تصنيف أكبر من 90% وكانت خوارزمية SVM الأعلى بدقة 96.77% وحساسية 87.93% وذلك مع السمات كاملة كما يوضح بالجدول (3-1)، وزادت دقة الخوارزميات مع التحقق المتقاطع لحوالي 94%. [32]

كانت نتائج هذه الدراسة واعدة حيث أجرت معالجة أولية جيدة للبيانات لكن استخدمت عدد قليل وغير متوازنة من العينات في الأصناف الثلاثة.

وفي دراسة أخرى قام ElZawawi وزملاؤه عام 2022 باقتراح إستراتيجية مختلطة من طرق تعلم الآلة تتضمن خوارزمية الغابة العشوائية مع طريقة الأسراب الجزئية في الأمثلة (RF-PSO) لاختيار أفضل

السمات ذات الصلة الوثيقة بمرض الزهايمر كون السمات دون صلة والمكررة ستسبب تدهور كبير عند تدريب المصنفات. [33]

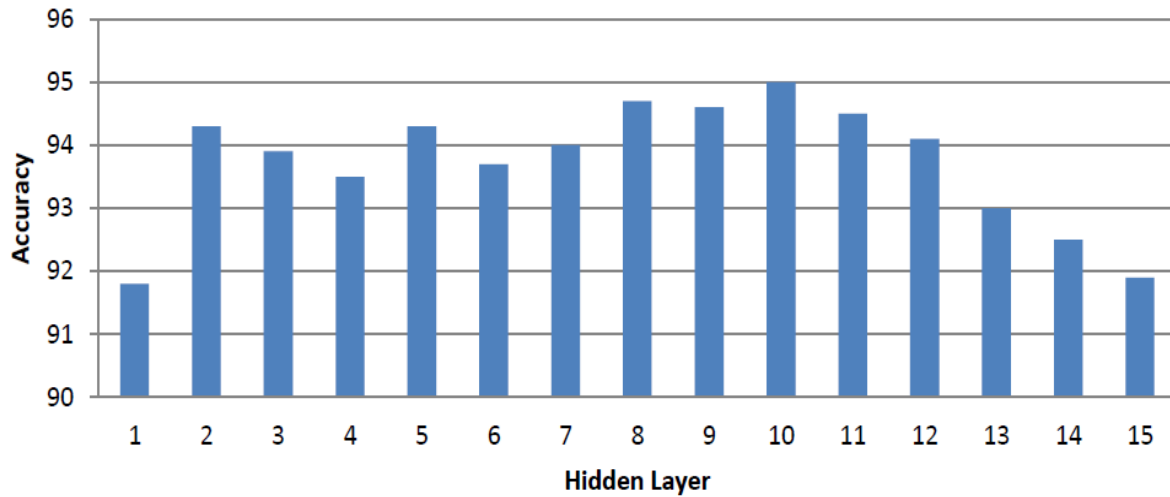
جمعت البيانات من قاعدة بيانات مبادرة التصوير العصبي لمرض الزهايمر ADNI، وذلك للتصنيف والتنبؤ بثلاث حالات: NL الإدراك الطبيعي، MCI ضعف الإدراك المعتدل، AD مرض الزهايمر، وحيث تم استخدام 45 سمة، و2700 سجلاً 900 في كل صنف. [33]

في البداية تم إعداد البيانات بحذف السمات التي تجاوزت نسبة القيم المفقودة فيها 60% وتم تعويض قيم الفقد في السمات المتبقية بطريقة المتوسط وبعضها باستخدام التقسيم لفئات، ثم تم تطبيق ثلاث طرق لاختيار أفضل السمات، وهي: Partial Swarm Optimization، Random Forest، و Random Forest Partial Swarm Optimizatio، وكانت دقة كل منها بالترتيب 89%، 86%، و95%، ويظهر الجدول (2-3) السمات التي تم اختيارها بكل من الطرق المذكور. [33]

الجدول (2-3) السمات المختارة في دراسة ElZawawi. [33]

Original Dataset	PSO	RF	RF-PSO
Baseline (19 attribute)	5	14	3
Socio (9 attribute)	1	2	1
MRI (7 attribute)	1	3	5
Neurological (9 attribute)	5	7	5

وكون طريقة RF-PSO كانت ذات الدقة الأعلى تم استخدام السمات التي اختارتها في تدريب شبكة عصبونية ذات الانتشار العكسي، وتابع الخطأ (Loss Function) ممثل بمربع الخطأ، وتابع التفعيل (Activation Function) هو Sigmoid، وتم زيادة عدد الطبقات المخفية من طبقة واحدة وحتى 15 طبقة، ويوضح الشكل (2-3) دقة الشبكة العصبونية حسب عدد الطبقات المخفية. [33]



الشكل (2-3) دقة الشبكة العصبونية حسب عدد الطبقات المخفية في دراسة ElZawawi [33]

تعد هذه الدراسة جيدة من حيث المنهجية المتبعة وموازنة العينات ولكن أغفلت ذكر أسماء السمات الناتجة عن طرق الاختيار بشكل صريح.

2-2-3 استخدام السجلات الصحية الإلكترونية للكشف عن مرض الزهايمر غير المشخص:

قام Stephan وزملاؤه عام 2021 بالعمل على إيجاد أساليب منخفضة التكلفة ومتوفرة بشكل دائم لفحص الأفراد والتنبؤ بخطر الإصابة بمرض الزهايمر أو الكشف المبكر عن الإصابة به، وذلك بالاستفادة من بيانات السجلات الصحية الإلكترونية، والحوسبة منخفضة التكلفة، والترجمة الإحصائية الحديثة، وخوارزميات تعلم الآلة حيث تم استخدام ثلاث خوارزميات Random Forest، Support Vector Machine، وAdaBoost، وكما يبين الجدول (3-3) دقة وحساسية كل من هذه الخوارزميات. [34]

الجدول (3-3) دقة وحساسية كل من هذه الخوارزميات في دراسة Stephan [34]

MODEL	RF	SVM	Adaboost
ACCURACY	0.8710	0.8754	0.9081
SENSITIVITY	0.1828	0.3071	0.1235

تم جمع البيانات من مركز العلوم الصحية بجامعة شمال تكساس، وكانت السجلات مرضية من 2016 وحتى 2019، لمرضى بعمر 65 سنة وأكبر، وحيث كان هناك 4995 مريض مميز (منهم 3950 غير مصاب بالزهايمر، و516 مصاب بمرض الزهايمر، و889 مصاب بالخرف (ناتج عن الزهايمر أو

غيره))، ويظهر الجدول (3-4) السمات العشرة الأعلى ارتباطاً بالإصابة بالزهايمر، والتي تم تحديدها باستخدام خوارزمية Random Forest، والمستخدم في تدريب الخوارزميات، والأهمية الإحصائية لها. [34]

الجدول (3-4) السمات العشرة الأعلى ارتباطاً بالإصابة بالزهايمر، والأهميتها الإحصائية في دراسة Stephan. [34]

الأهمية الإحصائية	الخصائص (الميزات)	
0.03366	معدل ضربات القلب	الفحص الصحي
0.05520	الوزن	
0.00927	معدل التنفس	
0.03214	ضغط الدم الانقباضي	
0.03214	ضغط الدم الانبساطي	
0.02867	درجة حرارة الجسم	
0.04997	كتلة الجسم	
0.02968	سكر الدم	
0.02048	علامات وأعراض لأمراض عصبية أخرى	الأمراض المصنفة
0.00927	الاكتئاب	
0.01274	عدد الزيارات للفحص الطبي	الديموغرافية

كانت الخطوات المتبعة في الدراسة جيدة لكن لم يتم موازنة العينات وبالإضافة إلى أن النماذج كانت ذات حساسية منخفضة.

وفي دراسة أخرى تهدف إلى التنبؤ بحدوث مرض الزهايمر باستخدام تعلم الآلة وبيانات صحية إدارية واسعة قام Park وزملاؤه عام 2020 باختبار إمكانية تنبؤ نماذج التعلم الآلة بحدوث مرض الزهايمر في المستقبل واكتشافه في وقت مبكر، وذلك بالاعتماد على المؤشرات الحيوية الموجودة في السجلات الطبية، وحيث تم استخدام خوارزميات Logistic Regression، Random Forest، Support Vector Machine. [35]

البيانات المستخدمة من قاعدة بيانات خدمة التأمين الصحي الوطنية الكورية المسجلة بين عامي 2002 و2010، وذلك بعدد عينات 40736 لأفراد أعمارهم فوق 65 سنة حيث 614 فرد مميز مع الإصابة

بمرض الزهايمر وفق تعريف زهايمر مؤكد، 2026 فرداً محتمل الإصابة بمرض الزهايمر وفق تعريف زهايمر محتمل، 38710 فرداً كبير في السن دون الإصابة بالزهايمر. [35]

وأظهرت خوارزمية Random Forest أفضل النتائج حيث كانت أعلى قيمة Accuracy وصلته إليها تساوي 82.30%، وبين الجدول (3-5) السمات العشرة الأعلى ارتباطاً بالزهايمر، والتي تم تحديدها باستخدام خوارزمية Logistic Regression، والمستخدم في تدريب الخوارزميات، وقيمة الـ P-value لها. [35]

الجدول (3-5) يبين السمات العشرة الأعلى ارتباطاً بالزهايمر، وقيمة الـ P-value لها في دراسة Park. [35]

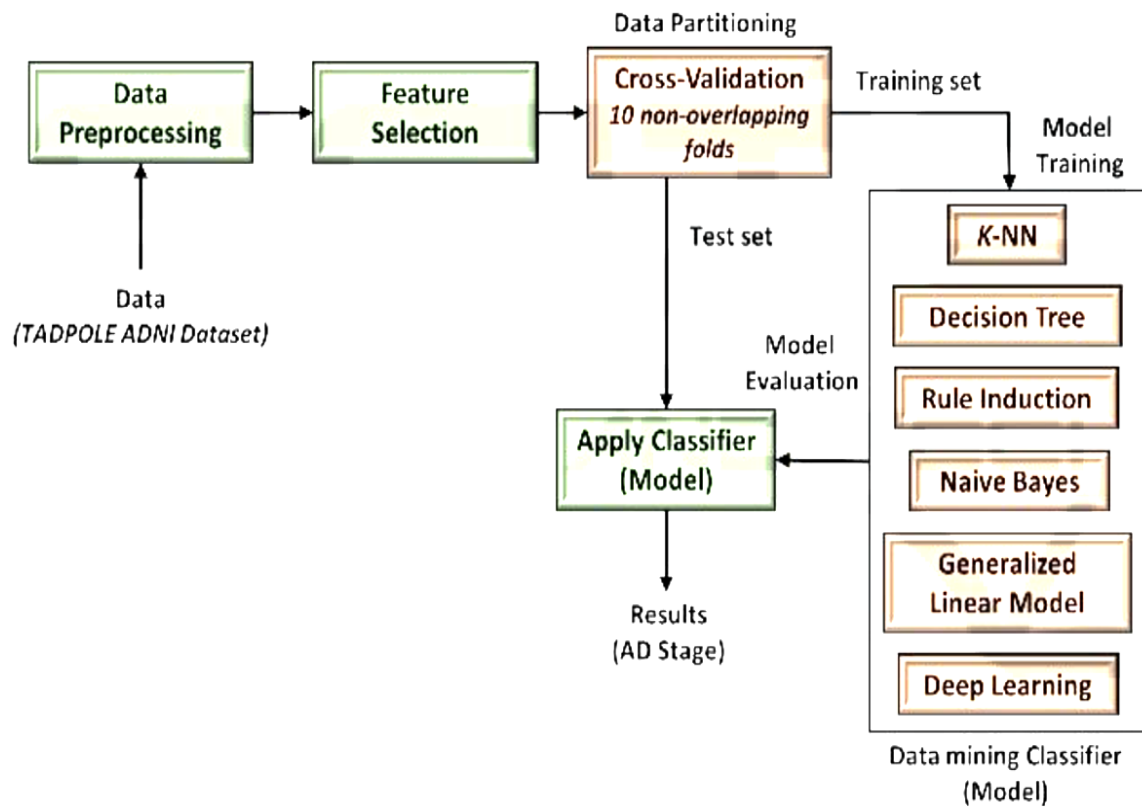
الخصائص (الميزات)	P-value
الفحص الصحي	0.001 >
الهيولوبين	0.001 >
بروتين في البول (الزلال)	0.001 >
الديموغرافية	0.001 >
العمر	0.001 >
الدواء المستخدم	0.001 >
زوتيبين (مضاد ذهان)	0.001 >
نيكومات سيترات (موسع للأوعية الدموية)	0.001 >
حمض التوفيناميك (مسكن الم)	0.001 >
ابيريزون هيدروكلوريد (مضاد تشنج)	0.001 >
الأمراض المصنفة	0.001 >
اضطرابات تنكسية عصبية أخرى	0.001 >
اضطرابات الأذن الخارجية	0.001 >
متلازمة الضائقة التنفسية	0.001 >

كانت دراسة جيدة لكن كان هناك تفاوت كبير في عدد العينات حيث عدد السليمين كان أكبر بكثير من باقي الأصناف مما قد يسبب انحياز في النماذج.

3-2-3 استخدام التعلم العميق للكشف عن مرض الزهايمر والتنبؤ المبكر به:

في دراسة هدفت إلى تصنيف مرض الزهايمر للكشف عنه وتحديد الأشخاص المعرضين لخطر الإصابة به قام Shahbaz وزملاؤه عام 2019 باستخلاص السمات الأكثر تميزاً لكل مرحلة من

مرحل مرض الزهايمر، وذلك انطلاقاً من كون الكشف المبكر عن مراحل مرض الزهايمر وتصنيفها مفيداً جداً في علاج أعراض المرض. ومن ناحية أخرى، تزايد استخدام الحاسوب وتسجيل بيانات المريض إلكترونياً في أقسام الرعاية الصحية بدلاً من النماذج الورقية مما يتيح عدد كبير من السجلات الصحية الإلكترونية (EHRs)، والتي تسمح باستخراج البيانات منها والاستفادة منها في تطبيق تقنيات التعلم بما يخدم المراكز الطبية والرعاية الصحية، ويرفع الجودة والإنتاجية، ويبين الشكل (3-3) مخطط العمل الدراسة. [36]



الشكل (3-3) مخطط عمل دراسة Shahbaz. [36]

تم تطبيق ست خوارزميات من تعلم الآلة: K-Nearest Neighbors، Decision Tree، Rule، Deep Learning، Generalized Linear Model، Naive Bayes، Induction، مجموعة فرعية من قاعدة بيانات مبادرة التصوير العصبي لمرض الزهايمر (ADNI)، وذلك للأصناف التالية: السليمين (CN)، ضعف الإدراك المعتدل المبكر والمتأخر (EMCI) (LMCI)، الشكوى الذاتية من الذاكرة (SMC)، ومرض الزهايمر (AD)، واستخدموا 2164 سجلاً للتدريب و 927 للاختبار، ويبين الجدول (3-6) عدد العينات في كل صنف. [36]

تم استخدام 27 سمة تتضمن سمات ديموغرافية واختبارات عصبية نفسية وسمات من صور الرنين المغناطيسي، ويظهر الجدول (3-6) دقة الخوارزميات المستخدمة. [36]

الجدول (3-6) دقة الخوارزميات المستخدمة في دراسة Shahbaz. [35]

Classifier	Accuracy (%)	
	Validation	Test
K-NN	73.10	43.26
Decision Tree	76.43	74.22
Rule Induction	92.47	69.69
Naive Bayes	79.44	74.65
Generalized Linear Model	92.75	88.24
Deep Learning	78.79	78.32

كانت الدراسة جيدة لكن لم تحدد طريقة اختيار السمات وكما أن عدد العينات يعتبر قليل بالنسبة إلى التعلم العميق.

وفي دراسة أخرى قام Lee وزملاؤه عام 2019 باستخدام التعلم العميق متعدد النماذج للتنبؤ بتطور مرض الزهايمر، وكان الهدف من الدراسة التنبؤ والكشف المبكر عن الزهايمر في مرحلة مبكرة قبل ظهور الأعراض السريرية، وحيث قاموا باستخراج البيانات الحيوية التكاملية المرتبطة بمرض الزهايمر وتصنيفها في 4 مجموعات، ومن ثم تم تدريب نموذج تنبؤي باستخدام كل مجموعة منها على حدة، وبعد ذلك تم دمج هذه النماذج الأربعة، واستخدامها لتدريب النموذج النهائي. [37]

تم جمع البيانات من قاعدة بيانات مبادرة التصوير العصبي لمرض الزهايمر، وذلك لأشخاص تتراوح أعمارهم بين 55 سنة و 91 سنة، وكان 415 أشخاص طبيعيين الإدراك (CN)، و 865 أشخاص لديهم ضعف الإدراك المعتدل (MCI)، و 338 مريض بالزهايمر (AD)، وكانت المجموعات الأربعة التي توصلوا إليها: [37]

- 1- المعلومات الديموغرافية (Demographic Information): العمر، الجنس، سنوات التعليم.
- 2- المؤشرات الحيوية للسائل الدماغي الشوكي (CFS): بروتين تاو، بروتين بيتا أميلويد.
- 3- الأداء المعرفي (Cognitive Performance): الذاكرة، الحالة العصبية النفسية، تنفيذ الأداء المطلوب.
- 4- المؤشرات الحيوية في صور الرنين المغناطيسي (MRI): حجم الحصين، سماكة القشرة الداخلية.

وكانت دقة النماذج الأربعة بشكل منفصل أقل من 78% وكان الأداء الأفضل للنموذج النهائي الناتج عن هذه المجموعات بدقة تقريبية 82%. [37]

كانت الدراسة جيدة ولكن لم يتم ذكر عدد السمات المستخدمة ولم توازن العينات بين الأصناف كما أن عدد العينات يعتبر قليل بالنسبة إلى التعلم العميق.

3-3 خاتمة:

تعددت الدراسات التي تعمل على تسخير تقنيات تعلم الآلة للوصول إلى التشخيص المبكر لمرض الزهايمر، لما في ذلك خدمة لهؤلاء المرضى الذين تظهر الأعراض السريرية عليهم بعد فترة طويلة من الإصابة، ومعظم الدراسات تؤكد على أن التشخيص المبكر قد يكون مفتاح العلاج.

الفصل الرابع:

منهج البحث وإجراءاته

4-1 مقدمة:

سيقدم هذا الفصل الجانب العملي من البحث، حيث سيتضمن منهجية البحث والخطوات العملية المتبعة لإتمام هذا البحث من تأمين قاعدة البيانات، ومعالجتها وتجهيزتها وتقسيمها، وبناء النماذج، وضبط محدداتها الفائقة، وتدريبها واختبارها، وتقييم أدائها للوصول إلى النموذج الأفضل دقة وكفاءة.

4-2 المنهجية:

تم العمل في هذا البحث وفق المخطط الموضح لمنهجية البحث في الشكل (4-1)، وتتألف هذه المنهجية من سبع خطوات:

حيث كانت الخطوة الأولى تأمين قاعدة البيانات التي ستبنى عليها باقي الخطوات فقد تم اختيار مجموعة بيانات ADNI للعمل عليها في هذا البحث.

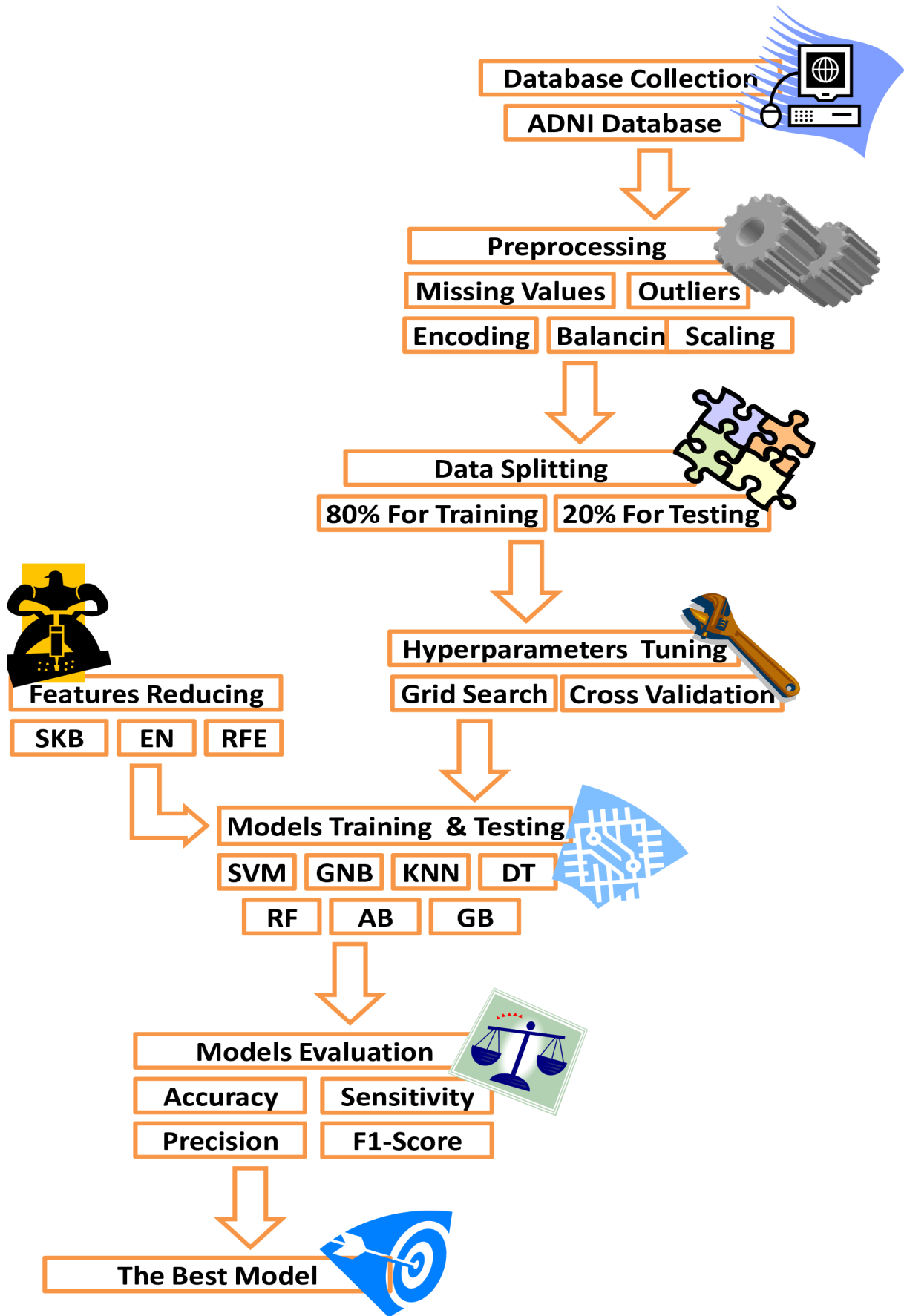
الخطوة الثانية هي من الخطوات الأساسية والضرورية للوصول للهدف المطلوب وتتضمن معالجة وتجهيز البيانات من خلال معالجة القيم المفقودة والشاذة، وبعض الإجراءات الإحصائية ومن ثم تحويل القيم الاسمية إلى رقمية، واختيار السمات، وأخيراً موازنة العينات بين الصفوف، وذلك بهدف رفع كفاءة العمل ودقة النتائج.

الخطوة الثالثة هي تقسيم مجموعة البيانات إلى مجموعتين من أجل تدريب النماذج واختبارها حيث قسمت إلى 80% من البيانات لتشكيل مجموعة التدريب و20% منها لتشكيل مجموعة الاختبار. وبالنسبة للخطوة الرابعة كانت بناء نماذج تعلم الآلة حيث تم بناء سبعة أنواع من نماذج تعلم الآلة ومن ثم ضبط المحددات الفائقة الخاصة بكل نموذج وذلك للوصول إلى أعلى كفاءة باستخدام البحث الشبكي (Grid Search) والتحقق المتقاطع (Cross Validation).

الخطوة الخامسة تدريب النماذج واختبارها باستخدام المحددات التي نتجت من عملية الضبط.

الخطوة السادسة هي تقييم أداء النماذج لتحديد النموذج الأفضل.

ومن ثم تم إضافة الخطوة السابعة وهي خطوة إضافية وذلك للوصول إلى أقل عدد ممكن من السمات مع أعلى دقة حيث تم إعادة تدريب واختبار وتقييم أداء النماذج لتحديد أفضل نموذج وأفضل مجموعة سمات تم الوصول لها.



الشكل (1-4) منهجية البحث المتبعة.

4-3 قاعدة البيانات:

تم العمل على قاعدة البيانات (Database) التي تم جمعها بواسطة مشروع ADNI (Alzheimer's Disease Neuroimaging Initiative) المتوفرة على الموقع التالي (adni.loni.usc.edu)، تم إطلاق مبادرة الصور العصبية لمرض الزهايمر (ADNI) في عام 2004 كشراكة عامة-خاصة بإشراف الباحث الرئيسي الدكتور Michael W. Weiner، وبتمويل من المعهد الوطني للشيخوخة وشركات الأدوية والمنظمات غير الربحية، وتعد مبادرة شاملة والأكبر من نوعها حيث تشمل أكثر من 50 موقعاً بحثياً في الولايات المتحدة وكندا.

وبشكل عام، مبادرة ADNI هي مشروع بحثي مبتكر قد قدم تقدماً كبيراً في فهم مرض الزهايمر من خلال تحديد علامات حيوية ودراسة تطور المرض، تعمل ADNI على البحث من أجل تشخيص وعلاج مرض الزهايمر، مما يؤدي في النهاية إلى تحسين حياة الملايين من الأشخاص المتأثرين بهذا الحالة المدمرة.

4-4 معالجة وتهيئة البيانات:

تعد مرحلة معالجة الأولوية للبيانات (Data Preprocessing) مهمة لتحسين جودة البيانات وتجنب الأخطاء، كما تساعد في تسهيل عملية التحليل وتحسين أداء النماذج فهي مرحلة أساسية للوصول إلى نتائج دقيقة.

4-4-1 القيم المفقودة:

تعتبر القيم المفقودة (Missing Values) من المشاكل الشائعة في قواعد البيانات نتيجة فقدان جزء من المعلومات وتكون عبارة عن خلايا فارغة أو تحتوي على كلمات مثل Nan وUnknown، وتحدث القيم المفقودة بسبب أخطاء في عملية إدخال البيانات أو عوامل خارجية أخرى مثلاً انسحاب المشاركين أو عدم الحصول على قيمة معينة لإحدى المتغيرات.

بعد حذف الأعمدة الإدارية وأعمدة خط البداية، تم معالجة القيم المفقودة على ثلاث خطوات: الخطوة الأولى تم حذف جميع أعمدة المتغيرات بشكل كامل والتي تحتوي في قيمها على أكثر من 50% قيم مفقودة، وذلك حرصاً على عدم تشويه البيانات وإفقادها مضمونها كون نسبة القيم المفقودة فيها كبيرة.

الخطوة الثانية حذف القيم المفقودة فقط من عمود متغير الخرج (التشخيص)، والتي كانت تبلغ 4500 قيمة مفقودة.

الخطوة الثالثة حذف القيم المفقودة فقط (الأسطر) من أعمدة المتغيرات المتبقية، وذلك لتبسيط قاعدة البيانات وتقليل من الضجيج، ويوضح الشكل (4-2) القيم المفقودة في أعمدة المتغيرات المتبقية في الخطوة الثالثة قبل وبعد حذفها.

DX	0	DX	0
AGE	7	AGE	0
PTGENDER	0	PTGENDER	0
PTEDUCAT	0	PTEDUCAT	0
APOE4	264	APOE4	0
CDRSB	608	CDRSB	0
ADAS11	600	ADAS11	0
ADAS13	693	ADAS13	0
ADASQ4	570	ADASQ4	0
MMSE	565	MMSE	0
RAVLT_immediate	680	RAVLT_immediate	0
RAVLT_learning	679	RAVLT_learning	0
RAVLT_forgetting	710	RAVLT_forgetting	0
LDELTOTAL	2571	LDELTOTAL	0
TRABSCOR	1008	TRABSCOR	0
FAQ	627	FAQ	0
Ventricles	3121	Ventricles	0
Hippocampus	3745	Hippocampus	0
WholeBrain	2901	WholeBrain	0
ICV	2601	ICV	0
mPACCdigit	562	mPACCdigit	0
mPACCtrailsB	557	mPACCtrailsB	0

After Delete
DX Missing
Values

الشكل (4-2) القيم المفقودة في أعمدة المتغيرات المتبقية في الخطوة الثالثة قبل وبعد حذفها.

وفي نهاية هذه الخطوة كان عدد العينات في عمود الخرج (التشخيص):

- 900 عينة في صنف مرض الزهايمر (AD).
- 2394 عينة في صنف ضعف الإدراك المعتدل (MCI).
- 2149 عينة في صنف الإدراك الطبيعي (NL).

4-4-2 القيم الشاذة (Outliers Values):

تعتبر عن القيم ذات الاختلافات الكبيرة عن باقي البيانات في قاعدة البيانات، ووجودها له تأثير سلبي على دقة تحليل البيانات والذي من الممكن أن يسبب تشوه في التنبؤات والنتائج، وتظهر القيم الشاذة نتيجة أخطاء عند جمع البيانات أو عند تفرغها، وقد تكون حقيقية لكن استثنائية بالنسبة لباقي قيم البيانات، ومعالجتها أمر أساسي لبناء نماذج أكثر دقة وأفضل أداء.

تم إزالة القيم المتطرفة في هذا البحث بتحديد الحدود الدنيا والعليا باستخدام تقنية الربيعيات (الربيع الأول والثالث والفرق بينهما)، وحذف جميع القيم الشاذة الواقعة خارج هذه الحدود في كل عمود من أعمدة المتغيرات المتبقية بعد مرحلة معالجة القيم المفقودة حسب الصفوف الثلاثة كلاً على حدة، وتوضح المعادلتين (1-4) و(2-4) العلاقتين المستخدمتين لتحديد الحدود الدنيا والعليا: [38]

$$Lower_{bound} = Q1 - 1.5 \times IQR \quad (1 - 4)$$

$$Upper_{bound} = Q3 + 1.5 \times IQR \quad (2 - 4)$$

Q1: الربيع الأول.

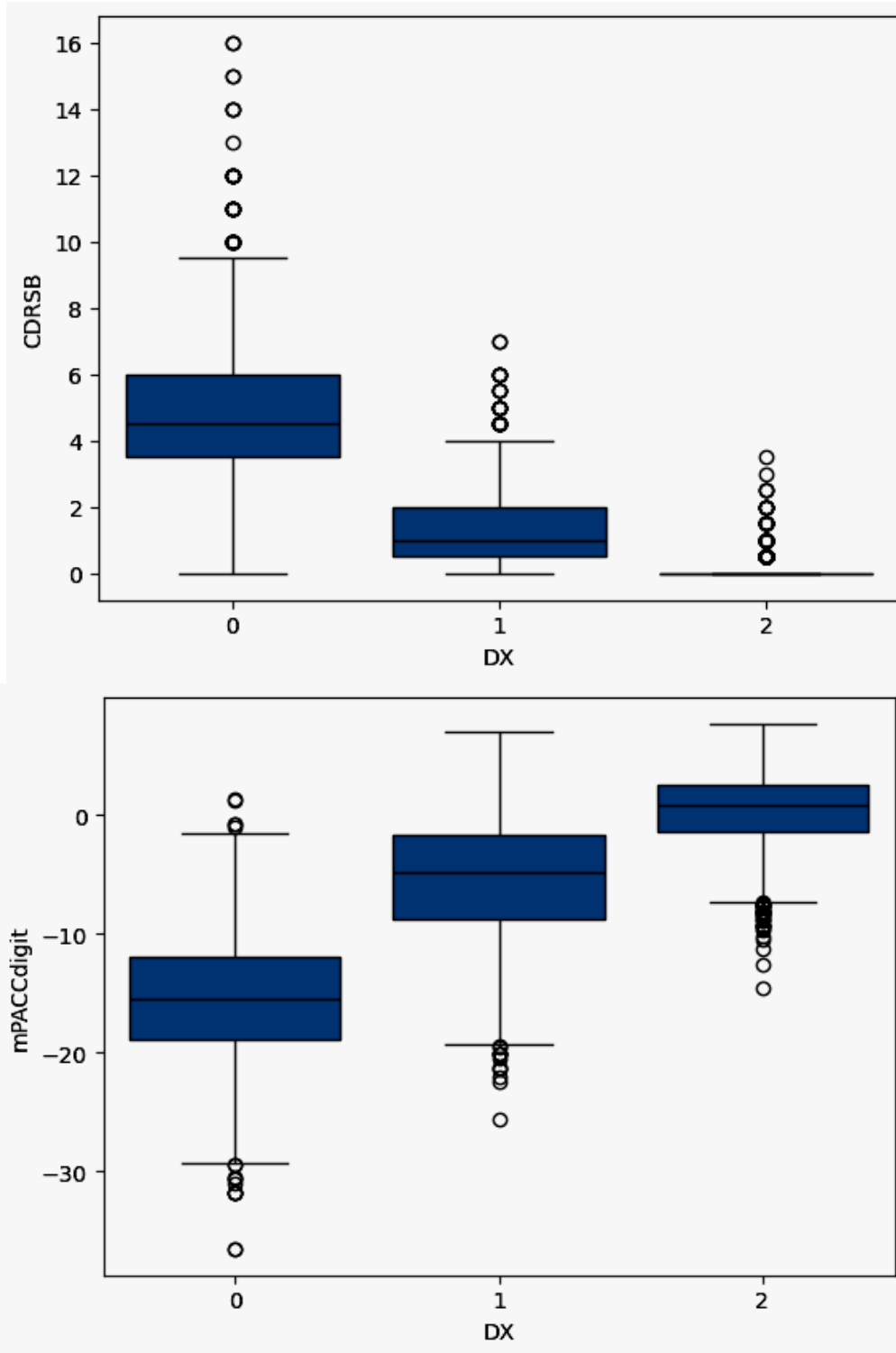
Q3: الربيع الثالث.

IQR: المدى بين الربيع الثالث والأول (Q3 - Q1).

يوضح الشكل (3-4) نقاط القيم الشاذة في مخططات الصندوق (Box Plots) لبعض المتغيرات.

وفي نهاية هذه الخطوة كان عدد العينات في عمود الخرج (التشخيص):

- 700 عينة في صنف مرض الزهايمر (AD).
- 1860 عينة في صنف ضعف الإدراك المعتدل (MCI).
- 1322 عينة في صنف الإدراك الطبيعي (NL).

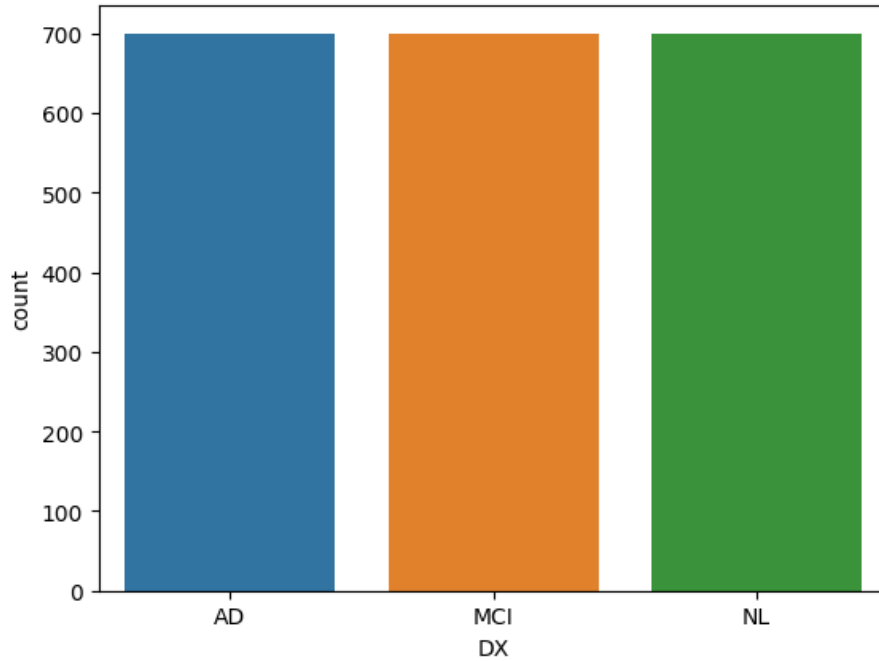


الشكل (4-3) مخططات الصندوق لبعض المتغيرات.

3-4-4 موازنة العينات:

تدريب النماذج بعدد عينات غير متوازن بين الأصناف يعطي نتائج قوية للصنف الذي فيه عدد عينات أكثر مقارنة بالنتائج الضعيفة في الأصناف ذوي عدد العينات الأقل، ولذلك تعد موازنة العينات

(Data Balancing) بين الأصناف من الخطوات المهمة لمنع تحيز وتفضيل النماذج أثناء التدريب للصنف ذو العينات الأكثر على باقي الأصناف، ويوضح الشكل (4-4) عدد العينات في كل صنف من أصناف الخرج الثلاثة بعد الموازنة، حيث تم تقليل عدد العينات واختيار عشوائي 700 عينة في كل من الصنف NL والصنف MCI وذلك بعد القيام بعملية خلط عشوائي، وذلك لكي يتوازن عدد عيناتهما مع عدد عينات الصنف AD لأنه يحتوي العدد الأقل من العينات (700 عينة).



الشكل (4-4) عدد العينات في كل صنف بعد الموازنة.

وفي نهاية هذه الخطوة كان عدد العينات في عمود الخرج (التشخيص):

- 700 عينة في صنف مرض الزهايمر (AD).
- 700 عينة في صنف ضعف الإدراك المعتدل (MCI).
- 700 عينة في صنف الإدراك الطبيعي (NL).

4-4-4 الترميز:

يطلق على عملية تحويل القيم الاسمية أو الفئوية إلى رقمية بالترميز (Encoding)، لأنه للتعامل مع البيانات ذات القيم الاسمية لابد من تحويلها إلى قيم رقمية أي للغة التي يفهمها الحاسوب وخوارزميات تعلم الآلة، حيث يسند رقم لكل اسم أو فئة، وهذا ما يساهم في جعل استخدامها أكثر فعالية وتسهيل عملية التدريب وتحسين أداء النماذج.

تم ترميز القيم الاسمية للخروج DX باستخدام طريقة Label Encoding حيث يتم الترميز وفق التسلسل الأبجدي كما يوضح الشكل (4-5).

AD: Alzheimer's Disease	→ (0)
MCI: Mild Cognitive Impairment	→ (1)
NL: Normal Cognitive	→ (2)

الشكل (4-5) Label Encoding للخروج DX.

4-4-5 تقسيم البيانات:

تقسيم مجموعة البيانات إلى مجموعتين من أجل تدريب النماذج واختبارها حيث قسمت إلى 80% من البيانات لتشكيل مجموعة التدريب و20% منها لتشكيل مجموعة الاختبار. في نهاية مرحلة المعالجة وصلنا إلى 21 سمة حيث يبين الجدول (4-1) السمات Features والخروج Label التي اعتمدت في تدريب النماذج واختبارها، وكما يبين توضيح لكل منها.

الجدول (4-1) السمات Features والخروج Label.

	Variables	Description
Label	DX	Diagnosis
Features	AGE	Age
	PTGENDER	Gender
	PTEDUCAT	Education Levels
	APOE4	Apolipoprotein E4 gene
	CDRSB	Clinical Dementia Rating – Sum of Boxes
	ADAS11	Alzheimer's Disease Assessment Scale (11 items)
	ADAS13	Alzheimer's Disease Assessment Scale (13 items)
	ADASQ4	Task 4 of ADAS11
	MMSE	Mini Mental State Examination
	RAVLT_immediate	Rey Auditory Verbal Learning Test – Immediate

RAVLT_learning	Rey Auditory Verbal Learning Test – Learning
RAVLT_forgetting	Rey Auditory Verbal Learning Test – Forgetting
LDELTOTAL	Delayed Total Recall
TRABSCOR	Trail Making Test
FAQ	Functional Assessment Questionnaire
mPACCdigit	Modified Preclinical Alzheimer Cognitive Composite with Digit
mPACCtrailsB	Modified Preclinical Alzheimer Cognitive Composite with Trails B
Ventricles	Ventricles Volume
Hippocampus	Hippocampus Volume
WholeBrain	Whole Brain Volume
ICV	Intracranial volume

حيث الخرج DX هو التشخيص النهائي للأشخاص، وبالنسبة للسّمات فإن السّمات الثلاثة الأولى هي متغيرات ديموغرافية، السّمة APOE4 هي متغير وراثي، السّمات من CDRSB إلى mPACCtrailsB هي اختبارات عصبية نفسية، والأربع سمات الأخيرة هي متغيرات تشريحية.

4-5 بناء النموذج:

تم بناء سبعة نماذج باستخدام تقنيات تعلم الآلة ML، وهي: Support Vector Machine، Naive Bayes، K-Nearest Neighbors، Decision Tree، Random Forest، AdaBoost، Gradient Boost.

وسنوضح فيما يلي المحددات الفائقة التي تم ضبطها لكل نموذج بهدف الوصول إلى نماذج نهائية ذات أفضل كفاءة وأداء.

4-5-1 آلة متجه الدعم:

يوضح الجدول (4-2) المحددات الفائقة والقيم التي أُسندت إليها للاختيار منها من أجل ضبط نموذج SVM للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (2-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج SVM.

Model	Hyperparameter	Range
SVM	kernel	"rbf", "sigmoid", "linear", "poly"
	C	0.001, 0.01, 0.1, 1, 5, 10, 100

حيث بداية تم ضبط النواة (Kernel) وهي تابع رياضي ينقل البيانات إلى فضاء ذو أبعاد أعلى أو أكثر تعقيداً لتسهيل عملية الفصل بين بيانات الفئات التي سيتم تصنيفها والتنبؤ بها، وتم اختيار النواة من ضمن خيارات النوى التالية: النواة الشعاعية (radial basis function)، النواة السينية (Sigmoid)، النواة الخطية (Linear)، النواة كثيرة الحدود (Polynomial).

وكما تم ضبط C هو معامل التحكم بالخطأ يعمل على الموازنة بين عرض الهامش Margin وخطأ التدريب من خلال العقوبة التي تفرضها على النقاط المصنفة بشكل خاطئ؛ أي أنه يعبر عن مدى تسامح نموذج SVM مع أخطاء التصنيف، وتم اختيار C من ضمن النطاق [0.001 – 100].

2-5-4 بايز البسيط:

يوضح الجدول (3-4) المحددات الفائقة والقيم التي أسندت إليها للاختيار منها من أجل ضبط نموذج NB للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (3-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج NB.

Model	Hyperparameter	Range
NB	Type	"Gaussian", "Multinomial", "Complement"
	priors	None, [0.1,]*3, [0.2,]*3, [0.3,]*3, [0.4,]*3, [0.5,]*3
	var_smoothing	0, 1e-3, 1e-6, 1e-9, 1e-12, 1e-15, 1e-18

حيث في البداية تم تحديد نوع Naive Bayes الذي سيتم استخدامه من بين ثلاث أنواع وهي GaussianNB، MultinomialNB، و ComplementNB، وبالاعتماد على الدقة الأعلى تم اختيار GaussianNaive Bayes.

بعد ذلك تم ضبط المحددات الفائقة الخاصة بهذا النوع، فأولاً تم ضبط Priors وهي تمثل الاحتمالات الأولية للفئات المختلفة حيث يتم إسناد قيم لها للحصول على معرفة مسبقة حول توزيع هذه الفئات، وقد تم اختيار Priors من ضمن نطاق القيم التالية [0.5 – 0].

وكما تم ضبط تسوية التباين (Variance smoothing) لتجنب حدوث مشكلة الاحتمالية الصفرية وذلك من أجل المحافظة على استقرار النموذج أثناء التدريب، وتم اختيار قيمة var_smoothing من ضمن النطاق [1e-18 – 0].

4-5-3 الجار الأقرب:

يوضح الجدول (4-4) المحددات الفائقة والقيم التي أسندت إليها للاختيار منها من أجل ضبط نموذج KNN للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (4-4) المحددات الفائقة ونطاق قيم اختيارها في نموذج KNN.

Model	Hyperparameter	Range
KNN	n_neighbors	[1, 2, 3, 4, 5, 6, 7, 8, 9, 10]
	Weights	"uniform ", "distance"
	P	1, 2

حيث تم ضبط (n-neighbors) وهو يمثل عدد الجيران الذي سيتم أخذهم بعين الاعتبار عند اتخاذ القرار لتصنيف عينة جديدة وذلك بحساب بعد العينة الجديدة المراد تصنيفها عنهم، وحيث تم اختيار عدد الجيران من ضمن نطاق القيم التالية [10 – 1].

وكما تم ضبط (Weights) وهو يحدد الأوزان المعطاة للجيران الأقرب، وقد تم اختيار الأوزان من بين خيارين: الأوزان متساوية (uniform)، و متغيرة بحسب البعد (distance).

بالإضافة إلى ضبط (P) وهو معامل المسافة الذي يحدد الطريقة المتبعة لقياس بعد العينة الجديدة المراد اختبارها عن العينات الجيران، حيث تم اختيار معامل المسافة من بين خيارين: 1 والذي يمثل مسافة مانهاتن، و2 والذي يمثل المسافة الإقليدية.

4-5-4 شجرة القرار:

يوضح الجدول (4-5) المحددات الفائقة والقيم التي أسندت إليها للاختيار منها من أجل ضبط نموذج DT للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (4-5) المحددات الفائقة ونطاق قيم اختيارها في نموذج DT.

Model	Hyperparameter	Range
DT	Criterion	"gini", "entropy"
	max_depth	1, 3, 5, 7, 9, 11, 13, 15, 17, 19, 21
	max_features	1, 3, 5, 7, 9, 11, 13, 15, 17, 19, 21
	min_samples_leaf	1, 3, 5, 7, 9, 11, 13, 15, 17, 19, 21

حيث تم ضبط المعيار (criterion) وهو الأسلوب الذي يحدد أفضل طريقة لتقسيم البيانات عند كل عقدة واختيار أفضل السمات للانقسام مما يساهم في تحسين دقة النموذج، وتم اختيار المعيار من ضمن المعيارين: Entropy، Gini impurity.

وكما تم ضبط العمق الأعظمي (max_depth) وهي تعبر عن أكبر عمق يمكن أن يصل إليه نموذج DT (أي الحد الأقصى لعدد المستويات (levels) أو التقسيمات (splits))، وتم اختيار العمق الأعظمي من ضمن النطاق [1 - 21].

وكذلك تم ضبط (min_samples_leaf) وهي تحدد أقل عدد ممكن من العينات اللازمة لتشكيل عقدة نهائية leaf node، وتم اختيارها من ضمن النطاق [1 - 21].

بالإضافة إلى ضبط (max_features) وهي تحدد عدد السمات المناسب للعثور على أفضل تقسيم في كل عقدة، وتم اختيار عدد السمات الأعظمي من ضمن النطاق [1 - 21].

4-5-5 الغابة العشوائية:

يوضح الجدول (4-6) المحددات الفائقة والقيم التي أسندت إليها للاختيار منها من أجل ضبط نموذج RF للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (4-6) المحددات الفائقة ونطاق قيم اختيارها في نموذج RF.

Model	Hyperparameter	Range
RF	Criterion	"gini", "entropy"
	n_estimators	10, 20, 30, 40, 50, 60, 70, 80, 90, 100, 110, 120, 130, 140, 150, 160, 170, 180, 190, 200
	max_depth	2, 4, 6, 8, 10, 12, 14, 16, 18, 20
	max_features	2, 4, 6, 8, 10, 12, 14, 16, 18, 20, 21

حيث بداية تم ضبط المعيار (criterion) وهو الأسلوب الذي يحدد أفضل الانقسامات عند بناء كل شجرة في النموذج واختيار أفضل السمات للاستخدام في كل انقسام، وذلك من خلال تقييم جودة الانقسامات مما يساهم في تحسين دقة وأداء النموذج، وتم اختيار المعيار من ضمن المعيارين: Gini Entropy، impurity.

بالإضافة إلى ضبط عدد الأشجار (n_estimators) وهي تحدد عدد أشجار القرار التي سوف يتم بناء نموذج RF عليها، وحيث ازدياد عددها يزيد من استقرار النموذج، وتم اختيار عدد أشجار القرار من ضمن النطاق [10 – 200].

وكما تم ضبط العمق الأعظمي (max_depth) والتي تحدد الحد الأقصى للعمق الذي يمكن أن تصل إليه الأشجار في النموذج (أي الحد الأقصى لعدد المستويات أو التقسيمات)، وتم اختيار العمق الأعظمي من ضمن النطاق [2 – 20].

بالإضافة إلى ضبط (max_features) وهي تحدد الحد الأقصى لعدد السمات المناسب لبناء كل شجرة وعقدها (أي التحكم بالعشوائية عند بناء الأشجار واختيار السمات)، وتم اختيار عدد السمات الأعظمي من ضمن النطاق [2 – 21].

4-5-6 التعزيز الكيفي:

يوضح الجدول (4-7) المحددات الفائقة والقيم التي أسندت إليها للاختيار منها من أجل ضبط نموذج AB للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (4-7) المحددات الفائقة ونطاق قيم اختيارها في نموذج AB.

Model	Hyper Parameter	Range
AB	n_estimators	10, 20, 30, 40, 50, 60, 70, 80, 90, 100, 110, 120, 130, 140, 150, 160, 170, 180, 190, 200
	Learning_rate	0.001, 0.01, 0.1, 1

حيث بداية تم ضبط (n_estimators) وهي تحدد عدد المتعلمين الضعفاء (weak learners) التي سيتم تدريبها على أوزان أخطاء بعضها البعض بالتتابع وبعد ذلك دمجها للوصول إلى نموذج AB النهائي (strong learner)، وتم اختيار عدد المتعلمين الضعفاء من ضمن النطاق [10 – 200]. وكما تم ضبط معدل تعلم النموذج (Learning_rate) الذي يحدد مدى تأثير كل نموذج ضعيف على النموذج النهائي، فحيث أن قيمته مهمة في تحقيق التوازن بين دقة النموذج وسرعة التعلم، وتم اختيار معدل التعلم من ضمن النطاق [0.001 – 1].

4-5-7 تعزيز التدرج:

يوضح الجدول (4-8) المحددات الفائقة والقيم التي أسندت إليها للاختيار منها من أجل ضبط نموذج GB للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (4-8) المحددات الفائقة ونطاق قيم اختيارها في نموذج GB

Model	Hyperparameter	Range
GB	n_estimators	10, 20, 30, 40, 50, 60, 70, 80, 90, 100, 110, 120, 130, 140, 150, 160, 170, 180, 190, 200
	max_depth	2, 4, 6, 8, 10, 12, 14, 16, 18, 20
	max_features	2, 4, 6, 8, 10, 12, 14, 16, 18, 20, 21
	learning_rate	0.001, 0.01, 0.1, 1
	Loss	"log_loss"

حيث بداية تم ضبط (n_estimators) وهي تحدد عدد المتعلمين الضعفاء (weak learners) التي سيتم تدريبها لتقليل تابع الخسارة لبعضها البعض بالتتابع وبعد ذلك دمجها للوصول إلى نموذج GB النهائي (strong learner)، وتم اختيار عدد المتعلمين الضعفاء من ضمن النطاق [10 – 200]. وكما تم ضبط (max_depth) والتي تحدد الحد الأقصى للعمق الذي يمكن أن تصل إليه النماذج الضعيفة، وتم اختيار العمق الأعظمي من ضمن النطاق [2 – 20]. وكما تم ضبط (max_features) وهي تحدد الحد الأقصى لعدد السمات المناسب لبناء كل نموذج ضعيف، وتم اختيار عدد السمات الأعظمي من ضمن النطاق [2 – 21]. بالإضافة إلى ضبط معدل تعلم النموذج (learning_rate) الذي يحدد مدى تأثير كل نموذج ضعيف على النموذج النهائي، فحيث أن قيمته مهمة في تحقيق التوازن بين دقة النموذج وسرعة التعلم، وتم اختيار معدل التعلم من ضمن النطاق [0.001 – 1].

4-6 تدريب النماذج وتقييمها:

بعد الانتهاء من الخطوات السابقة تم بناء النماذج (Models Building) و تم تدريب واختبار النماذج (Models Training & Testing) التي تم بنائها وذلك باستخدام السمات الناتجة من المراحل السابقة والتي عددها 21 سمة وفقاً لتقسيم البيانات المحدد 80/20 والخرج ذي الأصناف الثلاثة AD

و MCI و NL، ثم تم تقييم أداء وكفاءة هذه النماذج (Evaluation Models) باستخدام مقاييس تقييم الأداء الدقة والحساسية والإحكام ودرجة F1، وبناءً على هذه المقاييس تم تحديد أفضل النماذج (The Best Model).

4-7 تقليل السمات:

تم إضافة هذه الخطوة للعمل على تقليل عدد السمات التي تم الوصول إليها في المراحل السابقة، وذلك للوصول إلى المجموعة الفرعية للسمات التي تحقق التوازن بين عدد السمات الأقل ودقة النماذج في نفس الوقت، فالغرض الأساسي من تقليل السمات هو تقليل الأبعاد وزيادة دقة النموذج. تم استخدام ثلاث طرق تتبع أساليب اختيار سمات مختلفة للوصول إلى أقل عدد منها، ومن ثم اختبار مجموعة السمات المختارة من قبل كل طريقة على حدة في تدريب واختبار النماذج التي تم بناؤها، وذلك لتحديد أفضل مجموعة من السمات التي تحقق الأداء الأفضل للنماذج.

4-7-1 اختيار K الأفضل:

يوضح الجدول (4-9) المحددات الفائقة والقيم التي أسندت إليها من أجل ضبط خوارزمية SKB للوصول إلى أفضل كفاءة في اختيار مجموعة السمات.

الجدول (4-9) المحددات الفائقة وقيمها اختيارها في SKB.

Model	Hyperparameter	Value
SKB	Function	f_classif
	K	9

تم العمل على اختيار التابع الرياضي (الطريقة الإحصائية) التي سيتم اعتمدها في طريقة SKB لتقييم السمات لاختيار الأفضل منها، وبالنسبة لما يناسب هذا البحث تم اختيار التابع (Function) بطريقة f_classif وهي طريقة F-test الإحصائية، وتم اختيار قيمة K المسؤولة عن تحديد عدد السمات المراد اختيارها وكانت تساوي 9 حيث كان العدد الأفضل بالتجريب مقارنة مع قيم K تساوي 8 و 10.

4-7-2 الشبكة المرنة:

يوضح الجدول (4-10) المحددات الفائقة والقيم التي أسندت إليها للاختيار منها من أجل ضبط نموذج EN للوصول إلى أفضل كفاءة عن طريق الحصول على أفضل قيم محدّدات فائقة لهذا النموذج باستخدام البحث الشبكي والتحقق المتقاطع.

الجدول (4-10) المحددات الفائقة ونطاق قيم اختيارها في نموذج EN.

Model	Hyperparameter	Range
EN	Alph	0.001, 0.01, 0.1, 1, 10
	l1_ratio	0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9, 1

تم ضبط l1_ratio والذي يعبر عن نسبة Lasso المستخدمة في هذا البحث وباقي النسبة هو Ridge حيث كانت مساوية 0.1 أي 10% Lasso و 90% Ridge وتم اختيارها من ضمن النطاق [1 - 0.1]، وكما تم ضبط قيمة ألفا (Alph) والمسؤولة عن تحديد معامل العقوبة (Penalty) وكانت القيمة تساوي 0.1 وتم اختيارها من ضمن النطاق [10 - 0.001].

4-7-3 استبعاد السمات العودية:

يوضح الجدول (4-11) المحددات الفائقة والقيم التي أسندت إليها من أجل ضبط خوارزمية RFE للوصول إلى أفضل كفاءة في اختيار مجموعة السمات.

الجدول (4-11) المحددات الفائقة وقيمها اختيارها في RFE.

Model	Hyperparameter	Values
RFE	Model	Random Forest
	n_features_to_select	9

في هذه الطريقة يجب تحديد أحد نماذج تعلم الآلة لإعطاء أوزان للسمات ليتم الاختيار فيما بينها وعادة ما يكون الغلبة العشوائية أو الانحدار اللوجستي وكون نموذج RF كان أدائه من الأعلى في هذا البحث

فتم اختياره النموذج (Model) لهذه الطريقة وتم اختيار (n_features_to_select) المسؤولة عن تحديد عدد السمات المراد اختيارها وكانت تساوي 9 حيث كان العدد الأفضل بالتجريب مقارنة مع قيم K تساوي 8 و10.

8-4 خاتمة:

تناول هذا الفصل القسم العملي من البحث بكل مراحله وخطواته بدءاً من قاعدة البيانات ومن ثم معالجة وتهيئة هذه البيانات وتقسيمها لتدريب واختبار النماذج بعد ضبط محدداتها الفائقة إلى تقييم أداء وكفاءة هذه النماذج للوصول إلى أفضل نموذج منها، وكما تم توضيح خطوة تقليل السمات والطرق المتبعة للقيام بذلك للوصول إلى أفضل أقل عدد من السمات مع أفضل نموذج.

الفصل الخامس:

النتائج والمناقشة

1-5 مقدمة:

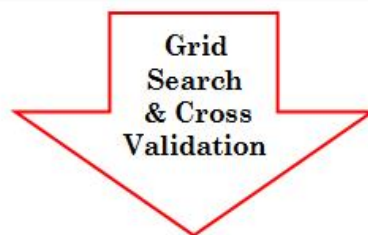
جذب انتشار الأمراض العصبية التنكسية مثل مرض الزهايمر اهتمام الباحثين في جامعات العالم، والذين يعملون على تطوير منهجيات ذات أداء عالي للتشخيص والعلاج والوقاية واكتشاف الأدوية المستهدفة، وذلك انطلاقاً من أهمية الكشف المبكر عن الزهايمر الذي من الممكن أن يكون مفتاح للعلاج الفعال، أو على الأقل التخطيط المستقبلي لعلاج مراحل هذا المرض، ورفع سوية الحياة والرعاية عند المرضى، سيقدم هذا الفصل نتائج البحث ومناقشتها.

2-5 المحددات الفائقة:

1-2-5 آلة متجه الدعم:

يوضح الشكل (1-5) القيم الأفضل للمحددات الفائقة لبناء نموذج آلة متجه الدعم، والتي تعطي أفضل كفاءة وأداء لهذا النموذج مقارنة مع باقي نطاق القيم المسندة للمحددات الفائقة التي اختير منها القيم الأفضل باستخدام البحث الشبكي والتحقق المتقاطع.

```
'kernel': ['rbf', 'linear', 'sigmoid', 'poly'],
'C': [0.001, 0.01, 0.1, 1, 10, 100],
```



```
The Best Hyperparameters: {'C': 1, 'kernel': 'linear'}
```

الشكل (1-5) القيم الأفضل للمحددات الفائقة لبناء نموذج SVM الأفضل.

وحيث تم اختيار النواة الخطية Linear Kernel لتكون هي نوع النواة Kernel لنموذج SVM، وبالنسبة إلى قيمة C فتم اختيار قيمة 1 لها وهي قيمة مقبولة كسماحية لضبط الأخطاء.

وهذه كانت أفضل قيم للمحددات الفائقة المختارة التي استخدمت في بناء هذا النموذج للوصول إلى نموذج SVM النهائي ذي الكفاءة والأداء الأفضل مقارنة مع نماذج SVM الأخرى المبنية باستخدام

نطاق القيم المسندة للمحددات الفائقة، وذلك كون النموذج الذي يبنى باستخدام القيم المختارة للمحددات الفائقة بواسطة البحث الشبكي والتحقق المتقاطع ستفضي إلى النموذج الأعلى دقة وكفاءة.

5-2-2 بايز البسيط الغاوسي:

يوضح الشكل (5-2) القيم الأفضل للمحددات الفائقة لضبط نموذج بايز البسيط الغاوسي، والتي تعطي أفضل كفاءة وأداء لهذا النموذج مقارنة مع باقي نطاق القيم المسندة للمحددات الفائقة التي اختير منها القيم الأفضل باستخدام البحث الشبكي والتحقق المتقاطع.

```
'priors': [None, [0.1,]* 3, [0.2,]* 3, [0.3,]* 3, [0.4,]* 3, [0.5,]* 3,],
'var_smoothing': [0, 1e-3, 1e-6, 1e-9, 1e-12, 1e-15, 1e-18]
```



The Best Hyperparameters: {'priors': None, 'var_smoothing': 1e-15}

الشكل (5-2) القيم الأفضل للمحددات الفائقة لبناء نموذج GNB الأفضل.

في البداية تم تحديد نوع بايز البسيط الذي سيتم استخدامه من بين ثلاث أنواع وهي GaussianNB، MultinomialNB، و ComplementNB، بالاعتماد على الدقة الأعلى تم اختيار بايز البسيط الغاوسي (GaussianNB).

بعد ذلك تم اختيار قيمة priors على أنها None أي لا يوجد قيم بدائية للاحتمالات الابتدائية، وأما بالنسبة إلى قيمة var_smoothing فكانت تساوي 1e-15.

وهذه كانت أفضل قيم للمحددات الفائقة المختارة التي استخدمت في بناء هذا النموذج للوصول إلى نموذج GNB النهائي ذي الكفاءة والأداء الأفضل مقارنة مع نماذج GNB الأخرى المبنية باستخدام نطاق القيم المسندة للمحددات الفائقة ، وذلك كون النموذج الذي يبنى باستخدام القيم المختارة للمحددات الفائقة بواسطة البحث الشبكي والتحقق المتقاطع ستفضي إلى النموذج الأعلى دقة وكفاءة.

5-2-3 الجار الأقرب:

يوضح الشكل (5-3) القيم الأفضل للمحددات الفائقة لضبط نموذج الجار الأقرب، والتي تعطي أفضل كفاءة وأداء لهذا النموذج مقارنة مع باقي نطاق القيم المسندة للمحددات الفائقة التي اختير منها القيم الأفضل باستخدام البحث الشبكي والتحقق المتقاطع.

```
'n_neighbors': [1, 2, 3, 4, 5, 6, 7, 8, 9, 10],
'weights': ['uniform', 'distance'],
'p': [1, 2],
```



The Best Hyperparameters: {'n_neighbors': 8, 'p': 1, 'weights': 'distance'}

الشكل (5-3) القيم الأفضل للمحددات الفائقة لبناء نموذج KNN الأفضل.

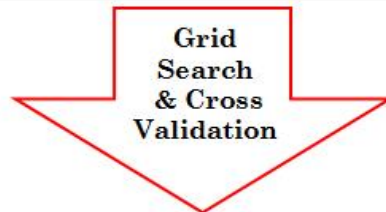
حيث كان عدد n-neighbors المختار هو 8، وبالنسبة إلى نوع Weights المختار كان من نوع distance أي الأوزان متغيرة بحسب البعد، وكان نوع P المختار هو 1 والذي يمثل حساب البعد بطريقة Manhattan.

وهذه كانت أفضل قيم للمحددات الفائقة المختارة التي استخدمت في بناء النموذج للوصول إلى نموذج KNN النهائي ذي الكفاءة والأداء الأفضل مقارنة مع نماذج KNN الأخرى المبنية باستخدام نطاق القيم المسندة للمحددات الفائقة، وذلك كون النموذج التي تبني باستخدام القيم المختارة بواسطة البحث الشبكي والتحقق المتقاطع ستفضي إلى النموذج الأعلى دقة وكفاءة.

5-2-4 شجرة القرار:

يوضح الشكل (5-4) القيم الأفضل للمحددات الفائقة لضبط نموذج شجرة القرار والتي تعطي أفضل كفاءة وأداء لهذا النموذج مقارنة مع باقي نطاق القيم المسندة للمحددات الفائقة التي اختير منها القيم الأفضل باستخدام البحث الشبكي والتحقق المتقاطع.

```
'criterion': ['gini', 'entropy'],
'max_depth': [3, 5, 7, 9, 11, 13, 15, 17, 19, 21],
'min_samples_leaf': [3, 5, 7, 9, 11, 13, 15, 17, 19, 21],
'max_features': [3, 5, 7, 9, 11, 13, 15, 17, 19, 21]
```



The Best Hyperparameters: {'criterion': 'gini', 'max_depth': 9, 'max_features': 9, 'min_samples_leaf': 3}

الشكل (5-4) القيم الأفضل للمحددات الفائقة لبناء نموذج DT الأفضل.

وحيث تم اختيار معيار criterion ليكون من نوع Gini لنموذج DT، وبالنسبة إلى قيمة max_depth فتم اختيار قيمة 9 لها، وكانت قيمة min_samples_leaf المختارة هي 3، في حين تم اختيار 9 لتكون قيمة max_features.

وهذه كانت أفضل قيم للمحددات الفائقة المختارة التي استخدمت في ضبط وتعديل النموذج للوصول إلى نموذج DT النهائي ذي الكفاءة والأداء الأفضل مقارنة مع نماذج DT الأخرى المبنية باستخدام نطاق القيم المسندة للمحددات الفائقة وذلك بحسب البحث الشبكي والتحقق المتقاطع، وذلك كون النموذج التي تبنى باستخدام القيم المختارة بواسطة البحث الشبكي والتحقق المتقاطع ستقضي إلى النموذج الأعلى دقة وكفاءة.

5-2-5 الغابة العشوائية:

يوضح الشكل (5-5) القيم الأفضل للمحددات الفائقة لضبط نموذج الغابة العشوائية والتي تعطي أفضل كفاءة وأداء لهذا النموذج مقارنة مع باقي نطاق القيم المسندة للمحددات الفائقة التي اختير منها القيم الأفضل باستخدام البحث الشبكي والتحقق المتقاطع.

```
'n_estimators': [10, 20, 30, 40, 50, 60, 70, 80, 90, 100, 110,
                  120, 130, 140, 150, 160, 170, 180, 190, 200],

'max_depth': [2, 4, 6, 8, 10, 12, 14, 16, 18, 20],

'max_features': [2, 4, 6, 8, 10, 12, 14, 16, 18, 20, 21],

'criterion': ['gini', 'entropy']
```



The Best Hyperparameters: {'criterion': 'gini', 'max_depth': 20, 'max_features': 8, 'n_estimators': 160}

الشكل (5-5) القيم الأفضل للمحددات الفائقة لبناء نموذج RF الأفضل.

وحيث تم اختيار معيار criterion ليكون من نوع Gini لنموذج RF، وبالنسبة إلى قيمة max_depth فتم اختيار قيمة 20 لها، في حين تم اختيار 8 لتكون قيمة max_features، وكانت قيمة n_estimators المختارة (عدد الأشجار) هي 160.

وهذه كانت أفضل قيم للمحددات الفائقة المختارة التي استخدمت في ضبط وتعديل النموذج للوصول إلى نموذج RF النهائي ذي الكفاءة والأداء الأفضل مقارنة مع نماذج RF الأخرى المبنية باستخدام نطاق القيم المسندة للمحددات الفائقة، وذلك كون النموذج الذي يبنى باستخدام القيم المختارة للمحددات الفائقة بواسطة البحث الشبكي والتحقق المتقاطع ستفضي إلى النموذج الأعلى دقة وكفاءة.

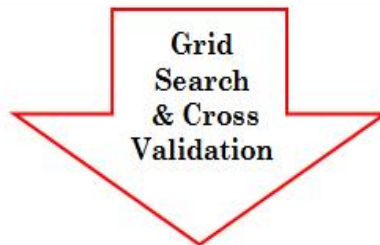
5-2-6 التعزيز الكيفي:

يوضح الشكل (5-6) القيم الأفضل للمحددات الفائقة لضبط نموذج التعزيز الكيفي والتي تعطي

أفضل كفاءة وأداء لهذا النموذج مقارنة مع باقي نطاق القيم المسندة للمحددات الفائقة التي اختير منها القيم الأفضل باستخدام البحث الشبكي والتحقق المتقاطع.

```
'n_estimators': [10, 20, 30, 40, 50, 60, 70, 80, 90, 100, 110,
                  120, 130, 140, 150, 160, 170, 180, 190, 200],

'learning_rate': [0.001, 0.01, 0.1, 1]
```



The Best Hyperparameters: {'learning_rate': 0.01, 'n_estimators': 40}

الشكل (5-6) القيم الأفضل للمحددات الفائقة لبناء نموذج AB الأفضل.

حيث كانت قيمة `n_estimators` (عدد المتعلمين الضعفاء) المختارة لنموذج AB هي 40، وبالنسبة إلى قيمة `learning_rate` فتم اختيار قيمة 0.01 إليها.

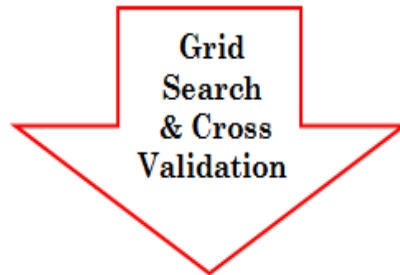
وهذه كانت أفضل قيم للمحددات الفائقة المختارة التي استخدمت في ضبط وتعديل النموذج للوصول إلى نموذج AB النهائي ذي الكفاءة والأداء الأفضل مقارنة مع نماذج AB الأخرى المبنية باستخدام نطاق القيم المسندة للمحددات الفائقة، ، وذلك كون النموذج الذي يبنى باستخدام القيم المختارة للمحددات الفائقة بواسطة البحث الشبكي والتحقق المتقاطع ستفضي إلى النموذج الأعلى دقة وكفاءة.

5-2-7 تعزيز التدرج:

يوضح الشكل (5-7) القيم الأفضل للمحددات الفائقة لضبط نموذج تعزيز التدرج والتي تعطي

أفضل كفاءة وأداء لهذا النموذج مقارنة مع باقي نطاق القيم المسندة للمحددات الفائقة التي اختير منها القيم الأفضل باستخدام البحث الشبكي والتحقق المتقاطع.

```
'n_estimators': [10, 20, 30, 40, 50, 60, 70, 80, 90, 100, 110,
                  120, 130, 140, 150, 160, 170, 180, 190, 200],
'learning_rate': [0.001, 0.01, 0.1, 1],
'max_depth': [2, 4, 6, 8, 10, 12, 14, 16, 18, 20],
'max_features': [2, 4, 6, 8, 10, 12, 14, 16, 18, 20, 21],
'loss': ['log_loss']
```



The Best Hyperparameters: {'learning_rate': 0.1, 'loss': 'log_loss', 'max_depth': 6, 'max_features': 8, 'n_estimators': 190}

الشكل (5-7) القيم الأفضل للمحددات الفائقة لبناء نموذج GB الأفضل.

حيث كانت قيمة `n_estimators` (عدد المتعلمين الضعفاء) المختارة لنموذج GB هي 190، وكانت قيمة `learning_rate` المختارة هي 0.1، وبالنسبة إلى قيمة `max_depth` فتم اختيار قيمة 6 إليها، في حين تم اختيار 8 لتكون قيمة `max_features`.

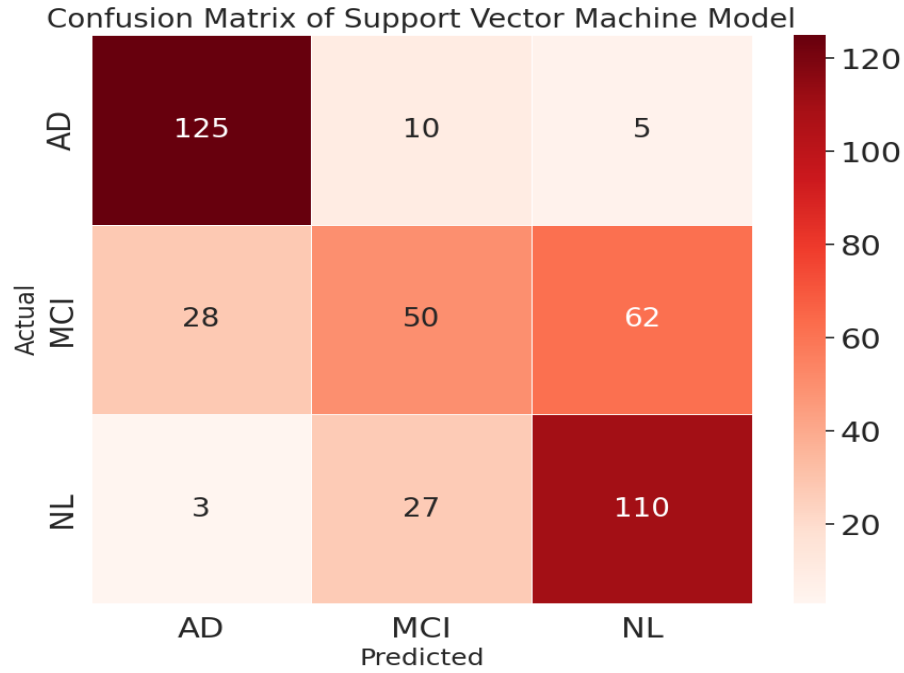
وهذه كانت أفضل قيم للمحددات الفائقة المختارة التي استخدمت في ضبط وتعديل النموذج للوصول إلى نموذج GB النهائي ذي الكفاءة والأداء الأفضل مقارنة مع نماذج GB الأخرى المبنية باستخدام نطاق القيم المسندة للمحددات الفائقة، وذلك كون النموذج الذي يبنى باستخدام القيم المختارة للمحددات الفائقة بواسطة البحث الشبكي والتحقق المتقاطع ستفضي إلى النموذج الأعلى دقة وكفاءة.

5-3 أداء النماذج مع كل السمات:

تم تقييم أداء النماذج النهائية التي تم الوصول إليها لكل صنف من أصناف الخرج من خلال المعايير التالية مصفوفة الارتباك وتقرير التصنيف وبالإضافة لتقييم أداء النماذج النهائية كاملاً بحساب المقاييس التالية الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل، وذلك بهدف الوصول إلى النموذج الأعلى أداء وكفاءة.

5-3-1 آلة متجه الدعم:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-8) عدد الحالات التي نجح وأخطأ نموذج SVM في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-8) مصفوفة الارتباك للنموذج SVM.

حيث يمكن ملاحظة أن نموذج SVM نجح في تصنيف 285 حالة من جميع عينات الاختبار، فمن عينات اختبار AD صنف 125 حالة AD صحيحة، و10 حالات على أنها MCI، و5 حالات على أنها NL، وبالنسبة إلى عينات اختبار NL صنف 110 حالة NL صحيحة، و3 حالات على أنها AD، و27 حالة على أنها MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 50 حالة MCI صحيحة، و28 حالة على أنها AD، و62 حالة على أنها NL، ويوضح الجدول (5-1) جميع حالات Positive و Negative للنموذج SVM.

الجدول (5-1) حالات Positive و Negative للنموذج SVM.

	True Positive	False Positive	True Negative	False Negative
AD	125	31	249	15
MCI	50	37	243	90
NL	110	67	213	30

وأما بالنسبة إلى تقرير التصنيف للنموذج SVM والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-9).

Classification Report of Support Vector Machine			
	precision	recall	f1-score
0	0.8013	0.8929	0.8446
1	0.5747	0.3571	0.4405
2	0.6215	0.7857	0.6940

الشكل (5-9) تقرير التصنيف للنموذج SVM.

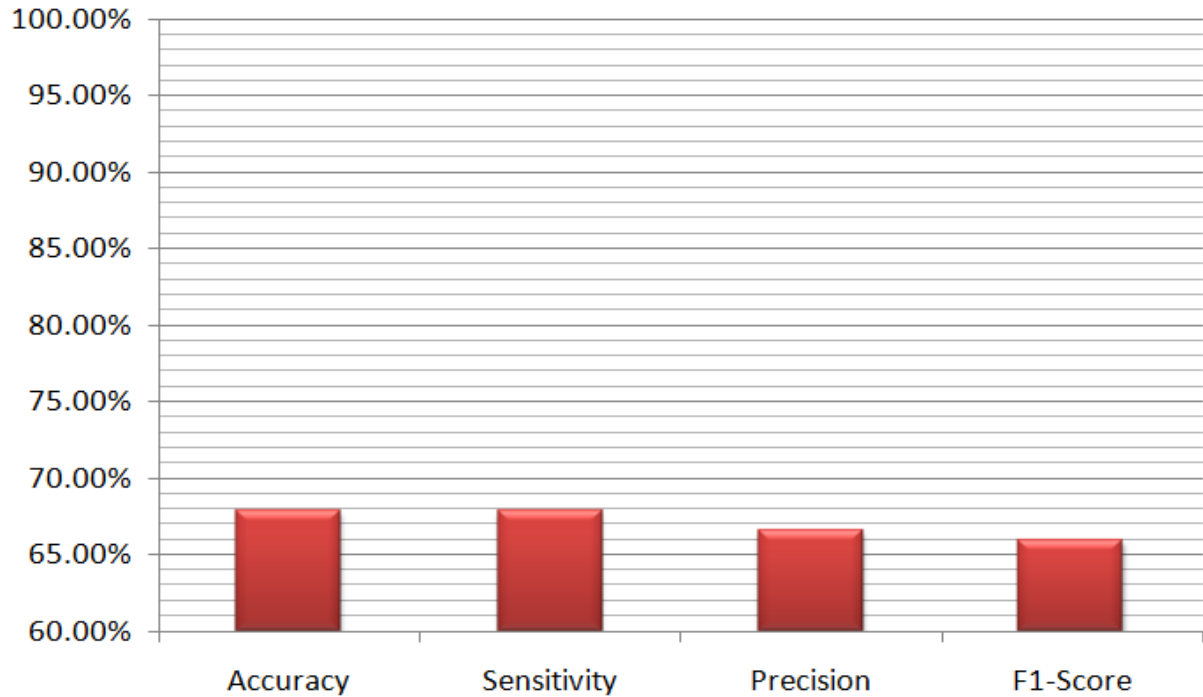
حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف AD (80.13%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD بشكل True Positive (125 حالة) وأقل عدد من حالات على أنها AD بشكل False Positive (31 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف MCI (57.47%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (50 حالة) وعدد من حالات على أنها MCI بشكل False Positive (37 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف NL كانت متوسطة بين القيمتين السابقتين (62.15%) كون النموذج صنف عدد حالات على أنها NL بشكل False Positive أكثر من AD و MCI (67 حالة) ولكنه صنف عدد من حالات على أنها NL بشكل True Positive أكثر من MCI (110 حالة).

أما بالنسبة إلى قيم الحساسية فكانت الأعلى أيضاً في الصنف AD (89.29%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD بشكل True Positive (125 حالة) وبأقل عدد من حالات

AD بشكل False Negative (15 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (35.71%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (50 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (90 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف NL كانت متوسطة بين القيمتين السابقتين (78.57%) كون النموذج صنف عدد من الحالات على أنها NL بشكل True Positive أقل من AD وأكثر من MCI (110 حالة) وعدد من الحالات NL بشكل False Negative أكثر من AD وأقل من MCI (30 حالة).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (84.46%) وأقل منها في الصنف NL (69.40%) وكانت قيمة درجة F1 الأقل في الصنف MCI (44.05%). في النهاية لابد من تقييم أداء وكفاءة نموذج SVM كاملاً حيث تم حساب كل من الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل ويوضح الشكل (5-10) قيم هذه المقاييس.

Performance Measures of Support Vector Machine Model

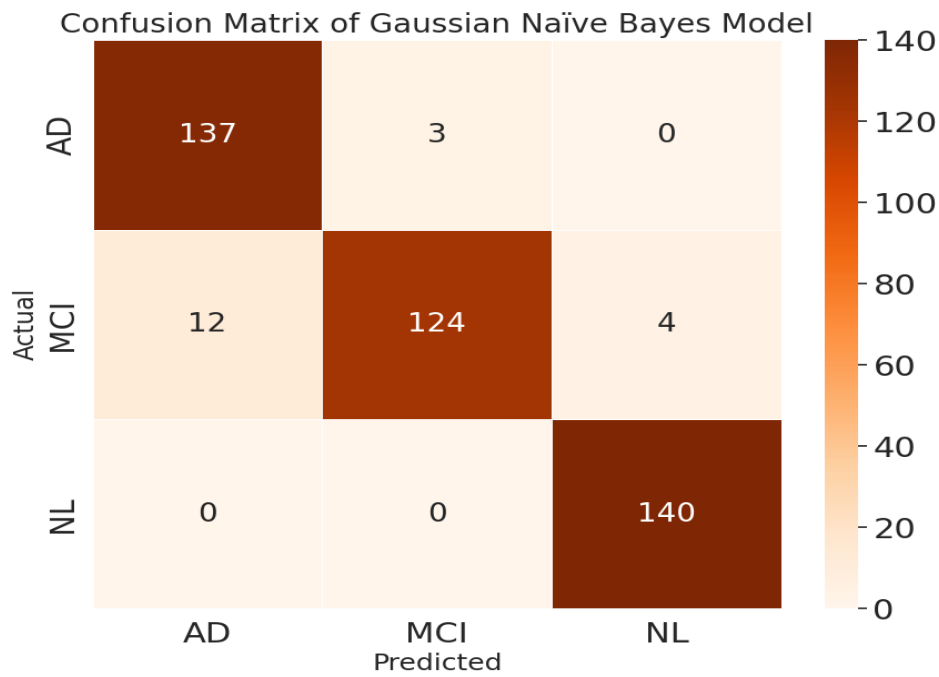


الشكل (5-10) مقاييس الأداء للنموذج SVM.

حيث تم حساب قيم هذه المقاييس للنموذج كنسبة مئوية حيث كانت قيمة الدقة تساوي 67.86%، وأما قيمة الحساسية بلغت 67.86%، وبالنسبة إلى قيمة الإحكام فكانت 66.58%، وبلغت قيمة درجة F1 65.97%.

5-3-2 بايز البسيط الغاوسي:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-11) عدد الحالات التي نجح وأخطأ نموذج GNB في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-11) مصفوفة الارتباك للنموذج GNB.

حيث يمكن ملاحظة أن نموذج GNB نجح في تصنيف 401 حالة من جميع عينات الاختبار ولم يخطئ النموذج في التصنيف بين الصنفين AD و NL، فمن عينات اختبار AD صنف 137 حالة AD صحيحة، و 3 حالات على أنها MCI، ولم يصنف أي منها على أنها NL، في حين بالنسبة إلى عينات اختبار NL صنف 140 حالة NL صحيحة، ولم يصنف أي منها على أنها AD أو MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 124 حالة MCI صحيحة، و 12 حالة على أنها AD، و 4 حالات على أنها NL، ويوضح الجدول (5-2) جميع حالات Positive و Negative للنموذج GNB.

الجدول (2-5) حالات Positive و Negative للنموذج GNB

	True Positive	False Positive	True Negative	False Negative
AD	137	12	268	3
MCI	124	3	277	16
NL	140	4	276	0

وأما بالنسبة إلى تقرير التصنيف للنموذج GNB والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-12).

Classification Report of Gaussian Naïve Bayes			
	precision	recall	f1-score
0	0.9195	0.9786	0.9481
1	0.9764	0.8857	0.9288
2	0.9722	1.0000	0.9859

الشكل (5-12) تقرير التصنيف للنموذج GNB.

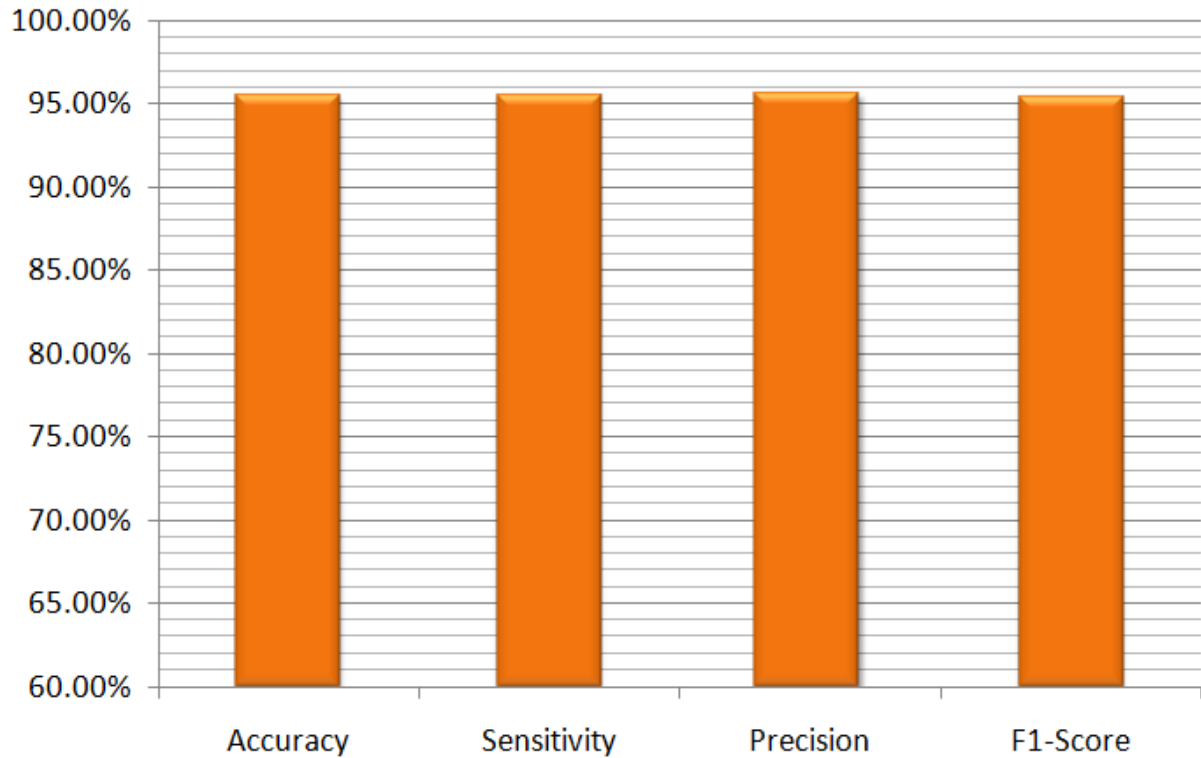
حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف MCI (97.64%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل False Positive (3 حالات) وكان عدد الحالات المصنفة على أنها MCI بشكل True Positive (124 حالات)، في حين كانت قيمة الإحكام الأقل في الصنف AD (91.95%) وذلك كون النموذج صنف أكثر عدد من حالات على أنها AD بشكل False Positive (12 حالة) وكان عدد الحالات المصنفة على أنها AD بشكل True Positive (137 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف NL كانت متوسطة بين القيمتين السابقتين (97.22%) كون النموذج صنف عدد حالات على أنها NL بشكل False Positive أقل من AD وأكثر من MCI (4 حالات) وكان عدد الحالات المصنفة على أنها NL بشكل True Positive (140 حالة).

أما بالنسبة إلى قيم الحساسية فكانت الأعلى في الصنف NL (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها NL بشكل True Positive (140 حالة) ولم يصنف أي من الحالات

NL بشكل False Negative (0 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (88.57%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (124 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (16 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف AD كانت متوسطة بين القيمتين السابقتين (97.86%) كون النموذج صنف عدد من الحالات على أنها AD بشكل True Positive أقل من NL وأكثر من MCI (137 حالة) وعدد من الحالات AD بشكل False Negative أكثر من NL وأقل من MCI (3 حالات).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف NL (98.59%) وأقل منها في الصنف AD (94.81%) وكانت قيمة درجة F1 الأقل في الصنف MCI (92.88%). في النهاية لابد من تقييم أداء وكفاءة نموذج GNB كاملاً حيث تم حساب كل من الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل ويوضح الشكل (5-13) قيم هذه المقاييس.

Performance Measures of Gaussian Naïve Bayes Model

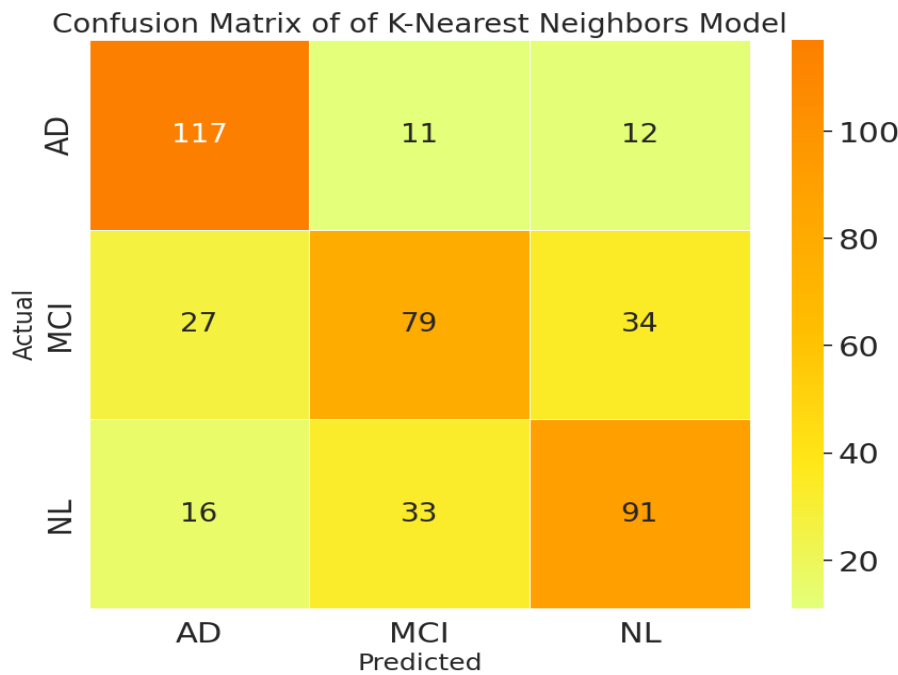


الشكل (5-13) مقاييس الأداء للنموذج GNB.

حيث تم حساب قيم هذه المقاييس للنموذج كنسبة مئوية حيث كانت قيمة الدقة تساوي 95.48%، وأما قيمة الحساسية بلغت 95.48%، وبالنسبة إلى قيمة الإحكام فكانت 95.60%، وبلغت قيمة درجة F1 95.43%.

3-3-5 الجار الأقرب:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-14) عدد الحالات التي نجح وأخطأ نموذج KNN في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-14) مصفوفة الارتباك للنموذج KNN.

حيث يمكن ملاحظة أن نموذج KNN نجح في تصنيف 287 حالة من جميع عينات الاختبار، فمن عينات اختبار AD صنف 117 حالة AD صحيحة، و11 حالة على أنها MCI، و12 حالة على أنها NL، وبالنسبة إلى عينات اختبار NL صنف 91 حالة NL صحيحة، و16 حالة على أنها AD، و33 حالة على أنها MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 79 حالة MCI صحيحة، و27 حالة على أنها AD، و34 حالة على أنها NL، ويوضح الجدول (5-3) جميع حالات Positive وNegative للنموذج KNN.

الجدول (3-5) حالات Positive و Negative للنموذج KNN.

	True Positive	False Positive	True Negative	False Negative
AD	117	43	237	23
MCI	79	44	236	61
NL	91	46	234	49

وأما بالنسبة إلى تقرير التصنيف للنموذج KNN والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-15).

Classification Report of K-Nearest Neighbors			
	precision	recall	f1-score
0	0.7312	0.8357	0.7800
1	0.6423	0.5643	0.6008
2	0.6642	0.6500	0.6570

الشكل (5-15) تقرير التصنيف للنموذج KNN.

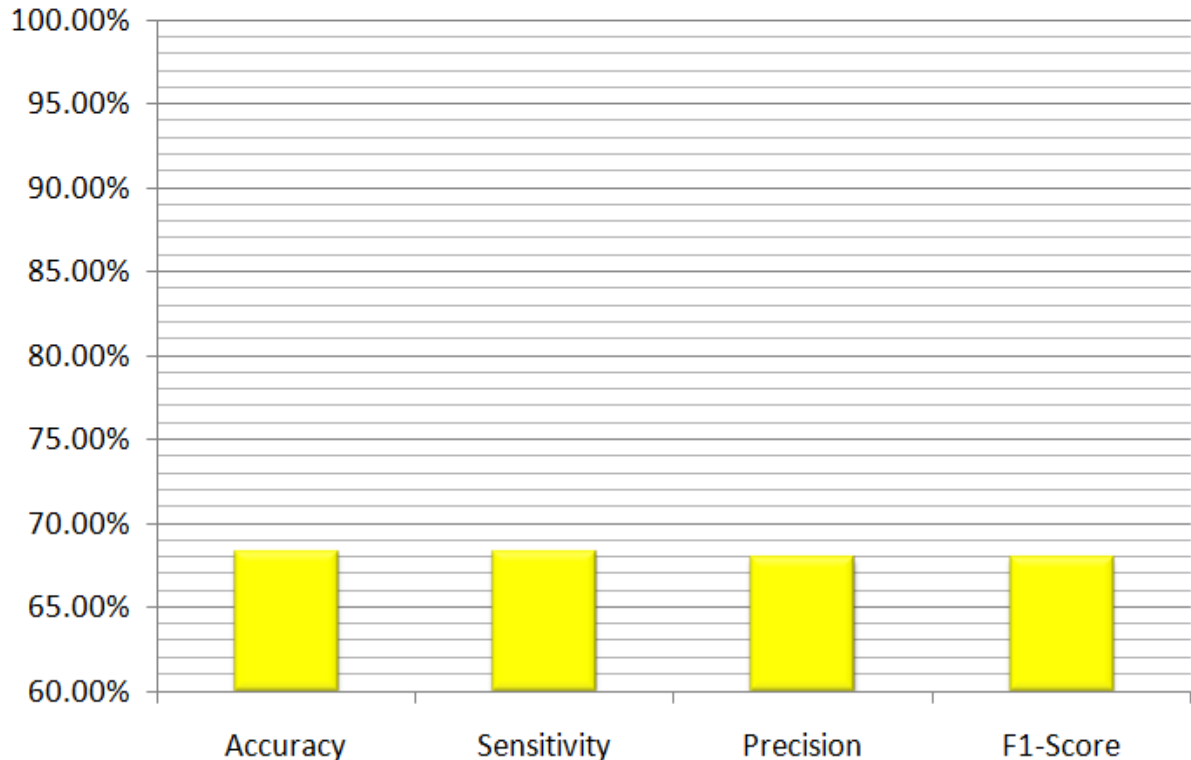
حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف AD (73.12%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD بشكل True Positive (117 حالة) وأقل عدد من حالات على أنها AD بشكل False Positive (43 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف MCI (64.23%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (79 حالة) وعدد من حالات على أنها MCI بشكل False Positive (44 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف NL كانت متوسطة بين القيمتين السابقتين (66.42%) كون النموذج صنف عدد حالات على أنها NL بشكل False Positive أكثر من AD و MCI (46 حالة) ولكنه صنف عدد من حالات على أنها NL بشكل True Positive أكثر من MCI (91 حالة).

أما بالنسبة إلى قيم الحساسية فكانت الأعلى أيضاً في الصنف AD (83.57%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD بشكل True Positive (117 حالة) وبأقل عدد من حالات AD بشكل False Negative (23 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI

(56.43%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (79 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (61 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف NL كانت متوسطة بين القيمتين السابقتين (65.00%) كون النموذج صنف عدد من الحالات على أنها NL بشكل True Positive أقل من AD وأكثر من MCI (91 حالة) وعدد من الحالات NL بشكل False Negative أكثر من AD وأقل من MCI (49 حالة).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (78.00%) وأقل منها في الصنف NL (65.70%) وكانت قيمة درجة F1 الأقل في الصنف MCI (60.08%). في النهاية لابد من تقييم أداء وكفاءة نموذج KNN كاملاً حيث تم حساب كل من الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل ويوضح الشكل (5-16) قيم هذه المقاييس.

Performance Measures of K-Nearest Neighbors Model

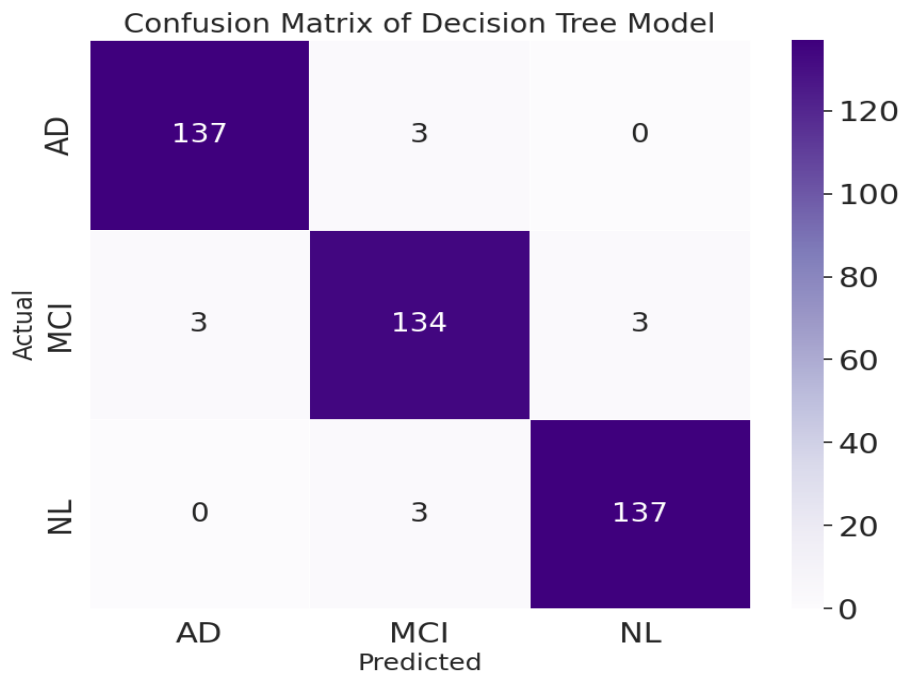


الشكل (5-16) مقاييس الأداء للنموذج KNN.

حيث تم حساب قيم هذه المقاييس للنموذج كنسبة مئوية حيث كانت قيمة الدقة تساوي 68.33%، وأما قيمة الحساسية بلغت 68.33%، وبالنسبة إلى قيمة الإحكام فكانت 67.93%، وبلغت قيمة درجة F1 67.93%.

5-3-4 شجرة القرار:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-17) عدد الحالات التي نجح وأخطأ نموذج DT في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-17) مصفوفة الارتباك للنموذج DT.

حيث يمكن ملاحظة أن نموذج DT نجح في تصنيف 408 حالة من جميع عينات الاختبار ولم يخطئ النموذج في التصنيف بين الصنفين AD و NL، فمن عينات اختبار AD صنف 137 حالة AD صحيحة، و 3 حالات على أنها MCI، ولم يصنف أي منها على أنها NL، وبالنسبة إلى عينات اختبار NL صنف 137 حالة NL صحيحة، و 3 حالات على أنها MCI، ولم يصنف أي منها على أنها AD، وأما بالنسبة إلى عينات اختبار MCI صنف 134 حالة MCI صحيحة، و 3 حالات على أنها AD، و 3 حالات على أنها NL، ويوضح الجدول (5-4) جميع حالات Positive و Negative للنموذج DT.

الجدول (4-5) حالات Positive و Negative للنموذج DT.

	True Positive	False Positive	True Negative	False Negative
AD	137	3	277	3
MCI	134	6	274	6
NL	137	3	277	3

وأما بالنسبة إلى تقرير التصنيف للنموذج DT والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-18).

Classification Report of Decision Tree Model:			
	precision	recall	f1-score
0	0.9786	0.9786	0.9786
1	0.9571	0.9571	0.9571
2	0.9786	0.9786	0.9786

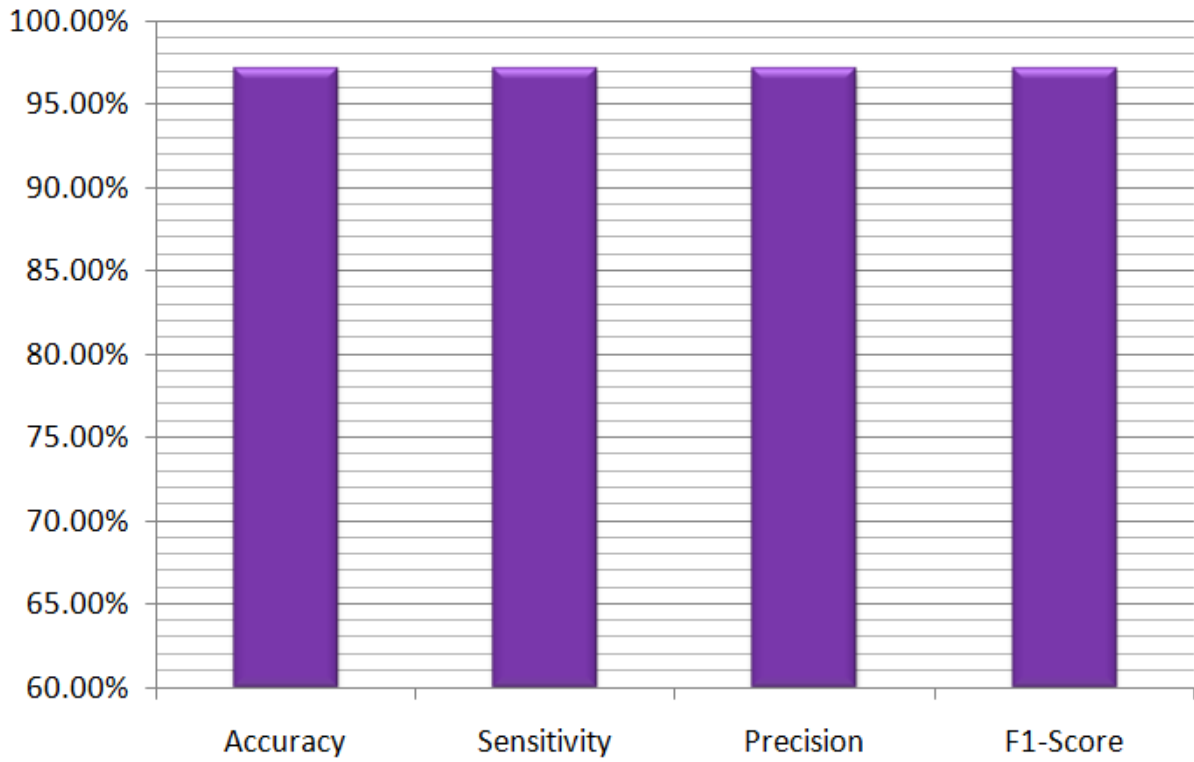
الشكل (5-18) تقرير التصنيف للنموذج DT.

حيث بالنسبة إلى قيم الإحكام كانت الأعلى ومتساوية في كل من الصنفين AD و NL (97.86%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD و NL بشكل True Positive (137 حالة) وأقل عدد من حالات على أنها AD و NL بشكل False Positive (3 حالات)، في حين كانت قيمة الإحكام الأقل في الصنف MCI (95.71%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (134 حالة) وبالإضافة إلى أكثر عدد من حالات على أنها MCI بشكل False Positive (6 حالات) مقارنة بكلا الصنفين AD و NL.

أما بالنسبة إلى قيم الحساسية فكانت الأعلى ومتساوية أيضاً في كل من الصنفين AD و NL (97.86%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD و NL بشكل True Positive (137 حالة) وأقل عدد من حالات AD و NL بشكل False Negative (3 حالات)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (95.71%) وذلك كون النموذج صنف أقل عدد من

حالات على أنها MCI بشكل True Positive (134 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (6 حالة) مقارنة بكلا الصنفين AD و NL. وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (97.86%) ومساوية لها في الصنف NL (97.86%) وكانت قيمة درجة F1 الأقل في الصنف MCI (95.71%). في النهاية لابد من تقييم أداء وكفاءة نموذج DT كاملاً حيث تم حساب كل من الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل ويوضح الشكل (5-19) قيم هذه المقاييس.

Performance Measures of Decision Tree Model

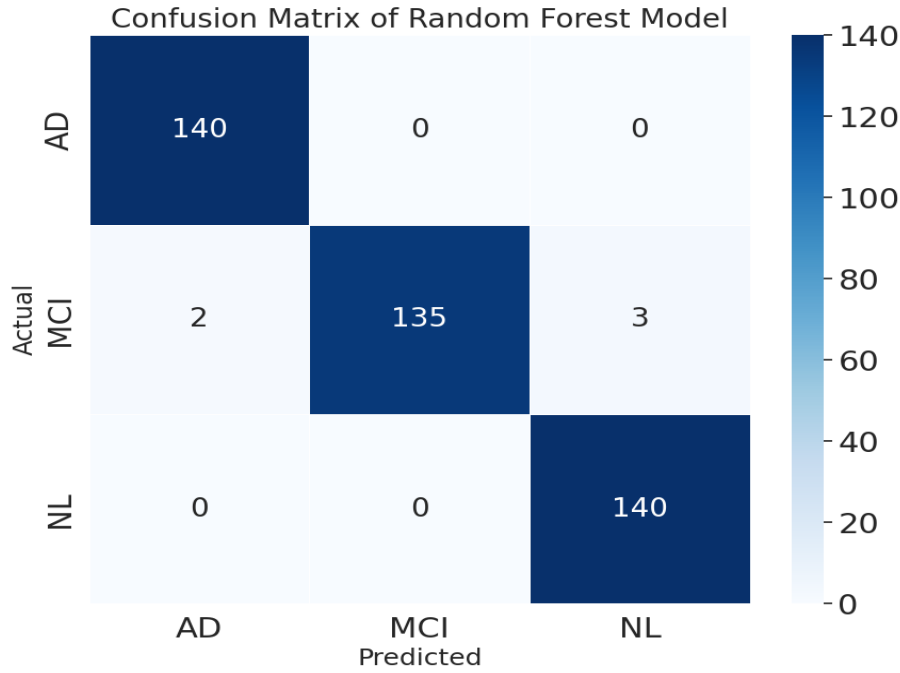


الشكل (5-19) مقاييس الأداء للنموذج DT.

حيث تم حساب قيم هذه المقاييس للنموذج كنسبة مئوية حيث كانت قيمة الدقة تساوي 97.14%، وأما قيمة الحساسية بلغت 97.14%، وبالنسبة إلى قيمة الإحكام فكانت 97.14%، وبلغت قيمة درجة F1 97.14%.

5-3-5 الغابة العشوائية:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-20) عدد الحالات التي نجح وأخطأ نموذج RF في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-20) مصفوفة الارتباك للنموذج RF.

حيث يمكن ملاحظة أن نموذج RF نجح في تصنيف 415 حالة من جميع عينات الاختبار وحيث نجح في تصنيف جميع عينات اختبار الصنفين AD و NL بشكل صحيح ولم يخطئ النموذج في التصنيف بينهما، فمن عينات اختبار AD صنف 140 حالة AD صحيحة، ولم يصنف أي منها على أنها MCI أو NL، وبالنسبة إلى عينات اختبار NL صنف 140 حالة NL صحيحة، ولم يصنف أي منها على أنها AD أو MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 135 حالة MCI صحيحة، و2 حالة على أنها AD، و3 حالات على أنها NL، ويوضح الجدول (5-5) جميع حالات Positive و Negative للنموذج RF.

الجدول (5-5) حالات Positive و Negative للنموذج RF.

	True Positive	False Positive	True Negative	False Negative
AD	140	2	278	0
MCI	135	0	280	5
NL	140	3	277	0

وأما بالنسبة إلى تقرير التصنيف للنموذج RF والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-21).

Classification Report of Random Forest Model:			
	precision	recall	f1-score
0	0.9859	1.0000	0.9929
1	1.0000	0.9643	0.9818
2	0.9790	1.0000	0.9894

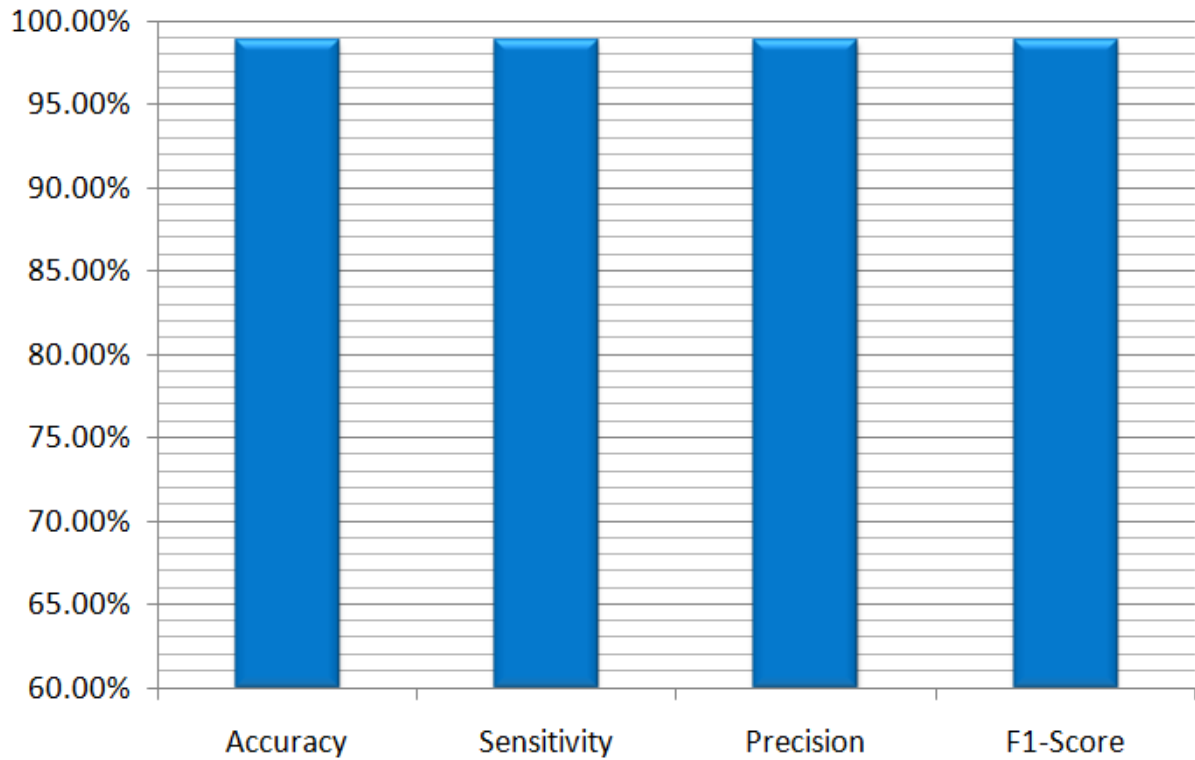
الشكل (5-21) تقرير التصنيف للنموذج RF.

حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف MCI (100.00%) وذلك كون النموذج لم يصنف أي من صنف حالات على أنها MCI بشكل False Positive (0 حالة) مع عدد من حالات على أنها MCI بشكل True Positive (135 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف NL (97.90%) وذلك كون النموذج صنف أكثر عدد من حالات على أنها NL بشكل False Positive (3 حالات) مع عدد من حالات على أنها NL بشكل True Positive (140 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف AD كانت متوسطة بين القيمتين السابقتين (98.59%) كون النموذج صنف عدد حالات على أنها AD بشكل False Positive أقل من NL وأكثر من MCI (2 حالة) مع عدد من حالات على أنها AD بشكل True Positive (140 حالة).

أما بالنسبة إلى قيم الحساسية فكانت الأعلى ومتساوية أيضاً في كل من الصنفين AD و NL (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD و NL بشكل True Positive (140 حالة) ولم يصنف أي من حالات AD و NL بشكل False Negative (0 حالة)،

في حين كانت قيمة الحساسية الأقل في الصنف MCI (96.43%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (135 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (5 حالة) مقارنة بكلا الصنفين AD و NL. وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (99.29%) وأقل منها في الصنف NL (98.94%) وكانت قيمة درجة F1 الأقل في الصنف MCI (98.18%). في النهاية لابد من تقييم أداء وكفاءة نموذج RF كاملاً حيث تم حساب كل من الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل ويوضح الشكل (5-22) قيم هذه المقاييس.

Performance Measures of Random Forest Model

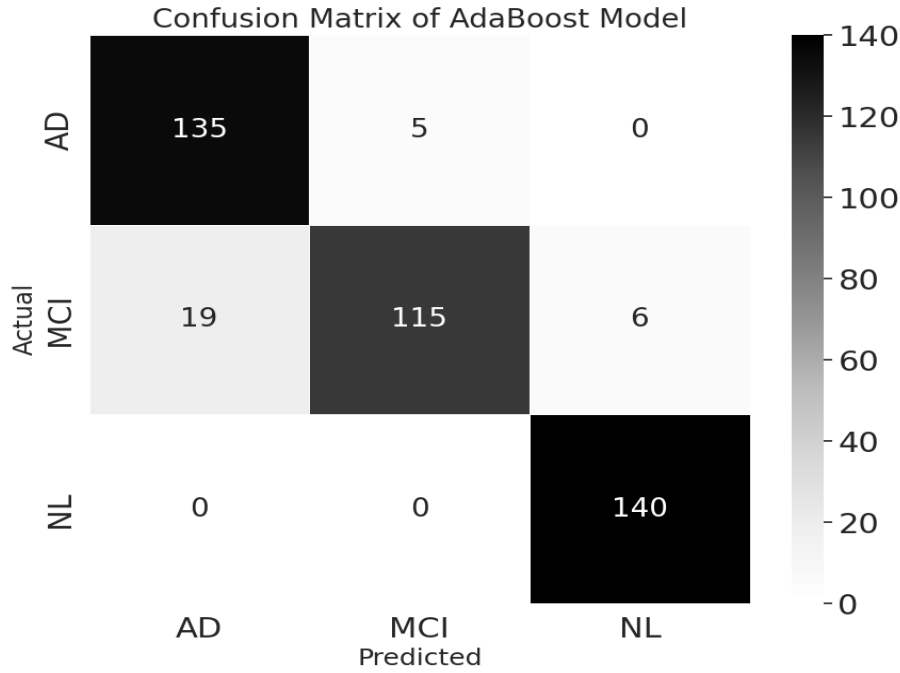


الشكل (5-22) مقاييس الأداء للنموذج RF.

حيث تم حساب قيم هذه المقاييس للنموذج كنسبة مئوية حيث كانت قيمة الدقة تساوي 98.81%، وأما قيمة الحساسية بلغت 98.81%، وبالنسبة إلى قيمة الإحكام فكانت 98.83%، وبلغت قيمة درجة F1 98.80%.

5-3-6 التعزيز الكيفي:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-23) عدد الحالات التي نجح وأخطأ نموذج AB في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-23) مصفوفة الارتباك للنموذج AB.

حيث يمكن ملاحظة أن نموذج AB نجح في تصنيف 390 حالة من جميع عينات الاختبار ولم يخطئ النموذج في التصنيف بين الصنفين AD و NL، فمن عينات اختبار AD صنف 135 حالة AD صحيحة، و 5 حالات على أنها MCI، ولم يصنف أي منها على أنها NL، في حين بالنسبة إلى عينات اختبار NL صنف 140 حالة NL صحيحة، ولم يصنف أي منها على أنها AD أو MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 115 حالة MCI صحيحة، و 19 حالة على أنها AD، و 6 حالات على أنها NL، ويوضح الجدول (5-6) جميع حالات Positive و Negative للنموذج AB.

الجدول (5-6) حالات Positive و Negative للنموذج AB.

	True Positive	False Positive	True Negative	False Negative
AD	135	19	261	5
MCI	115	5	275	25
NL	140	6	274	0

وأما بالنسبة إلى تقرير التصنيف للنموذج AB والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-24).

Classification Report of AdaBoost Model:			
	precision	recall	f1-score
0	0.8766	0.9643	0.9184
1	0.9583	0.8214	0.8846
2	0.9589	1.0000	0.9790

الشكل (5-24) تقرير التصنيف للنموذج AB.

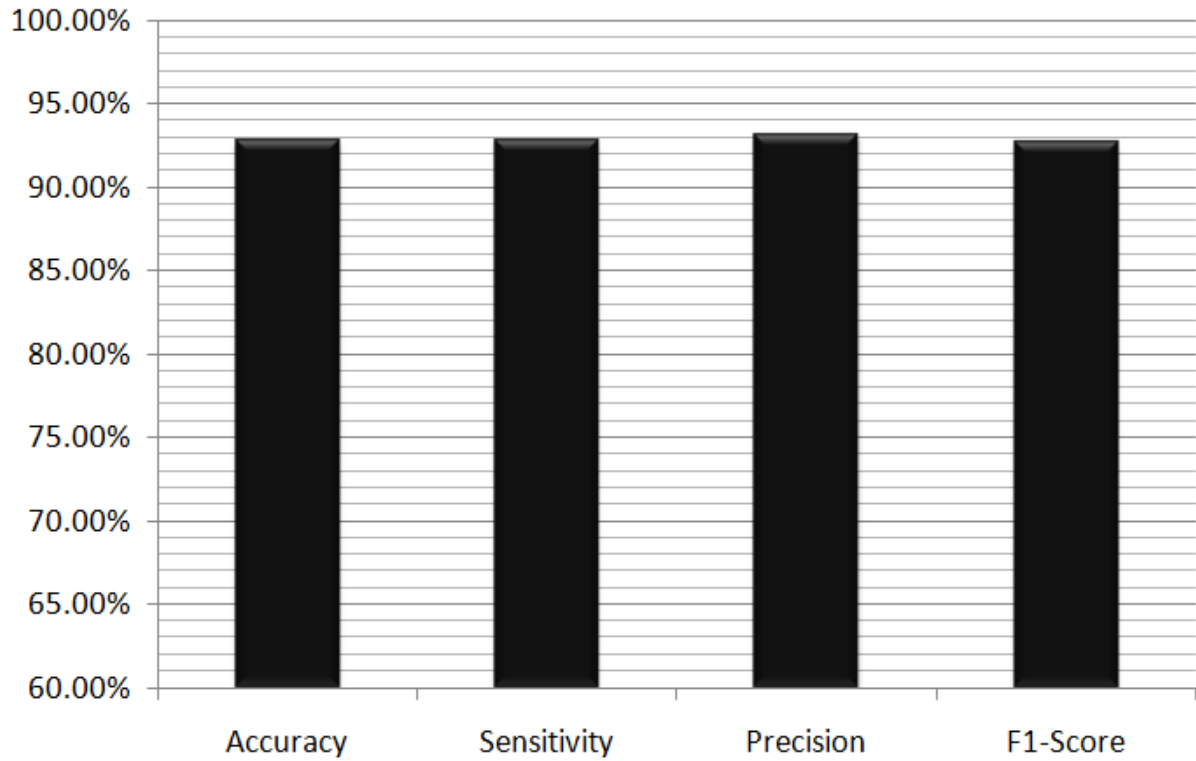
حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف NL (95.89%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها NL بشكل True Positive (140 حالة) وعدد من حالات على أنها NL بشكل False Positive (6 حالات)، في حين كانت قيمة الإحكام الأقل في الصنف AD (87.66%) وذلك كون النموذج صنف أكثر عدد من حالات على أنها AD بشكل False Positive (19 حالة) مع عدد من حالات على أنها AD بشكل True Positive (135 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف MCI كانت متوسطة بين القيمتين السابقتين (95.83%) كون النموذج صنف عدد حالات على أنها MCI بشكل False Positive أقل من AD و NL (5 حالات) ولكنه أيضاً صنف أقل عدد من حالات على أنها MCI بشكل True Positive (115 حالة).

أما بالنسبة إلى قيم الحساسية فكانت الأعلى أيضاً في الصنف NL (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها NL بشكل True Positive (140 حالة) ولم يصنف أي من

حالات NL بشكل False Negative (0 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (82.14%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (115 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (25 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف AD كانت متوسطة بين القيمتين السابقتين (96.43%) كون النموذج صنف عدد من الحالات على أنها AD بشكل True Positive أكثر من MCI وأقل من NL (135 حالة) وعدد من الحالات AD بشكل False Negative أقل من MCI وأكثر من NL (5 حالات).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف NL (97.90%) وأقل منها في الصنف AD (91.84%) وكانت قيمة درجة F1 الأقل في الصنف MCI (88.46%). في النهاية لابد من تقييم أداء وكفاءة نموذج AB كاملاً حيث تم حساب كل من الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل ويوضح الشكل (5-25) قيم هذه المقاييس.

Performance Measures of AdaBoost Model

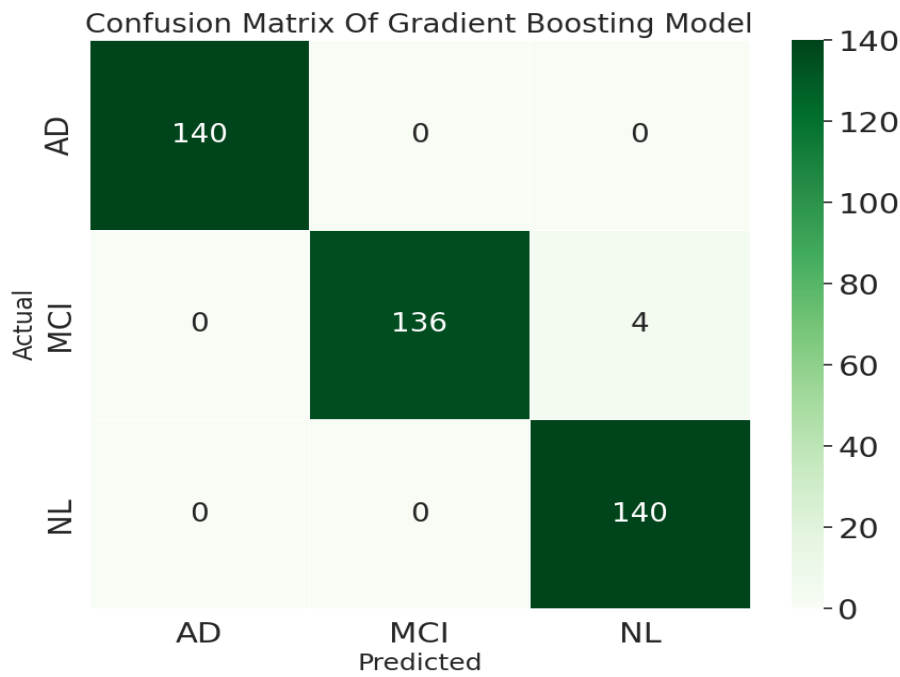


الشكل (5-25) مقاييس الأداء للنموذج AB.

حيث تم حساب قيم هذه المقاييس للنموذج كنسبة مئوية حيث كانت قيمة الدقة تساوي 92.86%، وأما قيمة الحساسية بلغت 92.86%، وبالنسبة إلى قيمة الإحكام فكانت 93.13%، وبلغت قيمة درجة F1 92.71%.

5-3-7 تعزيز التدرج:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-26) عدد الحالات التي نجح وأخطأ نموذج GB في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-26) مصفوفة الارتباك للنموذج GB.

حيث يمكن ملاحظة أن نموذج GB نجح في تصنيف 416 حالة من جميع عينات الاختبار وحيث نجح في تصنيف جميع عينات اختبار الصنفين AD و NL بشكل صحيح ولم يخطئ النموذج في التصنيف بينهما وكذلك لم يخطئ النموذج في التصنيف بين الصنفين AD أو MCI، فمن عينات اختبار AD صنف 140 حالة AD صحيحة، ولم يصنف أي منها على أنها MCI أو NL، وبالنسبة إلى عينات اختبار NL صنف 140 حالة NL صحيحة، ولم يصنف أي منها على أنها AD أو MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 136 حالة MCI صحيحة، و3 حالات على أنها NL، ولم يصنف أي منها على أنها AD، ويوضح الجدول (5-7) جميع حالات Positive و Negative للنموذج GB.

الجدول (5-7) حالات Positive و Negative للنموذج GB.

	True Positive	False Positive	True Negative	False Negative
AD	140	0	280	0
MCI	136	0	280	4
NL	140	4	276	0

وأما بالنسبة إلى تقرير التصنيف للنموذج GB والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-27).

Classification Report Of Gradient Boosting Model:			
	precision	recall	f1-score
0	1.0000	1.0000	1.0000
1	1.0000	0.9714	0.9855
2	0.9722	1.0000	0.9859

الشكل (5-27) تقرير التصنيف للنموذج GB.

حيث بالنسبة إلى قيم الإحكام كانت الأعلى ومتساوية في كل من الصنفين AD و MCI (100.00%) وذلك كون النموذج لم يصنف أي من حالات على أنها AD و MCI بشكل False Positive (0 حالة) مع عدد من حالات بشكل True Positive (140 حالة على أنها AD و 136 حالة على أنها MCI)، في حين كانت قيمة الإحكام الأقل في الصنف NL (97.22%) وذلك كون النموذج صنف أكثر عدد من حالات على أنها NL بشكل False Positive (4 حالات) مقارنة بكلا الصنفين AD و MCI مع عدد من حالات على أنها NL بشكل True Positive (140 حالة)

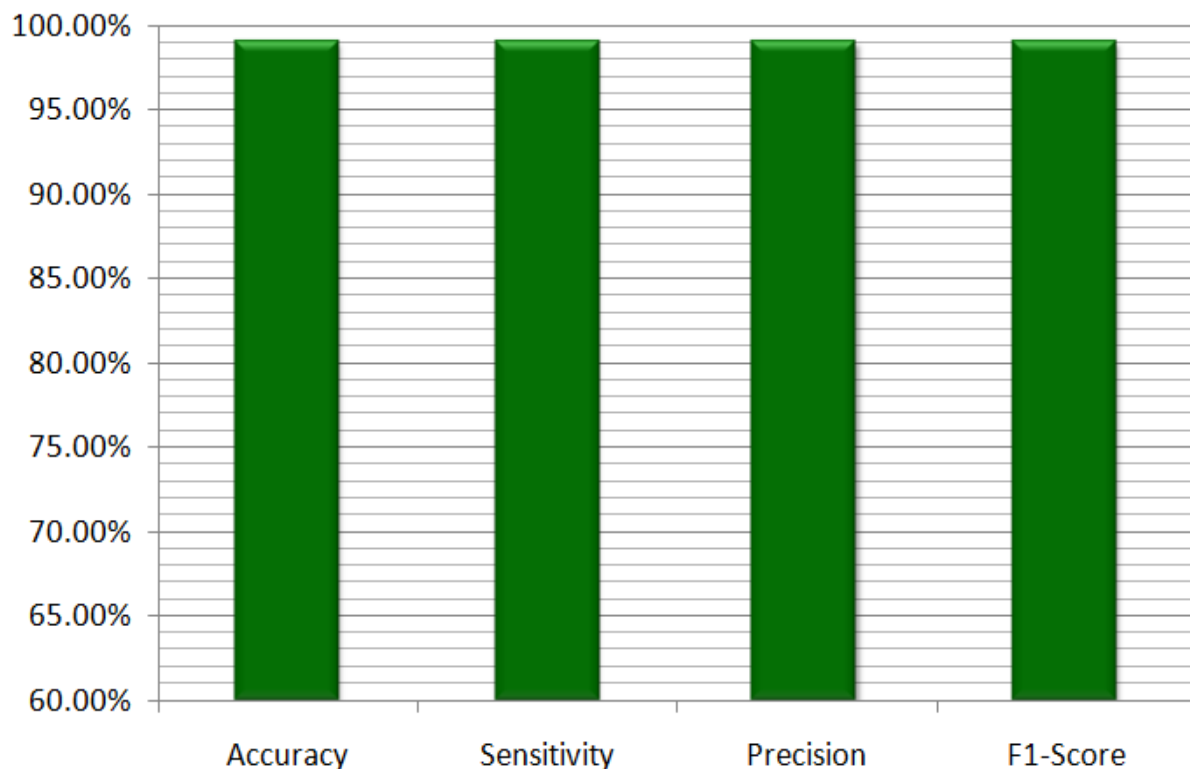
أما بالنسبة إلى قيم الحساسية فكانت الأعلى ومتساوية في كل من الصنفين AD و NL (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD و NL بشكل True Positive (140 حالة) وأقل عدد من حالات AD و NL بشكل False Negative (0 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (97.14%) وذلك كون النموذج صنف أقل عدد من حالات على

أنها MCI بشكل True Positive (136 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (4 حالة) مقارنة بكلا الصنفين AD و NL.

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (100.00%) وأقل منها في الصنف NL (98.59%) وكانت قيمة درجة F1 الأقل في الصنف MCI (98.55%).

في النهاية لابد من تقييم أداء وكفاءة نموذج GB كاملاً حيث تم حساب كل من الدقة، الحساسية، الإحكام، ودرجة F1 لنموذج ككل ويوضح الشكل (5-28) قيم هذه المقاييس.

Performance Measures of Gradient Boosting Model



الشكل (5-28) مقاييس الأداء للنموذج GB.

حيث تم حساب قيم هذه المقاييس للنموذج كنسبة مئوية حيث كانت قيمة الدقة تساوي 99.05%، وأما قيمة الحساسية بلغت 99.05%، وبالنسبة إلى قيمة الإحكام فكانت 99.07%، وبلغت قيمة درجة F1 99.05%.

5-3-8 مقارنة بين النماذج:

يظهر الجدول (5-8) ملخص عن مقاييس الأداء التي تم إيجادها لنموذج ككل لجميع النماذج التي تم استخدامها في هذه الدراسة بهدف المقارنة فيما بينها من أجل تحديد أفضل نموذج والذي يبدي أفضل أداء وكفاءة.

الجدول (5-8) مقارنة بين أداء النماذج المدربة باستخدام 21 سمة.

All Features	Model's Name	Accuracy	Sensitivity	Precision	F1-Score
	SVM	67.86%	67.86%	66.58%	65.97%
	NB	95.48%	95.48%	95.60%	95.43%
	KNN	68.33%	68.33%	67.93%	67.93%
	DT	97.14%	97.14%	97.14%	97.14%
	RF	98.81%	98.81%	98.83%	98.80%
	AB	92.86%	92.86%	93.13%	92.71%
	GB	99.05%	99.05%	99.07%	99.05%

حيث يمكن ملاحظة أن نموذج GB هو الأعلى قيمةً لجميع مقاييس الأداء وبالتالي هو النموذج الأفضل كفاءة وأداءً، ويعود ذلك إلى مجموعة المتعلمين الضعفاء التي تتدرب على أخطاء بعضها البعض لدمج في النموذج النهائي القوي، ويليه وتقريباً مساوي له نموذج RF حيث أيضاً جميع مقاييس أدائه كانت ذات القيم الأعلى وتقريباً مساوية لمقاييس أداء نموذج GB، وبالتالي يمكن اعتبار نموذج RF أيضاً من أفضل النماذج كفاءة وأداءً، ويعود ذلك إلى مجموعة أشجار القرار التي يتألف منها هذا النموذج واستقلالية كل منها والتي تعطي قرارها في التصنيف كل على حدا ليتم في النهاية اتخاذ القرار الذي اختير بشكل أكبر من مجموعة الأشجار كقرار نهائي.

5-4 أداء النماذج بعد تقليل السمات:

بعد الانتهاء من عرض نتائج النماذج مع 21 سمة ومناقشتها سيتم عرض نتائج النماذج بعد تقليل السمات ومناقشتها لكل من الطرق الثلاثة التي تم استخدامها في هذه الدراسة لتقليل السمات.

5-4-1 مع سمات طريقة اختيار K الأفضل:

يظهر الجدول (5-9) أفضل تسع سمات تم الوصول إليهم باستخدام طريقة اختيار K الأفضل SKB والتي سيتم إعادة تدريب جميع النماذج السابقة بها واختبارها وإيجاد مقاييس أداءها للوصول إلى النموذج الأفضل من حيث الكفاءة والأداء.

الجدول (5-9) أفضل تسع سمات بحسب SKB.

Top 9 SKB Features
CDRSB
ADAS11
ADAS13
ADASQ4
MMSE
LDELTOTAL
FAQ
mPACCdigit
mPACCtrailsB

يظهر الجدول (5-10) ملخص عن مقاييس الأداء التي تم إيجادها لنموذج ككل مع سمات SKB وذلك لجميع النماذج التي تم استخدامها في هذه الدراسة بهدف المقارنة فيما بينها من أجل تحديد أفضل نموذج والذي يبدي أفضل أداء وكفاءة.

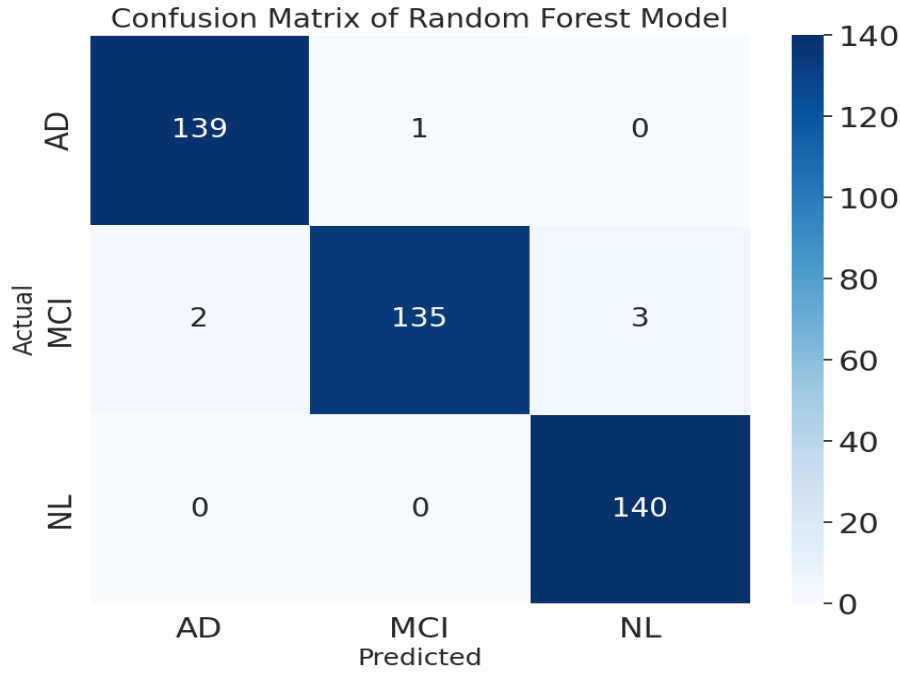
الجدول (5-10) مقارنة بين أداء النماذج المدربة باستخدام سمات SKB.

SKB Features	Model's Name	Accuracy	Sensitivity	Precision	F1-Score
	SVM	96.43%	96.43%	96.43%	96.43%
	GNB	95.71%	95.71%	95.76%	95.67%
	KNN	92.62%	92.62%	93.10%	92.44%
	DT	96.43%	96.43%	96.43%	96.41%
	RF	98.57%	98.57%	98.58%	98.57%
	AB	92.86%	92.86%	93.13%	92.73%
	GB	98.33%	98.33%	98.33%	98.35%

حيث يمكن ملاحظة أن نموذج RF هو الأعلى قيمةً لجميع مقاييس الأداء وبالتالي هو النموذج الأفضل كفاءة وأداء (حيث كانت القيم مساوية لقيم مقاييس أداء نموذج RF قبل تقليل السمات تقريباً مع اختلاف لم يتجاوز 0.24%)، فيما يليه وتقريباً مساوي له نموذج GB حيث أيضاً جميع مقاييس أداءه كانت ذات القيم الأعلى وتقريباً مساوية لمقاييس أداء نموذج RF، وبالتالي يمكن اعتبار نموذج GB أيضاً من أفضل النماذج كفاءة وأداء (حيث كانت القيم مساوية لقيم مقاييس أداء نموذج GB قبل تقليل السمات تقريباً مع اختلاف لم يتجاوز 0.71%)، وبمقارنة باقي النماذج مع النماذج قبل تقليل السمات يمكن ملاحظة تحسن مقاييس الأداء لكل من نموذج SVM ونموذج KNN مع محافظة باقي النماذج تقريباً على قيم مقاييس أداءها، وسيتم عرض نتائج كل من النموذج RF والنموذج GB ومناقشتها بشكل أكثر تفصيل كونهما النموذجين الأفضل في هذه الحالة أيضاً.

5-4-1-1 الغابة العشوائية مع سمات SKB:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-29) عدد الحالات التي نجح وأخطأ نموذج RF مع سمات SKB في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-29) مصفوفة الارتباك للنموذج RF مع سمات SKB.

حيث يمكن ملاحظة أن نموذج RF مع سمات SKB نجح في تصنيف 414 حالة من جميع عينات الاختبار وحيث نجح في تصنيف جميع عينات اختبار الصنف NL ولم يخطئ النموذج في التصنيف بين الصنفين AD و NL، فمن عينات اختبار AD صنف 139 حالة AD صحيحة، و 1 حالة على أنها MCI، ولم يصنف أي حالات على أنها NL، وبالنسبة إلى عينات اختبار NL صنف 140 حالة NL صحيحة، ولم يصنف أي منها على أنها AD أو MCI، وأما بالنسبة إلى عينات اختبار صنف MCI صنف 135 حالة MCI صحيحة، و 2 حالة على أنها AD، و 3 حالات على أنها NL، وبذلك تكون مصفوفة الارتباك للنموذج RF مع سمات SKB قريبة من مصفوفة الارتباك للنموذج RF قبل تقليل السمات مع اختلاف بسيط (1 حالة)، ويوضح الجدول (5-11) جميع حالات Positive و Negative للنموذج RF مع سمات SKB.

الجدول (5-11) حالات Positive و Negative للنموذج RF مع سمات SKB.

	True Positive	False Positive	True Negative	False Negative
AD	139	2	278	1
MCI	135	1	279	5
NL	140	3	277	0

وأما بالنسبة إلى تقرير التصنيف للنموذج RF مع سمات SKB والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-30).

Classification Report of Random Forest Model:			
	precision	recall	f1-score
0	0.9858	0.9929	0.9893
1	0.9926	0.9643	0.9783
2	0.9790	1.0000	0.9894

الشكل (5-30) تقرير التصنيف للنموذج RF مع سمات SKB.

حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف MCI (99.26%) وذلك كون النموذج من صنف أقل عدد من حالات على أنها MCI بشكل False Positive (1 حالة) مع عدد من حالات على أنها MCI بشكل True Positive (135 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف NL (97.90%) وذلك كون النموذج صنف أكثر عدد من حالات على أنها NL بشكل False Positive (3 حالات) مع عدد من حالات على أنها NL بشكل True Positive (140 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف AD كانت متوسطة بين القيمتين السابقتين (98.58%) كون النموذج صنف عدد حالات على أنها AD بشكل False Positive أكثر من MCI وأقل من NL (2 حالة) مع عدد من حالات على أنها AD بشكل True Positive (139 حالة).

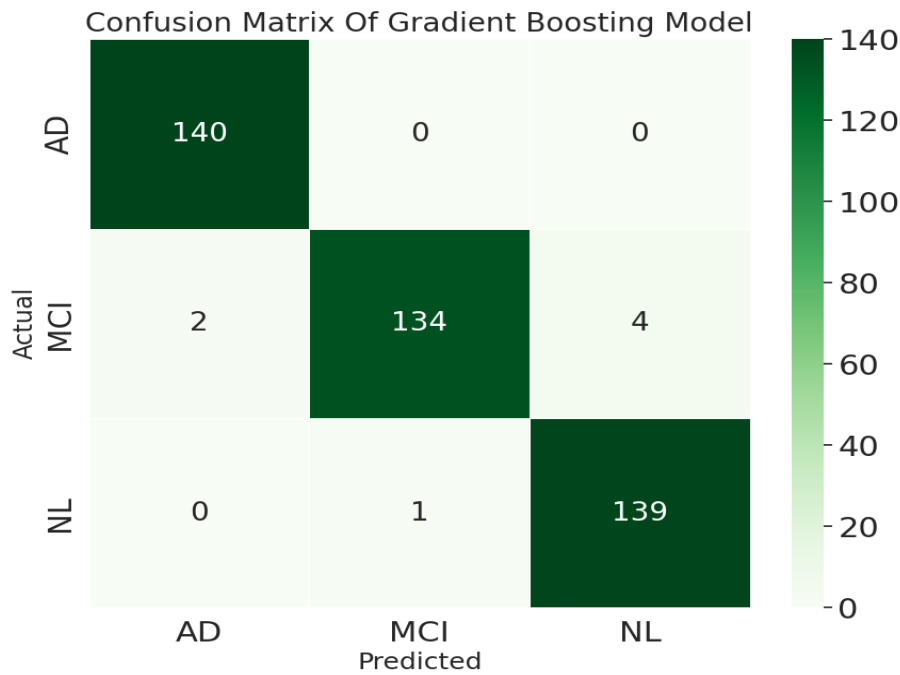
أما بالنسبة إلى قيم الحساسية فكانت الأعلى في الصنف NL (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها NL بشكل True Positive (140 حالة) ولم يصنف أي من حالات NL بشكل False Negative (0 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (96.43%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (135 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (5 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف AD كانت متوسطة بين القيمتين السابقتين (99.29%) كون النموذج صنف عدد من الحالات على أنها AD بشكل True Positive أكثر من MCI وأقل من

NL (139 حالة) وعدد من الحالات AD بشكل False Negative أقل من MCI وأكثر من NL (1 حالة).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف NL (98.94%) وأقل منها في الصنف AD (98.93%) وكانت قيمة درجة F1 الأقل في الصنف MCI (97.83%)، وبذلك تكون قيم تقرير التصنيف للنموذج RF مع سمات SKB قريبة من قيم تقرير التصنيف للنموذج RF قبل تقليل السمات.

5-4-1-2 تعزيز التدرج مع سمات SKB:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-31) عدد الحالات التي نجح وأخطأ نموذج GB مع سمات SKB في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-31) مصفوفة الارتباك للنموذج GB مع سمات SKB.

حيث يمكن ملاحظة أن نموذج GB مع سمات SKB نجح في تصنيف 413 حالة من جميع عينات الاختبار وحيث نجح في تصنيف جميع عينات اختبار الصنف AD بشكل صحيح ولم يخطئ النموذج في التصنيف بين الصنفين AD و NL، فمن عينات اختبار AD صنف 140 حالة AD صحيحة، ولم يصنف أي منها على أنها MCI أو NL، وبالنسبة إلى عينات اختبار NL صنف 139 حالة NL

صحيحة، و1 حالة على أنها MCI، ولم يصنف أي حالات منها على أنها AD، وأما بالنسبة إلى عينات اختبار MCI صنف 134 حالة MCI صحيحة، و4 حالات على أنها NL، و2 حالة على أنها AD، وبذلك تكون مصفوفة الارتباك للنموذج GB مع سمات SKB قريبة من مصفوفة الارتباك للنموذج GB قبل تقليل السمات مع اختلاف بسيط (3 حالات)، ويوضح الجدول (5-12) جميع حالات Positive و Negative للنموذج GB مع سمات SKB.

الجدول (5-12) حالات Positive و Negative للنموذج GB مع سمات SKB.

	True Positive	False Positive	True Negative	False Negative
AD	140	2	278	0
MCI	134	1	279	6
NL	139	4	276	1

وأما بالنسبة إلى تقرير التصنيف للنموذج GB مع سمات SKB والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-32).

Classification Report of Gradient Boosting Model:			
	precision	recall	f1-score
0	0.9859	1.0000	0.9929
1	0.9926	0.9571	0.9745
2	0.9720	0.9929	0.9823

الشكل (5-32) تقرير التصنيف للنموذج GB

حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف MCI (99.26%) وذلك كون النموذج من صنف 1 حالة على أنها MCI بشكل False Positive مع عدد من حالات على أنها MCI بشكل True Positive (134 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف NL (97.20%) وذلك كون النموذج صنف أكثر عدد من حالات على أنها NL بشكل False Positive (4 حالات) مع عدد من حالات على أنها NL بشكل True Positive (139 حالة)، وأما بالنسبة إلى قيمة الإحكام في

الصنف AD كانت متوسطة بين القيمتين السابقتين (98.59%) كون النموذج صنف عدد حالات على أنها AD بشكل False Positive أكثر من MCI وأقل من NL (2 حالة) مع عدد من حالات على أنها AD بشكل True Positive (140 حالة).

أما بالنسبة إلى قيم الحساسية فكانت الأعلى في الصنف AD (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها NL بشكل True Positive (140 حالة) ولم يصنف أي من حالات AD بشكل False Negative (0 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (95.71%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (134 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (6 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف NL كانت متوسطة بين القيمتين السابقتين (99.29%) كون النموذج صنف عدد من الحالات على أنها NL بشكل True Positive أكثر من MCI وأقل من AD (139 حالة) وعدد من الحالات NL بشكل False Negative أقل من MCI وأكثر من AD (1 حالة).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (99.29%) وأقل منها في الصنف NL (98.23%) وكانت قيمة درجة F1 الأقل في الصنف MCI (97.45%)، وبذلك تكون قيم تقرير التصنيف للنموذج GB مع سمات SKB قريبة من قيم تقرير التصنيف للنموذج GB قبل تقليل السمات.

5-4-2 مع سمات طريقة الشبكة المرنة:

يظهر الجدول (5-13) أفضل تسع سمات تم الوصول إليهم باستخدام طريقة الشبكة المرنة EN والتي سيتم إعادة تدريب جميع النماذج السابقة بها واختبارها وإيجاد مقاييس أدائها للوصول إلى النموذج الأفضل من حيث الكفاءة والأداء.

الجدول (5-13) أفضل تسع سمات بحسب EN.

Top 9 EN Features
CDRSB
ADASQ4
MMSE
RAVLT_immediate
RAVLT_learning
RAVLT_forgetting
LDELTOTAL
AGE
mPACCdigit

يظهر الجدول (5-14) ملخص عن مقاييس الأداء التي تم إيجادها لنموذج ككل مع سمات EN لجميع النماذج التي تم استخدامها في هذه الدراسة بهدف المقارنة فيما بينها من أجل تحديد أفضل نموذج والذي يبدي أفضل أداء وكفاءة

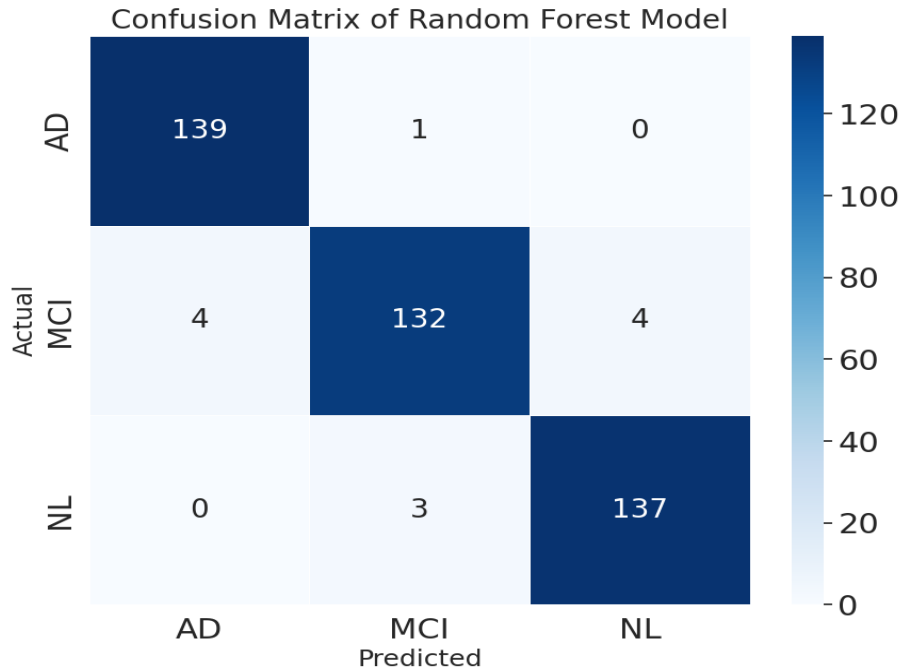
الجدول (5-14) مقارنة بين أداء النماذج المدربة باستخدام سمات EN.

EN Features	Model's Name	Accuracy	Sensitivity	Precision	F1-Score
	SVM	96.19%	96.19%	96.21%	96.16%
	NB	94.29%	94.29%	94.38%	94.21%
	KNN	88.57%	88.57%	88.58%	88.34%
	DT	95.00%	95.00%	95.03%	94.99%
	RF	97.14%	97.14%	97.14%	97.13%
	AB	92.86%	92.86%	93.13%	92.73%
	GB	97.38%	97.38%	97.38%	97.37%

حيث يمكن ملاحظة أن نموذج GB هو الأعلى قيمةً لجميع مقاييس الأداء وبالتالي هو النموذج الأفضل كفاءة وأداءً، فيما يليه وتقريباً مساوي له نموذج RF حيث أيضاً جميع مقاييس أداءه كانت ذات القيم الأعلى وتقريباً مساوية لمقاييس أداء نموذج GB، وبالتالي يمكن اعتبار نموذج RF أيضاً من أفضل النماذج كفاءة وأداءً (حيث انخفضت قيم مقاييس أداء النموذجين GB و RF قليلاً عن القيم قبل تقليل السمات مع اختلاف لم يتجاوز 1.67%)، وبالمقارنة مع النماذج قبل تقليل السمات فيمكن ملاحظة تحسن مقاييس الأداء لكل من نموذج SVM ونموذج KNN مع محافظة نموذج AB على مقاييس أدائها بالإضافة إلى انخفاض مقاييس أداء باقي النماذج، وسيتم عرض نتائج كل من نموذج GB ونموذج RF ومناقشتها بشكل أكثر تفصيل كونهما النموذجين الأفضل في هذه الحالة أيضاً.

5-4-2-1 الغابة العشوائية مع سمات EN:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-33) عدد الحالات التي نجح وأخطأ نموذج RF في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-33) مصفوفة الارتباك للنموذج RF مع سمات EN.

حيث يمكن ملاحظة أن نموذج RF مع سمات EN نجح في تصنيف 408 حالة من جميع عينات الاختبار ولم يخطئ النموذج في التصنيف بين كل من الصنفين AD و NL، فمن عينات اختبار AD صنف 139 حالة AD صحيحة، و 1 حالة على أنها MCI، ولم يصنف أي منها على أنها NL،

وبالنسبة إلى عينات اختبار NL صنف 137 حالة NL صحيحة، و3 حالات على أنها MCI، ولم يصنف أي منها على أنها AD، وأما بالنسبة إلى عينات اختبار MCI صنف 132 حالة MCI صحيحة، و4 حالات على أنها AD، و4 حالات على أنها NL، وبذلك توضح مصفوفة الارتباك للنموذج RF مع سمات SKB وجود 7 تصنيفات خاطئة زيادة عن مصفوفة الارتباك للنموذج RF قبل تقليل السمات، ويوضح الجدول (5-15) جميع حالات Positive و Negative للنموذج RF مع سمات EN.

الجدول (5-15) حالات Positive و Negative للنموذج RF مع سمات EN.

	True Positive	False Positive	True Negative	False Negative
AD	139	4	276	1
MCI	132	4	276	8
NL	137	4	276	3

وأما بالنسبة إلى تقرير التصنيف للنموذج RF مع سمات EN والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-34).

Classification Report of Random Forest Model:			
	precision	recall	f1-score
0	0.9720	0.9929	0.9823
1	0.9706	0.9429	0.9565
2	0.9716	0.9786	0.9751

الشكل (5-34) تقرير التصنيف للنموذج RF مع سمات EN.

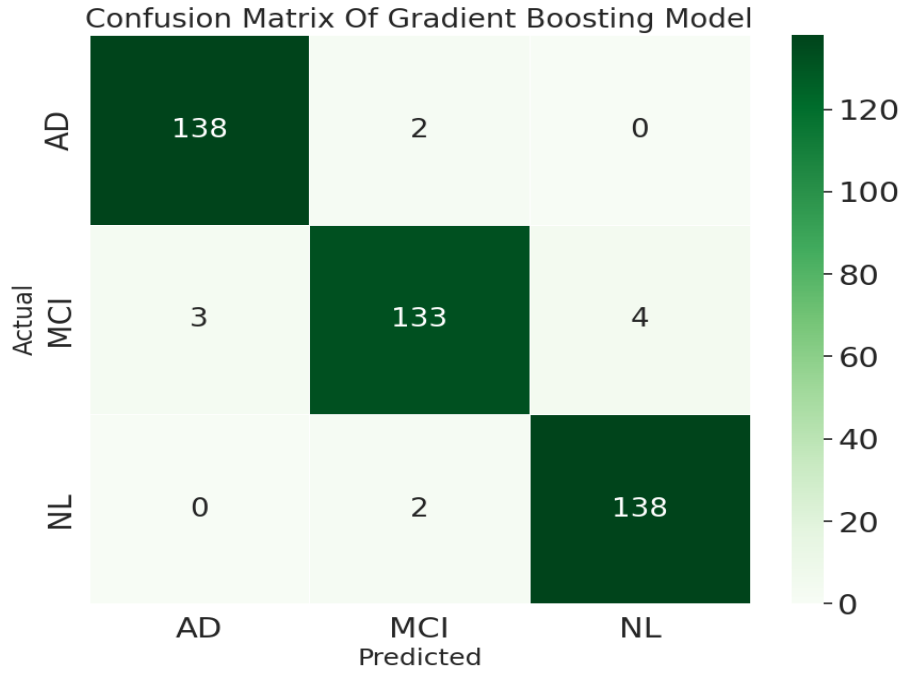
حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف AD (97.20%) وذلك كون النموذج صنف عدد من الحالات على أنها AD بشكل False Positive (4 حالات) مع أكبر عدد من حالات على أنها AD بشكل True Positive (139 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف MCI (97.06%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive

(132 حالات) مع عدد من الحالات على أنها MCI بشكل False Positive (4 حالات)، وأما بالنسبة إلى قيمة الإحكام في الصنف NL كانت متوسطة بين القيمتين السابقتين (97.16%) كون النموذج صنف عدد حالات على أنها NL بشكل False Positive نفس AD و MCI (4 حالة) ولكنه صنف عدد من حالات على أنها NL بشكل True Positive أقل من AD وأكثر من MCI (137 حالة). أما بالنسبة إلى قيم الحساسية فكانت الأعلى أيضاً في الصنف AD (99.29%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD بشكل True Positive (139 حالة) وأقل عدد من حالات AD بشكل False Negative (1 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (94.29%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (132 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (8 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف NL كانت متوسطة بين القيمتين السابقتين (97.86%) كون النموذج صنف عدد من الحالات على أنها NL بشكل True Positive أكثر من MCI وأقل من AD (137 حالة) وعدد من الحالات NL بشكل False Negative أقل من MCI وأكثر من AD (3 حالات).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (98.23%) وأقل منها في الصنف NL (97.51%) وكانت قيمة درجة F1 الأقل في الصنف MCI (95.65%)، وبذلك تكون قيم تقرير التصنيف للنموذج RF مع سمات EN أقل من قيم تقرير التصنيف للنموذج RF قبل تقليل السمات.

5-4-2 تعزيز التدرج مع سمات EN:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-35) عدد الحالات التي نجح وأخطأ نموذج GB مع سمات EN في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-35) مصفوفة الارتباك للنموذج GB مع سمات EN.

حيث يمكن ملاحظة أن نموذج GB مع سمات EN نجح في تصنيف 409 حالة من جميع عينات الاختبار ولم يخطئ النموذج في التصنيف بين كل من الصنفين AD و NL، فمن عينات اختبار AD صنف 138 حالة AD صحيحة، و 2 حالة على أنها MCI، ولم يصنف أي منها على أنها NL، وبالنسبة إلى عينات اختبار NL صنف 138 حالة NL صحيحة، و 2 حالات على أنها MCI، ولم يصنف أي منها على أنها AD، وأما بالنسبة إلى عينات اختبار MCI صنف 133 حالة MCI صحيحة، و 3 حالات على أنها AD، و 4 حالات على أنها NL، وبذلك توضح مصفوفة الارتباك للنموذج GB مع سمات EN وجود 7 تصنيفات خاطئة زيادة عن مصفوفة الارتباك للنموذج GB قبل تقليل السمات، ويوضح الجدول (5-16) جميع حالات Positive و Negative للنموذج GB مع سمات EN.

الجدول (5-16) حالات Positive و Negative للنموذج GB مع سمات EN.

	True Positive	False Positive	True Negative	False Negative
AD	138	3	277	2
MCI	136	4	276	7
NL	138	4	276	2

وأما بالنسبة إلى تقرير التصنيف للنموذج GB مع سمات EN والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-36).

Classification Report of Gradient Boosting Model:			
	precision	recall	f1-score
0	0.9787	0.9857	0.9822
1	0.9708	0.9500	0.9603
2	0.9718	0.9857	0.9787

الشكل (5-36) تقرير التصنيف للنموذج GB مع سمات EN.

حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف AD (97.87%) وذلك كون النموذج صنف أقل عدد من الحالات على أنها AD بشكل False Positive (3 حالات) مع أكبر عدد من حالات على أنها AD بشكل True Positive (138 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف MCI (97.08%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (136 حالات) مع عدد من الحالات على أنها MCI بشكل False Positive (4 حالات)، وأما بالنسبة إلى قيمة الإحكام في الصنف NL كانت متوسطة بين القيمتين السابقتين (97.18%) كون النموذج صنف عدد حالات على أنها NL بشكل True Positive نفس AD وأكثر من MCI (138 حالة) ولكنه صنف عدد من حالات على أنها NL بشكل False Positive أقل من AD وأكثر من MCI (4 حالات).

أما بالنسبة إلى قيم الحساسية فكانت الأعلى ومتساوية أيضاً في كل من الصنفين AD و NL (98.57%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD و NL بشكل True Positive (138 حالة) وأقل عدد من حالات AD و NL بشكل False Negative (2 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (95.00%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (133 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (7 حالة) مقارنة بكلا الصنفين AD و NL.

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (98.22%) وأقل منها في الصنف NL (97.87%) وكانت قيمة درجة F1 الأقل في الصنف MCI (96.03%) وبذلك تكون قيم تقرير التصنيف للنموذج GB مع سمات EN أقل من قيم تقرير التصنيف للنموذج GB قبل تقليل السمات.

3-4-5 مع سمات طريقة استبعاد السمات العودية:

يظهر الجدول (5-17) أفضل تسع سمات تم الوصول إليهم باستخدام طريقة استبعاد السمات العودية RFE والتي سيتم إعادة تدريب جميع النماذج السابقة بها واختبارها وإيجاد مقاييس أدائها للوصول إلى النموذج الأفضل من حيث الكفاءة والأداء.

الجدول (5-17) أفضل تسع سمات بحسب RFE.

Top 9 RFE Features
CDRSB
ADAS13
LDELTOTAL
TRABSCOR
FAQ
Ventricles
Hippocampus
mPACCdigit
mPACCtrailsB

يظهر الجدول (5-18) ملخص عن مقاييس الأداء التي تم إيجادها لنموذج ككل مع سمات RFE لجميع النماذج التي تم استخدامها في هذه الدراسة بهدف المقارنة فيما بينها من أجل تحديد أفضل نموذج والذي يبدي أفضل أداء وكفاءة.

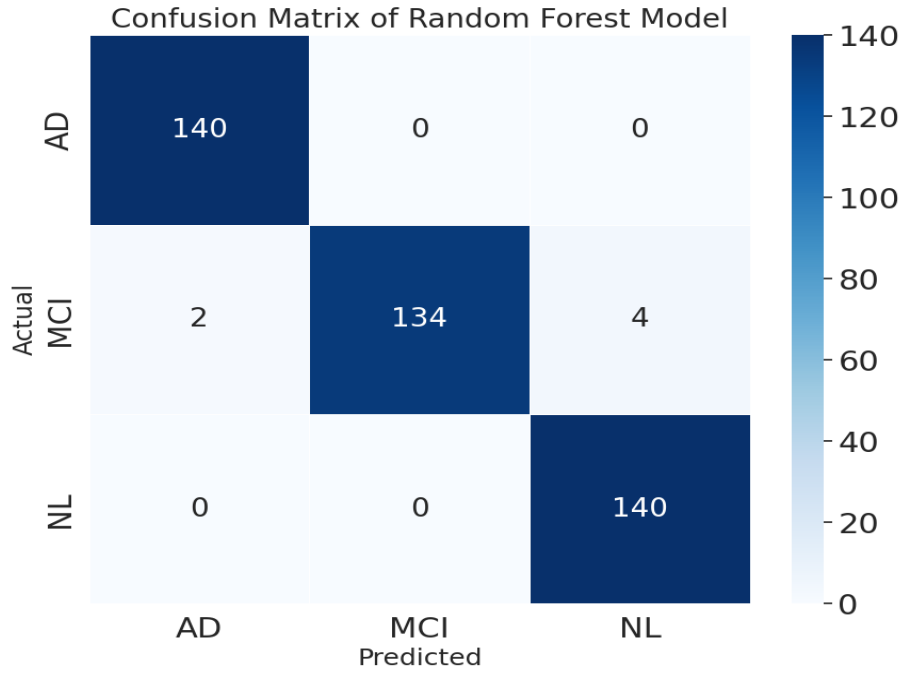
الجدول (5-18) مقارنة بين أداء النماذج المدربة باستخدام سمات RFE.

RFE Features	Model's Name	Accuracy	Sensitivity	Precision	F1-Score
	SVM	84.29%	84.29%	84.31%	84.27%
	NB	95.95%	95.95%	96.03%	95.91%
	KNN	72.62%	72.62%	72.47%	72.39%
	DT	96.19%	96.19%	96.19%	96.19%
	RF	98.57%	98.57%	98.60%	98.56%
	AB	92.86%	92.86%	93.13%	92.73%
	GB	99.05%	99.05%	99.06%	99.05%

حيث يمكن ملاحظة أن نموذج GB هو الأعلى قيمةً لجميع مقاييس الأداء وبالتالي هو النموذج الأفضل كفاءة وأداء (حيث كانت القيم مساوية لقيم مقاييس أداء نموذج GB قبل تقليل السمات تماماً)، فيما يليه وتقريباً مساوي له نموذج RF حيث أيضاً جميع مقاييس أداءه كانت ذات القيم الأعلى وتقريباً مساوية لمقاييس أداء نموذج GB، وبالتالي يمكن اعتبار نموذج RF أيضاً من أفضل النماذج كفاءة وأداء (حيث كانت القيم مساوية لقيم مقاييس أداء نموذج RF قبل تقليل السمات تقريباً مع اختلاف لم يتجاوز 0.24%)، وكما يمكن ملاحظة تحسن مقاييس الأداء لكل من النموذج SVM والنموذج KNN مع محافظة باقي النماذج تقريباً على قيم مقاييس أداءها، وسيتم عرض نتائج كل من النموذج RF والنموذج GB ومناقشتها بشكل أكثر تفصيل كونهما النموذجين الأفضل في هذه الحالة أيضاً.

5-4-3-1 الغابة العشوائية مع سمات RFE:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-37) عدد الحالات التي نجح وأخطأ نموذج RF مع سمات RFE في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-37) مصفوفة الارتباك للنموذج RF مع سمات RFE.

حيث يمكن ملاحظة أن نموذج RF مع سمات RFE نجح في تصنيف 414 حالة من جميع عينات الاختبار ولم يخطئ النموذج في التصنيف بين كل من الصنفين AD و NL، فمن عينات اختبار AD صنف 140 حالة AD صحيحة، ولم يصنف أي منها على أنها NL أو MCI، وبالنسبة إلى عينات اختبار NL صنف 140 حالة NL صحيحة، ولم يصنف أي منها على أنها AD أو MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 134 حالة MCI صحيحة، و2 حالة على أنها AD، و4 حالات على أنها NL، وبذلك توضح مصفوفة الارتباك للنموذج RF مع سمات RFE مساوية تقريباً مصفوفة الارتباك للنموذج RF قبل تقليل السمات مع اختلاف بسيط (1 حالة)، ويوضح الجدول (5-17) جميع حالات Positive و Negative للنموذج RF مع سمات RFE.

الجدول (5-19) حالات Positive و Negative للنموذج RF مع سمات RFE.

	True Positive	False Positive	True Negative	False Negative
AD	140	2	278	0
MCI	135	0	280	6
NL	140	4	276	0

وأما بالنسبة إلى تقرير التصنيف للنموذج RF مع سمات RFE والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-38).

Classification Report of Random Forest Model:			
	precision	recall	f1-score
0	0.9859	1.0000	0.9929
1	1.0000	0.9571	0.9781
2	0.9722	1.0000	0.9859

الشكل (5-38) تقرير التصنيف للنموذج RF مع سمات RFE.

حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف MCI (100.00%) وذلك كون النموذج لم يصنف أي من الحالات على أنها MCI بشكل False Positive مع عدد من حالات على أنها MCI بشكل True Positive (134 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف NL (97.22%) وذلك كون النموذج صنف أكبر عدد من الحالات على أنها NL بشكل False Positive (4 حالات) مع عدد من الحالات على أنها NL بشكل True Positive (140 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف AD كانت متوسطة بين القيمتين السابقتين (98.59%) كون النموذج صنف عدد حالات على أنها AD بشكل False Positive أقل NL وأكثر من MCI (2 حالة) مع عدد من حالات على أنها AD بشكل True Positive نفس NL وأكثر من MCI (140 حالة).

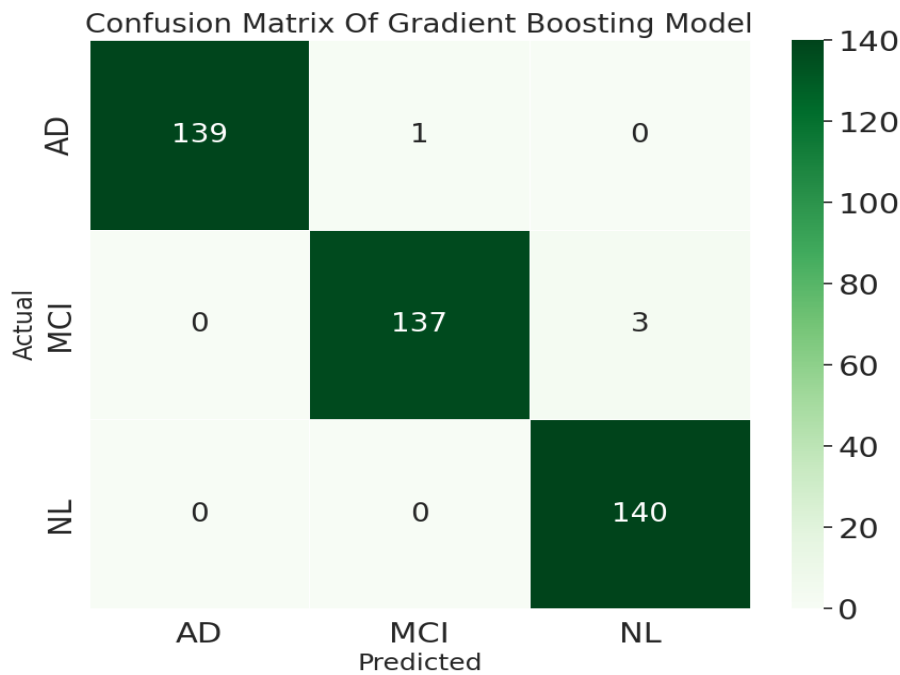
أما بالنسبة إلى قيم الحساسية فكانت الأعلى ومتساوية أيضاً في كل من الصنفين AD و NL (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها AD و NL بشكل True Positive (140 حالة) ولم يصنف أي من حالات AD و NL بشكل False Negative (0 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI (95.71%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (134 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (6 حالة).

وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (99.29%) وأقل منها في

الصنف NL (98.59%) وكانت قيمة درجة F1 الأقل في الصنف MCI (97.81%)، وبذلك تكون قيم تقرير التصنيف للنموذج RF مع سمات RFE مساوية تقريباً قيم تقرير التصنيف للنموذج RF قبل تقليل السمات.

5-4-3-2 تعزيز التدرج مع سمات RFE:

تظهر مصفوفة الارتباك الموضحة بالشكل (5-39) عدد الحالات التي نجح وأخطأ نموذج GB مع سمات RFE في تصنيفها من أصناف الخرج الثلاثة.



الشكل (5-39) مصفوفة الارتباك للنموذج GB مع سمات RFE.

حيث يمكن ملاحظة أن نموذج GB مع سمات RFE نجح في تصنيف 416 حالة من جميع عينات الاختبار ولم يخطئ النموذج في التصنيف بين كل من الصنفين AD و NL، فمن عينات اختبار AD صنف 139 حالة AD صحيحة، و 1 حالة على أنها MCI، ولم يصنف أي منها على أنها NL، وبالنسبة إلى عينات اختبار NL صنف 140 حالة NL صحيحة، ولم يصنف أي منها على أنها AD أو MCI، وأما بالنسبة إلى عينات اختبار MCI صنف 137 حالة MCI صحيحة، و 3 حالات على أنها NL، ولم يصنف أي منها على أنها AD، وبذلك توضح مصفوفة الارتباك لنموذج GB مع سمات RFE قريبة من مصفوفة الارتباك لنموذج GB قبل تقليل السمات حيث الاختلاف في تغير صنف خطأ واحد فقط، ويوضح الجدول (5-20) جميع حالات Positive و Negative للنموذج GB مع سمات RFE.

الجدول (5-20) حالات Positive و Negative للنموذج GB مع سمات RFE.

	True Positive	False Positive	True Negative	False Negative
AD	139	0	280	1
MCI	137	1	279	3
NL	140	3	277	0

وأما بالنسبة إلى تقرير التصنيف للنموذج GB مع سمات RFE والذي يعطي قيم مفصلة لمقاييس أداء هذا النموذج في كل صنف من أصناف الخرج الثلاثة، حيث يعطي قيمة كل من الإحكام والحساسية ودرجة F1، كما هو موضح في الشكل (5-40).

Classification Report of Gradient Boosting Model:			
	precision	recall	f1-score
0	1.0000	0.9929	0.9964
1	0.9928	0.9786	0.9856
2	0.9790	1.0000	0.9894

الشكل (5-40) تقرير التصنيف للنموذج GB مع سمات RFE.

حيث بالنسبة إلى قيم الإحكام كانت الأعلى في الصنف AD (100.00%) وذلك كون النموذج لم يصنف أي حالات على أنها AD بشكل False Positive مع عدد من حالات على أنها AD بشكل True Positive (139 حالة)، في حين كانت قيمة الإحكام الأقل في الصنف NL (97.90%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها NL بشكل False Positive (3 حالات) مع عدد من الحالات على أنها NL بشكل True Positive (140 حالة)، وأما بالنسبة إلى قيمة الإحكام في الصنف MCI كانت متوسطة بين القيمتين السابقتين (99.28%) كون النموذج صنف عدد حالات على أنها MCI بشكل False Positive أكثر من AD و أقل من NL (1 حالة) مع عدد من حالات على أنها MCI بشكل True Positive (137 حالة).

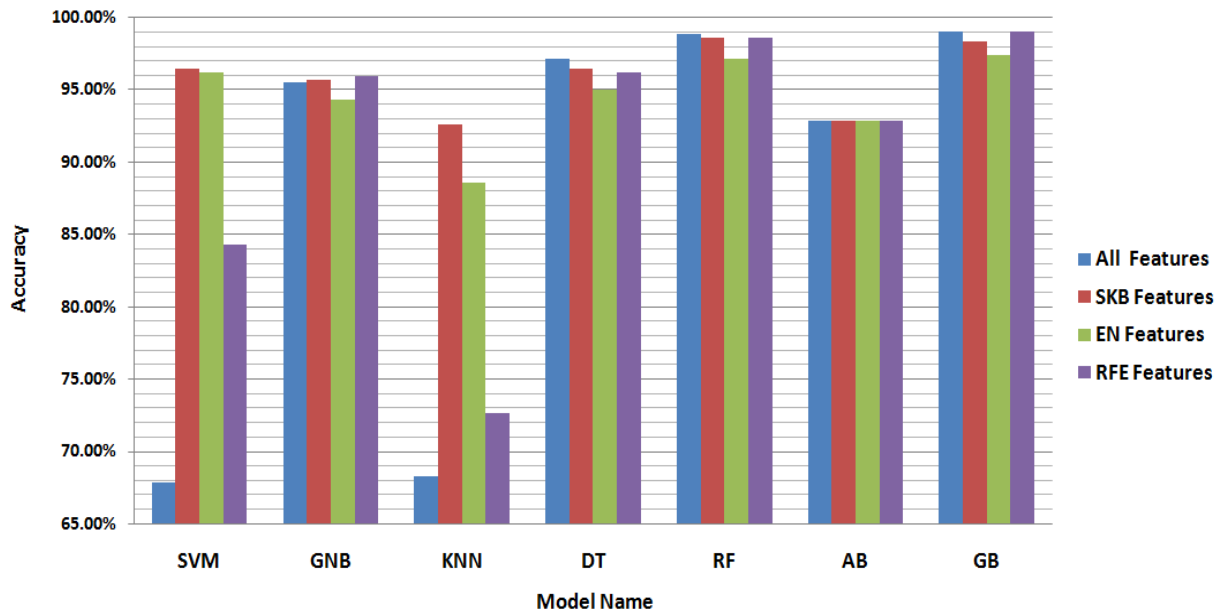
أما بالنسبة إلى قيم الحساسية فكانت الأعلى في الصنف NL (100.00%) وذلك كون النموذج صنف أكبر عدد من حالات على أنها NL بشكل True Positive (140 حالة) لم يصنف أي من حالات NL بشكل False Negative (0 حالة)، في حين كانت قيمة الحساسية الأقل في الصنف MCI

(97.86%) وذلك كون النموذج صنف أقل عدد من حالات على أنها MCI بشكل True Positive (137 حالة) وبالإضافة أنه صنف أكثر عدد من حالات MCI بشكل False Negative (3 حالة)، وأما بالنسبة إلى قيمة الحساسية في الصنف AD كانت متوسطة بين القيمتين السابقتين (99.92%) كون النموذج صنف عدد من الحالات على أنها AD بشكل True Positive أكثر من MCI وأقل من NL (139 حالة) وعدد من الحالات AD بشكل False Negative أقل من MCI وأكثر من NL (1 حالة).

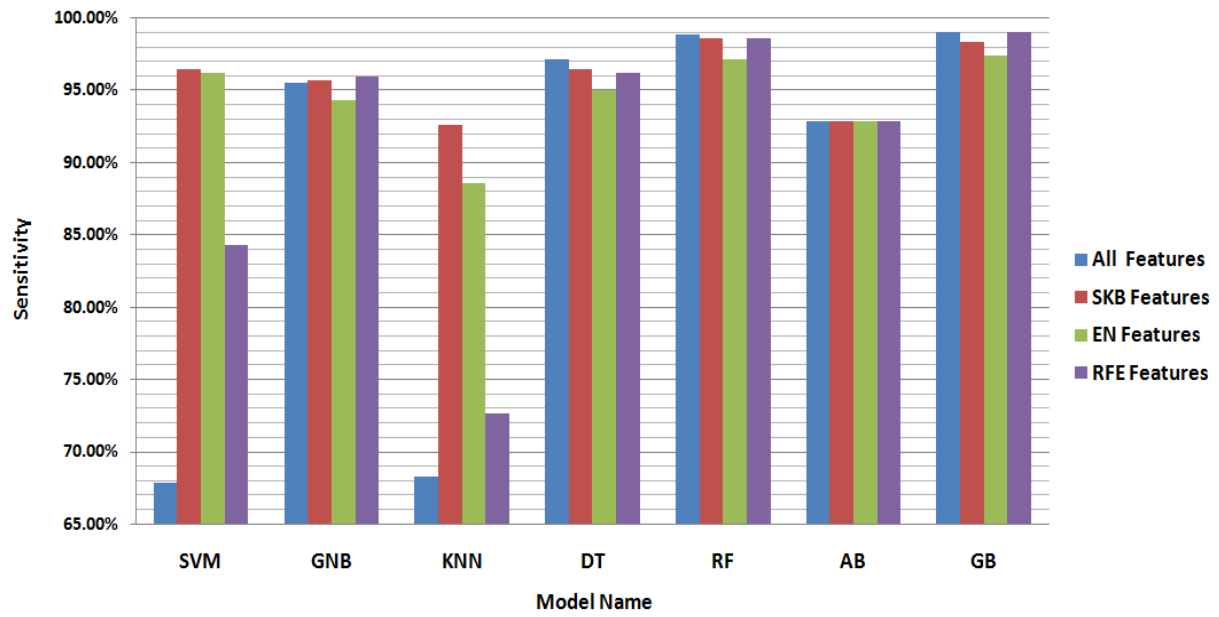
وأخيراً بالنسبة إلى قيم درجة F1 والتي تعتبر قيم شاملة لتحديد كفاءة النموذج في كل صنف وكونها تدمج بين المقياسين السابقين فكانت أعلى قيمة لدرجة F1 في الصنف AD (99.63%) وأقل منها في الصنف NL (98.94%) وكانت قيمة درجة F1 الأقل في الصنف MCI (98.56%)، وبذلك تكون قيم تقرير التصنيف للنموذج GB مع سمات RFE تقريباً مساوية إلى قيم تقرير التصنيف للنموذج GB قبل تقليل السمات.

5-5 مقارنة بين أداء جميع النماذج قبل وبعد تقليل السمات:

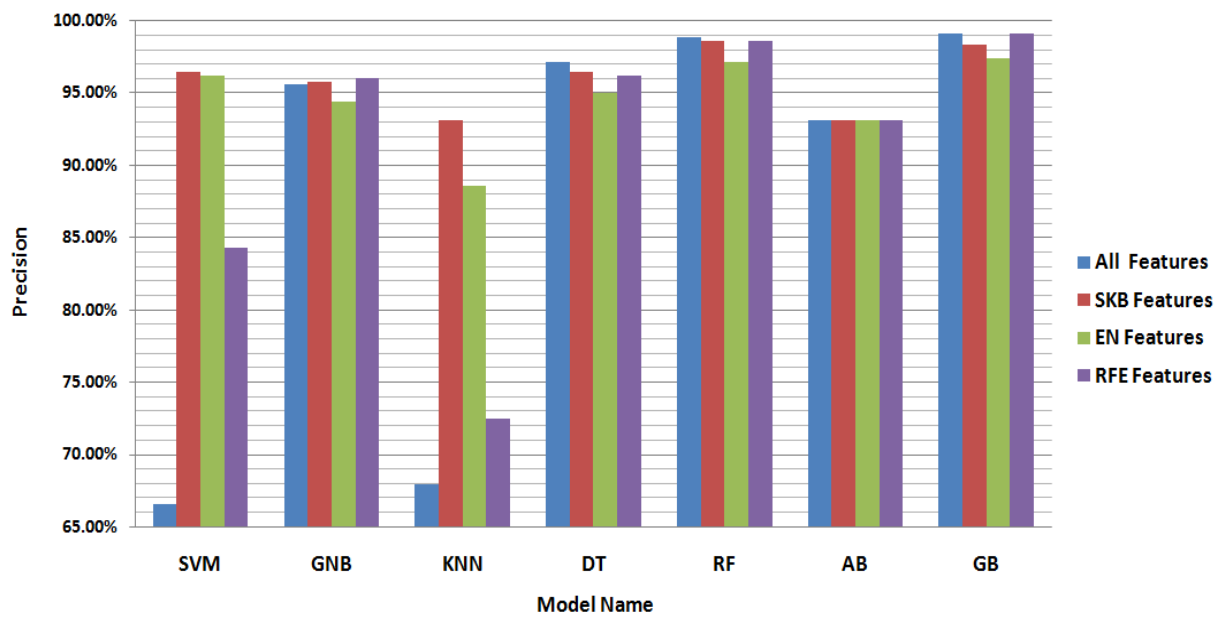
يظهر الشكل (5-41) قيم الدقة والشكل (5-42) قيم الحساسية والشكل (5-43) قيم الإحكام والشكل (5-44) قيم درجة F1 للنماذج قبل وبعد تقليل السمات، حيث نلاحظ تفوق نموذج GB ونموذج RF قبل تقليل السمات وبعدها.



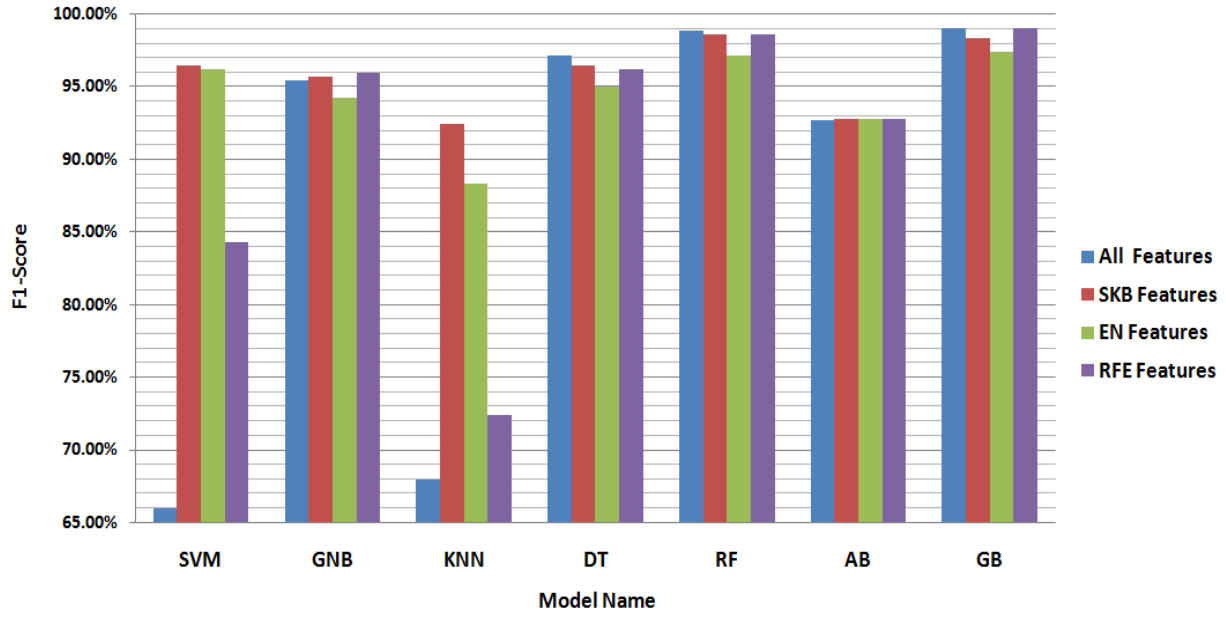
الشكل (5-41) قيم الدقة للنماذج قبل وبعد تقليل السمات.



الشكل (5-42) قيم الحساسية للنماذج قبل وبعد تقليل السمات.



الشكل (5-43) قيم الإحكام للنماذج قبل وبعد تقليل السمات.



الشكل (44-5) قيم درجة F1 للنماذج قبل وبعد تقليل السمات.

5-6 مقارنة مع الدراسات المرجعية:

تفوق هذا البحث على الدراسات السابقة بالنماذج التي تم بناؤها خصوصاً بكل من النموذجين GB و RF في كل من الحالتين قبل تقليل السمات وبعدها، وذلك نتيجة المنهجية المتبعة، ومعالجة البيانات بشكل جيد من التعامل مع كل من القيم المفقودة والقيم الشاذة والأهم موازنة عدد العينات في كل صنف الذي أغفلته أغلب الدراسات، وبالإضافة إلى ضبط المحددات الفائقة لكل نموذج باستخدام البحث الشبكي والتحقق المتقاطع، وكما لم تبين أغلب الدراسات آلية ضبط المحددات الفائقة وعملية التحقق المتقاطع، ويوضح الجدول (5-21) مقارنة هذا البحث مع الدراسات السابقة من حيث أفضل نموذج بناءً على قيمة كل من الدقة والحساسية.

الجدول (5-21) المقارنة هذا البحث مع الدراسات السابقة.

Study's Name	Year	Dataset	The Best Model	Accuracy & Sensitive
Predicting Alzheimer's disease progression using multi-modal deep learning approach	2019	ADNI	DL	Acc: 82.00%
Classification of Alzheimer's Disease using Machine Learning Techniques	2019	ADNI	GLM	Acc: 88.24%

Machine learning prediction of incidence of Alzheimer's disease using large-scale administrative health data	2020	NHIS	RF	Acc: 78.80%
Rule out Screening for Undiagnosed Dementia and Alzheimer's Disease Using an EHR Based Machine Learning Solution	2021	ADNI	AB	Acc: 90.81%
An Efficient Hybrid Approach For Diagnosis High Dimensional Data For Alzheimer's Diseases Using Machine Learning Algorithms	2022	ADNI	NN	Acc: 95.00%
Dementia Prediction Using Machine Learning	2023	OASIS	SVM	Acc: 96.77% Sen: 87.93%
Logistic random forest boosting technique for Alzheimer's diagnosis	2023	ADNI	LRFB	Acc: 76.91% Sen: 73.89%
This Study	2024	ADNI	RF	Acc: 98.81% Sen: 98.81%
			GB	Acc: 99.05% Sen: 99.05%

5-7 الخاتمة:

تعددت الدراسات التي تعمل على تسخير تقنيات تعلم الآلة للوصول إلى التشخيص المبكر لمرض الزهايمر، لما في ذلك خدمة لهؤلاء المرضى الذين تظهر الأعراض السريرية عليهم بعد فترة طويلة من الإصابة، ومعظم الدراسات تؤكد على أن التشخيص المبكر قد يكون مفتاح العلاج.

الفصل السادس:

الخاتمة

6-1 مقدمة:

يعد هذا الفصل ختام هذه الرسالة والذي يستعرض الاستنتاجات كملخص للنتائج والمناقشة وكذلك مميزات هذا البحث، ويطرح بعض التوصيات والمقترحات للعمل بها كخطوات مستقبلية.

6-2 الاستنتاجات:

من الاستنتاجات التي تم الوصول إليها:

- 1- يعد كل من نموذج GB ونموذج RF هما الأفضل أداء وكفاءة مقارنة مع باقي النماذج المستخدمة في هذا البحث، حيث أظهرتا أعلى قيم في كل الدقة والحساسية والإحكام ودرجة F1، لذلك يمكن اعتبار كل منهما أفضل نموذج.
- 2- طريقة SKB وطريقة RFE هما أفضل من طريقة EN في تقليل السمات.
- 3- تحسن أداء كل من النموذج SVM والنموذج KNN بشكل كبير بعد تقليل السمات.
- 4- بقاء كل من نموذج GB ونموذج RF هما الأفضل أداء وكفاءة مقارنة مع باقي النماذج المستخدمة في هذا البحث بعد تقليل السمات لذلك يمكن اعتبار كل منهما أفضل نموذج.
- 5- محافظة كل من نموذج GB ونموذج RF على نفس الأداء والكفاءة بعد تقليل السمات بكل من طريقة SKB وطريقة RFE مع الأداء والكفاءة قبل تقليل السمات.
- 6- طريقة SKB أعطت أفضل تسع سمات متعلقة فقط بالاختبارات العصبية النفسية (Neuropsychological Testing).

6-3 مميزات البحث:

من النقاط التي تميز بها هذا البحث:

- 1- المنهجية المقترحة للبحث.
- 2- ضبط المحددات الفائقة لكل من النماذج المستخدمة وذلك باستخدام كل من البحث الشبكي والتحقق المقاطع.
- 3- معالجة القيم الشاذة، موازنة عدد عينات مجموعة البيانات في كل صنف من أصناف الخرج.
- 4- استخدام عدة طرق لتقليل الخصائص.

5- عرض أداء كل نموذج بحسب كل صنف.

6- تفوق نموذج GB ونموذج RF في أدائهما على الدراسات المرجعية والمساوية تقريبا 99% لكل المقاييس.

4-6 التوصيات والمقترحات:

من التوصيات والمقترحات المستقبلية لهذا البحث:

- 1- نشر التوعية بشكل أكبر بمرض الزهايمر ومرحلة ضعف الإدراك المعتدل.
- 2- التشجيع على إطلاق مبادرة طبية لإجراء الفحوصات وجمع المعلومات لكبار السن وذلك لأهمية الاهتمام بالشيخوخة ولتشكيل قواعد بيانات محلية تفيد أبحاث الشيخوخة.
- 4- العمل على تقليل عدد السمات لأقل عدد ممكن.
- 5- العمل على السمات المتعلقة بالاختبارات النفسية العصبية (Neuropsychological Tests).
- 6- تجريب تقنيات تعلم الآلة الأخرى لبناء النموذج.

5-6 خاتمة:

عمل هذا البحث على بناء نموذج للتنبؤ والكشف المبكر عن الأشخاص المعرضين لخطر الإصابة بمرض الزهايمر وكانت نتائجه واعدة في الكشف عن الأصناف الثلاثة: مرض الزهايمر Alzheimer Disease (AD)، وضعف الإدراك المعتدل Mild Cognitive Impairment (MCI)، والإدراك الطبيعي Normal Cognitive (NL)، لذلك يجب العمل على تطويره مستقبلاً.

المراجع

- 1- Scheltens, P., Strooper, B., Kivipelto, M., Holstege, H., Ch  telat, G., Teunissen, Ch., Cummings, J., Wiesje van der Flier. (2021). Alzheimer's Disease. **Lancet**. 397. 1577-1590.
- 2- Alzheimer's Association. (2023). 2023 Alzheimer's disease facts and figures. **Alzheimer's & Dementia**. 17(3). 327  406.
- 3- Bondi, M. W., Edmonds, E. C., & Salmon, D. P. (2017). Alzheimer's disease: past, present, and future. **Journal of the International Neuropsychological Society**. 23(9-10). 818-831.
- 4- Hanseeuw, B., Betensky R., Jacobs H., Schultz A., Sepulcre, J., Becker, J., et al. (2019). Association of amyloid and tau with cognition in preclinical Alzheimer disease. **JAMA Neurol**. 76(8). 915-24.
- 5- Long, S., Benoist, C., Weidner, W. 2023. **World Alzheimer Report 2023: Reducing dementia risk: never too early, never too late**. England: Alzheimer's Disease International.
- 6- Geerts, H., Dacks, P. A., Devanarayan, V., Haas, M., Khachaturian, Z. S., Gordon, M. F., ... & Brain Health Modeling Initiative. (2016). Big data to smart data in Alzheimer's disease: The brain health modeling initiative to foster actionable knowledge. **Alzheimer's & Dementia**. 12(9). 1014-1021.
- 7- Breijyeh, Z., &Karaman, R. (2020). Comprehensive review on Alzheimer's disease: causes and treatment. **Molecules**. 25(24). 5789.
- 8- Murphy, K. P. (2012). **Machine learning: a probabilistic perspective**. MIT press.
- 9- Pisner, D. A., &Schnyer, D. M. (2020). Support vector machine. In Machine learning. **Academic Press**. 101-121.
- 10- Raschka, S. (2014). Naive bayes and text classification I-introduction and theory. **arXiv**. arXiv:1410.5329.
- 11- Kazmierska, J., &Malicki, J. (2008). Application of the Na  ve Bayesian Classifier to optimize treatment decisions. **Radiotherapy and Oncology**. 86(2). 211-216.
- 12- Mucherino, A., Papajorgji, P. J., Pardalos, P. M., Mucherino, A., Papajorgji, P. J., &Pardalos, P. M. (2009). K-nearest neighbor classification. **Data mining in agriculture**. 34. 83-106.
- 13- Shyrokykh, K., Girnyk, M., &Dellmuth, L. (2023). Short text classification with machine learning in the social sciences: The case of climate change on Twitter. **Plos one**. 18(9). e0290762.
- 14- Dutt, S., Chandramouli, S., & Dsa, A. (2018). **Machine learning**. India: Pearson Education.
- 15- Charbuty, B., & Abdulazeez, A. (2021). Classification based on decision tree algorithm for machine learning. **Journal of Applied Science and Technology Trends**. 2(1). 20-28

- 16- Priyanka, & Kumar, D. (2020). Decision tree classifier: a detailed survey. **International Journal of Information and Decision Sciences**. 12(3). 246-269.
- 17- Gaber, M. M., & Atwal, H. S. (2013). An entropy-based approach to enhancing Random Forests. **Intelligent Decision Technologies**, 7(4), 319-327.
- 18- Rastogi, R., & Shim, K. (2000). PUBLIC: A Decision Tree Classifier that Integrates Building and Pruning. **Data Mining and Knowledge Discovery**. 4(4). 315–344.
- 19- Parmar, A., Katariya, R., & Patel, V. (2019). **A review on random forest: An ensemble classifier**. In International conference on intelligent data communication technologies and internet of things (ICICI) 2018 (pp. 758-763). Springer International Publishing.
- 20- Sahour, H., Gholami, V., Torkaman, J., Vazifedan, M., & Saeedi, S. (2021). Random forest and extreme gradient boosting algorithms for streamflow modeling using vessel features and tree-rings. **Environmental Earth Sciences**. 80. 1-14.
- 21- Kim, T. H., Park, D. C., Woo, D. M., Jeong, T., & Min, S. Y. (2012). **Multi-class classifier-based adaboost algorithm**. In Intelligent Science and Intelligent Data Engineering: Second Sino-foreign-interchange Workshop, IScIDE 2011. China.
- 22- Khan, N. I., Mahmud, T., Islam, M. N., & Mustafina, S. N. (2020). **Prediction of cesarean childbirth using ensemble machine learning methods**. In Proceedings of the 22nd international conference on information integration and web-based applications & services. Thailand.
- 23- Natekin, A., & Knoll, A. (2013). Gradient boosting machines: a tutorial. **Frontiers in neurorobotics**. 7. Article 21.
- 24- Nhat-Duc, H., & Van-Duc, T. (2023). Comparison of histogram-based gradient boosting classification machine, random Forest, and deep convolutional neural network for pavement raveling severity classification. **Automation in Construction**. 148. 104767.
- 25- Grandini, M., Bagli, E., & Visani, G. (2020). Metrics for multi-class classification: an overview. **arXiv**. arXiv:2008.05756.
- 26- Markoulidakis, I., Kopsiaftis, G., Rallis, I., & Georgoulas, I. (2021, June). Multi-class confusion matrix reduction method and its application on net promoter score classification problem. In Proceedings of the **14th PErvasive Technologies Related to Assistive Environments Conference** (pp. 412-419).
- 27- Hossain, R., & Timmer, D. (2021). Machine learning model optimization with hyper parameter tuning approach. **Global Journal of Computer Science and Technology: D Neural & Artificial Intelligence**. 21(2). 31-39.
- 28- Ogunsanya, M., Isichei, J., & Desai, S. (2023). Grid search hyperparameter tuning in additive manufacturing processes. **Manufacturing Letters**. 35. 1031-1042.

- 29- Akinola, O., Ezugwu, A., Agushaka, J., Zitar, R., & Abualigah, L. (2022). Multiclass feature selection with metaheuristic optimization algorithms: a review. **Neural Computing and Applications**. 34(22). 19751-19790.
- 30- Kumar, V., & Minz, S. (2014). Feature selection: A literature review. **Smart Computing Review**. 4(3). 211-229.
- 31- Shastry, K. A., & Sattar, S. A. (2023). Logistic Random Forest Boosting Technique for Alzheimer's Diagnosis. **International Journal of Information Technology**. 15(3). 1719-1731.
- 32- Dhakal, S., Azam, S., Hasib, K. M., Karim, A., Jonkman, M., & Al Haque, A. F. (2023). Dementia Prediction Using Machine Learning. **Procedia Computer Science**. 219. 1297-1308.
- 33- ElZawawi, N. S., Saber, H. G., Hashem, M., & Gharib, T. (2022). An Efficient Hybrid Approach for Diagnosis High Dimensional Data for Alzheimer's Diseases Using Machine Learning Algorithms. **International Journal of Intelligent Computing and Information Sciences**. 22(2). 97-111.
- 34- Stephan, B., Julovich, D. A., Bracy, D., & Nguyen, J. (2021). Rule out Screening for Undiagnosed Dementia and Alzheimer's Disease Using an EHR Based Machine Learning Solution. **SMU Data Science Review**. 5(1). Article 5.
- 35- Park, J. H., Cho, H. E., Kim, J. H., Wall, M. M., Stern, Y., Lim, H., Yoo, S., Kim, H., & Cha, J. (2020). Machine Learning Prediction of Incidence of Alzheimer's Disease Using Large-Scale Administrative Health Data. **NPJ digital medicine**. 3(1). 46-53.
- 36- Shahbaz, M., Ali, S., Guergachi, A., Niazi, A., & Umer, A. (2019, July). Classification of Alzheimer's Disease using Machine Learning Techniques. **In Data**. 296-303.
- 37- Lee, G., Nho, K., Kang, B., Sohn, K. A., & Kim, D. (2019). Predicting Alzheimer's Disease Progression Using Multi-Modal Deep Learning Approach. **Scientific reports**. 9(1). 1952-1964.
- 38- Vinod, H. D. (2014). Matrix algebra topics in statistics and economics using R. **In Handbook of statistics**. (32, pp. 143-176). Elsevier.

Abstract

Alzheimer's disease is a neurodegenerative disease that worsens over time. Its clinical symptoms do not appear early but in the late stages. There are no drugs for this disease, but only an alleviation of its symptoms. Based on the importance of early detection of Alzheimer's disease in finding the appropriate treatment and allowing patients to choose to participate in clinical trials, this research worked to build a model for predicting and early detection of people at risk of developing Alzheimer's disease.

Seven Machine Learning Models, Support Vector Machine, Naive Bayes, Decision Tree, K-Nearest Neighbors, Random Forest, AdaBoost, and Gradient Boost, were used to detect the three classes: Alzheimer's Disease (AD), Mild Cognitive Impairment (MCI), and Normal Cognitive (NL).

The Alzheimer Disease Neuroimaging Initiative (ADNI) database was used, preprocessed, and initialized (missing and outlier values were processed, categorical features were encoded to numerical, and the number of samples in classes was balanced). Hyperparameters were tuned for all models using Grid Search and Cross Validation, then trained, tested, and evaluated to determine the best model. Features were reduced using Three methods, and models were retrained to reach the top nine features with the best model.

Both Gradient Boost and Random Forest outperformed the other models before and after feature reduction and were the best performing and efficient with RF model accuracy and sensitivity values equal to 98.81% and GB model accuracy and sensitivity values equal to 99.05%. This indicates promising results for using these models to support clinical decision-making.

Keywords: Alzheimer's Disease, Mild Cognitive Impairment, Early Detection of Alzheimer's, Machine Learning, Hyperparameters, Support Vector Machine, Naive Bayes, K-Nearest Neighbors, Decision Tree, Random Forest, AdaBoost, Gradient Boosting.

Syrian Arab Republic

Damascus University

Faculty of Mechanical & Electrical Engineering

Department of Biomedical Engineering



Early prediction and detection of Alzheimer's disease using Machine Learning

A dissertation submitted in partial fulfillment of the requirements for the degree
of Master in Biomedical Engineering

By:

Eng. George Riashi

Supervisor:

Prof. Rasha Massoud

2024 - 2025