



Validación de un algoritmo de clasificación para la identificación de interacciones farmacológicas

Validation of a classification algorithm for identifying pharmacological interactions

Colmenares-Guillén Luis Enrique

Benemérita Universidad Autónoma de Puebla

Facultad de Ciencias de la Computación

Correo: lecolme@gmail.com

<https://orcid.org/0000-0002-9921-8813>

Carrillo-Ruiz Maya

Benemérita Universidad Autónoma de Puebla

Facultad de Ciencias de la Computación

Correo: crrllrzmy@gmail.com

<https://orcid.org/0000-0001-6152-456X>

Morales-Murillo Victor Giovanni

Benemérita Universidad Autónoma de Puebla

Facultad de Ciencias de la Computación

Correo: vg055@hotmail.com

<https://orcid.org/0000-0002-6786-9232>

López y López José Gustavo

Benemérita Universidad Autónoma de Puebla

Facultad de Ciencias Químicas

Correo: jose.lopez@correo.buap.mx

<https://orcid.org/0000-0003-4119-3833>

Resumen

La interacción farmacológica es la modificación del efecto de un fármaco por la acción de otro fármaco. En México, se ha incrementado año con año el número de reacciones no deseadas de medicamentos provocadas por interacciones farmacológicas. En este artículo se realizó un análisis de los diccionarios médicos iDoctus México, PLM México y Vademécum. Con este análisis, se desarrolló un corpus con 540 interacciones farmacológicas que pueden presentarse en los medicamentos más frecuentes del Hospital Universitario de Puebla. Se generó un modelo de clasificación en la plataforma Weka utilizando un algoritmo Naïve Bayes que predice la posibilidad de una interacción farmacológica, clasificada en leve, moderada o grave. Se realizaron las pruebas con el algoritmo Naïve Bayes utilizando el método de validación cruzada con 10 pliegues. Posteriormente, las pruebas de validación, se compararon con el resultado obtenido con el algoritmo Random Forest, utilizando nuevamente el método de validación cruzada con 10 pliegues. Los resultados en el algoritmo Naïve Bayes son en precisión 79.1%, en recuerdo 75.7% y en F-measure 74.8%. En el algoritmo Random Forest son en precisión 51.7%, en recuerdo 51.7% y en F-measure 50.7%. Finalmente, los resultados obtenidos beneficiarían a la farmacovigilancia para aproximarse a predecir nuevas interacciones farmacológicas antes de la comercialización de un medicamento.

Descriptores: Interacción farmacológica, farmacovigilancia, aprendizaje automático, clasificación, validación.

Abstract

The pharmacological interaction is the modification of the effect of one drug by the action of another. In Mexico, the number of undesired drug reactions caused by pharmacological interactions has increasing year in year out. In this article an analysis of medical dictionaries such as iDoctus México, PLM México and Vademécum was carried out. With this analysis, a corpus with 540 pharmacological interactions that can occur in the most frequent medications at the Hospital Universitario de Puebla was developed. A classification model was generated on the Weka platform using a Naïve Bayes algorithm, which predicts the possibility of a pharmacological interaction classified as mild, moderate or severe. The tests with the Naïve Bayes algorithm were performed using the cross-validation method with 10 folds. Subsequently, the validation tests were compared with the results obtained by the Random Forest algorithm, again using the cross validation method with 10 folds. The results of the Naïve Bayes algorithm are 79.1%, in precision; 75.7% in recall, and 74.8%, in F-measure. With regard to the Random Forest algorithm, the data is as follows: 51.7%, in precision; 51.7% in recall, and 50.7%, in F-measure. Finally, the results obtained would benefit pharmacovigilance to approximate predicting new pharmacological interactions, prior to marketing a drug.

Keywords: Pharmacological Interaction, Pharmacovigilance, Learning Machine, Classifier, Validation.

INTRODUCCIÓN

Los medicamentos representan un riesgo para la salud de la población, por tal motivo, es imprescindible evaluar el riesgo/beneficio para cada uno de ellos. El Centro Nacional de Farmacovigilancia (CNFV) de México demuestra que desde 1995 al 2014, se incrementó el número de notificaciones de reacciones no deseadas de medicamentos provocadas por interacciones farmacológicas y que la industria farmacéutica presentó el mayor número de esas notificaciones (COFEPRIS, 2014). La industria farmacéutica en México ha participado con 1.30% en el PIB, esta industria se encuentra dentro de los 15 principales mercados a nivel mundial y esto, significó ser el segundo mercado más importante para América latina (Caso, 2011). Por estas razones, el sector farmacéutico, en conjunto con el sector salud, son fundamentales para México (Cómite de Competitividad Centro de Estudios Sociales y de Opinión Pública, 2010).

La Interacción Farmacológica (IF) es la modificación del efecto de un fármaco por la acción de otro fármaco, además, puede causar efectos leves, moderados o graves; o hasta provocar la muerte de un ser humano. Por tal motivo, existe una ciencia que se llama farmacovigilancia que investiga, vigila y evalúa la información sobre los efectos de los medicamentos para identificar información de las reacciones no deseadas de medicamentos y prevenir los daños en los pacientes (Secretaría de Salud, 2014).

Este trabajo tiene como objetivo establecer un proceso tecnológico de soporte, a la farmacovigilancia con base en los Sistemas de Reconocimiento de Patrones (SRP) para identificar la gravedad de la IF, antes y después de la comercialización de un medicamento y así, aproximarse a predecir nuevas interacciones farmacológicas.

Los SRP son procesos complejos que se han implementado en aplicaciones de reconocimiento de voz, identificación de huellas dactilares, reconocimiento óptico de caracteres, la identificación de secuencias de ADN, reconocimiento de rostros, clasificación de correos spam, entre otras (Duda *et al.*, 2001; Hernández *et al.*, 2004). Entonces, es muy importante construir SRP para problemas específicos que requieran el uso de estas técnicas. El problema que se plantea en este trabajo es la clasificación de la gravedad de una IF que presenta tres clases: leve, moderada o grave (Idoctus México, 2016).

Se propone generar un Modelo de Clasificación (MC) que utilice un algoritmo Naïve Bayes (NB) que se fundamenta en características o atributos condicional-

mente independientes, porque la IF presenta ese tipo de características. Es decir, en un paciente la gravedad de una IF se manifiesta por el consumo de dos o más fármacos; la dosis de cada fármaco; el consumo de alimentos; y aspectos físicos del paciente como su altura, peso, edad, metabolismo e historial clínico (Secretaría de Salud, 2014). Este trabajo de investigación se enfoca en la IF del tipo fármaco-fármaco.

El MC se generó en Weka (Waikato Environment for Knowledge Analysis) porque es una plataforma de software para el aprendizaje automático y la minería de datos que permite utilizar diferentes técnicas de clasificación para poder evaluar el rendimiento de diferentes algoritmos (Han *et al.*, 2005).

El MC se realizó empleando validación cruzada a 10 pliegues, comúnmente utilizada en procesos de clasificación (Han *et al.*, 2012). Este método de validación cruzada divide el corpus en 10 grupos en donde se usa un grupo para el entrenamiento y 9 para pruebas, se calcula el error y se repite el proceso 10 veces cambiando el conjunto de entrenamiento (Han *et al.*, 2012).

En la propuesta del modelo de clasificación del algoritmo Naïve Bayes se comparó este con el algoritmo Random Forest con la finalidad de discutir los resultados de cada uno de ellos.

MATERIALES Y MÉTODOS

En este trabajo se desarrolló un modelo de clasificación. Para ello se construyó un corpus utilizando diccionarios médicos como iDoctus México (Idoctus México, 2016), PLM México (PLM México, 2017) y Vademécum (VADEMECUM, 2011) en formato electrónico, que contienen información de las interacciones farmacológicas de México. Estos recursos digitales facilitan la creación de corpus de interacciones farmacológicas. El modelo que se construyó utiliza el aprendizaje supervisado (Manning y Schütze, 1999).

Siguiendo las etapas: colección de datos, selección de características, selección del modelo matemático, entrenamiento, pruebas y complejidad computacional (Duda *et al.*, 2001).

El esquema general del modelo de clasificación se muestra en la Figura 1, en el que se describen las 5 etapas del ciclo de diseño que se utilizaron para construir el MC. En la primera etapa de colección de datos se recolecta información de muestras de interacciones farmacológicas leves, moderadas o graves para desarrollar un corpus. En la segunda etapa de selección de características se seleccionan los nombres de dos fármacos que forman una IF, en la tercera etapa de selección del modelo matemático se analiza precisamente el modelo ma-

temático Naïve Bayes, en la cuarta etapa de entrenamiento se genera el Modelo de Entrenamiento (ME) con base en el modelo matemático NB y el corpus y en la quinta etapa de evaluación se realizan las pruebas del clasificador.

COLECCIÓN DE DATOS

Se analizaron diferentes diccionarios médicos como PLM, VADEMECUM, MICROMEDEX e IDOCTUS México, que fueron recomendados por un especialista farmacéutico, el Dr. José Gustavo López de la Facultad de Ciencias Químicas. Se seleccionó IDOCTUS México por su usabilidad; su organización de la información y porque está basado en la base de datos farmacológica del Consejo General de Colegios Oficiales Farmacéuticos (Idoctus Mexico, 2016). Posteriormente, se recolectaron del IDOCTUS, las interacciones farmacológicas presentadas en los medicamentos más frecuentes del Hospital Universitario de Puebla. De los 360 fármacos se obtuvo un corpus con 540 IFs, donde 180 IFs pertenecen a cada clase: leve, moderada o grave.

Es importante el procesamiento del lenguaje natural porque influye en los resultados del MC. Por este motivo, se realizó un pre-procesamiento al corpus de IFs porque su datos pueden ser incompletos o inconsistentes, lo que ocasiona la extracción de patrones poco útiles (Zhang *et al.*, 2003). Se aplicó un pre-procesamiento al corpus a través de un algoritmo desarrollado con AWK (Aho *et al.*, 1979), que realiza la tarea de corrección y eliminación de los datos incompletos e inconsistentes del corpus, convirtiendo letras mayúsculas en minúsculas, borrando los signos de puntuación y acentos (Aguirre y Edmonds, 2007). En la Tabla 1 se observa el diseño del corpus de IFs que presenta cinco columnas. En la primera columna se indica el número de IF que es una asignación de un identificador a la interacción farmacológica; en la segunda columna se indica la etiqueta que es la clase a la que pertenece la IF (leve,

moderada o grave); en la columna tres y cuatro se encuentran los fármacos que forman a la IF; y en la última columna se ubica el efecto de la IF en el humano. Se diseñó de esta forma el corpus de IFs para facilitar el reconocimiento de patrones, evitar información irrelevante y realizar la extracción de características en el MC. En la Figura 2 se observa un resultado pre-procesado del corpus.

SELECCIÓN DE CARACTERÍSTICAS

La selección de características es un campo de investigación importante en los SRP, en el aprendizaje automático y en la minería de datos. La selección de características ha demostrado tanto en la teoría como en la práctica que es eficaz para mejorar la eficiencia del aprendizaje y aumentar la precisión predictiva del MC, su objetivo es elegir un subconjunto de datos de entrada eliminando características que son irrelevantes o que no tienen información predictiva del corpus de IFs (Ramaswami y Bhaskaran, 2009).

En este trabajo se seleccionaron tres características o atributos del corpus de IFs: el fármaco 1, el fármaco 2 y la gravedad (leve, moderada o grave). La razón de su selección fue que estas características constituyen a una IF del tipo fármaco-fármaco (Secretaría de Salud, 2014) y cumplen con los requerimientos de la selección de características, en donde los datos deben ser invariantes y únicos (Duda *et al.*, 2001).

SELECCIÓN DEL MODELO MATEMÁTICO

La selección del modelo matemático es Naïve Bayes porque una IF presenta características independientes, es decir, la IF no se presenta únicamente por la combinación de dos o más fármacos, sino también por la dosis de cada fármaco, por el consumo de alimentos, el historial clínico del paciente, por los aspectos físicos del paciente como su edad, altura, peso, metabolismo y

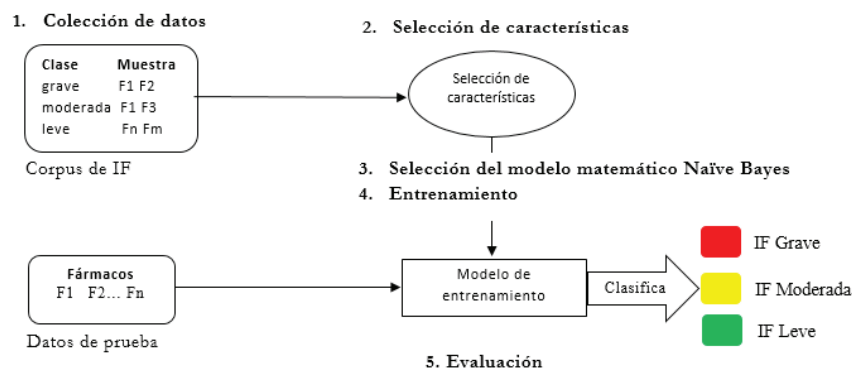


Figura 1. Diseño del modelo de clasificación

Tabla 1. Diseño del corpus de interacciones farmacológicas

If #	Etiqueta	Fármaco 1	Fármaco 2	Efecto
IF 1#	Grave	amiodarona	acenocumarol	posible potenciación del efecto anticoagulante peligro de hemorragias
IF ...#		...	...	...
IF 180#	Grave	levomepromazina	fosfenitoina	
IF ...#		...	...	...
IF 360#	Moderada	levomepromazina	asenapina	...
IF ...#		...	...	...
IF 540#	Leve	ketorolaco	Valsartan	...

IF 1# grave	amiodarona	acenocumarol	posible potenciación del efecto anticoagulante peligro de hemorragias
IF 2# grave	amiodarona	adenosina	la administración conjunta de adenosina con antiarrítmicos de clase Ia disipiramida hidroquinidina procainamida o iii a
IF 3# grave	amiodarona	citalopram	la administración conjunta de citalopram o escitalopram con fármacos susceptibles de prolongar el intervalo qt tales co
IF 4# grave	amiodarona	claritromicina	la administración conjunta de macrólidos en general ej azitromicina claritromicina eritromicina telitromicina etc con
IF 5# grave	amiodarona	cobicistat	posible acumulación del sustrato del cyp3a4 con riesgo de aumento de su toxicidad
IF 6# grave	amiodarona	daclatasvir	posible riesgo de bradicardia sintomática grave y potencialmente mortal cuando ledipasvirsofosbuvir o sofosbuvir combin
IF 7# grave	amiodarona	darunavir	posible de las concentraciones plasmáticas del antiarrítmico pudiendo conducir a efectos tóxicos graves y potencialmen
IF 8# grave	amiodarona	digoxina	posible incremento de los niveles de digoxina con el consiguiente riesgo de intoxicación
IF 9# grave	amiodarona	dihidroartemisinina	la administración conjunta de piperquinadhidroartemisinina por su posible acción sobre el intervalo qt con ot
IF 10# grave	amiodarona	eritromicina	la administración conjunta de macrólidos en general ej azitromicina claritromicina eritromicina telitromicina etc con
IF 11# grave	amiodarona	escitalopram	la administración conjunta de citalopram o escitalopram con fármacos susceptibles de prolongar el intervalo qt tales co
IF 12# grave	amiodarona	esparfloxacino	posible potenciación de la toxicidad
IF 13# grave	amiodarona	fenitoina	posible aumento de los niveles plasmáticos de fenitoina riesgo de intoxicación
IF 14# grave	amiodarona	fenitoina antiarrítmico	posible aumento de los niveles plasmáticos de fenitoina riesgo de intoxicación
IF 15# grave	amiodarona	fidaxomicina	posible aumento de las concentraciones plasmáticas de fidaxomicina con posible aumento de la toxicidad sistémica ej mol
IF 16# grave	amiodarona	ingolimod	riesgo de disminución excesiva de la frecuencia cardíaca
IF 17# grave	amiodarona	fosamprenavir	posible de las concentraciones plasmáticas del antiarrítmico pudiendo conducir a efectos tóxicos graves y potencialmen
IF 18# grave	amiodarona	fosfenitoina	posible aumento de los niveles plasmáticos de fenitoina riesgo de intoxicación
IF 19# grave	amiodarona	hidroxizina	la administración conjunta de hidroxizina con fármacos susceptibles de prolongar el intervalo qt puede aumentar el ries
IF 20# grave	amiodarona	indinavir	posible de las concentraciones plasmáticas del antiarrítmico pudiendo conducir a efectos tóxicos graves y potencialmen

Figura 2. Corpus pre-procesado de interacciones farmacológicas con un total de 540 interacciones, 180 muestras por cada clase: leve, moderada o grave

otros factores (COFEPRIS, 2014). Este trabajo se enfoca en la IF del tipo fármaco-fármaco (Secretaría de Salud, 2014). Además, Zhang menciona que, NB es uno de los más eficientes y efectivos algoritmos de aprendizaje inductivo para el desarrollo de aprendizaje automático y de minería de datos (Zhang, 2005). Adicionalmente, es uno de los clasificadores más utilizados por su simplicidad y rapidez (Witten y Frank, 2005).

NB se fundamenta en el Teorema de Bayes que consiste en la probabilidad condicional. Sin embargo, NB asume ingenuamente que las características no son condicionales y son independientes, esto se demostró en el trabajo de Domingos y Pazzani (1997), obteniendo como resultado la fórmula de NB. A diferencia de otros modelos matemáticos, en este modelo se contabilizan las frecuencias de ocurrencias de las características seleccionadas y no se realizan búsquedas de hipótesis (Morales, 2012).

$$V_{NB} = \arg \max_{V_j \in V} (P(V_j) \prod_i P(a_i | V_j))$$

donde

$P(V_j)$ = Frecuencia de las clases leve, moderada y grave

$P(a_i | V_j)$ = Frecuencia de los datos observados, que son los fármacos

V_{NB} = Argumento Máximo de las probabilidades leve, moderada o grave

En el trabajo de Ruiz, se menciona que para mejorar la precisión numérica de NB, generalmente se implementa la propiedad de logaritmo, esto permite obtener mejores resultados en la fórmula de NB, porque evita que las estimaciones de probabilidades se aproximen a cero. La fórmula de NB con la implementación de la propiedad de logaritmo es $V_{NB} = \arg \max_{V_j \in V} (\log P(V_j) + \sum_i \log P(a_i | V_j))$ (Ruiz, 2012).

ENTRENAMIENTO

En la fase de entrenamiento se estima $P(a_i | V_j)$ de la fórmula de NB en el Corpus de Entrenamiento (CE) de IFs, en el cual se implementa un suavizado Laplaciano *adding one*, que consiste en sumar un uno a todas las

probabilidades y que las probabilidades no se aproximen a cero. El suavizado Laplaciano se realiza para obtener un espacio de probabilidad a eventos que no son visibles, esto significa la reducción del margen de error del algoritmo NB y una probabilidad uniforme (Manning y Schütze, 1999). En la siguiente fórmula se estiman las probabilidades de los fármacos del CE implementando el suavizado Laplaciano.

$$P(\text{Fármaco} | \text{Clase}) = \frac{F_k + 1}{F + |\text{vocabulario}|}$$

donde

F_k = Frecuencia de un fármaco que pertenece a la misma clase

F = Frecuencia del total de fármacos que pertenecen a la misma clase

Vocabulario = Número de fármacos de todas las clases del CE

Los resultados de las estimaciones de las probabilidades de los fármacos por cada una de sus clases son almacenados en una estructura de datos llamada tabla hash o Diccionario de Palabras (DP). Este DP es el modelo de entrenamiento que se genera con base en el CE previamente pre-procesado, su diseño se presenta en la Tabla 2.

PRUEBAS

En esta fase se utilizó la plataforma Weka que consta de 3 partes. En la primera parte se generó un archivo arff que contiene tres atributos: fármaco 1, fármaco 2 y gravedad. En el mismo archivo, se agregaron las 540 instancias de interacciones farmacológicas. En la Tabla 3 se muestran los datos de interacciones farmacológicas en un archivo arff que se utilizó en la plataforma Weka.

En la segunda parte, se introduce el archivo de datos en la plataforma Weka y se selecciona el algoritmo

de clasificación NB. Se ejecuta el mismo utilizando validación cruzada a 10 pliegues.

El método de validación cruzada es una técnica utilizada para modelos de clasificación. La validación cruzada se implementó con un valor de K igual a 10 (Hernández *et al.*, 2004). En la Figura 3, se muestra el proceso de validación cruzada, en donde el corpus de IFs se divide en 10 grupos, en el primer pliegue K=1 se selecciona el grupo 1 para pruebas y se seleccionan del grupo 2 al 10 para entrenamiento, posteriormente en el segundo pliegue K=2 se selecciona el grupo 2 para pruebas y se seleccionan del grupo 3 al 10 y el grupo 1 para entrenamiento. En cada pliegue, se dividió el corpus en 90% para entrenamiento y 10% para pruebas. Este proceso se repite hasta la pliegue K = 10.

Finalmente, en cada pliegue se evalúan los resultados y se promedian de los 10 pliegues para obtener el resultado de las métricas de precisión, recuerdo y F-measure (Hernández *et al.*, 2004).

EVALUACIÓN

Es importante analizar las métricas de clasificación que existen en los sistemas de aprendizaje automático para evaluar y comprender el rendimiento de un modelo de clasificación (Lever *et al.*, 2016). En este trabajo se evalúa el MC por las siguientes métricas: precisión, recuerdo y F-measure (Computer Science Department Cornell University, 2003). En la Figura 4, se muestra una matriz de confusión para n clases (Manligues, 2017), con esta matriz se calcularon las métricas de evaluación.

El número Total de Falsos Negativos (TFN), Total de Falsos Positivos (TFP), Total de Verdaderos Negativos (TVN) y Total de Verdaderos Positivos (TVP) para cada clase i se calculan con las siguientes fórmulas (Manligues, 2017).

Tabla 2. Diseño del modelo de entrenamiento del algoritmo de clasificación NB

Tabla Hash		Valores	
Llave (Key)	Grave	Moderada	Leve
Fármaco 1	P(F1 Grave)	P(F1 Moderada)	P(F1 Leve)
Fármaco 2	P(F2 Grave)	P(F2 Moderada)	P(F2 Leve)
Fármaco 3	P(F3 Grave)	P(F3 Moderada)	P(F3 Leve)
.....			
.....			
Fármaco 360	P(F360 Grave)	P(F360 Moderada)	P(F360 Leve)

Tabla 3. Archivo arff en la plataforma Weka

@relation interacciones.symbolic

@attribute farmaco1 {amiodarona, levomepromazina, ketorolaco, piridoxina, dopamina, amikacina, parecoxib, meclozina, metamizol, ibuprofeno, altizida, telaprevir, ajo, cimetidina, oxazepam, metadona, efavirenz, ergotamina, misoprostol, ... }

@attribute farmaco2 {amiodarona, levomepromazina, ketorolaco, piridoxina, dopamina, amikacina, parecoxib, meclozina, metamizol, ibuprofeno, altizida, telaprevir, ajo, cimetidina, oxazepam, metadona, efavirenz, ergotamina, misoprostol, ... }

@attribute gravedad {grave, moderada, leve}

@data

amiodarona,acenocumarol,grave

amiodarona,adenosina,grave

amiodarona,citalopram,grave

amiodarona,claritromicina,grave

amiodarona,cobicistat,grave

...

...

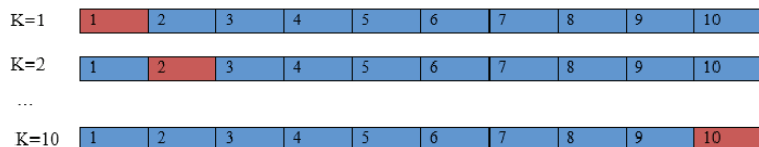


Figura 3. Validación cruzada con K = 10 pliegues

		Predicted Number			
		Class 1	Class 2	...	Class n
Actual Number	Class 1	x_{11}	x_{12}	...	x_{1n}
	Class 2	x_{21}	x_{22}	...	x_{2n}
	.	.	.	.	.
	Class n	x_{n1}	x_{n2}	...	x_{nn}

Figura 4. Matriz de confusión para n clases (Manligues, 2017)

$$TFN = \sum_{j=1}^n x_{ij} \quad TFP = \sum_{j=1}^n x_{ij}$$

$$TVN = \sum_{j=1}^n \sum_{k=1, k \neq i}^n x_{jk} \quad TVP_{all} = \sum_{j=1}^n x_{ij}$$

(TVN) Total de Verdaderos Negativos = número total de interacciones farmacológicas correctamente rechazadas a su clase.

donde

- (TVP) Total de Verdaderos Positivos = número total de interacciones farmacológicas asignadas correctamente a su clase (leve, moderada o grave).
- (TFP) Total de Falsos Positivos = número total de interacciones farmacológicas asignadas incorrectamente a su clase.
- (TFN) Total de Falsos Negativos = número total de interacciones farmacológicas incorrectamente rechazadas a su clase.

Finalmente, se calcula la precisión, el recuerdo y el F-measure para cada clase con las siguientes fórmulas (Manligues, 2017).

$$Precisión_i = \frac{TVP_{all}}{TVP_{all} + TFP_i}$$

$$Recuerdo_i = \frac{TVP_{all}}{TVP_{all} + TFN_i}$$

$$F-measure_i = \frac{2 * (Precisión * Recuerdo)}{Precisión + Recuerdo}$$

ANÁLISIS DE RESULTADOS

El primer resultado es un corpus con 540 instancias de interacciones farmacológicas etiquetadas por su tipo de gravedad como leve, moderada o grave.

Los resultados que se presentaron en la evaluación y los experimentos del modelo de clasificación con el algoritmo Naïve Bayes son los siguientes. En el experimento con el método de validación cruzada con 10 iteraciones se obtuvo una precisión de 79.1%, un recuerdo de 75.7% y un F-measure de 74.8%. Los resultados por cada clase leve, moderada o grave se presentan en la Tabla 4.

El algoritmo NB se comparó con el algoritmo Random Forest. Los resultados del MC con el algoritmo Random Forest son en precisión 51.7%, en recuerdo un 51.7% y en F-measure 50.7%. Los resultados por cada clase se presentan en la Tabla 5.

Se observa que al utilizar el método de validación cruzada con 10 pliegues, el algoritmo NB presentó mejores resultados que el algoritmo de Random Forest.

CONCLUSIONES Y TRABAJO FUTURO

En este artículo se desarrolló un corpus de interacciones farmacológicas clasificadas por su tipo de gravedad (leve, moderada o grave) usando el diccionario médico de IDOCTUS, en donde se obtuvieron las 540 interacciones fármacológicas de los fármacos más utilizados

en el Hospital Universitario de Puebla. Se generó un modelo de clasificación utilizando el algoritmo Naïve Bayes porque la IF presenta atributos o características independientes. Además, se realizaron las pruebas para evaluar el MC utilizando un método de experimentos llamado validación cruzada con un K igual a 10 pliegues. Se analizaron los resultados y se comparó con el mismo método de validación cruzada el algoritmo Naïve Bayes con el algoritmo Random Forest. En la comparación, el algoritmo NB presentó mejores resultados.

En el trabajo de Kukreja, se presenta una comparación de un algoritmo NB con el algoritmo Random Forest utilizando el método de validación cruzada, en donde, Random Forest presentó un mejor rendimiento en conjuntos de datos con múltiples clases y el algoritmo NB presentó mejores resultados en un conjunto de datos con 4 clases (Kukreja *et al.*, 2012).

En este trabajo, el algoritmo NB presentó mejores resultados que Random Forest porque se identificaron solo tres clases en la IF: leve, moderada o grave. Además, la IF presentó características independientes por lo cual esta investigación se enfocó en el modelo matemático Naïve Bayes.

El trabajo futuro es incluir nuevos atributos al corpus de IF para actualizar el MC, los atributos que se pueden agregar son: la dosis de los fármacos, los aspectos físicos de un paciente: altura, peso, edad y metabolismo y el historial clínico de un paciente.

Tabla 4. Resultados de NB con validación cruzada a 10 pliegues

Clase	Precisión	Recuerdo	F-measure
Grave	92.9%	51.1%	65.9%
Moderada	68.4%	90%	77.7%
Leve	76%	86.1%	80.7%
Total	79.1%	75.7%	74.8%

Tabla 5. Resultados del algoritmo Random Forest con validación cruzada con 10 pliegues

Clase	Precisión	Recuerdo	F-measure
Grave	45.5%	28.3%	34.9%
Moderada	42.7%	57.2%	48.9%
Leve	66.8%	69.4%	68.1%
Total	51.7%	51.7%	50.7%

Finalmente, este MC brinda una herramienta para automatizar la identificación de la gravedad de las interacciones farmacológicas en las áreas de ciencias químicas, farmacología y farmacovigilancia. Adicionalmente, permite aproximarse a descubrir nuevas interacciones farmacológicas para evitar dañar la salud de los pacientes con tratamientos farmacológicos.

AGRADECIMIENTOS

Los autores agradecen el financiamiento parcial por la Vicerrectoría de Investigación y Estudios de Posgrado (VIEP) de la Benemérita Universidad Autónoma de Puebla (BUAP). Además, se agradece el apoyo del laboratorio de sistemas en tiempo real de la Facultad de Ciencias de la Computación en conjunto del laboratorio de farmacia clínica de la Facultad de Ciencias Químicas de la BUAP.

REFERENCIAS

- Aguirre, E. y Edmonds, P. (2007). *Word Sense Disambiguation*. 1a ed. Netherlands: Springer. 1-28.
- Aho, A., Kernighan, B., Weinberger, P. (1979). Awk-APattern Scanning and Processing Language, Software. *Practice and Experience*, 9(4): 267-279.
- Caso, P. (2011). *Esquema regulatorio de medicamentos en México: Oportunidades y retos*. México: COFEPRIS.
- COFEPRIS. 7° Boletín Informativo (2014). *Farmacovigilancia y Tecnovigilancia*. México: COFEPRIS.
- Comité de Competitividad Centro de Estudios Sociales y de Opinión Pública. *Situación del sector farmacéutico en México* (2010). 1a ed., México: Centro de Estudios Sociales y de Opinión Pública. 83-85.
- Computer Science Department of Cornell University. (2003). *Performance Measures for Machine Learning*. Estados Unidos de América: Cornell University.
- Domingos, P. y Pazzani, M. (1997). On the optimality of the simple bayesian classifier under Zero-One Loss. *Machine Learning*, 29(2): 103-130.
- Duda, R., Hart, P., Stork, D. (2001). *Pattern classification*. 2a ed. California: Ricoh California Research Center.
- Eikan, C. (2011). *Evaluating classifiers*. San Diego: Universidad de California.
- Han, J., Kamber, M., Pei, J. (2005). *DATA MINING: Concepts and Techniques*. 2a ed. Estados Unidos: Elsevier.
- Han J., Kamber, M., Pei, J. (2012). *DATA MINING: Concepts and Techniques*, 3a ed. Estados Unidos: Elsevier.
- Hernández, O., Ramírez, Q., Ferri, R. (2004). *Introducción a la minería de datos*. 1a ed. Madrid: Pearson educación.
- Idoctus México. Interacciones. Recuperado el 25 de julio de 2016 de <http://mx.idoctus.com/interacciones>
- Kukreja, M., Johnston, S., Stafford, P. (2012). Comparative study of classification algorithms for immunosignaturing data. *BMC Bioinformatic*, 13(139): 1-15.
- Lever, J., Krzywinski, M., Altman, N. (2016). Points of significance: Classification evaluation. *Nature Methods*, 13: 541-542.
- Manning, C. y Schütze, H. (1999). *Foundations of statistical natural language processing*. 2a ed. Estados Unidos de América: Massachusetts Institute of Technology.
- Manligues, G. Generalized confusion matrix for multiple clases, (2017). Recuperado el 21 de mayo 2017 de https://www.researchgate.net/publication/310799885_Generalized_Confusion_Matrix_for_Multiple_Classes
- Morales, E. *Aprendizaje Bayesiano*. (2012). Recuperado el 11 de septiembre 2016 de <https://ccc.inaoep.mx/~emorales/Cursos/NvoAprend/Acetatos/bayes.pdf>
- PLM México. Medicamentos PLM. Recuperado el 10 de enero de 2017 de <http://www.medicamentosplm.com>
- Ramaswami, M. y Bhaskaran, R.A. (2009). Study on feature selection techniques in educational. *Data Mining Journal of Computing*, 1(1): 7-11.
- Ruiz, J. (2016). Aprendizaje de modelos probabilísticos 2012. Recuperado el 13 de Septiembre de 2016 de <https://www.cs.us.es/cursos/iati-2012/temas/tema-10.pdf>
- Secretaría de Salud. *FARMACOPEA de los Estados Unidos Mexicanos*. 5a ed. México: Comisión Permanente de la Farmacopea de los Estados Unidos Mexicanos.
- VADMECUM. *Vademecum Internacional 11*. (2011). España: MEDICOM.
- Witten I. y Frank, Eibe (2005). *Data Mining: Practical Machine Learning Tools and Techniques*. 2a ed. California.
- Zhang, H. (2005). Exploring conditions for the optimality of naive Bayes. *International Journal of Pattern Recognition and Artificial Intelligence*, 19(2): 183-198.
- Zhang, S., Zhang, C., Yang, Q. (2003). Data preparation for data. *Applied Artificial Intelligence*, 17(5-6): 375-381.