

DESARROLLO DE MODELOS DE APRENDIZAJE AUTOMÁTICO APLICADOS A LA PATOLOGÍA DE ALZHEIMER

CLASIFICADOR DE GRADO DE DEMENCIA BASADO EN APRENDIZAJE AUTOMÁTICO SUPERVISADO.



Estudiante: Barros Tobar Ricardo Josué.
Director: Dr Aguiar Pontes Josafá De Jesús.



ESCUELA
POLITÉCNICA
NACIONAL

PLANTEAMIENTO DEL TRABAJO

El Alzheimer es un trastorno que causa el deterioro de las células cerebrales, que en principio se lo confunde con el envejecimiento natural, en trabajos previos se han desarrollado modelos para clasificar a un paciente dentro de las etapas del deterioro cognitivo derivados de la patología, en este TIC se desarrolla varios modelos de aprendizaje automático supervisado con el fin de observar cual tiene mejor rendimiento para el conjunto de datos.

Table 3. Confusion Matrix of given subjects TND*: True Non-Demented; TD*: True Demented; TC*: True Converted; PND*: Predict Non-Demented; PD*: Predict Demented, and PC*: Predict Converted.

	TND	TD	TC	precision
PND	43	14	10	64.18%
PD	8	27	1	75.00%
PC	2	0	0	0.00%
Recall	81.13%	65.85%	0.00%	0.00%

- Objetivo General:

Implementar modelos de aprendizaje automático supervisado con la capacidad necesaria para identificar el grado de demencia derivado del Alzheimer.

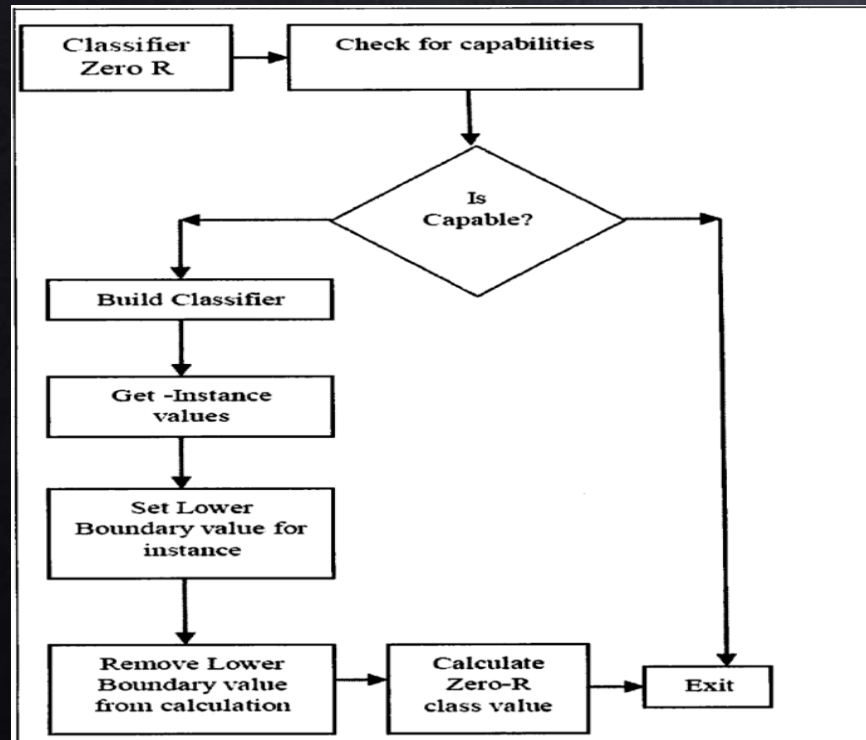
- Objetivos Específicos:

- Preparar los conjuntos de datos necesarios para llevar a cabo esta investigación.
- Desarrollar modelos estadísticos de aprendizaje supervisado para identificar el grado de demencia derivado del Alzheimer.
- Analizar los resultados obtenidos a partir de la calidad predictiva de los modelos desarrollados.
- Determinar cuál es la modelo más eficiente dada la base de datos estudiada.



METODOLOGÍA

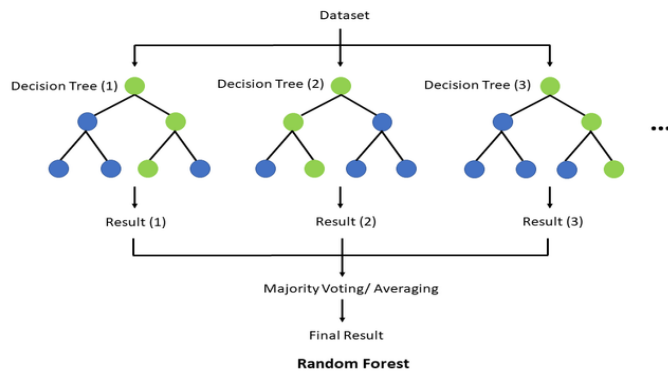
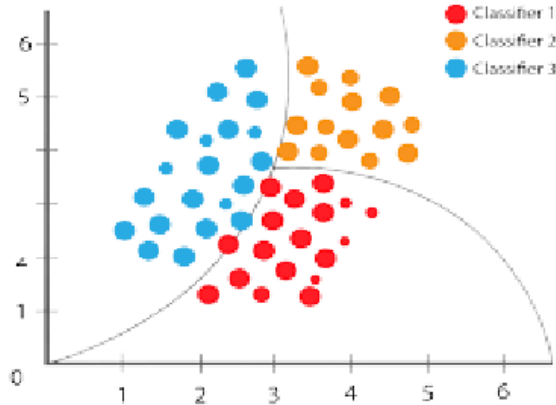
El establecer los baselines y metas a superar se realizó con la ayuda de la herramienta de análisis de datos Weka y el algoritmo Zero Rule teniendo así una referencia a superar para cada modelo desarrollado.



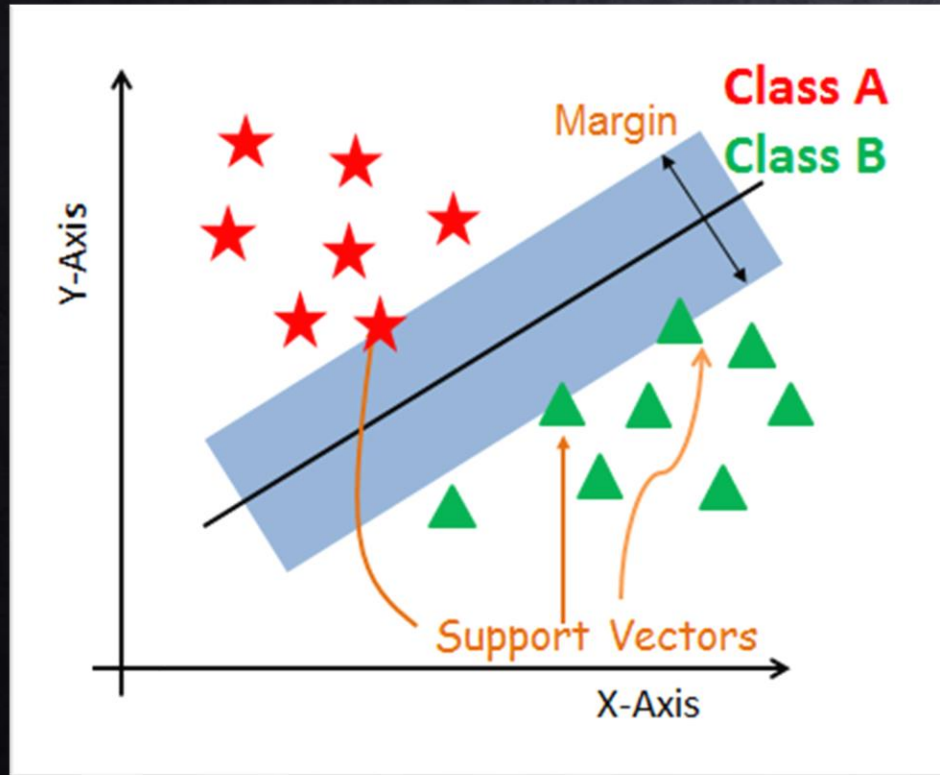
MODELOS SELECCIONADOS

El conjunto de datos a trabajar posee 375 instancias por lo que se opta por emplear modelos estadísticos de clasificación como lo son:

- Modelo Naïve bayes classifier: Modelo basado en el teorema de Bayes, que asume la independencia entre variables para que calculando la probabilidad de que una instancia pertenezca a una clase, dado un conjunto de características, y elige la clase con la mayor probabilidad.
- Modelo Random Forrest Classifier: Modelo de aprendizaje automático basado en la construcción de múltiples árboles de decisión que clasifica los datos promediando las predicciones de estos árboles, lo que mejora la precisión y reduce el riesgo de sobreajuste.



MODELOS SELECCIONADOS



- Support Vector Machine (SVM): Modelo de clasificación que se basa en encontrar el hiperplano óptimo que separa a los datos en diferentes clases, tratando de maximizar el margen entre las clases más cercanas (vectores de soporte) y el hiperplano, SVM determina en qué lado del hiperplano se encuentra una instancia y la asigna a una clase específica para esto puede utilizar diferentes funciones kernel para manejar datos no linealmente separables.

METODOLOGÍA

Una vez se realiza un análisis al conjunto de datos se determine en borrar variables poco significativas para eliminar posible ruido para los modelos. De manera general se estableció tomar el 20% de la data para testeo y el 80% restante para entrenamiento con 5 carpetas para realizar la validación cruzada con una semilla de tamaño 36, además de calcular métricas de evaluación como lo son el Accuracy, Recall, P-Value y obtener la matriz de confusión correspondiente.

CONFUSION MATRIX FOR MULTI-CLASS CLASSIFIER

P R E D I C T E D	ACTUAL			
		Actual_Red	Actual_Blue	Actual_Green

RESULTADOS

El baseline empleando el algoritmo Zero Rule tanto en la herramienta Weka como en el cuaderno de desarrollo para los modelos, quedo establecido en un 52% para la clasificación de las distintas instancias de la variable objetivo “Group”, además, se estableció los nombres de clase 0, 1 y 2, para representar los grupos de la variable quedando así:

- 0 representa al grupo Converted.
- 1 representa al grupo Demented.
- 2 representa al grupo Nondemented.

Classification Report:				
	precision	recall	f1-score	support
0	0.00	0.00	0.00	7
1	0.00	0.00	0.00	29
2	0.52	1.00	0.68	39
accuracy			0.52	75
macro avg	0.17	0.33	0.23	75
weighted avg	0.27	0.52	0.36	75
Confusion Matrix:				
[[0 0 7]				
[0 0 29]				
[0 0 39]]				
Test Accuracy: 0.5200				
Cross-validated Accuracy: 0.5094				

RESULTADOS

El modelo de Naïve Bayes Classifier obtuvo los siguientes resultados en la herramienta Weka, mismos que fueron superados en el modelo desarrollado.

Resultados obtenidos en Weka

test options

☐ Use training set

☐ Supplied test set

☒ Cross-validation

☐ Percentage split

Set...

Folds

5

%

66

More options...

Nom) Group

Start

Stop

result list (right-click for options)

5:25:06 - rules.ZeroR

5:25:22 - functions.LibSVM

5:32:19 - bayes.NaiveBayes

5:32:33 - trees.RandomForest

Classifier output

0.574MMSE=0.432nWBV-0.359Age-0.369EDUC+0.318CDR...

mean

-0.1031

0.0903

0.1725

std. dev.

0.4747

0.6373

0.4419

weight sum

190

146

37

precision

0.0097

0.0097

0.0097

Time taken to build model: 0 seconds

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances

321

86.059 %

Incorrectly Classified Instances

52

13.941 %

Kappa statistic

0.7475

Mean absolute error

0.1667

Root mean squared error

0.2761

Relative absolute error

43.2457 %

Root relative squared error

62.9289 %

Total Number of Instances

373

=== Detailed Accuracy By Class ===

TP Rate

FP Rate

Precision

Recall

F-Measure

MCC

ROC Area

PRC Area

Class

0.963

0.180

0.847

0.963

0.901

0.793

0.947

0.919

Nondemented

0.884

0.057

0.908

0.884

0.896

0.831

0.972

0.955

Demented

0.243

0.018

0.600

0.243

0.346

0.343

0.783

0.432

Converted

Weighted Avg.

0.861

0.116

0.847

0.861

0.844

0.763

0.940

0.885

=== Confusion Matrix ===

a

b

c

<-- classified as

183

7

0

a = Nondemented

11

129

6

b = Demented

22

6

9

c = Converted

Resultados del modelo desarrollado

Cross-validated Accuracy: 0.8713				
Permutation test p-value: 0.0099				
Classification Report:				
	precision	recall	f1-score	support
0	0.00	0.00	0.00	7
1	0.93	0.97	0.95	29
2	0.86	0.97	0.92	39
accuracy			0.88	75
macro avg	0.60	0.65	0.62	75
weighted avg	0.81	0.88	0.84	75
Confusion Matrix:				
[[0 1 6]				
[1 28 0]				
[0 1 38]]				
Test Accuracy: 0.8800				

RESULTADOS

Random Forrest Classificator obtuvo los siguientes resultados en la herramienta Weka, mismos que fueron superados en el modelo desarrollado, en ambos casos se usó 100 arboles de decisión.

Resultados obtenidos en Weka

Test options

☐ Use training set

☐ Supplied test set

☒ Cross-validation

☐ Percentage split

Set...

Folds 5

% 66

More options...

(Nom) Group

StartStop

Result list (right-click for options)

09:25:06 - rules.ZeroR

09:25:22 - functions.LibSVM

09:32:19 - bayes.NaiveBayes

09:32:33 - trees.RandomForest

Classifier output

=== Classifier model (full training set) ===

RandomForest

Bagging with 100 iterations and base learner

weka.classifiers.trees.RandomTree -K 0 -M 1.0 -V 0.001 -S 1 -do-not-check-capabilities

Time taken to build model: 0.24 seconds

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances33188.7399 %

Incorrectly Classified Instances4211.2601 %

Kappa statistic0.7967

Mean absolute error0.1545

Root mean squared error0.2522

Relative absolute error40.0814 %

Root relative squared error57.4694 %

Total Number of Instances373

=== Detailed Accuracy By Class ===

	TP Rate	FP Rate	Precision	Recall	F-Measure	MCC	ROC Area	PRC Area	Class
	0.974	0.109	0.902	0.974	0.937	0.869	0.963	0.931	Nondemented
	0.945	0.075	0.890	0.945	0.917	0.862	0.981	0.967	Demented
	0.216	0.015	0.615	0.216	0.320	0.328	0.769	0.414	Converted
Weighted Avg.	0.887	0.086	0.869	0.887	0.868	0.812	0.951	0.894	

=== Confusion Matrix ===

a	b	c	<-- classified as
185	5	0	a = Nondemented
3	138	5	b = Demented
17	12	8	c = Converted

Resultados del modelo desarrollado

Cross-validated Accuracy: 0.9035

Permutation test p-value: 0.0099

Classification Report:

	precision	recall	f1-score	support
0	1.00	0.14	0.25	7
1	0.97	1.00	0.98	29
2	0.86	0.97	0.92	39
accuracy			0.91	75
macro avg	0.94	0.71	0.72	75
weighted avg	0.92	0.91	0.88	75

Confusion Matrix:

[[1 0 6]

[0 29 0]

[0 1 38]]

Test Accuracy: 0.9067

RESULTADOS

El SVM tiene un antecedente para el conjunto de datos de donde se recuperó el empleo del kernel lineal, con la herramienta Weka se obtuvo otro resultado, estos antecedentes permitieron desarrollar un modelo superior en cuanto al valor de Accuracy

Resultados obtenidos en Weka

Classifier	Choose	RandomForest	-P 100 -I 100 -num-splits 1 -K 0 -M 1.0 -V 0.001 -S 1
Test options	☐ Use training set		
	☐ Supplied test set	Set...	
	☒ Cross-validation	Folds	5
	☐ Percentage split	%	66
	More options...		
(Num) Group			
	Start	Stop	
Result list (right-click for options)			
09:25:06 - rules.ZeroR			
09:25:22 - functions.LibSVM			
09:32:19 - bayes.NaiveBayes			
09:32:33 - trees.RandomForest			

Classifier output

-0.581CDR-0.571nMBV-0.439Age-0.3018MBE+0.1578E8...

0.574MBSE-0.432nMBV-0.399Age-0.369EDUC+0.318CDR...

Group

Test mode: 5-fold cross-validation

=== Classifier model (full training set) ===

LibSVM wrapper, original code by Yasser El-Manzalawy (= WEBSVM)

Time taken to build model: 0.04 seconds

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances 332 89.008 %

Incorrectly Classified Instances 41 10.992 %

Kappa statistic 0.8017

Mean absolute error 0.0733

Root mean squared error 0.2707

Relative absolute error 15.0047 %

Root relative squared error 61.6973 %

Total Number of Instances 373

=== Detailed Accuracy By Class ===

	TP Rate	FP Rate	Precision	Recall	F-Measure	MCC	ROC Area	PRC Area	Class
	0.989	0.098	0.913	0.989	0.949	0.896	0.946	0.908	NonDemented
	0.945	0.066	0.902	0.945	0.923	0.872	0.940	0.874	Demented
	0.162	0.024	0.429	0.162	0.235	0.218	0.569	0.153	Converted
Weighted Avg.	0.890	0.078	0.860	0.890	0.868	0.819	0.906	0.820	

=== Confusion Matrix ===

a	b	c	<-- classified as
188	2	0	a = NonDemented
0	138	8	b = Demented
18	13	6	c = Converted

Resultados del modelo desarrollado

Cross-validated Accuracy: 0.8927					
Classification Report:					
	precision	recall	f1-score	support	
0	0.00	0.00	0.00		7
1	0.94	1.00	0.97		29
2	0.86	0.97	0.92		39
accuracy			0.89		75
macro avg	0.60	0.66	0.63		75
weighted avg	0.81	0.89	0.85		75
Confusion Matrix:					
[[0 1 6]					
[0 29 0]					
[0 1 38]]					
Test Accuracy: 0.8933					
T-statistic: -0.0558					
P-value: 0.9582					

Resultados del antecedente

Table 3. <u>Confusion Matrix</u> of given subjects TND*: True Non-Demented; TD*: True Demented; TC*: True Converted; PND*: Predict Non-Demented; PD*: Predict Demented, and PC*: Predict Converted.				
	TND	TD	TC	precision
PND	43	14	10	64.18%
PD	8	27	1	75.00%
PC	2	0	0	0.00%
Recall	81.13%	65.85%	0.00%	0.00%

CONCLUSIONES

- El modelo Naive Bayes es particularmente útil cuando se tienen muchas características independientes en el caso de estudio donde a pesar de la fuerte suposición de independencia entre características, sin embargo, resultó ser el modelo con menor precisión de los modelos desarrollados teniendo un valor de 0.88 es decir que sus predicciones son correctas en un 88% de los casos.
- La demencia producida por Alzheimer afecta principalmente a personas mayores, y debido a la falta de una cura definitiva, profesionales de la salud y científicos de la computación realizan investigaciones para identificar características relevantes que permitan predecir este deterioro. La floresta randómica de Árboles de Decisión logró los mejores resultados con valores de rendimiento eficientes dentro de los modelos desarrollados.

CONCLUSIONES

- La poca influencia y aparición de la clase “Converted” la volvió un problema para la clasificación de estas características afectado el rendimiento de los modelos, pese a que se ha realizado una eliminación de variables no influyentes, la correlación entre las variables restantes podría no haber sido completamente optimizada.
- El modelo SVM con Kernel lineal, adecuado para clases linealmente separables, mostró el mejor ajuste en los datos, con una diferencia mínima de 0.0006 entre Test Accuracy y Cross-validated Accuracy. Aunque es efectivo, se suele utilizar para problemas con grandes cantidades de datos debido a su conexión con técnicas de Deep Learning.

GRACIAS POR SU ATENCIÓN

:D