

Millán Solarte, Julio César; Cerezo, Edinson Caicedo

Article

Modelos para otorgamiento y seguimiento en la gestión del riesgo de crédito

Revista de Métodos Cuantitativos para la Economía y la Empresa

Provided in Cooperation with:

Universidad Pablo de Olavide, Sevilla

Suggested Citation: Millán Solarte, Julio César; Cerezo, Edinson Caicedo (2018) : Modelos para otorgamiento y seguimiento en la gestión del riesgo de crédito, Revista de Métodos Cuantitativos para la Economía y la Empresa, ISSN 1886-516X, Universidad Pablo de Olavide, Sevilla, Vol. 25, pp. 23-41

This Version is available at:

<https://hdl.handle.net/10419/195396>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



<https://creativecommons.org/licenses/by-sa/4.0/>



UNIVERSIDAD
PABLO DE OLAVIDE
SEVILLA



REVISTA DE MÉTODOS CUANTITATIVOS PARA LA
ECONOMÍA Y LA EMPRESA (25). Páginas 23–41.
Junio de 2018. ISSN: 1886-516X. D.L: SE-2927-06.
www.upo.es/revistas/index.php/RevMetCuant/article/view/2370

Modelos para otorgamiento y seguimiento en la gestión de riesgo de crédito

MILLÁN SOLARTE, JULIO CÉSAR

Departamento de Contabilidad y Finanzas, Facultad de Administración
Universidad del Valle, Cali (Colombia)

Correo electrónico: julio.millan@correounivalle.edu.co

CAICEDO CEREZO, EDINSON

Departamento de Contabilidad y Finanzas, Facultad de Administración
Universidad del Valle, Cali (Colombia)

Correo electrónico: edinson.caicedo@correounivalle.edu.co

RESUMEN

Esta investigación muestra la aplicación y desempeño de tres modelos para la clasificación de solicitantes de créditos: el modelo de análisis discriminante, el de regresión logística y el de redes neuronales; técnicas empleadas por las instituciones financieras en el cálculo del *scoring* de crédito. Los resultados obtenidos muestran un mejor desempeño del modelo de redes neuronales en comparación con el de regresión logística y análisis discriminante, logrando una tasa de aciertos en la clasificación del 86.9%. Para los tres modelos se emplearon catorce variables que informan sobre las características socioeconómicas del prestatario y sobre las características propias de la operación crediticia. En el ámbito de la gestión financiera, este resultado es importante dado que puede complementarse con el cálculo de la probabilidad de incumplimiento, con los montos expuestos en cada operación de crédito y con la tasa de recuperación de la entidad para establecer el valor de las pérdidas esperadas a nivel individual y a nivel del portafolio de créditos de la entidad.

Palabras claves: *scoring* de crédito; riesgo de crédito; probabilidad de incumplimiento; análisis discriminante; regresión logística; redes neuronales.

Clasificación JEL: C14; C45; C51; D14.

MSC2010: 91G40; 62M45; 62G08; 91G70; 91B82.

Models for Granting and Tracking in Credit Risk Management

ABSTRACT

This research shows the application and performance of three models for the classification of credit applicants: discriminant analysis, logistic regression and neural networks; techniques used by financial institutions for the calculation of credit scoring. The results show a better performance of the neural network model compared to logistic regression and discriminant analysis, achieving a success rate of 86.9% in the classification. For the three models, fourteen variables were used to inform about applicant's socioeconomic characteristics and those of the credit operation. In the area of credit risk management, this result is relevant since it can be complemented by the calculation of default probability, the exposure at default and the recovery rate of the entity to establish the value of expected losses at both the individual level and the whole credit portfolio of the entity.

Keywords: Credit scoring; credit risk; default probability; discriminant analysis; logistic regression; neural networks.

JEL classification: C14; C45; C51; D14.

MSC2010: 91G40; 62M45; 62G08; 91G70; 91B82.



1. Introducción

El tema de riesgo en el ámbito financiero siempre será cuestión de discusión por varios factores propios a las implicaciones que tiene sobre el desempeño financiero de las actividades de inversión o de financiación, que realiza una entidad o un individuo.

Para posibilitar su estudio, generalmente el riesgo financiero se ha clasificado en riesgo de mercado, riesgo de crédito y riesgo operativo (McNeil et al., 2005).

El riesgo de crédito, según McNeil et al. (2005), se define como las pérdidas originadas por el incumplimiento en las obligaciones contraídas (por ejemplo créditos, bonos), incumplimiento generado por varios factores, entre los que se pueden mencionar los movimientos bruscos en el mercado de activos financieros, situaciones de iliquidez, imposibilidad de ejecutar garantías o cobros. De Lara Haro (2005) define el riesgo de crédito como la pérdida potencial que se registra con motivo del incumplimiento de una contraparte en una transacción financiera (o en alguno de los términos y condiciones de la transacción). También se concibe como un deterioro en la calidad crediticia de la contraparte o en la garantía o colateral pactado originalmente.

En este sentido, las decisiones relativas al otorgamiento y seguimiento de créditos son vitales para cualquier tipo de institución financiera, ya que se pueden originar grandes pérdidas financieras generadas por el retraso, o no pago de las obligaciones (Yu et al, 2008). Un buen número de estas instituciones, emplean el juicio y la experiencia de los analistas de crédito en la selección de sus clientes, es decir se realiza el estudio de cada solicitud por separado, otras instituciones emplean un sistema de calificación crediticia (scoring de crédito), o realizan ambas tareas para aprobar o rechazar una solicitud y en relación con el producto solicitado.

Generalmente, en un sistema de scoring de crédito se revisan y evalúan los datos del solicitante que informan sobre su situación financiera, historial de pagos y antigüedad en el empleo, entre otras variables, información que permitirá distinguir entre un "buen" y un "mal" solicitante (Hand & Henley, 1997), el proceso se realiza tomando una muestra del historial de clientes (Thomas et al., 2004), por lo general, los modelos de scoring de crédito se ocupan de dos clases de crédito, préstamos de consumo y préstamos comerciales (Thomas et al., 2002). En el campo de la gestión del riesgo de crédito, se han aplicado muchos modelos y algoritmos que buscan brindar apoyo en la calificación crediticia, incluyendo técnicas estadística, algoritmo genético y redes neuronales (Yu et al., 2008).

Para estimar el nivel de riesgo de los prestatarios, se asigna una probabilidad de default (PD) diferente para cada uno de ellos, este un indicador ampliamente empleado en las instituciones financieras. La PD indica que una contraparte determinada no podrá cumplir con sus obligaciones. La errónea estimación de la PD conduce a calificaciones no adecuadas, precios incorrectos de los instrumentos financieros, subestimación del colateral, entre otros eventos. La probabilidad de incumplimiento es también un parámetro utilizado en el cálculo del capital económico o del capital regulatorio bajo el enfoque de Basilea II en una institución financiera.

Esta investigación tiene como propósito principal emplear las variables que se utilizan con mayor frecuencia en los sistemas de scoring de crédito (calificación y clasificación), analizando el desempeño de cuatro técnicas: El análisis discriminante, la regresión logística, las redes neuronales y un modelo híbrido, luego con la información obtenida se podría estimar las probabilidad de incumplir en los pagos, como una función de la calificación calculada, insumo útil en la estimación de la pérdida esperada por exposición al riesgo de crédito y en la determinación del capital necesario para la cobertura de dichas pérdidas.

2. Revisión de la literatura

2.1 Modelos de Scoring

El credit scoring es un método estadístico para estimar la probabilidad de incumplimiento (default) de un prestatario, usando su información histórica y estadística para obtener un indicador que permita distinguir los buenos deudores de los malos deudores. Los modelos de scoring de crédito son empleados para evaluar el riesgo de crédito a nivel individual de un deudor o persona que este solicitando un crédito.

Hand & Henley (1997), los definen como los métodos estadísticos utilizados para clasificar a solicitantes de un crédito o a quienes son ya clientes de la entidad evaluadora en las categorías “bueno” y “malo”.

Kiefer & Larson (2006) indican que los modelos de scoring son empleados para clasificar a los solicitantes del crédito con base en su desempeño esperado, igualmente analizan las ventajas de utilizar modelos paramétricos y no paramétricos empleando un modelo estilizado simple.

Gutiérrez (2007) indica que los modelos de scoring son muy importantes en el proceso de gestión del crédito, su trabajo busca explicar en detalle la construcción, composición y operatividad de estos modelos, emplean para ilustración una base de datos del sistema financiero argentino.

La información que arroja el scoring, aunada a la disponible en la solicitud del crédito permite el análisis del solicitante para tomar la decisión de otorgar o no el crédito. El objetivo de un modelo de scoring de crédito es en esencia realizar la distinción entre un buen solicitante de un crédito (bajo riesgo de default) y un mal solicitante (alto riesgo de default).

2.2 Análisis discriminante

Este método se utiliza para realizar la clasificación de diferentes individuos en grupos, a partir de los valores de las variables observadas sobre los individuos objeto de categorización. Cada individuo pertenecerá a un solo grupo. En la gestión del riesgo de crédito específicamente se busca detectar si el solicitante de un préstamo pertenecerá en el futuro al grupo de personas que devuelven el crédito o por el contrario pertenecerá al grupo de personas que no lo hacen.

En su desarrollo el análisis discriminante se emplea sobre una población previamente dividida en grupos (para riesgo de crédito generalmente se emplean dos grupos, grupo cumplidos en el pago del crédito y grupo incumplidos en el pago del crédito), el análisis discriminante encuentra una función que permite, con un grado determinado de certeza explicar esta división de los grupos, lo que se conoce generalmente como visión explicativa, obtenida esta función se puede emplear para clasificar a nuevos individuos en alguno de los grupos en los cuales se encuentra dividida la población, lo que se ha denominado como visión predictiva.

Se busca estimar la relación entre una única variable dependiente no métrica (categórica) y un conjunto de variables independientes métricas así:

$$Y_i = X_1 + X_2 + X_3 + \dots + X_n, \quad (1)$$

donde Y_i es una variable no métrica (categórica) y X_i (con $i = 1$ hasta n) es una variable métrica. De acuerdo con Lee et al. (2002), la expresión general del análisis discriminante es:

$$Z = \alpha + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \dots + \beta_n X_n, \quad (2)$$

donde Z es el score (puntuación zeta) discriminante, α es el término intercepto y β representa el coeficiente respectivo en la combinación lineal de las variables explicativas X_i .

Cuando a los solicitantes del crédito se les quiere clasificar en dos grupos (cumplido e incumplido), bastará con una función discriminante Z , pero si se les quiere clasificar en tres grupos, harán falta dos funciones discriminantes. En general serán necesarias $K-1$ funciones discriminantes donde K es el número de grupos en que se divide la población.

El objetivo en este análisis es obtener una función discriminante con alto poder de clasificación y conocer qué variables determinan la pertenencia a cada grupo. Uno de los métodos empleados es el de inclusión por pasos, en este procedimiento (paso a paso) puede ingresar, y también ser excluida cualquier variable independiente, que cumpla o no, con los siguientes requisitos: Valor requerido del estadístico F , tenga un valor más pequeño del estadístico Lambda de Wilks (este indicador para un conjunto de variables independientes mide las desviaciones dentro de cada grupo respecto a las desviaciones totales) y que el nivel de tolerancia sea lo más cercano a uno, la tolerancia es un indicador de la asociación lineal entre las variables independientes.

El otro método utilizado se denomina introducción conjunta de variables independientes.

Efectuado el análisis discriminante, se verificará que la función discriminante sea significativa y la bondad del ajuste sea aceptable, después de ejecutada esta tarea, el interés se colocará en interpretar los resultados, para ello se deben examinar las funciones discriminantes obtenidas y establecer la importancia relativa de cada variable independiente en el momento de realizar la discriminación entre los grupos. Los métodos generalmente empleados para ello son:

- i) Análisis de los coeficientes estandarizados de las funciones discriminantes de Fisher
- ii) Análisis de la matriz de estructura y
- iii) Análisis del estadístico F univariante.

Una utilidad importante del análisis es la clasificación de los deudores en diferentes categorías de riesgo, esto se puede llevar a cabo empleando las funciones discriminantes de Fisher, para obtener las probabilidades de impago, esta probabilidad se utiliza para indicar la pertenencia a alguna de las categorías definidas, con este procedimiento se puede diseñar un sistema interno de rating estructurado en diferentes grados de riesgo.

2.3 Regresión logística

El modelo de regresión logística permite estimar la probabilidad que el cliente de una institución financiera incumpla en sus pagos, el cálculo de la probabilidad de incumplimiento para cada crédito se puede realizar haciendo uso del modelo de regresión logística, tal como lo plantea Gujarati (2003), este modelo pertenece al grupo de modelos de elección binaria, la función de distribución acumulada en la que se basa el modelo es la función logística. Wiginton (1980) fue uno de los primeros en publicar resultados de scoring de crédito empleando regresión logística.

Si se consideran los clientes de una entidad financiera, a los cuales se le ha otorgado un crédito y denotando como Y la variable que representa el cumplimiento o incumplimiento del cliente; Y toma el valor de 1 cuando el cliente incumple y 0 cuando el cliente no incumple.

La probabilidad P_i de incumplimiento de una persona se puede notar como:

$$P_i = P(Y_i = 1) \quad (3)$$

La probabilidad de incumplimiento de un cliente se puede explicar en función de un conjunto de factores $X = (X_1, X_2, \dots, X_p)$, en este sentido la probabilidad P_i está determinada por la siguiente ecuación:

$$P_i = E(Y_i = 1 | X) = \frac{\text{Exp}(\beta X)}{1 + \text{Exp}(\beta X)} \quad (4)$$

La anterior ecuación representa la función de distribución logística, se puede verificar que si X está comprendida entre $(-\infty, \infty)$, P_i está en el intervalo $(0,1)$ además de que P_i está relacionada en forma no lineal con X_i .

Los coeficientes asociados a los factores X , esto es, los β_i , están representando la contribución del factor X_i en la explicación de la probabilidad de incumplimiento de la persona analizada. Si β_i es menor que cero, entonces se menciona que a medida que se incrementa el valor del factor X_i la probabilidad de incumplimiento disminuye; si β_i es mayor que cero indica que un incremento del valor del factor X_i , incrementaría la probabilidad de incumplimiento del cliente.

La regresión logística ha sido ampliamente utilizada en la estimación de las probabilidades de incumplimiento por ejemplo en los trabajos de Xiao et al. (2006), Tsai et al. (2009), Yap et al. (2011), Lawrence & Arshadi (1995) y Hamdi & Mestiri (2014). Otras investigaciones han comparado los resultados obtenidos en el scoring de crédito para predecir el incumplimiento en el pago de obligaciones, empleando un modelo de regresión logística, con los de un modelo de regresión logística difusa por ejemplo en Sohn et al. (2016), igualmente con otras técnicas de obtención de scoring como lo muestran Akkoç (2012), Kiruthika & Dilsha (2015).

2.4 Redes neuronales

Una red neuronal es un sistema que recibe entradas numéricas y genera salidas de uno o más valores numéricos. Estas redes son un intento de crear redes que funcionen de una manera muy similar al cerebro humano, utilizando componentes que se comportan como el cerebro. En el cerebro humano las señales electrónicas se llevan a una neurona por un gran número de dendritas, entonces tiene lugar la conversión de señales a pulsos eléctricos, enviados en un axón a un número de sinapsis, que transfieren ideas o información a las dendritas de otras neuronas. Por tanto, una neurona puede enviar o recibir una señal hacia o desde otras neuronas. Entonces una red neuronal consiste en elementos cada uno de los cuales recibe un número de entradas y genera una sola salida.

Lee et al. (2002), indican que las características esenciales de una red neural son los nodos (organizados en capas), la arquitectura de la red, que describe la conexión entre los nodos, y el algoritmo utilizado para encontrar los valores de los parámetros de la red (pesos). Las capas de una red pueden ser:

- De entrada, conformadas por las neuronas que introducen las pautas de entrada en la red, en estas neuronas no se realiza procesamiento.
- Ocultas (intermedias) conformadas por las neuronas cuyas entradas vienen de capas anteriores y sus salidas pasaran a neuronas de capas siguientes.
- De salida, formadas por neuronas que sus valores indican la salida de toda la red, para el scoring de crédito la variable de salida indicaría la probabilidad de impago.

Para efectos de clasificación, se utiliza la red neuronal denominada perceptron multicapa (MLP). La red de trabajo está formada por una capa de entrada, una o más capas ocultas, y una capa de salida, esto es lo que generalmente se conoce como arquitectura de la red; cada una de estas capas contiene varias neuronas. Cada neurona procesa los valores de entrada (atributos) y genera un valor de salida que es transmitido a las neuronas de la capa siguiente. Cada neurona en la capa de entrada (con subíndices $i=1, 2, 3, \dots, n$) entrega el valor de una variable predictor (una característica) a partir de un vector x . Cuando se obtiene la distinción entre default y no default, entonces el valor de la neurona de salida es adecuado.

2.5 Árboles de decisión

Este método se encuentra clasificado dentro de las técnicas no paramétricas de cálculo de scoring de crédito, esto significa que no utilizan supuestos de distribución iniciales. Henley & Hand (1996) los ubican dentro de los métodos de partición recursiva, la finalidad de una partición recursiva es

dividir un conjunto de observaciones-datos en conjuntos disjuntos con el objetivo de aumentar la homogeneidad, cuando la partición es originada por la evaluación de una condición que tiene solamente dos alternativas, por ejemplo cumplido o incumplido, se conoce como partición recursiva binaria. En esta técnica se emplean reglas de clasificación hasta conseguir una categorización final, en donde se obtienen tanto los posibles resultados de un evento como la probabilidad de ocurrencia.

El objetivo de los árboles de decisión es predecir o clasificar una variable objetivo dependiente a través de la combinación y partición de variables independientes. Este método es muy útil en la selección de variables significativas entre una gran cantidad de variables. El resultado de un árbol (nodos finales) indica las variables independientes que se encuentran más fuertemente relacionadas (directa o inversa) con la variable dependiente, además indica los rangos de la variable dependiente en los cuales existe una mayor concentración del valor de la variable independiente (segmentación), permitiendo también la estratificación, predicción y exploración.

El método en su desarrollo implementa nodos (que corresponden a cada grupo o subgrupo de casos) que pueden ser: nodo raíz (es el primer nodo que representa la muestra completa y en el que se describen las diferentes categorías de la variable de interés), nodos filiales-rama (nodos que se crean cuando se particiona un nodo) y nodo terminal-rama (que corresponde al último nodo). Cuando se construyen modelos de scoring de crédito, las probabilidades de incumplimiento siempre se van a obtener en los nodos terminales, esto se realiza identificando el número de créditos incumplidos en un nodo final respecto del número de créditos que están en cada uno de los nodos finales, lo que significa que se tendrían tantas probabilidades de incumplimiento como nodos terminales hayan. Algunos algoritmos utilizados para realizar el estudio entre la variable dependiente y la independiente son CHAID (Chi-square Automatic Interaction Detection), CART (Classification And Regression Trees).

Los estudios de Baesens et al. (2003), Abdou & Pointon (2011), Brown & Mues (2012), realizan un análisis de datos en donde comparan los resultados obtenidos en scoring de crédito empleando varios métodos de clasificación, entre ellos los árboles de decisión.

2.6 Modelos Híbridos

Los modelos hasta aquí presentados, han sido utilizados en forma satisfactoria en cuanto al manejo de los datos y resultados obtenidos, las investigaciones y desarrollos relacionados en la gestión del riesgo financiero y en particular del riesgo de crédito, han buscado la construcción de nuevos métodos que propendan por una mejor evaluación y desempeño en torno a la predicción y calificación que se busca con ellos, son varias las investigaciones que sobre este aspecto se han llevado a cabo, un direccionamiento en este sentido ha sido la construcción de modelos híbridos, lo cuales combinan técnicas que pueden analizarse en forma separada o conjunta. En general, un modelo híbrido se basa en combinar técnicas de agrupamiento y clasificación, en Lenard et al. (1998), se propone un modelo que emplea la técnica de clustering (clasificación) como método de aprendizaje no supervisada, el cual es entrenado inicialmente y su salida utilizada posteriormente como entrada en una técnica de agrupamiento, buscando mejorar el resultado final. Steinberg & Cardell (1998) plantean una técnica para ayudar a resolver problemas de clasificación de varias tipologías, el trabajo lo elaboran mediante la hibridación de dos métodos de clasificación típicos, árboles de decisión y regresión logística, herramientas útiles en la minería de datos, detallan la manera de implementar este enfoque desde el punto de vista teórico y práctico. Lee et al. (2002), establecen un modelo híbrido entre redes neuronales artificiales (ANN) y análisis discriminante (DA), utilizan las redes neuronales tipo back propagation (BPN), encuentran que el modelo híbrido converge más rápido hacia una solución que una ANN tradicional, en primera instancia el DA es utilizado para determinar las variables apropiadas que alimentarán una ANN. En el trabajo de Hsieh (2005) empleando la minería de datos se construye un modelo híbrido que junta las redes neuronales artificiales y el método de clustering, el modelo resultado es empleado para obtener el

scoring de crédito. El estudio elaborado por Huanh, Chen & Wang (2007) lleva a cabo una combinación entre la técnica de máquinas de soporte vectorial (SVM) y los algoritmos genéticos (GA), para emplearlo en una base de datos del repositorio UCI, los resultados del modelo se comparan con los obtenidos por las dos técnicas utilizadas individualmente y también con el empleo de redes neuronales y árboles de decisión; Espín-García & Rodríguez-Caballero (2013) elaboran un modelo híbrido combinado árboles de decisión y regresión logística empleado para obtener la clasificación de clientes sin referencias de crédito en una entidad financiera en México.

Otros métodos empleados para para obtener puntuaciones crediticias son el método del vecino más cercano (Fix & Hodges1952, Chatterjee & Barcun, 1970; Henley & Hand,1996; Hastie et al. 2005), el método de sistemas automáticos de soporte vectorial (Vapnik, 1998), método de lógica difusa (Hoffmann et al.2007; Malhotra & Malhotra,1999; Tang & Chi,2005).

2.7 Probabilidad de incumplimiento

El evento de incumplimiento en una operación crediticia es objeto de asignación de una probabilidad de ocurrencia, la cual puede analizarse a nivel del acreditado, o deudor. De acuerdo con Crosbie (1997) citado por Elizondo & Altman (2003) los elementos que deben considerarse en el análisis del riesgo de crédito individual son:

- i) La probabilidad de incumplimiento obtenida como la frecuencia relativa con la que ocurre el evento en el que la contraparte no cumpla con las obligaciones contractuales e incumpla en el compromiso contraído.
- ii) La tasa de recuperación, referida a la proporción de la deuda que podrá ser recuperada cuando que la contraparte ha entrado en incumplimiento.
- iii) La migración del crédito, que indica el grado con que la calidad o calificación del crédito puede mejorar o deteriorarse.

3. Metodología

Este análisis de carácter cuantitativo, tiene como objetivo estudiar los datos y variables que caracterizan los solicitantes de crédito con el fin de obtener una metodología de clasificación como buen o mal cliente (0 y 1) en función de la probabilidad de cumplimiento con la obligación crediticia contraída, empleando para ello técnicas estadísticas como el análisis discriminante y la regresión logística, también se utilizan las redes neuronales, clasificadas como técnicas de inteligencia artificial y un modelo híbrido que combina los árboles de decisión con la regresión logística.

Los resultados obtenidos pueden emplearse para estimar la probabilidad de incumplimiento (PD) de cada uno de los créditos analizados. Con el propósito de evaluar el desempeño de los modelos elaborados se construye la curva de operaciones características (ROC) y la matriz de confusión o tabla de clasificación.

3.1 Datos y variables

Se empleará una base de datos conformada por 673 registros de clientes de una entidad financiera en Colombia, que no hacen parte del buró de créditos central, los datos fueron obtenidos en un periodo de doce meses entre los años 2014 y 2015, el conjunto de datos brinda información de los créditos otorgados a personas naturales, permitiendo conocer las características del deudor y de la operación crediticia. Las variables independientes son similares a las utilizadas en los trabajos de Avery et al. (2004), Quintana (2005), Boj et al. (2009), Rayo et al. (2010) y Villano (2013).

En la tabla 1 se muestra las variables empleadas en cada uno de los modelos y referidas a cada solicitante del crédito.

Tabla 1. Variables utilizadas

Nombre de la Variable	Nombre de la Variable
Estado actual	Préstamo
Género	Plazo
Edad	Ingreso/ deuda
Actividad económica	Línea crédito
Estado civil	Tasa de interés
Tipo vivienda	Garantía
Número personas a cargo	

Fuente: Elaboración propia

3.2 Resultados empleando análisis discriminante

El propósito principal de un análisis discriminante es predecir la pertenencia a un grupo determinado con base en una combinación lineal de un conjunto variables métricas. El procedimiento comienza con un conjunto de observaciones donde se conoce tanto la pertenencia al grupo y los valores de las variables. El resultado final del procedimiento es un modelo que permite la predicción de la pertenencia al grupo cuando se conocen sólo las variables independientes.

Para esta investigación se utilizó el método de inclusión por pasos con el fin de conocer la significancia individual de cada variable en la función discriminante, lo que permite la construcción de una función que emplee las variables más útiles en la clasificación y obtener la contribución individual de cada variable al modelo discriminante, en la tabla 2 se observa como finalmente fueron seleccionadas cuatro variables (clasificadoras) las cuales como se mencionó en la sección 2.2 cumplen con los valores requeridos de los estadísticos Lambda de Wilks y F, para este análisis tienen significancia estadística.

Tabla 2. Variables introducidas/excluidas

Paso	Variables Introducidas	Lambda de Wilks							
		Estadístico	gl1	gl2	gl3	F exacta			
						Estadístico	gl1	gl2	Sig.
1	Vivienda	,948	1	1	671,000	36,733	1	671,000	,000
2	Garantía	,910	2	1	671,000	33,071	2	670,000	,000
3	Estado	,904	3	1	671,000	23,592	3	669,000	,000
4	Genero	,898	4	1	671,000	18,948	4	668,000	,000

Fuente: Elaboración propia

En la tabla 3 se muestra el resultado de las variables seleccionadas en el procedimiento paso a paso y sus respectivos coeficientes, los cuales se emplearán en la función discriminante.

Con los valores anteriores se obtiene el score para cada cliente de la entidad, como se indica en la expresión (2).

Tabla 3. Coeficientes de la función canónica discriminante

Variable	Función
	1
Genei	,524
Ecivi	,323
Tvivi	1,659
Garani	1,329
(Constant)	-2,405

Fuente: Elaboración propia

En la tabla 4, se muestra que el modelo discriminante construido es capaz de clasificar a 401 solicitantes como buenos clientes (Def = 0) de 548 buenos solicitantes, por lo tanto, tiene un 73.2% de precisión en la clasificación para el grupo buenos clientes. Por otro lado, el mismo modelo discriminante es capaz de clasificar 81 malos solicitantes como malos clientes (Def =1) de 125 malos solicitantes. Por lo tanto, tiene un 64.8% de precisión de clasificación para el grupo malos clientes. Luego el modelo es capaz de obtener una precisión en la clasificación del 71,6% en los grupos combinados.

Tabla 4. Resultados de la clasificación Análisis Discriminante

		Def	Grupo de pertenencia pronosticado		Total
			0	1	
Original	Recuento	0	401	147	548
		1	44	81	125
	%	0	73,2	26,8	100,0
		1	35,2	64,8	100,0
Validación cruzada	Recuento	0	401	147	548
		1	44	81	125
	%	0	73,2	26,8	100,0
		1	35,2	64,8	100,0

Fuente: Elaboración propia

3.3 Resultados empleando regresión logística

En el modelo de regresión logística la probabilidad de incumplimiento puede obtenerse en función de un conjunto de variables explicativas X_i . El modelo tiene la siguiente estructura, utilizando las variables descritas en la tabla 1.

$$Y_i = \ln[(P_i|1 - P_i)] = \beta_1 + \beta_2 Gene_i + \beta_3 Edad_i + \beta_4 Act_i + \beta_5 Eciv_i + \beta_6 Tviv_i + \beta_7 Npcar_i + \beta_8 Pres_i + \beta_9 Nper_i + \beta_{10} Rind_i + \beta_{11} Lcre_i + \beta_{12} Tint_i + \beta_{13} Garan_i + \mu_i \quad (5)$$

En el desarrollo de la regresión logística se obtiene un modelo estadísticamente significativo con una buena capacidad de clasificación, en la tabla 5, se muestra que el modelo logístico construido es capaz de clasificar a 531 solicitantes como buenos clientes (Def = 0) de 548 buenos solicitantes, por lo tanto, tiene un 96.9% de precisión en la clasificación para el grupo buenos clientes. Por otro lado, el mismo modelo logístico es capaz de clasificar 14 malos solicitantes como malos clientes (Def =1) de 125 malos solicitantes. Por lo tanto, tiene un 11.2% de precisión de clasificación para el grupo

malos clientes. Luego el modelo es capaz de obtener una precisión en la clasificación del 81,0% para ambos grupos.

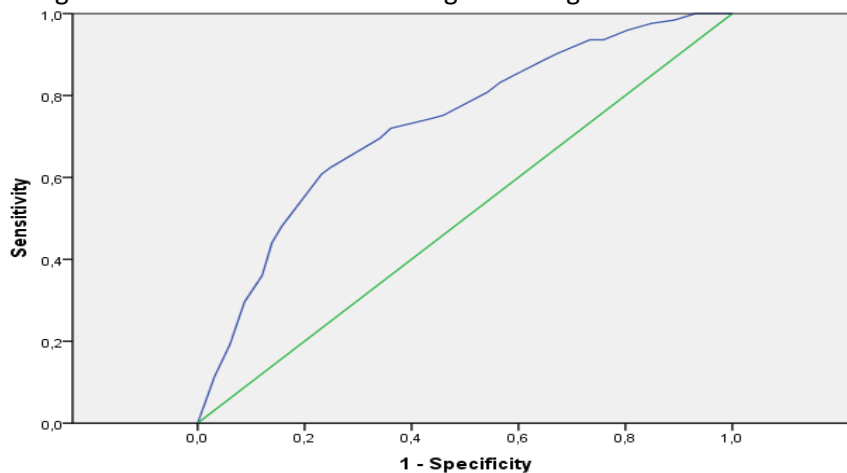
Tabla 5. Tabla de clasificación Regresión Logística

Observado		Pronosticado		
		Def Previo		Porcentaje correcto
		0	1	
Paso 4	Def Previo 0	531	17	96,9
	1	111	14	11,2
Porcentaje global				81,0

Fuente: Elaboración propia

El modelo logra un buen ajuste a los datos como se observa en la curva de operaciones características (figura 1) con un área de 73%, resultado que se puede ver en la Tabla 6, la curva se encuentra alejada de la línea de referencia indicando una buena representación de los datos por parte del modelo.

Figura 1. Curva ROC Modelo de Regresión Logística



Fuente: Elaboración propia

Tabla 6. Área bajo la curva modelo Regresión Logística

Área	Error típ.	Sig. asintótica	Intervalo de confianza asintótico al 95%	
			Límite inferior	Límite superior
,730	,025	,000	,682	,779

Fuente: Elaboración propia

3.4 Resultados empleando redes neuronales

En el modelo de redes neuronales que se muestra en la tabla 7 de clasificación, se observan los resultados prácticos de la utilización de la red perceptron multicapa (MLP). En general, el 90.1% de los casos de entrenamiento se clasificaron correctamente, lo que corresponde al 9.9% de incorrectos mostrado en la tabla resumen del modelo. Un modelo mejor debe identificar correctamente un mayor porcentaje de los casos.

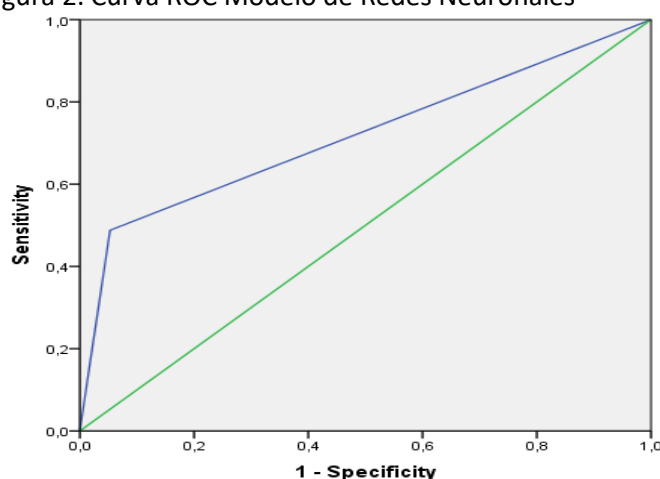
Tabla 7. Tabla de Clasificación redes neuronales

Muestra	Observado	Pronosticado		Porcentaje correcto
		0	1	
Entrenamiento	0	386	9	97.7%
	1	39	49	55.7%
	Porcentaje global	88.0%	12.0%	90,1%
Reserva	0	133	20	86.9%
	1	25	12	32.4%
	Porcentaje global	83.2%	16.8%	76.3%

Fuente: Elaboración propia

El modelo logra un buen ajuste a los datos, se ve reflejado en los valores de la curva ROC (figura 2) con un área de 71.8%, como se puede ver en la tabla 8.

Figura 2. Curva ROC Modelo de Redes Neuronales



Fuente: Elaboración propia

Tabla 8. Área bajo la curva modelo Redes Neuronales

Área	Error típ.	Sig. asintótica	Intervalo de confianza asintótico al 95%	
			Límite inferior	Límite superior
,718	,030	,000	,660	,775

Fuente: Elaboración propia

3.5 Resultados empleando modelo híbrido

Con el objetivo de mejorar la clasificación de los clientes analizados (los que ya poseen créditos) y realizar una propuesta para mejorar la clasificación de nuevos solicitantes, se construyó un modelo híbrido combinando los métodos árboles de decisión y regresión logística, siguiendo los trabajos de Steinberg & Cardell (1998); Espín-García & Rodríguez-Caballero (2013).

Inicialmente se elaboró un modelo de árboles de decisión, el resumen del modelo construido, empleando el método de crecimiento CHAID, se presenta en el Anexo 1. Las variables independientes seleccionadas para el análisis fueron Vivienda, Garantía, Género. Los resultados del modelo anterior y las probabilidades obtenidas en cada nodo terminal, se pueden observar en la estructura del mismo presentada en el Anexo 2. El árbol contiene nueve nodos, de los cuales cinco son nodos terminales, con la información provista por este método, se elaboró un modelo de regresión logística, en la tabla 9, se puede ver como el modelo híbrido clasifica a 548 solicitantes

como buenos clientes (Def = 0). En términos porcentuales el modelo tiene una precisión del 81.4% para ambos grupos (buenos clientes y malos clientes).

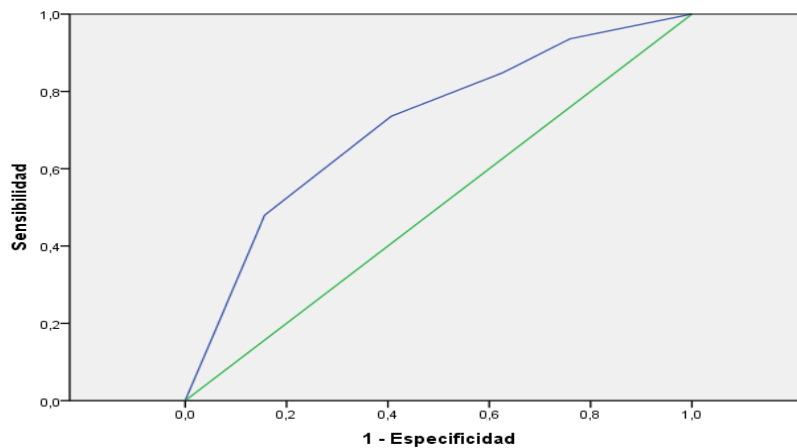
Tabla 9. Tabla de clasificación Modelo Híbrido

Observado			Pronosticado		
			Def Previo		Porcentaje correcto
			0	1	
Paso 2	Def Previo	0	548	0	100,0
		1	125	0	,0
	Porcentaje global				81,4

Fuente: Elaboración propia

El ajuste del modelo a los datos es adecuado, se observa en la curva ROC (figura 3) cubriendo un área de 71.5%, información consignada en la tabla 10.

Figura 3. Curva ROC Modelo Híbrido



Fuente: Elaboración propia

Tabla 10. Área bajo la curva modelo Híbrido

Área	Error típ.	Sig. asintótica	Intervalo de confianza asintótico al 95%	
			Límite	
			Límite inferior	superior
,715	,026	,000	,665	,765

Fuente: Elaboración propia

3.6 Comparación de los modelos

Una vez obtenidos los resultados de los cuatro modelos se puede condensar la información para realizar la comparación de la capacidad predictiva de cada técnica como se muestra en la tabla 11:

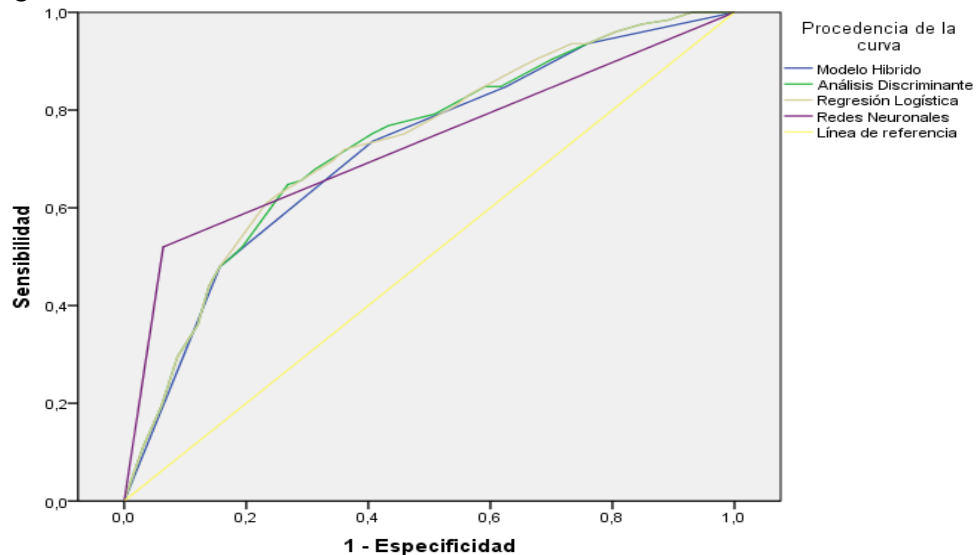
Tabla 11. Comparación capacidad predictiva de los cuatro modelos

Modelo	Buenos aceptados	Buenos rechazados	Malos aceptados	Malos rechazados	Tasa de éxitos
Análisis discriminante	401	147	44	81	71.6%
Regresión logística	531	17	111	14	81.0%
Redes neuronales	386	9	39	49	86.9%
Híbrido	548	0	125	0	81.4%

Fuente: Elaboración propia

Se puede apreciar el mejor desempeño del modelo de redes neuronales en la clasificación de los clientes solicitantes de crédito, en tanto que el análisis discriminante presenta la más baja tasa de aciertos en la clasificación, la segunda mejor opción para calificar los clientes es el modelo híbrido construido en la combinación de árboles de decisión y regresión logística, este modelo logra un mejor clasificación que el modelo de regresión logística, cuando éste se utilizó individualmente. Con el indicador de la curva ROC (figura 4) se puede realizar un análisis similar de forma conjunta para los cuatro modelos como se observa en la tabla 12.

Figura 4. Curva ROC de los cuatro modelos



Fuente: Elaboración propia

Tabla 12. Área bajo la curva de los cuatro modelos

Variables resultado de contraste	Área	Error típ	Sig. asintótica	Intervalo de confianza asintótico al 95%	
				Límite inferior	Límite superior
Modelo Híbrido	,715	,026	,000	,665	,765
Análisis discriminante	,729	,025	,000	,681	,778
Regresión Logística	,730	,025	,000	,682	,779
Redes Neuronales	,718	,030	,000	,660	,775

Fuente: Elaboración propia

Se aprecia que el modelo de regresión logística cubre un área un tanto mayor que la cubierta por los modelos de redes neuronales, análisis discriminante e híbrido, en la clasificación de los prestatarios.

4. Conclusiones

El estudio realizado ha permitido desarrollar cuatro modelos uno mediante el procedimiento perceptrón multicapa, para construir una red neuronal y pronosticar la probabilidad de que un cliente tenga mora en un crédito para clasificarlo como buen o mal cliente, las otras técnicas empleadas en el estudio son el modelo de regresión logística, el modelo de análisis discriminante y el modelo híbrido árboles de decisión-regresión logística. Una opción alterna a las técnicas de clasificación lo constituye este último modelo permitiendo juntar una técnica paramétrica (regresión logística) con una no paramétrica (árboles de decisión). Los resultados de los modelos son comparables y podrían confrontarse con otras técnicas (lógica difusa, algoritmo genético, máquinas de soporte vectorial). Los indicadores de bondad muestran que se puede estar bastante

seguro que los datos no contienen relaciones que no puedan capturar estos modelos, y por lo tanto, pueden utilizarse para seguir investigando la relación entre las variable dependiente e independientes.

La curva ROC y la tabla de clasificación (matriz de confusión) son herramientas muy útiles para medir la efectividad de un modelo de scoring. La curva ROC permite saber qué tan bien separados están los buenos deudores de los malos. La matriz de confusión permite conocer qué tan efectivo es el modelo en el momento de calificar y clasificar.

El scoring de crédito puede utilizarse para controlar la selección de riesgos (clientes), gestionar las probables pérdidas por exposición al riesgo de crédito, evaluar nuevas solicitudes de préstamos, mejorar el tiempo de procesamiento en la aprobación de préstamos, asegurar que los criterios de crédito existentes son relevantes y aplicados de manera consistente, entre otros aspectos.

Los resultados obtenidos constituye una forma importante de apoyo a la dirección de las instituciones financieras en donde se empleen estos modelos para presupuestar la pérdida esperada, la pérdida no esperada, realizar las provisiones económicas necesarias e informar sobre el capital económico necesario para hacer frente a estos eventos, permitiendo que la entidad siga funcionando, es decir constituyen herramientas de ayuda en la gestión financiera. El análisis financiero se puede complementar con la obtención de medidas como el valor en riesgo (VaR y TVaR) y modelar la distribución de las probables pérdidas.

5. Referencias Bibliográficas

- Abdou, H. A., & Pointon, J. (2011). Credit scoring, statistical techniques and evaluation criteria: A review of the literature. *Intelligent Systems in Accounting, Finance and Management*, 18(2-3), 59-88.
- Akkoç, S. (2012). An empirical comparison of conventional techniques, neural networks and the three stage hybrid Adaptive Neuro Fuzzy Inference System (ANFIS) model for credit scoring analysis: The case of Turkish credit card data. *European Journal of Operational Research*, 222(1), 168-178.
- Avery, R. B., Calem, P. S., & Canner, G. B. (2004). Consumer credit scoring: Do situational circumstances matter? *Journal of Banking & Finance*, 28(4), 835-856.
- Baesens, B., Van Gestel, T., Viaene, S., Stepanova, M., Suykens, J., & Vanthienen, J. (2003). Benchmarking state-of-the-art classification algorithms for credit scoring. *Journal of the Operational Research Society*, 54(6), 627-635.
- Boj, E., Claramunt, M. M., Esteve, A., & Fortiana, J. (2009). *Criterio de selección de modelo en credit scoring. Aplicación del análisis discriminante basado en distancias*. Artículo presentado en Anales del Instituto de Actuarios Españoles.
- Brown, I., & Mues, C. (2012). An experimental comparison of classification algorithms for imbalanced credit scoring data sets. *Expert Systems with Applications*, 39(3), 3446-3453.
- Chatterjee, S., & Barcun, S. (1970). A nonparametric approach to credit screening. *Journal of the American statistical Association*, 65(329), 150-154.
- De Lara Haro, A. (2005). *Medición y control de riesgos financieros*: Editorial Limusa.
- Elizondo, A., & Altman, E. I. (2003). *Medición integral del riesgo de crédito*: Editorial Limusa.
- Espin-García, O., & Rodríguez-Caballero, C. V. (2013). Metodología para un scoring de clientes sin referencias crediticias. *Cuadernos de Economía*, 32(59), 137-162.
- Fix, E., & Hodges Jr, J. L. (1952). Discriminatory analysis-nonparametric discrimination: Small sample performance: DTIC Document.
- Gujarati, D. N. (2003). *Basic Econometrics*. 4th: New York: McGraw-Hill.
- Gutierrez, G. M. A. (2007). Credit scoring models: what, how, when and for what purposes. *Munich Personal RePEc Archive Paper*, 16377.

- Hamdi, M., & Mestiri, S. (2014). Bankruptcy prediction for Tunisian firms: An application of semi-parametric logistic regression and neural networks approach. *Economics Bulletin*, 34(1), 133-143.
- Hand, D. J., & Henley, W. E. (1997). Statistical Classification Methods in Consumer Credit Scoring: A Review. *Journal of the Royal Statistical Society. Series A (Statistics in Society)*, 160(3), 523-541. doi: 10.2307/2983268
- Hastie, T., Tibshirani, R., Friedman, J., & Franklin, J. (2005). The elements of statistical learning: data mining, inference and prediction. *The Mathematical Intelligencer*, 27(2), 83-85.
- Henley, W., & Hand, D. J. (1996). A k-nearest-neighbour classifier for assessing consumer credit risk. *The Statistician*, 77-95.
- Hoffmann, F., Baesens, B., Mues, C., Van Gestel, T., & Vanthienen, J. (2007). Inferring descriptive and approximate fuzzy rules for credit scoring using evolutionary algorithms. *European Journal of Operational Research*, 177(1), 540-555.
- Hsieh, N.-C. (2005). Hybrid mining approach in the design of credit scoring models. *Expert Systems with Applications*, 28(4), 655-665.
- Huang, C.-L., Chen, M.-C., & Wang, C.-J. (2007). Credit scoring with a data mining approach based on support vector machines. *Expert Systems with Applications*, 33(4), 847-856.
- Kiefer, N. M., & Larson, C. E. (2006). Specification and informational issues in credit scoring. *International Journal of Statistics and Management Systems*, 1, 152-178.
- Kiruthika, & Dilsha, M. (2015). A Neural Network Approach for Microfinance Credit Scoring. *Journal of Statistics and Management Systems*, 18(1-2), 121-138.
- Lawrence, E. C., & Arshadi, N. (1995). A Multinomial Logit Analysis of Problem Loan Resolution Choices in Banking. *Journal of Money, Credit and Banking*, 27(1), 202-216. doi: 10.2307/2077859
- Lee, T.-S., Chiu, C.-C., Lu, C.-J., & Chen, I. F. (2002). Credit scoring using the hybrid neural discriminant technique. *Expert Systems with Applications*, 23(3), 245-254.
- Lenard, M. J., Madey, G. R., & Alam, P. (1998). The design and validation of a hybrid information system for the auditor's going concern decision. *Journal of Management Information Systems*, 14(4), 219-237.
- Malhotra, R., & Malhotra, D. (1999). Fuzzy systems and neuro-computing in credit approval. *Journal of lending and Credit Risk Management*, 81, 24-27.
- McNeil, A., Frey, R., & Embrechts, P. (2005). Quantitative risk management: Concepts, techniques and tools. *Princeton Series in Finance, Princeton*.
- Quintana, M. J. M., Gallego, A. G., & Pascual, M. E. V. (2005). Aplicación del análisis discriminante y regresión logística en el estudio de la morosidad en las entidades financieras: comparación de resultados. *Pecunia: Revista de la Facultad de Ciencias Económicas y Empresariales, Universidad de León*(1), 175-199.
- Rayo Cantón, S., Lara Rubio, J., & Camino Blasco, D. (2010). Un modelo de Credit Scoring para instituciones de microfinanzas en el marco de Basilea II. *Journal of Economics, Finance and Administrative Science*, 15(28), 89-124.
- Sohn, S. Y., Kim, D. H., & Yoon, J. H. (2016). Technology credit scoring model with fuzzy logistic regression. *Applied Soft Computing*, 43, 150-158.
- Steinberg, D., & Cardell, N. S. (1998). The hybrid CART-Logit model in classification and data mining. *Salford Systems White Paper*.
- Tang, T.-C., & Chi, L.-C. (2005). Predicting multilateral trade credit risks: comparisons of Logit and Fuzzy Logic models using ROC curve analysis. *Expert Systems with Applications*, 28(3), 547-556.
- Thomas, L., Edelman, D., & Crook, J. (2002). Credit scoring & its applications, Society for Industrial Mathematics: Philadelphia.
- Thomas, L. C., Edelman, D. B., & Crook, J. N. (2004). *Readings in credit scoring: foundations, developments, and aims*: Oxford University Press on Demand.

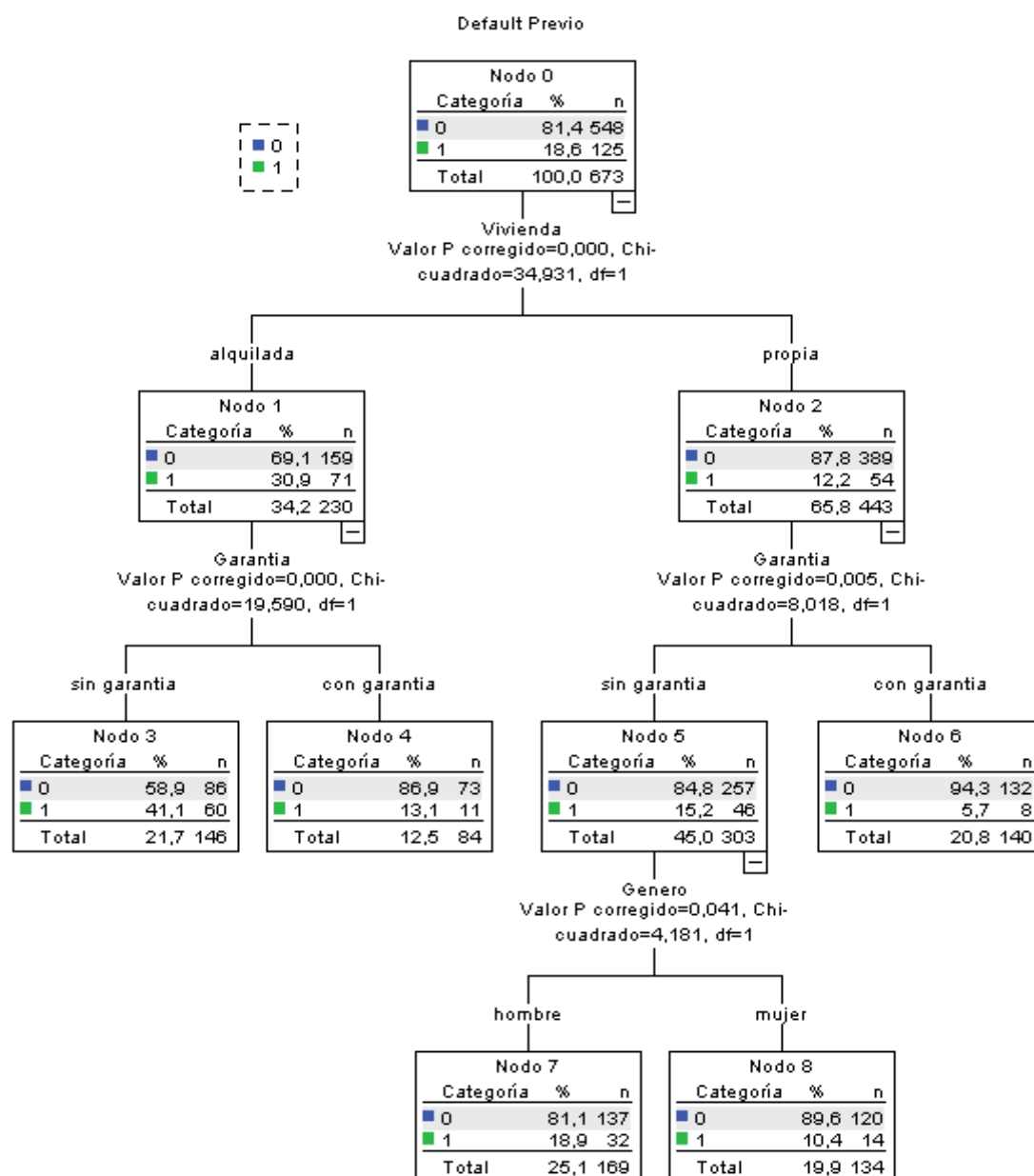
- Tsai, M.-C., Lin, S.-P., Cheng, C.-C., & Lin, Y.-P. (2009). The consumer loan default predicting model – An application of DEA–DA and neural network. *Expert Systems with Applications*, 36(9), 11682-11690.
- Vapnik, V. (1998). The support vector method of function estimation *Nonlinear Modeling* (pp. 55-85): Springer.
- Villano, F. E. S. (2013). Cuantificación del riesgo de incumplimiento en créditos de libre inversión: un ejercicio econométrico para una entidad bancaria del municipio de Popayán, Colombia. *Estudios Gerenciales*, 29(129), 416-427.
- Wiginton, J. C. (1980). A note on the comparison of logit and discriminant models of consumer credit behavior. *Journal of Financial and Quantitative Analysis*, 15(03), 757-770.
- Xiao, W., Zhao, Q., & Fei, Q. (2006). A comparative study of data mining methods in consumer loans credit scoring management. *Journal of Systems Science and Systems Engineering*, 15(4), 419-435.
- Yap, B. W., Ong, S. H., & Husain, N. H. M. (2011). Using data mining to improve assessment of credit worthiness via credit scoring models. *Expert Systems with Applications*, 38(10), 13274-13283.
- Yu, L., Wang, S., Lai, K. K., & Zhou, L. (2008). *BioInspired Credit Risk Analysis*: Springer.

Anexo 1. Resumen del modelo Árboles de decisión

Especificaciones	Método de crecimiento	CHAID	
	Variable dependiente	Default Previo	
	Variables independientes	Genero, Edad, Actividad económica, Estado, Vivienda, Persona a cargo, Préstamo, Periodos, Relación, Línea, Tasa, Garantía	
	Validación	Ninguna	
	Máxima profundidad de árbol		3
	Mínimo de casos en un nodo filial		100
	Mínimo de casos en un nodo parental		50
Resultados	Variables independientes incluidas	Vivienda, Garantía, Genero	
	Número de nodos		9
	Número de nodos terminales		5
	Profundidad		3

Fuente: Elaboración propia

Anexo 2. Estructura modelo Árboles de decisión



Fuente: Elaboración propia