

TRABAJO DE INVESTIGACIÓN FINAL

Sistemas de credit scoring basados en machine learning y datos alternativos:

Evaluación de su aplicabilidad en el sistema financiero argentino a partir de la construcción de un dataset sintético

Autores:

Juan Ignacio Bárbaro

Nehuen Oscar Fernández

Carrera:

Licenciatura en Finanzas

Tutor: Juan Y Caligiuri Alessandro

Año: 2026

A nuestras familias, por el apoyo constante y la paciencia a lo largo de toda la carrera. A nuestro tutor, por la guía técnica y el compromiso con este trabajo de investigación. Y a la Universidad Argentina de la Empresa (UADE), por brindarnos el espacio y las herramientas para nuestra formación profesional.

Índice

Abstract.....	5
Resumen	5
Palabras clave:.....	6
Introducción	6
1 Planteamiento del problema.....	8
2 Justificación de la investigación	9
3 Preguntas de investigación	10
4 Objetivo General	11
5 Objetivos Particulares	11
6 Hipótesis	11
7 Marco Teórico	11
7.1 Riesgo crediticio y concepto de default.....	11
7.2 Indicadores de riesgo crediticio: PD, LGD, EAD.....	12
7.3 Evaluación del desempeño y métricas de contraste	13
7.4 La matriz de confusión: estructura y dimensiones del error de clasificación....	13
7.5 El AUC-ROC como métrica central de evaluación en credit scoring	14
7.6 El reporte de clasificación: Precision, Recall y F1-Score	15
7.7 Modelos tradicionales de credit scoring.....	17
7.8 Modelos estadísticos clásicos.....	17
7.9 Utilización de variables tradicionales en credit scoring.....	18
7.10 Limitaciones de los modelos tradicionales	18
7.11 Machine learning aplicado al riesgo crediticio	19
8 Principales algoritmos aplicados al credit scoring	19
8.1 Classification and Regression Trees (CART).....	19
8.2 Random Forest	20
8.3 Gradient Boosting	20
8.4 XGBoost (eXtreme Gradient Boosting).....	21
8.5 LightGBM (Light Gradient Boosting Machine).....	21
9 Datos alternativos en credit scoring	22
10 Ventajas frente a los modelos tradicionales	23
11 Inclusión financiera y contexto argentino	23

12	Justificación del uso de datasets sintéticos en investigación sobre credit scoring	25
13	Validez metodológica y aplicación al contexto argentino	25
14	Metodología	26
14.1	Enfoque metodológico	26
14.2	Construcción del dataset sintético	27
14.3	Construcción de la variable objetivo	28
14.4	Entrenamiento de los modelos	29
14.5	Estrategia de análisis	29
15	Resultados de los modelos	30
15.1	Regresión Logística	30
15.2	Resultados del modelo LightGBM	31
	Comparación entre ambos modelos	33
16	Interpretación de la importancia de las variables	34
17	Limitaciones.....	37
18	Conclusión.....	38
19	Futuras líneas de investigación	39
20	Bibliografía	41
21	ANEXOS	47
	Anexo 1: Datos Bloomberg	47
	Anexo 2: Detalle de variables.....	49
	Anexo 3: Código Python	58

Abstract

The rigidity of traditional credit evaluation systems systematically excludes informal workers and individuals with little or no credit history. In Argentina, where labor informality exceeds 36%, this limitation represents a significant structural barrier to financial inclusion. Within this context, the present study evaluates the use of machine learning models driven by alternative variables, such as transactional behavior patterns in digital wallets and mobile device segmentation telemetry, as a tool to improve credit risk assessment and expand access to credit.

Given the unavailability of real microdata, a synthetic dataset is constructed and calibrated using official and industry-specific Argentine sources, including the Central Bank of Argentina (BCRA), the National Institute of Statistics and Census (INDEC), the Economic Commission for Latin America and the Caribbean (ECLAC), and the Argentine Fintech Chamber. Based on this dataset, a machine learning model, specifically LightGBM, is trained and compared against a logistic regression model as its traditional counterpart, with the objective of assessing the differences in predictive performance. Model performance is evaluated using metrics such as AUC-ROC, the Kolmogorov–Smirnov (KS) statistic, and the F1-score.

The results indicate that alternative variables possess significant predictive power with respect to default risk and that ensemble methods are capable of capturing non-linear relationships that traditional linear approaches fail to represent. In particular, LightGBM demonstrates superior performance relative to the baseline model across most of the evaluation metrics considered. Overall, these findings support the viability of machine learning and alternative data as a robust approach to credit assessment for unbanked and underbanked populations in Argentina.

Resumen

La rigidez de los sistemas tradicionales de evaluación crediticia excluye de manera sistemática a trabajadores informales y a individuos sin historial crediticio. En Argentina, donde la informalidad laboral supera el 36%, este sesgo constituye una barrera estructural relevante para la inclusión financiera. En este contexto, la presente investigación evalúa el uso de modelos de machine learning alimentados con variables alternativas, como patrones de comportamiento transaccional en billeteras digitales y

telemetría de segmentación por gama de dispositivo móvil, como herramienta para mejorar la estimación del riesgo crediticio y ampliar el acceso al crédito.

Dado que no se dispone de microdatos reales, se construye un dataset sintético calibrado a partir de fuentes oficiales y sectoriales argentinas (BCRA, INDEC, CEPAL y Cámara Argentina Fintech). Sobre esta base se entrenan y comparan un modelo de machine learning, específicamente LightGBM, contra un modelos de regresión logística como contraparte tradicional, con el objetivo de contrastar los resultados obtenidos. Dicho desempeño se evalúa mediante métricas como AUC-ROC, estadístico KS, F1-Score.

Los resultados evidencian que las variables alternativas poseen poder predictivo significativo sobre el riesgo de default, y que los modelos de ensamble capturan relaciones no lineales que los enfoques lineales tradicionales no logran representar. En particular, LightGBM muestra un desempeño superior respecto del modelo base en la mayoría de las métricas evaluadas. En conjunto, estos hallazgos respaldan la viabilidad del uso de machine learning y datos alternativos como una alternativa robusta para la evaluación crediticia de poblaciones no bancarizadas o con historial limitado en Argentina.

Palabras clave:

Riesgo Crediticio, scoring, Machine Learning, Variables alternativas, LightGBM, regresión logística, dataset sintético, inclusión financiera

Introducción

El acceso al crédito constituye uno de los principales mecanismos de movilidad económica, desarrollo productivo e inclusión financiera, al permitir que personas y empresas financien proyectos de inversión, consumo y crecimiento (Thomas et al. 2017). En este escenario, la adecuada evaluación del riesgo crediticio resulta indispensable para que las entidades financieras puedan asignar eficientemente sus recursos y preservar la calidad de sus carteras de préstamos.

Durante las últimas décadas, los sistemas de credit scoring han evolucionado desde modelos estadísticos tradicionales, como la regresión logística y el análisis discriminante lineal, hacia herramientas basadas en técnicas de machine learning, capaces de

procesar grandes volúmenes de información e identificar patrones complejos de comportamiento financiero (Lessmann et al. 2015). Mientras los modelos tradicionales dependen del historial bancario formal, los modelos de machine learning han demostrado capacidad para incorporar fuentes de información alternativas, como el comportamiento digital, historial de pagos en servicios y alquileres, uso de billeteras virtuales, que permiten evaluar a personas sin trazabilidad en el sistema financiero convencional (Berg. et al., 2020).

En Argentina, esta problemática adquiere una relevancia particular debido a las características propias del mercado laboral y del sistema financiero. Según el Instituto Nacional de Estadística y Censos (INDEC, 2025), el 36,7% de los asalariados urbanos se desempeña en condiciones de informalidad laboral, mientras que el crédito al sector privado representa aproximadamente el 13% del Producto Bruto Interno, uno de los niveles más bajos de América Latina (Banco Mundial, 2024). Estas condiciones limitan el acceso al financiamiento formal de un amplio segmento de la población que, si bien desarrolla actividades económicas y genera ingresos, no dispone del historial bancario requerido por los modelos tradicionales de evaluación crediticia.

A fin de ilustrar con datos de mercado contemporáneos el perfil de riesgo del sistema bancario formal argentino, se consultó la plataforma Bloomberg Finance L.P. (21 de abril de 2026) para cuatro entidades de referencia del sector: BBVA Argentina (BBAR), Banco Macro (BMA), Grupo Financiero Galicia (GGAL) y Banco Hipotecario (BHIP). Los modelos de riesgo de Bloomberg estiman para estas entidades probabilidades de impago a un año que oscilan entre el 0,13% (Banco Macro) y el 0,56% (Banco Hipotecario), con ratios de préstamos en situación irregular (NPL) que van del 0,76% (Galicia) al 3,92% (Hipotecario). (Ref. Anexo 1). Estos indicadores corresponden exclusivamente al segmento formal y regulado del sistema financiero, el único con trazabilidad suficiente para ser modelado por los burós de crédito convencionales y contrastan marcadamente con las tasas de mora estimadas para el segmento informal y fintech, que según reports el BCRA en su informe sobre proveedores no financieros de crédito del segundo semestre de 2025 supera el 30%. Esta brecha entre los perfiles de riesgo del segmento formal y el informal es precisamente la que los sistemas tradicionales de scoring no pueden capturar, y la que esta investigación busca abordar mediante la incorporación de variables alternativas en modelos de machine learning.

Frente a este escenario, el crecimiento de la digitalización financiera ha generado nuevas fuentes de información capaces de complementar los mecanismos tradicionales de evaluación del riesgo. Variables como el comportamiento de pago de servicios públicos y alquileres, el uso de billeteras virtuales, la frecuencia de utilización de aplicaciones financieras y otros registros de comportamiento digital han comenzado a incorporarse como datos alternativos en modelos de credit scoring desarrollados mediante técnicas de machine learning. Diversos estudios sostienen que estas variables contienen información relevante sobre la capacidad y la voluntad de pago de los individuos, permitiendo mejorar la evaluación crediticia de personas con escaso historial financiero formal (Jagtiani & Lemieux, 2019; Berg et al., 2020).

En este contexto, la presente investigación analiza la aplicabilidad de modelos de machine learning basados en variables alternativas para la evaluación del riesgo crediticio de personas sin historial financiero formal en el sistema financiero argentino. Dado que el acceso a bases de datos crediticias reales se encuentra restringido por cuestiones regulatorias, comerciales y de protección de datos personales, conforme a la Ley N° 25.326 de Protección de Datos Personales, el estudio propone la construcción de un dataset sintético calibrado mediante estadísticas provenientes de organismos oficiales y antecedentes académicos, con el propósito de evaluar el potencial predictivo de estas variables y contribuir al desarrollo de metodologías aplicables al contexto argentino.

1 Planteamiento del problema

La finalidad de este estudio es analizar si los modelos de machine learning alimentados con variables alternativas pueden predecir el riesgo de default crediticio con mayor precisión que los modelos estadísticos tradicionales para la población excluida del sistema financiero formal en Argentina.

Las investigaciones existentes documentan que los modelos de machine learning superan a los modelos lineales tradicionales en la predicción del default crediticio cuando los datos exhiben relaciones no lineales y efectos de interacción entre variables (Lessmann et al., 2015). Sin embargo, esta evidencia se concentra principalmente en mercados desarrollados con alta bancarización, donde los datos alternativos complementan un historial crediticio preexistente. En Argentina, donde más de un tercio de la fuerza laboral carece de ese historial, la pregunta de si los datos alternativos

pueden complementarlo, permanece sin respuesta empírica local. Esta es la carencia central que la presente investigación busca solventar.

En este contexto, la tesis de grado de UADE titulada *"Modelos Predictivos de Riesgo Crediticio: Comparación entre Enfoques Tradicionales y de Machine Learning"* (2025) estableció una base conceptual y metodológica relevante, pero reconoció como limitación la ausencia de variables alternativas de comportamiento financiero informal y digital en su análisis. El presente trabajo se propone continuar y ampliar esa línea de investigación, incorporando ese conjunto de variables y evaluando si su inclusión produce una mejora medible en la capacidad predictiva de los modelos de machine learning sobre la regresión logística en el contexto argentino. Ante la imposibilidad de acceder a microdatos financieros reales, la investigación construye un dataset sintético calibrado con parámetros provenientes del BCRA, el INDEC, la CEPAL y la Cámara Argentina Fintech.

2 Justificación de la investigación

La presente investigación se justifica por la necesidad de generar evidencia empírica sobre la aplicabilidad de modelos de machine learning basados en variables alternativas para la evaluación del riesgo crediticio en el contexto argentino, donde la disponibilidad de datos constituye una de las principales limitaciones para el desarrollo de investigaciones de este tipo.

Desde el punto de vista de su conveniencia y relevancia social, el estudio aborda una problemática de amplio impacto: el 36,7% de los asalariados urbanos argentinos trabaja informalmente (INDEC, 2025) y queda excluido del crédito formal no por incapacidad de pago sino por ausencia de trazabilidad bancaria. Si los modelos propuestos demuestran ser predictivamente superiores, sus resultados pueden orientar a instituciones financieras, empresas fintech y organismos regulatorios hacia sistemas de evaluación más inclusivos, beneficiando directamente a millones de personas que hoy acceden al financiamiento únicamente a través del mercado informal, a tasas significativamente más elevadas. En cuanto a sus implicaciones prácticas, la investigación ofrece un framework metodológico replicable: el dataset sintético construido y calibrado con fuentes oficiales locales puede ser adoptado por otras instituciones o investigadores para evaluar modelos de scoring sin necesidad de acceder a datos reales restringidos, resolviendo así uno de los obstáculos operativos más frecuentes en la investigación aplicada al

crédito en Argentina, tal como señalan Bussmann et al. (2021), quienes identifican las restricciones de privacidad y disponibilidad de datos como el principal obstáculo para este tipo de investigación.

Desde el punto de vista de su valor teórico, la investigación llena un vacío en la literatura: la evidencia disponible sobre la superioridad predictiva de los modelos de machine learning con datos alternativos se concentra en mercados desarrollados (Berg et al., 2020), y existe escasa investigación empírica sobre su aplicabilidad en economías con alta informalidad estructural como la argentina. Al extender la línea de trabajo iniciada por la tesis precedente de UADE (2025) mediante la incorporación de variables alternativas, este estudio aporta nueva evidencia sobre qué variables y qué modelos son más adecuados para evaluar poblaciones *thin-file* en el contexto local, con posibilidad de generalización a otros mercados latinoamericanos con características similares. Finalmente, en términos de utilidad metodológica, la construcción del dataset sintético mediante parámetros calibrados con fuentes primarias argentinas constituye en sí misma una contribución instrumental: define un procedimiento sistemático para simular poblaciones crediticias locales en ausencia de microdatos reales, sugiriendo cómo estudiar más adecuadamente este tipo de poblaciones y sentando un antecedente metodológico replicable para investigaciones futuras en el campo del riesgo crediticio en Argentina.

3 Preguntas de investigación

¿Cuáles son los criterios necesarios para construir un dataset sintético representativo que incorpore variables alternativas destinadas al análisis del riesgo crediticio en ausencia de datos reales?

¿En qué medida un dataset sintético construido a partir de variables alternativas permite analizar la aplicabilidad de modelos de machine learning para la evaluación del riesgo crediticio en el sistema financiero argentino?

¿Qué tipo de variables alternativas resultan más adecuadas para representar el comportamiento crediticio de individuos sin historial financiero formal?

4 Objetivo General

Adaptar sistemas de credit scoring basados en machine learning y datos alternativos al sistema crediticio argentino, a partir de la construcción de un dataset sintético representativo del mercado local y el contraste del poder discriminativo de modelos de machine learning frente a la regresión logística tradicional.

5 Objetivos Particulares

Desarrollar un dataset sintético, basado en parámetros realistas que incluya variables alternativas y represente el comportamiento financiero de individuos con y sin historial crediticio.

Evaluar la aplicabilidad de los modelos de machine learning para scoring crediticio, contemplando restricciones de datos debido a las regulaciones vigentes de Argentina.

6 Hipótesis

Dado el carácter descriptivo y exploratorio de la investigación, no se plantea una hipótesis formal a ser contrastada empíricamente. En su lugar, el trabajo se basa en el supuesto de que la incorporación de datos alternativos permite representar de manera más amplia el comportamiento financiero de individuos sin historial crediticio, lo cual será analizado a través del modelo propuesto.

7 Marco Teórico

7.1 Riesgo crediticio y concepto de default

El riesgo crediticio constituye uno de los conceptos fundamentales de la teoría financiera moderna y representa la probabilidad de que un deudor incumpla total o parcialmente las obligaciones pactadas. Formalmente, Bessis (2015) lo define con la siguiente afirmación: “Credit risk is the risk of losses due to borrowers defaults or deterioration of credit standing” [El riesgo de crédito es el riesgo de pérdidas derivadas del default de los prestatarios o del deterioro de su calidad crediticia.] (Bessis, 2015, p. 5). Donde deja en claro que no implica únicamente pérdidas por incumplimiento (parcial o total), sino

también por el aumento de la probabilidad de ocurrencia del mismo (deterioro de la calidad crediticia).

De acuerdo con el Comité de Basilea (2006), el riesgo crediticio representa históricamente la fuente de pérdidas más significativa para las instituciones bancarias a nivel global, siendo su adecuada medición un requisito indispensable para la solvencia del sistema (Basel Committee on Banking Supervision, 2006, International Convergence of Capital Measurement, BIS)

El concepto de incumplimiento crediticio, denominado habitualmente *default*, refiere al evento en que un deudor deja de cumplir con los terminos comprometidos en un contrato de crédito. La definición operativa varía según el marco regulatorio, si bien los estándares de Basilea II y III establecen como referencia un atraso superior a noventa días. Siddiqi (2012) precisa que "La definición de *incumplimiento (bad)* suele establecerse en 90 días de mora o más, aunque este umbral varía según la institución y el producto financiero" (Siddiqi, 2012, Credit Risk Scorecards, 2ª ed., Wiley, p. 9)

7.2 Indicadores de riesgo crediticio: PD, LGD, EAD

La cuantificación del riesgo crediticio descansa sobre un conjunto de indicadores técnicos. El más relevante para el presente trabajo es la Probabilidad de Incumplimiento (*Probability of Default*, PD), que Bessis (2015) define como "La probabilidad de que el prestatario incurra en incumplimiento dentro de un horizonte de un año" (Bessis, 2015, Risk Management in Banking, 4ª ed., Wiley, p. 21).

Complementariamente, la Pérdida Dado el Incumplimiento (LGD) cuantifica la proporción de la exposición que la institución perdería en caso de default, y la Exposición al Momento del Incumplimiento (EAD) refleja el saldo de la deuda en ese momento. La combinación de estos tres parámetros permite calcular la Pérdida Esperada ($EL = PD \times LGD \times EAD$), base de los requerimientos de capital bajo el enfoque IRB de Basilea II (BCBS, 2006, International Convergence of Capital Measurement, BIS, p. 52).

Una vez definidos los principales indicadores utilizados para cuantificar el riesgo crediticio, resulta necesario analizar cómo se evalúa la capacidad de los modelos para estimar dichas variables de manera precisa. En otras palabras, no basta con identificar qué dimensiones del riesgo deben ser consideradas; también es fundamental contar con herramientas que permitan determinar si las predicciones generadas por un modelo reflejan adecuadamente el comportamiento observado en la realidad. En este contexto,

adquieren relevancia las métricas de evaluación y contraste, las cuales permiten medir el desempeño predictivo de los sistemas de scoring y analizar los distintos tipos de errores que pueden surgir durante el proceso de clasificación crediticia.

7.3 Evaluación del desempeño y métricas de contraste

Entre las distintas herramientas utilizadas para evaluar el desempeño de un modelo de clasificación, la matriz de confusión ocupa un lugar central debido a que permite descomponer los resultados obtenidos en categorías específicas de aciertos y errores. Esta representación constituye el punto de partida para el cálculo de múltiples métricas de desempeño y facilita la comprensión de las implicancias que cada tipo de error puede generar en la toma de decisiones crediticias.

7.4 La matriz de confusión: estructura y dimensiones del error de clasificación

La matriz de confusión es la herramienta fundamental que sistematiza los cuatro resultados posibles de un clasificador binario. Hastie, Tibshirani y Friedman (2009) la definen en el contexto del aprendizaje automático como:

A confusion matrix is a cross-classification of the predicted and actual class labels. It summarizes how often the classifier correctly or incorrectly assigns each observation to its true class. For a two-class problem, the confusion matrix has four cells: true positives (TP), false positives (FP), true negatives (TN), and false negatives (FN). Accuracy, defined as the proportion of correct classifications, is an inadequate measure of performance for skewed distributions, as a classifier that always predicts the majority class will achieve high accuracy while providing no useful information. [Una matriz de confusión es una tabulación cruzada de las clases predichas y las clases reales. Resume con qué frecuencia el clasificador asigna correcta o incorrectamente cada observación a su clase verdadera. En un problema de clasificación binaria, la matriz de confusión contiene cuatro celdas: verdaderos positivos (TP), falsos positivos (FP), verdaderos negativos (TN) y falsos negativos (FN). La exactitud (*accuracy*), definida como la proporción de clasificaciones correctas, constituye una medida de desempeño inadecuada en el caso de distribuciones sesgadas, ya que un clasificador que siempre predice

la clase mayoritaria puede alcanzar una elevada exactitud sin proporcionar información útil] (Hastie et al., 2009, p. 301)

Bahnsen (2014) señala dentro del contexto crediticio, las interpretaciones económicas y sociales de los cuatro cuadrantes. Los Verdaderos Negativos (TN) son aprobaciones correctas de solicitantes solventes, generan ingresos. Los Verdaderos Positivos (TP) son rechazos correctos de futuros defaultadores, evitan pérdidas. Los Falsos Negativos (FN) son aprobaciones incorrectas de individuos que luego incumplen: su costo es $\text{Costo}(\text{FN}) \approx \text{EAD} \times \text{LGD}$. Los Falsos Positivos (FP) son rechazos injustos de solicitantes solventes. (Bahnsen et al., 2014)

Si bien la matriz de confusión constituye una herramienta fundamental para comprender la distribución de aciertos y errores en un modelo de clasificación, su capacidad descriptiva resulta limitada cuando se busca realizar comparaciones entre distintos modelos o evaluar su desempeño bajo diferentes umbrales de decisión. Esta limitación adquiere especial relevancia en el ámbito del riesgo crediticio, donde los eventos de incumplimiento suelen representar una proporción reducida del total de observaciones, generando conjuntos de datos desbalanceados. En consecuencia, surge la necesidad de emplear métricas que permitan evaluar la capacidad discriminatoria de un modelo de manera más integral e independiente de un umbral específico de clasificación. Entre estas métricas, el área bajo la curva ROC (AUC-ROC) se ha consolidado como uno de los estándares más utilizados en la industria financiera y en la literatura académica para medir la capacidad de un modelo de distinguir entre clientes con diferente nivel de riesgo crediticio.

7.5 El AUC-ROC como métrica central de evaluación en credit scoring

El Área Bajo la Curva ROC (AUC-ROC) es la métrica de evaluación más ampliamente utilizada en credit scoring. Fawcett (2006) establece su definición y su interpretación probabilística en el trabajo que fijó el estándar metodológico para su uso en machine learning. De acuerdo con el autor, las curvas ROC (*Receiver Operating Characteristic*) constituyen representaciones bidimensionales que contraponen la tasa de verdaderos positivos en las ordenadas frente a la tasa de falsos positivos en las abscisas. Este gráfico ilustra el balance entre los aciertos del modelo (beneficios) y sus errores (costos).

En este contexto, el área bajo la curva (AUC) representa la probabilidad de que una instancia positiva seleccionada al azar reciba una puntuación mayor por parte del clasificador que una instancia negativa elegida de la misma forma. Esta propiedad, que es aplicable a cualquier métrica de desempeño que no dependa de un umbral fijo, cobra especial importancia en el ámbito del *credit scoring*; esto se debe a que permite medir la capacidad de ordenamiento (*rank-ordering*) del modelo sin verse afectada por el punto de corte (*cutoff*) determinado para la aprobación del crédito.

En términos prácticos, el AUC-ROC responde a: si tomamos al azar un individuo que defaulteará y otro que no, ¿qué probabilidad tiene el modelo de asignarle un score de riesgo mayor al primero? Un modelo perfectamente aleatorio obtiene $AUC=0,50$; un modelo perfecto obtiene $AUC=1,0$. Lessmann et al. (2015) establecen la escala de referencia práctica: $AUC < 0,60$ indica ausencia de poder discriminativo; $0,70-0,80$ indica discriminación aceptable; $0,80-0,90$ indica discriminación buena.

7.6 El reporte de clasificación: Precision, Recall y F1-Score

Mientras que el AUC-ROC evalúa la capacidad de ordenamiento del modelo de manera agregada e independiente del umbral de clasificación, el reporte de clasificación responde a una pregunta distinta: una vez que el modelo emite una decisión binaria concreta, aprobar o rechazar, sobre un umbral de corte determinado, ¿qué tan acertadas son esas decisiones sobre la clase de interés? Sokolova y Lapalme (2009) destacan que estas métricas son sensibles a cambios en la matriz de confusión y, a diferencia de la exactitud (*accuracy*), no son invariantes ante el desbalance de clases.

La matriz de confusión es una herramienta que evidencia el desempeño del modelo dividiendo las predicciones en verdaderos positivos (TP), verdaderos negativos (TN), falsos positivos (FP) y falsos negativos (FN). En múltiples ocasiones no sólo resulta relevante cuántos errores comete un modelo, sino también el tipo de error que genera. (Kuhn & Johnson, 2013)

Predicted	Observed	
	Event	Nonevent
Event	<i>TP</i>	<i>FP</i>
Nonevent	<i>FN</i>	<i>TN</i>

Ilustración 1: Extraída del libro Applied Predictive Modeling.

(Kuhn & Johnson, 2013).

La primera pregunta que se hace ese analista es: de todas las personas a las que rechacé por considerarlas riesgosas, ¿cuántas realmente lo eran? Esa pregunta la responde la Precision (Sokolova & Lapalme, 2009):

$$\text{Precisión} = \frac{TP}{TP + FP}$$

Ecuación 1: Formula Precision

Si la Precision es baja, significa que el analista, o el modelo, está rechazando a muchas personas que en realidad habrían pagado puntualmente, generando exactamente el costo de exclusión financiera. (Bahnsen, Aouada y Ottersten 2014).

La segunda pregunta es la inversa: de todas las personas que efectivamente terminaron incumpliendo, ¿cuántas había detectado a tiempo? Esa es la pregunta del Recall (Sokolova & Lapalme, 2009):

$$\text{Recall} = \frac{TP}{TP + FN}$$

Ecuación 2: Formula Recall

Un Recall bajo cuenta una historia distinta: el modelo dejó pasar solicitantes que terminaron incumpliendo, generando pérdidas crediticias directas que Bessis (2015) cuantifica como $EAD \times LGD$.

La integración de las dos métricas anteriores da lugar al F1-Score, definido como la media armónica entre Precision y Recall. A diferencia de una media aritmética convencional, la media armónica castiga con mayor severidad los valores bajos, de modo que el F1-Score únicamente alcanza puntuaciones elevadas cuando tanto la Precision como el Recall son, a su vez, individualmente altos (Géron, 2019).

$$F1 = 2 * \frac{\text{Precisión} * \text{Recall}}{\text{Precisión} + \text{Recall}}$$

Ecuación 3: Formula F1

Una vez definido el concepto de riesgo crediticio y las métricas utilizadas para cuantificarlo, resulta necesario analizar cómo las instituciones financieras traducen estos conceptos en decisiones concretas de otorgamiento de crédito. En este contexto surgen los sistemas de credit scoring, herramientas diseñadas para estimar la probabilidad de incumplimiento a partir de la información disponible sobre los solicitantes.

7.7 Modelos tradicionales de credit scoring

El credit scoring es el proceso mediante el cual una institución financiera estima el riesgo asociado a una solicitud de crédito a partir de la información disponible sobre el solicitante. Thomas, Edelman y Crook (2002) define el credit scoring como "Un conjunto de modelos de decisión y de las técnicas subyacentes que asisten a los prestamistas en la concesión de crédito al consumo" (Thomas, Edelman & Crook, 2002, Credit Scoring and Its Applications, SIAM, p. 1). En este proceso, información de diversa naturaleza se transforma en un indicador que permite clasificar a los solicitantes según su probabilidad de incumplimiento (Thomas et al., 2002).

Los primeros antecedentes del credit scoring moderno se remontan a los trabajos de Durand (1941). Posteriormente, Altman (1968) desarrolló el modelo Z-Score, basado en análisis discriminante lineal, que demostró la utilidad de los modelos cuantitativos para predecir la insolvencia empresarial. Aunque fue diseñado para empresas, este enfoque sentó las bases metodológicas para el desarrollo de los sistemas modernos de evaluación del riesgo crediticio.

7.8 Modelos estadísticos clásicos

Los modelos estadísticos más ampliamente utilizados en el credit scoring tradicional son la regresión logística (logit) y el modelo probit. Hosmer y Lemeshow (2000) describen la regresión logística como "El método preferido para el análisis de datos que involucran un resultado binario." (Hosmer & Lemeshow, 2000, p. 1), destacando que su principal ventaja reside en que los coeficientes pueden expresarse como razones de probabilidad

relativa (odds ratios), lo que facilita la comunicación con audiencias no estadísticas y la justificación de decisiones ante organismos regulatorios.

El modelo expresa la probabilidad de default como transformación sigmoide de una combinación lineal de predictoras, garantizando valores en el intervalo $[0, 1]$, propiedad que Hosmer y Lemeshow (2000) destacan como ventaja técnica fundamental.

Otro enfoque clásico es el análisis discriminante lineal de Fisher, si bien históricamente relevante en las aplicaciones iniciales de FICO (Hand & Henley, 1997), ha perdido terreno progresivamente. Hand y Henley (1997) explican esta obsolescencia señalando que el discriminante requiere normalidad multivariada y homocedasticidad, que raramente se cumplen en datos crediticios reales.

7.9 Utilización de variables tradicionales en credit scoring

Los modelos tradicionales de scoring se nutren de un conjunto relativamente estable de variables predictoras, organizadas habitualmente en las categorías del framework FICO: historial de pagos, nivel de endeudamiento, antigüedad del historial crediticio, mezcla de tipos de crédito y solicitudes recientes. La correcta selección de características es fundamental para asegurar la potencia predictiva de un sistema de puntuación. De acuerdo con Thomas et al. (2002), omitir variables de riesgo significativas impide que el *scorecard* discrimine de manera óptima entre perfiles de deudores. En sentido opuesto, la inclusión de predictores irrelevantes o distorsionados propicia el sobreajuste del algoritmo a la muestra histórica, reduciendo su eficiencia y estabilidad al evaluar a futuros solicitantes. En consecuencia, las variables seleccionadas delimitan estrictamente el alcance y la precisión de la decisión crediticia.

Aunque las variables tradicionales continúan siendo el principal insumo de los modelos de credit scoring, su eficacia se reduce cuando los solicitantes carecen de historial crediticio suficiente. Esta situación evidencia las principales limitaciones de los enfoques convencionales.

7.10 Limitaciones de los modelos tradicionales

Los modelos tradicionales exhiben un desempeño marcadamente inferior cuando se aplican a poblaciones sin historial crediticio formal, colectivo denominado *thin-file* o

credit-invisible. Jagtiani y Lemieux (2019) confirman empíricamente, mediante datos de LendingClub, que "El uso de datos alternativos ha permitido que algunos prestatarios que, según criterios tradicionales, habrían sido clasificados como de alto riesgo (*subprime*), sean asignados a categorías crediticias superiores". (Jagtiani & Lemieux, 2019, p. 1018).

En Argentina, donde el 36,7% de los asalariados trabaja informalmente (INDEC, 2025), esta limitación implica la exclusión sistemática de millones de potenciales prestatarios, no por ausencia de capacidad de pago sino por ausencia de información en las fuentes de datos convencionales

7.11 Machine learning aplicado al riesgo crediticio

Entre los algoritmos de Machine Learning [ML] más utilizados en credit scoring se destacan los árboles de decisión, los métodos de ensemble y las redes neuronales. Sin embargo, previo a analizar los modelos en específico, es pertinente hacer la distinción entre algoritmos de aprendizaje supervisado y no supervisado. Según lo hacen Hastie et al. (2009) el aprendizaje supervisado opera con datasets que contienen tanto variables predictoras como una variable objetivo conocida, que permite guiar el proceso de aprendizaje y que es la que se buscará predecir. En cambio, en el aprendizaje no supervisado no tenemos variables resultado, sino que se presentan características y el objetivo es encontrar asociaciones y patrones entre las mismas. En el dominio del credit scoring, el aprendizaje supervisado constituye el paradigma dominante, dado que el objetivo es predecir un evento binario conocido históricamente, el incumplimiento crediticio, a partir de características observables del solicitante. Los algoritmos que se presentan a continuación pertenecen todos a la categoría supervisada.

8 Principales algoritmos aplicados al credit scoring

8.1 Classification and Regression Trees (CART)

Breiman et al. (1984) desarrollaron los árboles de clasificación y regresión (CART), una técnica que segmenta los datos mediante particiones sucesivas para generar reglas de decisión fácilmente interpretables. Su principal ventaja radica en la transparencia del

proceso de clasificación, ya que cada rama del árbol representa una secuencia explícita de decisiones. Sin embargo, la literatura señala que los árboles individuales suelen presentar problemas de sobreajuste y alta varianza, lo que reduce su capacidad para generalizar adecuadamente sobre observaciones no utilizadas en el entrenamiento (Hastie et al., 2009; Müller & Guido, 2016). Esta limitación constituyó uno de los principales motivos para el desarrollo de los métodos de ensamble, que combinan múltiples árboles con el objetivo de mejorar la estabilidad y el desempeño predictivo de los modelos (James et al., 2013).

8.2 Random Forest

Como respuesta a las limitaciones de los árboles de decisión individuales, Breiman (2001) desarrolló el algoritmo Random Forest, un método de ensamble que combina múltiples árboles para obtener predicciones más robustas y estables. El modelo utiliza técnicas de *bagging* y selección aleatoria de variables, lo que reduce la correlación entre los árboles y disminuye el riesgo de sobreajuste presente en las estructuras individuales (Breiman, 1996; Breiman, 2001). Su principal ventaja radica en la mejora de la capacidad de generalización y del desempeño predictivo respecto de un único árbol de decisión. Sin embargo, al basarse en la agregación de árboles independientes, puede presentar limitaciones para capturar interacciones complejas entre variables, lo que impulsó el desarrollo posterior de métodos de *gradient boosting* capaces de modelar relaciones más sofisticadas (Hastie et al., 2009).

8.3 Gradient Boosting

A diferencia de Random Forest, que construye múltiples árboles de manera independiente para reducir la varianza, Gradient Boosting genera los árboles de forma secuencial, de modo que cada nuevo árbol busca corregir los errores cometidos por los anteriores (Friedman, 2001). Este enfoque permite reducir el sesgo del modelo e identificar relaciones no lineales e interacciones complejas entre las variables, características frecuentes en los problemas de riesgo crediticio.

Friedman (2001) desarrolló este algoritmo bajo un esquema de optimización basado en el gradiente de una función de pérdida, incorporando una tasa de aprendizaje (*learning rate*) que controla la contribución de cada árbol y ayuda a reducir el riesgo de

sobreajuste. Asimismo, Hastie et al. (2009) destacan que la posibilidad de utilizar distintas funciones de pérdida otorga al método una elevada flexibilidad para adaptarse a diferentes problemas de clasificación, entre ellos la predicción del riesgo crediticio.

8.4 XGBoost (eXtreme Gradient Boosting)

La evolución del paradigma secuencial hacia la arquitectura Extreme Gradient Boosting (XGBoost), introducida por Chen y Guestrin (2016), perfeccionó el esquema tradicional de Friedman al integrar una función de pérdida regularizada cuya optimización se instrumenta a través de una expansión de Taylor. Estas optimizaciones matemáticas permiten controlar de forma nativa la complejidad de los árboles y acelerar la convergencia del sistema al incorporar la información de la curvatura del error. En el plano operativo, el algoritmo agiliza la ejecución mediante el procesamiento en paralelo de las variables y la gestión automática de los datos faltantes. A pesar de estas optimizaciones de software, el modelo encuentra un límite crítico ante grandes volúmenes de información. La necesidad de ordenar y examinar exhaustivamente cada valor numérico continuo para identificar el punto de corte óptimo genera un cuello de botella computacional severo ante datasets masivos.

8.5 LightGBM (Light Gradient Boosting Machine)

Los hallazgos empíricos de Lessmann et al. (2015) confirmaron la superioridad consistente de los métodos de ensamble frente a las técnicas lineales convencionales, situando a estas variantes secuenciales en la frontera de eficiencia de la industria financiera. La solvencia analítica de estas arquitecturas descansa en su capacidad para capturar relaciones no lineales complejas y asimilar aquellas señales débiles de comportamiento económico que los modelos paramétricos tradicionales descartan por completo como ruido informativo.

Sin embargo como se detalla en la sección previa, correspondiente al modelo XGBoost, la viabilidad operativa de estos sistemas encontró un límite crítico ante datasets masivos debido a la carga computacional requerida. La superación de este cuello de botella computacional se materializó con la introducción de LightGBM, diseñado por Ke et al. (2017). A diferencia del ordenamiento y análisis de corte exacto de variables continuas propio de XGBoost, esta arquitectura sustituye el procesamiento de valores individuales por un algoritmo basado en histogramas discretos. El mapeo de características se

condensa en Intervalos acotados. Dicha discretización reduce el consumo de memoria y acelera los tiempos de ejecución. En el plano estructural el modelo complementa esta ventaja, desestimando el desarrollo simétrico por niveles o depth-wise adoptado por las arquitecturas de ensamble convencionales y en su lugar adopta una estrategia de crecimiento por hojas o leaf-wise, la cual elige la bifurcación que maximiza la reducción del error global sin la obligatoriedad de balancear simétricamente cada nivel del árbol. Lo que, por sí mismo, ya implica la optimización de recursos y la aceleración del procesamiento. Sin embargo, esto, se termina de consolidar formalmente mediante dos técnicas de optimización complementarias. Para disminuir el volumen de observaciones sin degradar la precisión, el algoritmo ejecuta un muestreo unilateral basado en gradientes o GOSS, el cual preserva las instancias con grandes errores y realiza un muestreo aleatorio sobre aquellas con gradientes pequeños. En paralelo, para gestionar la alta dispersión y la presencia de variables cualitativas, el modelo aplica el empaquetado de características exclusivas o EFB. Este segundo mecanismo fusiona atributos mutuamente excluyentes en vectores densos individuales; esta fusión disminuye la dimensionalidad del espacio de búsqueda sin resignar capacidad predictiva. Por ello, debido a su eficiencia en términos de velocidad y bajo costos de procesamiento y su probada robustez estadística para mitigar el sobreajuste, sin perder capacidad para detectar relaciones complejas, este algoritmo constituye la herramienta metodológica seleccionada para la fase de evaluación cuantitativa y aplicabilidad del presente trabajo de investigación.

9 Datos alternativos en credit scoring

Los datos alternativos refieren a cualquier fuente de información sobre el comportamiento de un individuo que no forma parte de los registros crediticios mantenidos por los burós de crédito convencionales.

Jagtiani y Lemieux (2019), a partir del análisis de datos de LendingClub correspondientes al período 2007-2015, encontraron que los prestamistas fintech han reducido progresivamente su dependencia de los puntajes FICO tradicionales al incorporar fuentes de datos alternativas en sus evaluaciones crediticias. Los autores observaron que esta información adicional permite identificar de manera más precisa la solvencia de ciertos solicitantes que, bajo los criterios convencionales, habrían sido clasificados como

prestatarios de alto riesgo. Como consecuencia, algunos de estos individuos pueden acceder a categorías crediticias más favorables y obtener financiamiento en mejores condiciones. Los hallazgos sugieren que los datos alternativos aportan información complementaria sobre la calidad crediticia que no es capturada por los sistemas tradicionales de información crediticia, contribuyendo así a una mayor inclusión financiera de segmentos previamente desatendidos (Jagtiani & Lemieux, 2019).

10 Ventajas frente a los modelos tradicionales

Berg et al. (2020), en un estudio empírico sobre ML y datos alternativos para credit scoring, cuantifican con precisión las ventajas de los modelos de ML sobre los sistemas basados en burós de crédito: En su investigación, evidencian que el uso de la huella digital ofrece una perspectiva informativa que enriquece los registros tradicionales, mejorando la selección de perfiles y disminuyendo el riesgo de impago. Desde una perspectiva métrica, un modelo estimado puramente con variables digitales simples logra un incremento de 5.3 puntos porcentuales en el AUC en comparación con el *score* de crédito habitual. Asimismo, al combinar ambos conjuntos de datos, la capacidad de discriminación (AUC) se eleva en 6.8 puntos porcentuales frente al uso exclusivo del buró. Con ello, los autores concluyen que el entorno digital captura dimensiones críticas de la capacidad y comportamiento de pago del consumidor que no están correlacionadas con las bases de datos de crédito habituales.

11 Inclusión financiera y contexto argentino

La profundidad del sistema financiero argentino, medida como el ratio de crédito al sector privado sobre el PBI, se ubica en torno al 13%, significativamente por debajo del promedio latinoamericano y muy alejada de los niveles de economías desarrolladas (FMI, 2023, *Argentina: Financial Sector Assessment*, IMF Country Report No. 23/168). Esta baja profundidad refleja tanto restricciones de oferta, imitaciones en los sistemas de evaluación crediticia para poblaciones informales, como factores de demanda vinculados a la desconfianza histórica en el sistema financiero.

La tasa de informalidad laboral es el condicionante estructural más relevante de la exclusión crediticia. La tasa de informalidad laboral constituye uno de los principales

factores estructurales que restringen el acceso al crédito formal en Argentina. En este sentido, el INDEC (2025) informa que el 36,7 % de los asalariados se desempeña en condiciones de informalidad. De acuerdo con la OECD (2025), la magnitud de este fenómeno limita la inclusión financiera, ya que los trabajadores informales suelen carecer de la documentación requerida por los modelos tradicionales de evaluación crediticia, como recibos de sueldo, contratos laborales o registros de aportes a la seguridad social. Como consecuencia, muchas personas quedan excluidas del sistema financiero formal independientemente de su verdadera capacidad de pago. La organización también destaca que esta exclusión contribuye a perpetuar desigualdades económicas y sociales, y señala que el uso de datos alternativos y técnicas de aprendizaje automático aplicadas al *credit scoring* constituye una de las alternativas más prometedoras para ampliar el acceso al crédito de estos sectores.

La convergencia de alta informalidad, baja profundidad financiera y desconfianza histórica crea un escenario en el que los datos alternativos y los modelos de ML tienen un potencial transformador especialmente elevado: al ampliar el conjunto de señales evaluables más allá del historial crediticio bancario, estos sistemas pueden hacer visible, en términos de scoring, a una porción significativa de la población que los mecanismos tradicionales mantienen excluida. Gambacorta, Huang, Qiu y Wang (2019) demuestran que el impacto de los modelos de machine learning alimentados con datos no tradicionales es particularmente significativo en mercados emergentes, donde la proporción de población sin historial crediticio formal es estructuralmente mayor que en economías desarrolladas, ampliando así la base de clientes evaluables sin necesidad de que estos cuenten con un historial bancario previo.

No obstante, la posibilidad de ampliar el acceso al crédito mediante la utilización de datos alternativos y modelos avanzados de evaluación crediticia plantea desafíos que exceden las cuestiones puramente técnicas. La incorporación de nuevas fuentes de información requiere que las entidades financieras accedan, almacenen y procesen datos que, en muchos casos, no formaban parte de los esquemas tradicionales de evaluación del riesgo. En consecuencia, el desarrollo de estas herramientas debe analizarse en conjunto con los marcos regulatorios vigentes y las normas de protección de datos personales, ya que dichos factores condicionan tanto la disponibilidad de la información

como las posibilidades de implementación efectiva de estos modelos en el sistema financiero.

12 Justificación del uso de datasets sintéticos en investigación sobre credit scoring

Los datasets sintéticos son conjuntos de datos generados artificialmente mediante procedimientos estadísticos o algorítmicos que reproducen las propiedades de los datos reales sin contener información de individuos específicos. Jordon et al. (2022) destacan que los datos sintéticos permiten preservar la privacidad, mejorar la reproducibilidad y modelar escenarios hipotéticos no observables en datos reales.

En el campo del credit scoring, la investigación empírica enfrenta una paradoja estructural: los datos más valiosos para el avance científico están en manos de las instituciones financieras, con incentivos limitados para compartirlos. A diferencia de otras áreas donde existen repositorios públicos de referencia, el campo del credit scoring dispone de un número reducido de datasets públicos (como son el German Credit Dataset, Australian Credit Dataset, entre otros) con características muy distintas del sistema financiero argentino. Esta escasez impone restricciones metodológicas severas a la investigación académica local y justifica el recurso a metodologías de simulación.

13 Validez metodológica y aplicación al contexto argentino

Si bien los datasets sintéticos surgen como una alternativa metodológica frente a las limitaciones de acceso a información crediticia real, su utilización plantea interrogantes respecto de la capacidad de estos datos para reproducir adecuadamente los patrones y relaciones observados en la realidad. En este sentido, la utilidad de un conjunto de datos sintético no depende únicamente de su disponibilidad, sino también de su capacidad para representar de manera consistente las características de la población bajo estudio y preservar las relaciones estadísticas relevantes para el fenómeno analizado. Por ello, resulta necesario examinar los criterios que sustentan la validez metodológica de este enfoque y analizar su potencial aplicación en investigaciones orientadas al sistema financiero argentino.

Walonoski et al. (2018) sostienen que la validez de los datos sintéticos depende de que reproduzcan adecuadamente las características estadísticas de los datos reales, permitan desarrollar modelos predictivos con un desempeño comparable y reduzcan el riesgo de identificación de individuos. El cumplimiento de estos criterios respalda su utilización en investigaciones aplicadas.

En el contexto de la presente investigación, el uso de datasets sintéticos se justifica por la convergencia de tres factores: la ausencia de microdatos públicos representativos del sistema financiero argentino; las restricciones regulatorias sobre el uso de datos reales en investigación académica; y la necesidad de simular condiciones macroeconómicas particulares del entorno local. La estrategia de calibración adoptada, parametrizar las distribuciones a partir de estadísticas agregadas del BCRA, INDEC y fuentes de la industria fintech argentina, confiere al enfoque una fundamentación metodológica sólida. Los datos sintéticos representan en esta investigación no una limitación sino una solución deliberada que permite avanzar en el conocimiento sin comprometer la privacidad ni depender de accesos institucionales restringidos.

14 Metodología

14.1 Enfoque metodológico

La presente investigación adopta un enfoque cuantitativo, con enfoque exploratorio, orientado a evaluar el desempeño de técnicas de aprendizaje automático aplicadas al riesgo crediticio en el contexto argentino. El objetivo consiste en comparar la capacidad predictiva de un modelo estadístico tradicional, la Regresión Logística, y un algoritmo de ensamble basado en árboles de decisión, LightGBM, utilizando un conjunto de datos sintético construido para representar las principales características del mercado de crédito local.

La elección de un enfoque cuantitativo responde a la necesidad de medir objetivamente los resultados obtenidos en ambos modelos mediante indicadores estadísticos de desempeño. Sin embargo, la investigación enfrenta una limitación relevante: la imposibilidad de acceder a bases de datos individuales reales debido al secreto bancario, las restricciones regulatorias y la inexistencia de repositorios públicos que contengan información crediticia detallada de personas físicas.

Frente a esta limitación, se desarrolló un Proceso Generador de Datos (Data Generating Process, DGP) capaz de construir un dataset sintético calibrado a partir de información proveniente de fuentes oficiales, principalmente el Instituto Nacional de Estadística y Censos (INDEC), el Banco Central de la República Argentina (BCRA) y la Cámara Argentina Fintech. Este procedimiento permite reproducir relaciones económicas y financieras observadas en la realidad sin contar con la información correspondiente a individuos reales.

El dataset sintético constituye el insumo utilizado para entrenar y evaluar ambos modelos predictivos. De esta manera, la comparación entre algoritmos se realiza bajo idénticas condiciones de información, permitiendo analizar las ventajas y limitaciones de cada metodología en la predicción del incumplimiento crediticio.

14.2 Construcción del dataset sintético

El conjunto de datos fue generado íntegramente mediante un Proceso Generador de Datos (DGP) desarrollado en Python. Su diseño busca reproducir las principales características socioeconómicas y financieras de la población argentina potencialmente sujeta a crédito formal. El algoritmo genera inicialmente las variables exógenas mediante distribuciones probabilísticas previamente parametrizadas y, posteriormente, construye el resto de las variables respetando las relaciones de dependencia e interacción definidas en el modelo. De este modo, cada observación representa un individuo con características consistentes desde el punto de vista económico y estadístico, preservando la coherencia entre variables como formalidad laboral, ingresos, gastos, ahorro, endeudamiento y comportamiento financiero. Finalmente, a partir de estas variables se calcula la probabilidad de incumplimiento y se asigna la variable objetivo (*default*), completando la generación del conjunto de datos utilizado para el entrenamiento y evaluación de los modelos.

La población de referencia corresponde a personas mayores de 18 años residentes en Argentina con posibilidad de acceder al sistema financiero. El dataset final contiene 100.000 observaciones, cantidad que permite entrenar modelos de machine learning con suficiente robustez estadística y representar adecuadamente la heterogeneidad de la población analizada.

Las variables incorporadas pueden agruparse en cuatro grandes dimensiones:

Variables demográficas: edad, provincia de residencia, nivel educativo y carga familiar.

Variables laborales: condición de formalidad laboral, antigüedad en el empleo, ingresos y volatilidad de ingresos.

Variables financieras: gastos mensuales, ahorro líquido, relación deuda-ingreso (Debt Service Ratio), atrasos créditos recientes, uso de crédito informal y posesión de tarjeta de crédito.

Variables de comportamiento digital: utilización de billeteras digitales, uso de efectivo e índice de calidad del dispositivo móvil.

La lógica para la modelización de cada variable, su descripción, tipo de dato, la fuente utilizada para su parametrización y las principales relaciones con el resto de las variables; Se encuentran en la tabla “Modelización de variables” del Anexo 2. Allí mismo podrá acceder, también, al código desarrollado en Python, (Anexo 3), con el propósito de facilitar la reproducibilidad del estudio sin sobrecargar el desarrollo metodológico.

Finalmente, se incorpora la variable objetivo default, correspondiente al incumplimiento crediticio, cuya construcción se realizó mediante un mecanismo probabilístico basado en un score latente.

14.3 Construcción de la variable objetivo

La variable de incumplimiento (default) constituye la variable dependiente del estudio y representa el evento que los modelos intentan predecir.

Su generación no se realizó mediante una asignación aleatoria, sino a través de un procedimiento secuencial compuesto por tres etapas.

En primer lugar, se construyó un score latente que integra variables económicas, laborales, financieras y digitales mediante relaciones lineales y no lineales. Posteriormente, dicho score fue transformado en una probabilidad de incumplimiento utilizando una función sigmoide. Finalmente, la asignación de la variable binaria de default se realizó mediante un ensayo de Bernoulli, preservando la naturaleza probabilística propia del riesgo crediticio.

Una vez obtenidos los resultados para la variable objetivo, se le aplicó un procedimiento iterativo de sesgo acumulativo calibración con el objetivo de que la tasa global de incumplimiento convergiera hacia el valor definido durante la parametrización.

La descripción detallada del algoritmo utilizado para construir el score latente y las funciones implementadas puede consultarse en el Anexo 3.

14.4 Entrenamiento de los modelos

Al contar con dataset de la muestra ya completamente estructurado. Este fue dividido en un conjunto de entrenamiento (80%) y un conjunto de prueba (20%), preservando la proporción de casos de incumplimiento en ambos subconjuntos.

Sobre el conjunto de entrenamiento se estimaron dos modelos predictivos: Regresión Logística, utilizada como modelo base (baseline) y LightGBM, algoritmo de Gradient Boosting basado en árboles de decisión.

En el caso de LightGBM se aplicó validación cruzada de cinco particiones (5-fold cross validation) durante el proceso de entrenamiento, con el fin de reducir el riesgo de sobreajuste y obtener una estimación más robusta del desempeño del modelo.

Una vez entrenados, ambos algoritmos fueron evaluados exclusivamente sobre el conjunto de prueba, con observaciones no utilizadas durante el entrenamiento.

El código correspondiente al entrenamiento, validación y evaluación de ambos modelos, nuevamente se detalla en el Anexo 3.

14.5 Estrategia de análisis

La estrategia de análisis se desarrolló en cuatro etapas sucesivas.

En primer lugar, se verificó que el dataset sintético reprodujera adecuadamente los parámetros utilizados durante su construcción, contrastando las principales distribuciones con la información proveniente de las fuentes oficiales empleadas en la parametrización.

Posteriormente, se realizó un análisis exploratorio mediante matrices de correlación con el objetivo de identificar inconsistencias estadísticas y validar que las relaciones observadas entre las variables fueran consistentes con la teoría económica y la evidencia empírica.

En una tercera etapa se entrenaron los modelos predictivos sobre el conjunto de entrenamiento y se evaluó su desempeño utilizando el conjunto de prueba.

Finalmente, se comparó el rendimiento de ambos modelos mediante distintas métricas de clasificación. El indicador principal utilizado fue el Área Bajo la Curva ROC (AUC-ROC), complementado con el coeficiente de Gini, el estadístico Kolmogorov-Smirnov (KS), matrices de confusión, precisión, recall y F1-score.

Adicionalmente, con el propósito de interpretar el comportamiento del modelo LightGBM, se empleó el método SHAP (SHapley Additive exPlanations), que permite cuantificar la contribución individual de cada variable sobre la probabilidad estimada de incumplimiento y analizar las relaciones no lineales capturadas por el algoritmo de ensamble.

15 Resultados de los modelos

15.1 Regresión Logística

La Regresión Logística obtuvo un AUC-ROC de 0,8093, un coeficiente de Gini de 0,6186 y un estadístico KS de 0,4800. Estos resultados evidencian una adecuada capacidad para discriminar entre clientes con y sin probabilidad de incumplimiento. La matriz de confusión muestra que el modelo clasificó correctamente 11.780 clientes No Default y 2.535 clientes Default, mientras que registró 5.017 falsos positivos y 668 falsos negativos, sobre un conjunto de prueba de 20.000 observaciones.

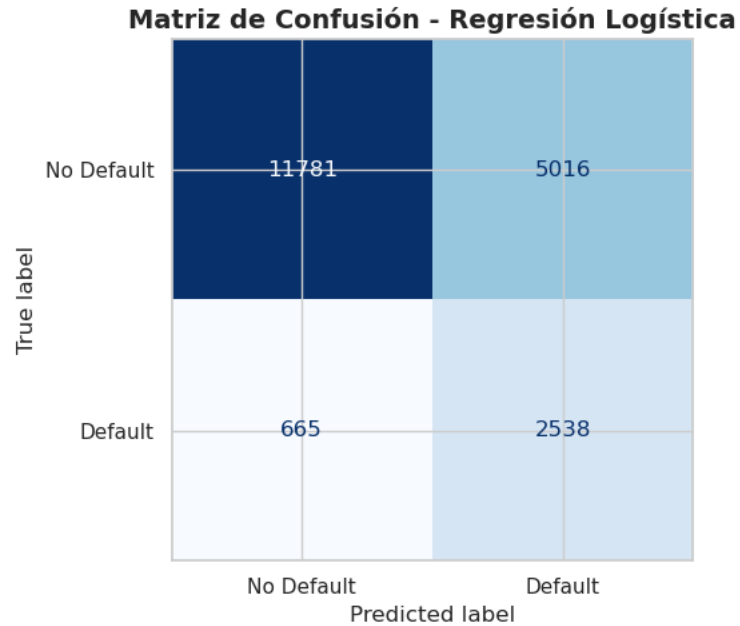


Ilustración 2: Extraída código Python – Anexo 3

El reporte de clasificación indica que la clase No Default alcanzó una precisión de 0,95, un recall de 0,70 y un F1-Score de 0,81, mientras que la clase Default obtuvo una precisión de 0,34, un recall de 0,79 y un F1-Score de 0,47. Estos resultados muestran que el modelo identifica correctamente una elevada proporción de clientes incumplidores, aunque a costa de clasificar como riesgosos a un número considerable de clientes solventes. Este comportamiento es esperable en modelos de evaluación del riesgo crediticio que priorizan la detección del default, ya que los costos de un falso negativo (no detectar un incumplidor) superan los costos de un falso positivo (rechazar un cliente solvente).

index	Precision	Recall	F1-Score
No Default (0)	0.95	0.70	0.81
Default (1)	0.34	0.79	0.47
AUC-ROC			0.8093

Ilustración 3: Extraída código Python – Anexo 3

15.2 Resultados del modelo LightGBM

El modelo LightGBM obtuvo un AUC-ROC de 0,8263, un coeficiente de Gini de 0,6526 y un estadístico KS de 0,5100. Estos resultados evidencian una buena capacidad para

discriminar entre clientes con y sin probabilidad de incumplimiento. La matriz de confusión muestra que el modelo clasificó correctamente 12.419 clientes No Default y 2.460 clientes Default, mientras que registró 4.378 falsos positivos y 743 falsos negativos, sobre el mismo conjunto de prueba de 20.000 observaciones.

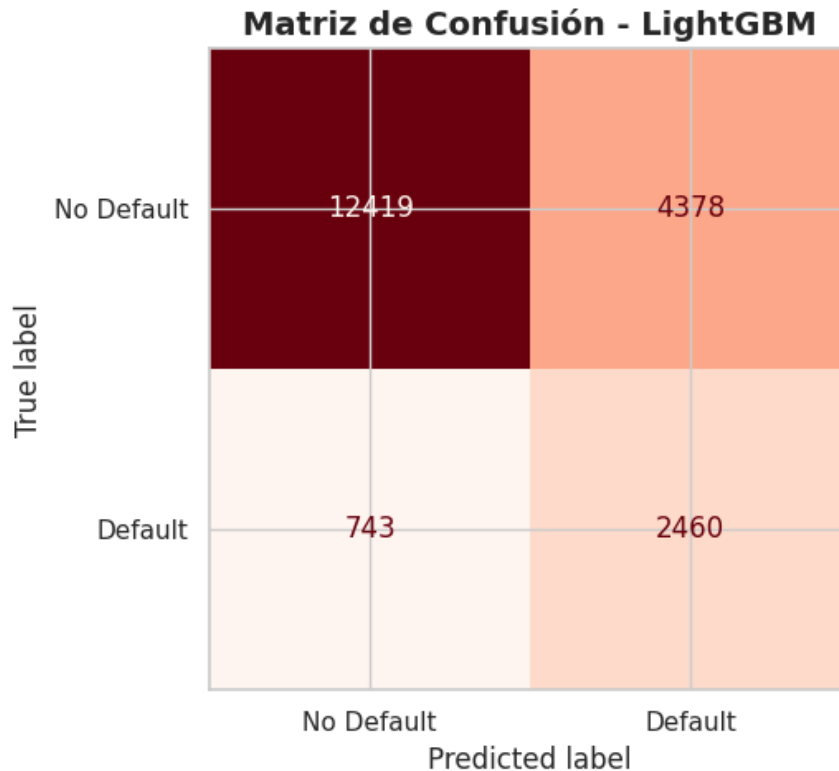


Ilustración 4: Extraída código Python – Anexo 3

El reporte de clasificación indica que la clase No Default alcanzó una precisión de 0,94, un recall de 0,74 y un F1-Score de 0,83, mientras que la clase Default obtuvo una precisión de 0,36, un recall de 0,77 y un F1-Score de 0,49. El modelo mantiene una elevada capacidad para identificar clientes con riesgo de incumplimiento, al tiempo que reduce la cantidad de falsos positivos respecto de la Regresión Logística, permitiendo clasificar correctamente un mayor número de clientes solventes. Las métricas obtenidas reflejan un desempeño predictivo sólido y evidencian la capacidad del modelo para capturar relaciones más complejas entre las variables alternativas del dataset sintético.

	Precision	Recall	F1-Score
No Default (0)	0.94	0.74	0.83
Default (1)	0.36	0.77	0.49
AUC-ROC			0.8263

Ilustración 5: Extraída código Python – Anexo 3

15.3 Comparación entre ambos modelos

La comparación entre ambos modelos evidencia que tanto la Regresión Logística como LightGBM presentan un desempeño satisfactorio para la evaluación del riesgo crediticio sobre el dataset sintético desarrollado. Sin embargo, LightGBM obtuvo mejores resultados en las principales métricas de desempeño.

LightGBM alcanzó un AUC-ROC de 0,8263 frente al 0,8093 de la Regresión Logística, representando una mejora de 0,017 puntos (1,7%). El coeficiente de Gini también mejoró: 0,6526 frente a 0,6186 (3,4 puntos de ganancia). El estadístico KS mostró incremento de 0,4800 a 0,5100 (3,0 puntos). En la clase No Default, LightGBM mejoró el recall (0,74 frente a 0,70) y el F1-Score (0,83 frente a 0,81), permitiendo identificar correctamente un mayor número de clientes solventes y reduciendo falsos positivos de 5.017 a 4.378 (639 menos). Para la clase Default, mantuvo un nivel de recall similar (0,77 frente a 0,79) y mejoró la precisión (0,36 frente a 0,34) y el F1-Score (0,49 frente a 0,47), reflejando un mejor equilibrio entre la detección de incumplimientos y la reducción de falsos positivos.

Estos resultados indican que el modelo basado en machine learning aprovecha de manera más eficiente la información contenida en las variables alternativas del dataset sintético, obteniendo una mejora consistente respecto de la Regresión Logística. Aunque las diferencias son moderadas, respaldan la utilización de algoritmos de ensamble como una herramienta complementaria para la evaluación del riesgo crediticio de personas sin historial financiero formal en el contexto del sistema financiero argentino. La mejora en Gini (3,4 puntos) es particularmente significativa porque captura ganancia de información sobre el baseline, sugiriendo que LightGBM captura no-linealidades que Logística no puede modelar.

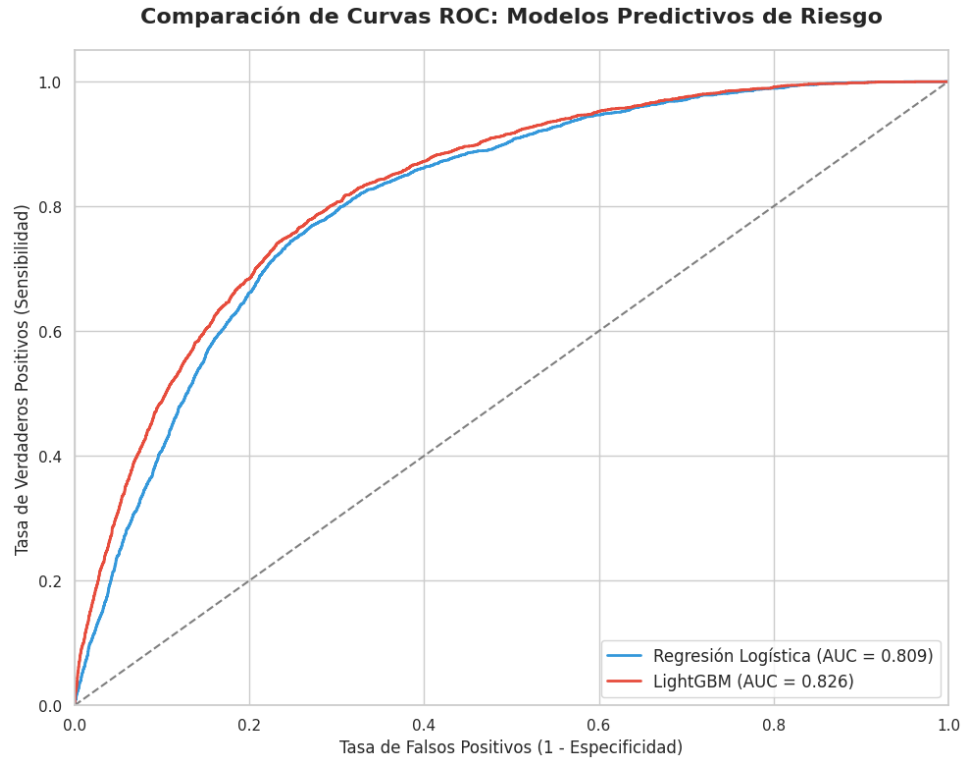


Ilustración 6: Extraída código Python – Anexo 3

16 Interpretación de la importancia de las variables

Como se mencionó anteriormente en la metodología es importante realizar una descomposición de la varianza del modelo no paramétrico. La cual, en nuestro caso, se realizó mediante el proceso de valores SHAP, este permite examinar cómo impacta cada variable predictora en la probabilidad latente de default. El gráfico de resumen revela una jerarquía clara de riesgo y, fundamentalmente, confirma la presencia de dinámicas no lineales y efectos de interacción que la regresión logística tradicional no puede asimilar debido a su rigidez matemática.

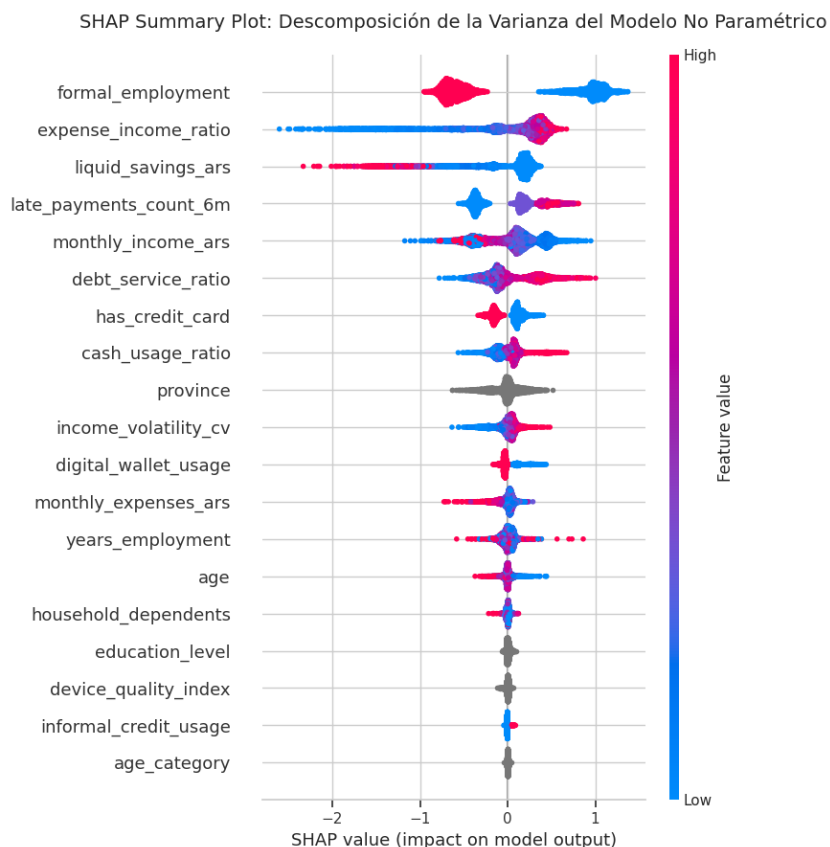


Ilustración 7: Gráfico resumen de análisis SHAP – Anexo 3

La condición laboral del deudor, representada por la variable `formal_employment`, es el principal vector de discriminación de riesgo en el sistema. El agrupamiento de la muestra exhibe un comportamiento binario muy marcado. Los registros asociados a la informalidad (valores bajos en azul) presentan valores SHAP positivos y elevados, lo que significa que la ausencia de un empleo registrado desplaza de forma directa el score hacia la zona de default. Por el contrario, la presencia de un empleo en blanco (valores altos en rojo) opera a la izquierda del gráfico como un fuerte mitigador del riesgo. Este comportamiento valida la hipótesis central sobre la fractura estructural del mercado de trabajo argentino, donde el sector informal concentra los mayores ratios de vulnerabilidad financiera.

El ratio entre gastos e ingresos (`expense_income_ratio`) ocupa el segundo lugar en importancia y es un claro ejemplo de la asimetría que capturan los modelos basados en árboles. La distribución muestra una extensa cola azul hacia el extremo izquierdo del gráfico, con valores SHAP inferiores a -2.0. Esto indica que mantener una estructura de gastos fijos muy baja respecto a los ingresos mensuales genera un blindaje crítico y no

lineal contra el impago. A medida que el ratio de consumo se eleva (puntos fucsias y rojos), el impacto se traslada hacia el cuadrante derecho, incrementando la probabilidad de default a medida que el deudor pierde capacidad de absorción ante shocks económicos.

El análisis de los activos líquidos (`liquid_savings_ars`) complementa esta dinámica de flujos. La acumulación de reservas en pesos (puntos rojos) reduce el riesgo de manera consistente, ubicándose del lado izquierdo. Sin embargo, la mayor densidad de la muestra se concentra en un bloque azul del lado derecho, lo que demuestra que la falta de ahorros líquidos (`liquid_savings_ars`= 0) funciona como un detonante del incumplimiento. Esto se asocia directamente con la variable de comportamiento transaccional `late_payments_count_6m`, donde el historial de mora reciente exhibe una penalización escalonada sobre el ahorro. La ausencia de atrasos (puntos azules) mitiga el riesgo, pero la acumulación de incidentes de pago en los últimos seis meses (puntos rojos) empuja el score con fuerza hacia la derecha. LightGBM procesa esta variable bajo un criterio de saturación marginal decreciente, donde los primeros atrasos aportan mayor sorpresa informativa que los subsiguientes.

Las variables alternativas introducidas para identificar al segmento sub-bancarizado muestran un comportamiento predictivo claro y justifican la adopción de metodologías de aprendizaje automático. El uso de efectivo (`cash_usage_ratio`) y la adopción de billeteras digitales (`digital_wallet_usage`) operan de manera contrapuesta y asimétrica. Una alta dependencia del dinero físico (puntos rojos en `cash_usage_ratio`) eleva el riesgo de default, ya que la falta de trazabilidad aísla al deudor de los circuitos de asistencia financiera formal. Como contrapartida, la adopción de billeteras virtuales (puntos rojos en `digital_wallet_usage`) se extiende hacia la izquierda. Esto demuestra que la digitalización de los cobros y pagos en el sector informal actúa como una señal alternativa de inclusión y estabilidad financiera. LightGBM aprovecha esta interacción para identificar subgrupos de menor riesgo dentro de la informalidad, quiebre no lineal que la regresión logística tradicional licúa.

La volatilidad de los ingresos (`income_volatility_cv`) y la carga de deuda fáctica (`debt_service_ratio`) se posicionan como aceleradores directos del default cuando toman valores elevados (puntos rojos a la derecha). Esta dispersión confirma que la inestabilidad en los flujos de caja mensuales impide una planificación financiera regular, potenciando el impacto de la carga de deuda sobre el saldo de consumo. Finalmente,

variables demográficas como la edad (age) y la composición del hogar (household_dependents) muestran impactos moderados pero consistentes con la teoría del capital humano. La juventud (puntos azules en age) y una mayor cantidad de familiares a cargo (puntos rojos en household_dependents) se desplazan hacia la derecha, reflejando que la menor estabilidad laboral temprana y las mayores exigencias presupuestarias fijas de los hogares vulnerables erosionan de forma constante la capacidad de repago del crédito.

La correspondencia entre la importancia del gráfico SHAP y las reglas del Proceso Generador de Datos (PGD) confirma la validez del diseño metodológico, al tiempo que dota de interpretabilidad al algoritmo ante las exigencias de transparencia del sistema financiero. La jerarquía de las variables principales refleja que las ecuaciones condicionales se trasladaron sin distorsiones al espacio de entrenamiento. Esto asegura que LightGBM opere de forma auditable y responda de manera directa a la estructura macroeconómica inyectada en el origen

17 Limitaciones

La principal limitación de esta investigación radica en la utilización de un dataset sintético en lugar de información crediticia real. Si bien el Proceso Generador de Datos fue calibrado utilizando parámetros provenientes de fuentes oficiales y diseñado para reproducir relaciones económicas observables, los resultados deben interpretarse dentro del contexto de la muestra generada. Asimismo, al incorporar relaciones no lineales entre las variables, es posible que modelos como LightGBM, capaces de capturar este tipo de patrones, presenten un desempeño relativamente superior respecto de modelos lineales.

Por otra parte, aunque LightGBM obtuvo mejores resultados que la Regresión Logística en todas las métricas analizadas, las diferencias observadas fueron moderadas. Esto indica que, si bien el algoritmo de machine learning mostró un mejor desempeño, la ganancia predictiva sobre el modelo tradicional no fue sustancial.

Finalmente, la implementación práctica de este tipo de modelos requiere considerar aspectos regulatorios vinculados al uso de datos alternativos y a la protección de datos personales, procurando garantizar el cumplimiento del marco normativo vigente y la transparencia de los procesos de decisión automatizados.

18 Conclusión

La presente investigación tuvo como objetivo analizar la aplicabilidad de modelos de machine learning y datos alternativos para la evaluación del riesgo crediticio en individuos con escaso o nulo historial financiero formal en Argentina. Los resultados obtenidos permiten concluir que este enfoque constituye una alternativa viable para complementar los sistemas tradicionales de credit scoring, especialmente en un contexto donde la informalidad laboral alcanza al 36,7% de los asalariados y limita el acceso al crédito de una proporción significativa de la población.

La construcción de un dataset sintético calibrado a partir de información del INDEC, BCRA, CEPAL y la Cámara Argentina Fintech permitió evaluar el comportamiento de distintas arquitecturas predictivas bajo condiciones representativas del contexto argentino. Esta construcción se sustentó en tres criterios centrales: la calibración de los parámetros contra fuentes oficiales verificables, la preservación de las relaciones económicas y teóricas entre variables, y la validación posterior de esa consistencia mediante el contraste de los resultados agregados contra los datos reales utilizados en la calibración. Sobre esta base, tanto la Regresión Logística como LightGBM demostraron una capacidad discriminatoria satisfactoria, obteniendo valores de AUC-ROC superiores a 0,80. No obstante, el modelo LightGBM presentó un desempeño superior, alcanzando un AUC-ROC de 0,8263 frente a 0,8093 de la Regresión Logística, además de mejoras en los indicadores Gini y KS, lo que evidencia una mayor capacidad para capturar patrones complejos de riesgo crediticio. En este sentido, el dataset sintético permitió analizar la aplicabilidad de los modelos de machine learning de manera acotada pero informativa: acotada, porque los resultados dependen de la estructura de relaciones definida en el proceso generador de datos y no de comportamiento observado; e informativa, porque replicó condiciones plausibles del mercado argentino y permitió contrastar de forma controlada el desempeño relativo de ambas arquitecturas bajo idénticas condiciones de información.

Los resultados también sugieren que las variables alternativas incorporadas poseen contenido informativo relevante para la predicción del incumplimiento. De acuerdo con el análisis SHAP, las variables alternativas más adecuadas para representar el comportamiento crediticio de individuos sin historial financiero formal fueron la formalidad

laboral, el ratio gasto-ingreso, los ahorros líquidos, el historial reciente de atrasos y el uso de billeteras digitales combinado con la calidad del dispositivo móvil, en tanto lograron diferenciar perfiles de riesgo incluso dentro del segmento informal, donde los modelos tradicionales carecen de señales. Asimismo, el mejor desempeño de LightGBM respalda la existencia de relaciones no lineales e interacciones entre estas variables que los modelos lineales tradicionales no logran representar completamente.

Estos hallazgos contribuyen a la literatura sobre inclusión financiera al proporcionar evidencia de que los datos alternativos pueden ampliar la información disponible para evaluar a individuos históricamente excluidos de los mecanismos convencionales de análisis crediticio. En este sentido, el aporte principal del trabajo no radica únicamente en la mejora de las métricas predictivas, sino en la posibilidad de desarrollar herramientas de evaluación más adecuadas para poblaciones con limitada trazabilidad financiera.

Sin embargo, los resultados deben interpretarse considerando las limitaciones del estudio. La investigación se desarrolló sobre un dataset sintético, por lo que futuras investigaciones deberían contrastar estos hallazgos utilizando bases de datos reales anonimizadas y profundizar el análisis de aspectos vinculados a explicabilidad, sesgos algorítmicos y regulación del uso de datos alternativos en procesos de decisión crediticia.

En definitiva, la evidencia obtenida indica que la combinación de machine learning y datos alternativos puede mejorar la evaluación del riesgo crediticio en contextos de alta informalidad como el argentino. Si bien aún existen desafíos regulatorios, operativos y metodológicos para su implementación, los resultados sugieren que estas herramientas poseen el potencial de contribuir simultáneamente a una mejor gestión del riesgo por parte de las entidades financieras y a una mayor inclusión financiera de sectores actualmente excluidos del crédito formal.

19 Futuras líneas de investigación

A partir de los resultados obtenidos, futuras investigaciones podrían profundizar el análisis mediante la incorporación de nuevas fuentes de datos alternativas y el estudio del desempeño de los modelos bajo diferentes escenarios económicos. Asimismo, resulta de interés avanzar en el desarrollo de metodologías que fortalezcan la

interpretabilidad y permitan evaluar potenciales sesgos algorítmicos, contribuyendo a la construcción de sistemas de evaluación crediticia más robustos, transparentes y acordes con las exigencias del sistema financiero.

20 **Bibliografía**

Adeline, A., Choiniere-Crevecoeur, I., Fonseca, R. & Michaud, P.C. (2019). *Income Volatility, Health and Well-Being*. IZA Discussion Paper No. 12823. Institute of Labor Economics (IZA). https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3503773

Akerlof, G.A. (1970). *The Market for 'Lemons': Quality Uncertainty and the Market Mechanism*. The Quarterly Journal of Economics. Vol. 84, No. 3, pp. 488–500. MIT Press / Oxford University Press. <https://doi.org/10.2307/1879431>

Altman, E.I. (1968). *Financial Ratios, Discriminant Analysis and the Prediction of Corporate Bankruptcy*. The Journal of Finance. Vol. 23, No. 4, pp. 589–609. American Finance Association / Wiley. <https://doi.org/10.2307/2978933>

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. Vol. 54, No. 6, pp. 627–635. Palgrave Macmillan. <https://doi.org/10.1057/palgrave.jors.2601545>

Bahnsen, A.C., Aouada, D. & Ottersten, B. (2014). *Example-Dependent Cost-Sensitive Logistic Regression for Credit Scoring*. Proceedings of the 13th IEEE International Conference on Machine Learning and Applications (ICMLA). pp. 263–269. IEEE. <https://doi.org/10.1109/ICMLA.2014.48>

Banco Central de la República Argentina (BCRA). (2026). *Informe sobre Bancos*. Buenos Aires: BCRA. Publicación mensual. <https://www.bcra.gob.ar/publicaciones/informe-sobre-bancos-marzo-de-2026/>

Banco Central de la República Argentina (BCRA). (2026). *Informe de Proveedores No Financieros de Crédito*, Buenos Aires: BCRA. Publicación semestral. <https://www.bcra.gob.ar/publicaciones/informe-de-proveedores-no-financieros-de-credito-junio-de-2026/>

Basel Committee on Banking Supervision (BCBS). (2006). *International Convergence of Capital Measurement and Capital Standards: A Revised Framework*. Bank for International Settlements. Basilea, Suiza. <https://www.bis.org/publ/bcbs128.htm>

Berg, T., Burg, V., Gombović, A. & Puri, M. (2020). *On the Rise of FinTechs: Credit Scoring Using Digital Footprints*. The Review of Financial Studies. Vol. 33, No. 7, pp. 2845–2897. Oxford University Press. <https://doi.org/10.1093/rfs/hhz099>

Bessis, J. (2015). *Risk Management in Banking*. 4^a ed. John Wiley & Sons. Hoboken, NJ. ISBN: 978-1-118-66021-8

Breiman, L. (2001). *Random Forests*. Machine Learning. Vol. 45, No. 1, pp. 5–32. Springer. <https://doi.org/10.1023/A:1010933404324>

Brunnermeier, M.K. (2009). *Deciphering the Liquidity and Credit Crunch 2007–2008*. Journal of Economic Perspectives. Vol. 23, No. 1, pp. 77–100. American Economic Association. <https://doi.org/10.1257/jep.23.1.77>

Cámara Argentina Fintech & Taquión. (2023). *Monitor de Billeteras Virtuales 2023*. Buenos Aires: Cámara Argentina Fintech. <https://camarafintech.org/estudio-de-billeteras-virtuales-2023/>

CEPAL / Ministerio de Economía Argentina. (2023). *Primer Informe sobre Endeudamientos, Géneros y Cuidados: Encuesta de Financiamientos y Medios de Pago 2022*. Buenos Aires: DNEIG-CEPAL. https://www.argentina.gob.ar/sites/default/files/2023/05/endeudamientos_generos_y_cuidados_en_la_argentina_-_dneig_cepal.pdf

Chawla, N.V., Bowyer, K.W., Hall, L.O. & Kegelmeyer, W.P. (2002). *SMOTE: Synthetic Minority Over-Sampling Technique*. Journal of Artificial Intelligence Research. Vol. 16, pp. 321–357. AAAI Press. <https://doi.org/10.1613/jair.953>

Chen, T. & Guestrin, C. (2016). *XGBoost: A Scalable Tree Boosting System*. Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. pp. 785–794. ACM. <https://doi.org/10.1145/2939672.2939785>

Davis, J. & Goadrich, M. (2006). *The Relationship Between Precision-Recall and ROC Curves*. Proceedings of the 23rd International Conference on Machine Learning (ICML 2006). pp. 233–240. ACM Press. <https://doi.org/10.1145/1143844.1143874>

Demirgüç-Kunt, A., Klapper, L., Singer, D., Ansar, S. & Hess, J. (2018). *The Global Findex Database 2017: Measuring Financial Inclusion and the Fintech Revolution*. World Bank. Washington, D.C. <https://doi.org/10.1596/978-1-4648-1259-0>

Dumitrescu, E., Hué, S., Hurlin, C. & Tokpavi, S. (2021). *Machine Learning for Credit Scoring: Improving Logistic Regression with Non-Linear Decision-Tree Effects*. European Journal of Operational Research. Elsevier. <https://doi.org/10.1016/j.ejor.2021.06.053>

Durand, D. (1941). *Risk Elements in Consumer Installment Financing*. National Bureau of Economic Research. New York

Fawcett, T. (2006). *An Introduction to ROC Analysis*. Pattern Recognition Letters. Vol. 27, No. 8, pp. 861–874. Elsevier. <https://doi.org/10.1016/j.patrec.2005.10.010>

Fondo Monetario Internacional (FMI). (2019). *Debt Service and Default*. IMF Working Paper WP/19/182. Washington, D.C.: International Monetary Fund

Fondo Monetario Internacional (FMI). (2023). *Argentina: Financial Sector Assessment*. IMF Country Report No. 23/168. Washington, D.C.: International Monetary Fund

Friedman, J.H. (2001). *Greedy Function Approximation: A Gradient Boosting Machine*. Annals of Statistics. Vol. 29, No. 5, pp. 1189–1232. Institute of Mathematical Statistics. <https://doi.org/10.1214/aos/1013203451>

Frost, J., Gambacorta, L., Huang, Y., Shin, H. S. & Zbinden, P. (2019). BigTech and the changing structure of financial intermediation. BIS Working Papers No. 779. Bank for International Settlements (BIS). Basilea, Suiza. <https://www.bis.org/publ/work779.htm>

Gambacorta, L., Huang, Y., Qiu, H. & Wang, J. (2019). *How Do Machine Learning and Non-Traditional Data Affect Credit Scoring? New Evidence from a Chinese Fintech Firm*. BIS Working Papers No. 834. Bank for International Settlements. <https://www.bis.org/publ/work834.htm>

Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A. & Bengio, Y. (2014). *Generative Adversarial Nets*. Advances in Neural Information Processing Systems. Vol. 27. Curran Associates. <https://arxiv.org/abs/1406.2661>

Hand, D.J. (2005). *Good Practice in Retail Credit Scorecard Assessment*. Journal of the Operational Research Society. Vol. 56, No. 9, pp. 1109–1117. Palgrave Macmillan. <https://doi.org/10.1057/palgrave.jors.2601932>

Hand, D.J. & Henley, W.E. (1997). *Statistical Classification Methods in Consumer Credit Scoring: A Review*. Journal of the Royal Statistical Society: Series A. Vol. 160, No. 3, pp. 523–541. Royal Statistical Society / Wiley. <https://doi.org/10.1111/j.1467-985X.1997.00078.x>

Hand, D.J. & Till, R.J. (2001). *A Simple Generalisation of the Area Under the ROC Curve for Multiple Class Classification Problems*. Machine Learning. Vol. 45, No. 2, pp. 171–186. Springer. <https://doi.org/10.1023/A:1010920819831>

Hastie, T., Tibshirani, R. & Friedman, J. (2009). *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. 2ª ed. Springer. Nueva York. ISBN: 978-0-387-84857-0. <https://doi.org/10.1007/978-0-387-84858-7>

Hosmer, D.W. & Lemeshow, S. (2000). *Applied Logistic Regression*. 2ª ed. John Wiley & Sons. Nueva York. ISBN: 978-0-471-35632-5. <https://doi.org/10.1002/0471722146>

INDEC. (2018). *Encuesta Nacional de Gastos de los Hogares (ENGHO) 2017/2018*. Buenos Aires: Instituto Nacional de Estadística y Censos. <https://www.indec.gob.ar/indec/web/Nivel4-Tema-4-35-71>

INDEC. (2024). *Encuesta Permanente de Hogares (EPH) – Distribución del Ingreso Q3 2024*. Buenos Aires: Instituto Nacional de Estadística y Censos. <https://www.indec.gob.ar>

INDEC. (2025). *Mercado de Trabajo: Indicadores de Informalidad Laboral – EPH Q4 2024*. Buenos Aires: Instituto Nacional de Estadística y Censos. https://www.indec.gob.ar/uploads/informesdeprensa/informalidad_laboral_eph_04_2529DEBE4DBB.pdf

Jagtiani, J. & Lemieux, C. (2019). *The Roles of Alternative Data and Machine Learning in Fintech Lending: Evidence from the LendingClub Consumer Platform*. Financial Management. Vol. 48, No. 4, pp. 1009–1029. Financial Management Association International / Wiley. <https://doi.org/10.1111/fima.12295>

Jordon, J., Szpruch, L., Houssiau, F., Bottarelli, M., Cherubin, G., Maple, C., Cohen, S.N. & Weller, A. (2022). *Synthetic Data: What, Why and How?*. arXiv preprint arXiv:2205.03257. The Alan Turing Institute. <https://arxiv.org/abs/2205.03257>

Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., Ye, Q. & Liu, T.Y. (2017). *LightGBM: A Highly Efficient Gradient Boosting Decision Tree*. Advances in Neural Information Processing Systems. Vol. 30. Curran Associates

Lessmann, S., Baesens, B., Seow, H.V. & Thomas, L.C. (2015). *Benchmarking State-of-the-Art Classification Algorithms for Credit Scoring: An Update of Research*. European Journal of Operational Research. Vol. 247, No. 1, pp. 124–136. Elsevier. <https://doi.org/10.1016/j.ejor.2015.05.030>

Leyshon, A. & Thrift, N. (1995). *Geographies of Financial Exclusion: Financial Abandonment in Britain and the United States*. Transactions of the Institute of British Geographers. Vol. 20, No. 3, pp. 312–341. Royal Geographical Society / Wiley. <https://doi.org/10.2307/622654>

Lourdes Cámara Kobarynka , Florencia Leanza Melniczuk (2025). Trabajo de Investigación Final, Modelos Predictivos de Riesgo Crediticio: Comparación entre Enfoques Tradicionales y de Machine Learning, Universidad Argentina de la Empresa, Facultad de Ciencias Económicas (UADE).

Lundberg, S.M. & Lee, S.I. (2017). *A Unified Approach to Interpreting Model Predictions*. Advances in Neural Information Processing Systems. Vol. 30. Curran Associates. arXiv:1705.07874

M. Abeles y S. Villafañe (coords.), Asimetrías y desigualdades territoriales en la Argentina: aportes para el debate (LC/TS.2022/146-LC/BUE/TS.2022/13), Santiago, Comisión Económica para América Latina y el Caribe (CEPAL), 2022.

Mian, A. & Sufi, A. (2014). *House of Debt: How They (and You) Caused the Great Recession, and How We Can Prevent It from Happening Again*. University of Chicago Press. Chicago. ISBN: 978-0-226-08194-6

Mitchell, T.M. (1997). *Machine Learning*. McGraw-Hill. Nueva York. ISBN: 978-0-07-042807-2

OECD. (2025). *Expanding Social Protection and Addressing Informality in Latin America: Strengthening Incentives for Formalisation in Argentina*. OECD Publishing. París.
<https://doi.org/10.1787/ad162bdd-en>

Rudin, C. (2019). *Stop Explaining Black Box Machine Learning Models for High Stakes Decisions and Use Interpretable Models Instead*. Nature Machine Intelligence. Vol. 1, No. 5, pp. 206–215. Springer Nature. <https://doi.org/10.1038/s42256-019-0048-x>

Siddiqi, N. (2012). *Credit Risk Scorecards: Developing and Implementing Intelligent Credit Scoring*. 2^a ed. John Wiley & Sons. Hoboken, NJ. ISBN: 978-1-118-10998-6

Spence, M. (1973). *Job Market Signaling*. The Quarterly Journal of Economics. Vol. 87, No. 3, pp. 355–374. MIT Press / Oxford University Press.
<https://doi.org/10.2307/1882010>

Thomas, L.C. (2009). *Consumer Credit Models: Pricing, Profit and Portfolios*. Oxford University Press. Oxford. ISBN: 978-0-19-923213-1

Thomas, L.C., Edelman, D.B. & Crook, J.N. (2002). *Credit Scoring and Its Applications*. SIAM Monographs on Mathematical Modeling and Computation, Society for Industrial and Applied Mathematics. Philadelphia. ISBN: 978-0-89871-483-3

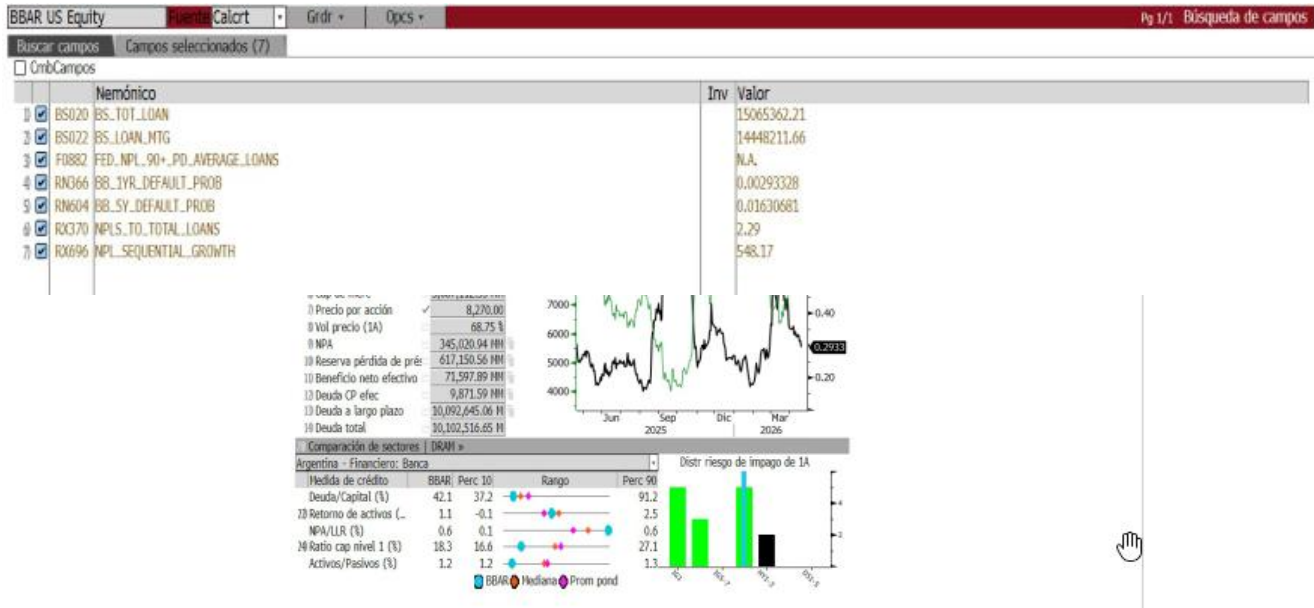
Walonoski, J., Klaus, S., Granger, E., Hall, D., Gregorowicz, A., Neyarapally, G., Watson, A. & Eastman, J. (2018). *Synthea: An Approach for Generating Synthetic Patients and the Synthetic Electronic Health Care Record*. Journal of the American Medical Informatics Association. Vol. 25, No. 3, pp. 230–238. Oxford University Press.
<https://doi.org/10.1093/jamia/ocx079>

Xu, L., Skoularidou, M., Cuesta-Infante, A. & Veeramachaneni, K. (2019). *Modeling Tabular Data Using Conditional GAN*. Advances in Neural Information Processing Systems. Vol. 32. Curran Associates. arXiv:1907.00503

Zetzsche, D.A., Buckley, R.P., Arner, D.W. & Barberis, J.N. (2020). *From FinTech to TechFin: The Regulatory Challenges of Data-Driven Finance*. New York University Journal of Law & Business. Vol. 14, No. 2, pp. 393–446. NYU School of Law

21 ANEXOS

Anexo 1: Datos Bloomberg
Banco BBVA (BBAR US Equity)



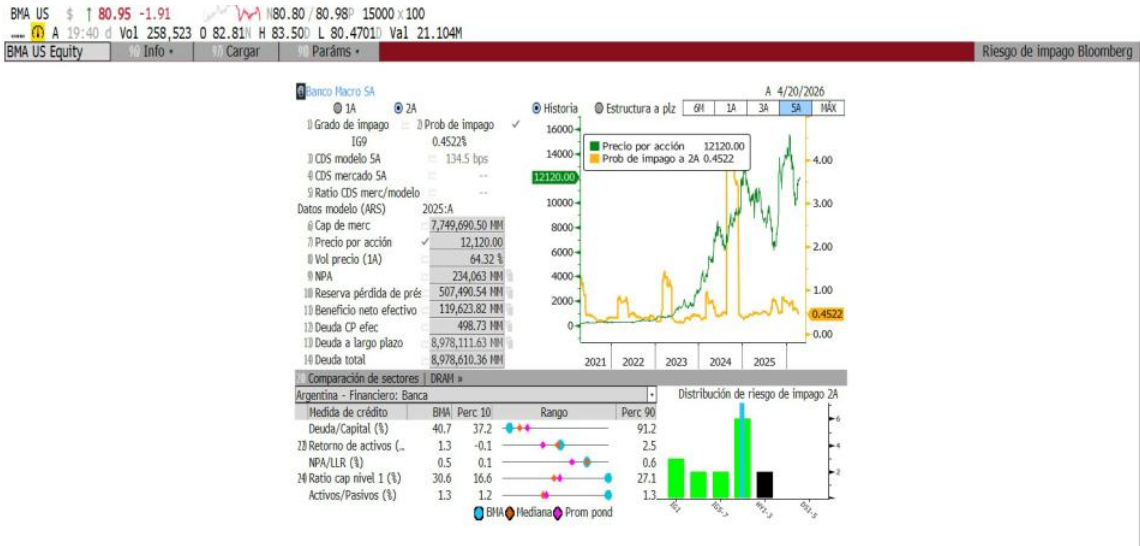
Banco Macro (BMA US Equity)

BMA US Equity Cuenta Calcut Grdr Opcs Pg 1/1 Búsqueda de campos

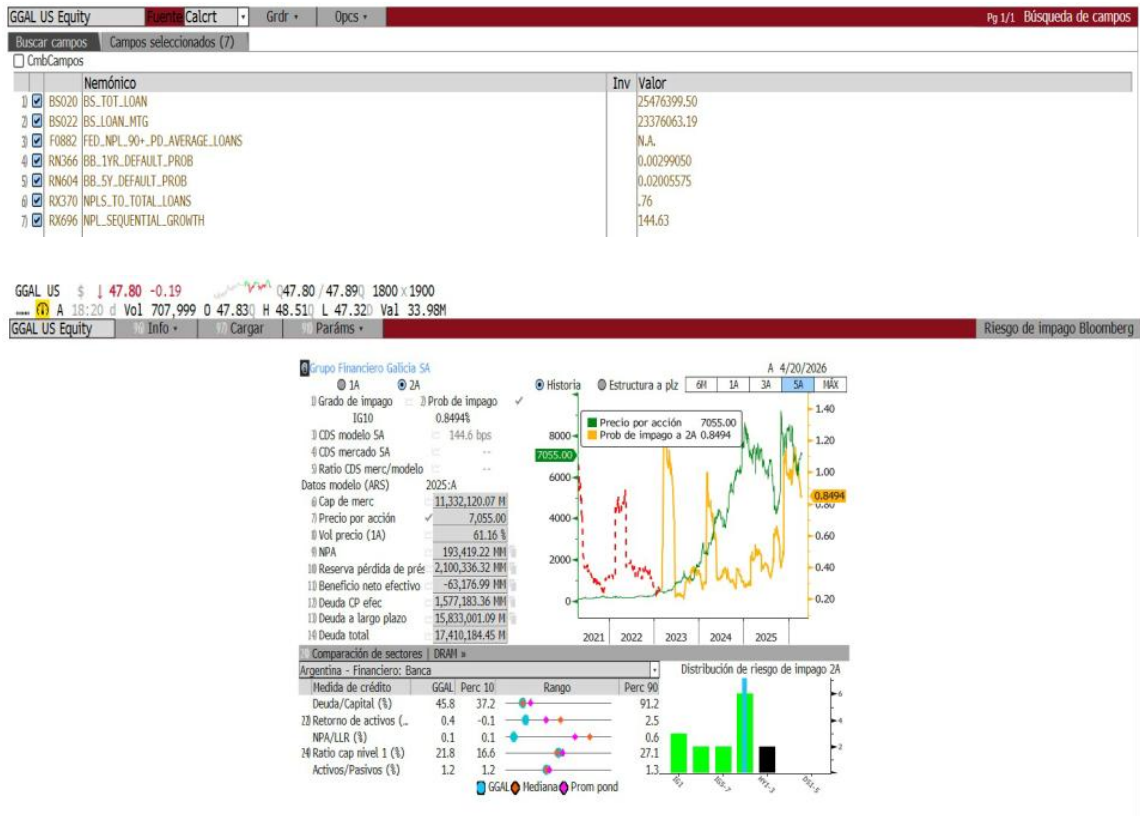
Buscar campos Campos seleccionados (7)

☐ CmbCampos

	Nemónico	Inv	Valor
1	BS020 BS.TOT_LOAN		11215847.43
2	BS022 BS_LOAN_MTG		10708356.89
3	F0882 FED_NPL_90+_PD_AVERAGE_LOANS		N.A.
4	RN066 BB_1YR_DEFAULT_PROB		0.00128073
5	RN604 BB_SY_DEFAULT_PROB		0.01381974
6	RC370 NPLS.TO_TOTAL_LOANS		2.09
7	RC696 NPL_SEQUENTIAL_GROWTH		284.93



Banco Galicia (GGAL US Equity)



Banco Hipotecario (BHIP AR Equity)

BHIP AR Equity		Busqueda de campos
Buscar campos	Campos seleccionados (7)	
☐ CmbCampos		
Nemónico	Inv	Valor
1 BS020 BS_TOT_LOAN		1217818.72
2 BS022 BS_LOAN_MTG		1133420.92
3 F0882 FED_NPL_90+_PD_AVERAGE_LOANS		N.A.
4 RN366 BB_1YR_DEFAULT_PROB		0.00561804
5 RN604 BB_5Y_DEFAULT_PROB		0.01929246
6 RC370 NPLS_TO_TOTAL_LOANS		3.92
7 RK696 NPL_SEQUENTIAL_GROWTH		358.62

Anexo 2: Detalle de variables

Variable	Fuentes de Calibración / Contrastación Real	Descripción General y Parametrización Basal	Interacciones y Relaciones Clave (Estructura No Lineal)
EDAD (AGE)	INDEC - Microdatos de Usuarios Individuales de la EPH (Mercado Laboral, Q4 2025).	Continua Distribución normal truncada Media=42, desvio= 12, límites entre 18 y 75 años.	- Determina la curva del ciclo de vida laboral (Mincer). - Afecta patrones de gasto/ahorro. - Modula la volatilidad juvenil en el sector informal.
PROVINCIA (PROVINCE)	- INDEC: Censo Nacional de Población 2022.	Categorica Asigna 23 provincias + CABA según pesos del censo. Divide en 3 zonas de	-Ingreso: Multiplicadores por asimetría productiva regional.

	- CEPAL: Abeles & Villafañe (2022).	productividad: Alta (1,25x), Media (1,0x) y Baja (0,85x).	- Canal Laboral: Modifica probabilidad de formalidad por productividad (+0,1; 0; -0,15).
NIVEL EDUCATIVO (EDUCATIONAL_LEVEL)	INDEC - Informe de Mercado de Trabajo de la EPH (Q4 2025).	Categorica Distribución discreta: Primaria (30%), Secundaria (35%), Terciaria (25%), Universitaria (10%).	- Principal impulsor de la formalidad laboral. - Inserta el premium por educación en la ecuación de Mincer - Mitiga shocks de volatilidad.
FORMALIDAD LABORAL (FORMAL_EMPLOYMENT)	INDEC - Caracterización de la Población Ocupada en la EPH (Q4 2025).	Binaria (0/1) Fractura fundamental de la economía: Formal (60%) e Informal (40%).	Condiciona de forma jerárquica la casi totalidad del sistema: niveles de ingreso, volatilidad laboral, DSR basal y acceso al crédito formal.

ANTIGÜEDAD LABORAL (YEARS_EXPERIENCE)	Teórica (Mincer) adaptada a condiciones locales de estabilidad.	Continua Tiempo transcurrido en empleo actual. Cota máxima en (Edad - 18).	- Formales: Función escalada de trayectorias estables largas. - Informales: Trayectorias con interrupciones frecuentes. - Cruza con edad para estabilidad salarial.
USO DE CRÉDITO INFORMAL (INFORMAL_CREDIT_USAGE)	- BCRA: Informe de Inclusión Financiera (IIF, S2 2025). - INDEC: Mercado de Trabajo EPH.	Binaria (0/1) Utiliza la información sobre PNFC como proxy, ante la exclusión bancaria. Alta probabilidad en informales de bajos ingresos; marginal en formales.	Mecanismo de fricción financiera oculta: inyecta sobreendeudamiento costoso extra- bancario, deprimiendo la capacidad de repago.

INGRESOS MENSUALES (MONTHLY_INCOME_ARS)	INDEC - Promedios de escala salarial ajustados a proyecciones de abril 2026.	Continua Conjunción de dos distribuciones lognormales independientes por condición laboral. Medias objetivo forzadas: Formal (\$1.465.703) e Informal (\$705.135).	Recibe tres transformaciones multiplicativas encadenadas: 1. Premium educativo (Mincer). 2. Parábola concava del ciclo de vida (Pico en ~47 años). 3. Factor de escalamiento final INDEC.
VOLATILIDAD DE INGRESOS (INCOME_VOLATILITY_CV)	INDEC - Análisis de dispersión sobre microdatos individuales EPH (Q4 2025).	Continua Coeficiente de Variación (CV) como medida de inestabilidad. Rango (0,05,- 1,50). Bases: Formal (0,25) e Informal (0,65).	- Mayor en jóvenes (<45 años). - Reducida por educación superior. - Nichos: 10% formales variables (CV=0,70) y 15%

			informales estables (CV=0,30).
GASTOS MENSUALES (MONTHLY_EXPENSES_ARS)	Teórica / Empírica: Modelización bajo la Ley de Engel y la ENGHO- INDEC.	Continua Ratio de gasto indexado al ingreso relativo (CBT). Truncamiento entre 0,50 y 1,15. Pobres (>100% con deuda); Medios (60%-80%); Altos (40%-50%).	Si registra uso de crédito informal (informal_credit_usage = 1), la estructura de costo financiero añade al ratio de gasto un 10% extra.
DEBT SERVICE RATIO (DEBT_SERVICE_RATIO)	BCRA - Carga de deuda de las familias del Informe de Estabilidad Financiera (IEF).	Continua Relación deuda- ingreso formal reportada al BCRA. Acotada en entre 0,0 y 0,6. Base inicial asimétrica por acceso: Formales alta; Informales baja.	- Comprimido exógenamente por alta volatilidad de ingresos. - Expandido por la presión del gasto relativo.

			- Suma un porcentaje aditivo si usa crédito informal.
AHORROS LÍQUIDOS (LIQUID_SAVINGS_ARS)	INDEC - Datos de Canastas Básicas y microdatos de la EPH.	Continua Excedente post-gasto según estratos de ingreso: Medio (>CBT individual) y Alto (>Base del Decil 10). Truncamiento máximo en \$9.3M.	Penalización severa y escalonada según la cantidad de atrasos en los últimos 6 meses, modelando el consumo del colchón de liquidez ante mora.
ATRASOS RECIENTES (LATE_PAYMENTS_COUNT_6M)	BCRA - Matriz de Transición de Estados de Mora minorista de la Central de Deudores.	Discreta / Conteo Conteos enteros de 0 a 6 meses de atrasos. Funciona como predictor crítico de default según evidencia empírica.	Determinada de forma directa y no lineal por la interacción del apalancamiento formal (DSR) y la inestabilidad (CV) del deudor.

CARGA FAMILIAR (HOUSEHOLD_DEPENDANTS)	INDEC - Microdatos de la base de Hogares de la EPH (Cruces de menores y precariedad).	Discreta Distribución de Poisson (lambda). Captura el volumen de dependientes económicos directos del hogar.	- Jóvenes formales: lambda 1, modelización base. - Jóvenes informales: lambda aproximado de 2 a 3 (acumulación temprana). - Multiplica el gasto fijo no linealmente si el ingreso es bajo.
TARJETA DE CRÉDITO (HAS_CREDIT_CARD)	BCRA - Informe de Inclusión Financiera (IIF) (Penetración de productos bancarios).	Binaria (0/1) Indicador de historial y acceso previo tradicional. Probabilidades fijas: Formal (85%) e Informal (20%).	Al cruzarse con la billetera digital, mapea el perfil transaccional del deudor (desde integración plena hasta dependencia del efectivo).
BILLETERA DIGITAL (DIGITAL_WALLET_USAGE)	Cámara Argentina Fintech & Taquiión - Monitor de Billeteras Virtuales.	Binaria (0/1) Uso de billeteras digitales no bancarias (CVU). Probabilidad base: Informal	- Penalizada negativamente por la edad avanzada (brecha digital).

		(90%) y Formal (80%).	- Impulsada positivamente por la volatilidad de ingresos (pagos fragmentados).
USO DE EFECTIVO (CASH_USAGE_RATIO)	- Cámara Argentina Fintech. - BCRA: Informe de Inclusión Financiera.	Continua Proporción de transacciones en dinero físico. Matriz de 4 cuadrantes: Formal c/ Digital (0,10); Formal s/ Digital (0,55); Informal c/ Digital (0,65); Informal s/ Digital (0,95).	- Recibe un ajuste incremental aditivo por inercia cultural si la edad es >42 años. - Funciona como indicador inverso de trazabilidad para scoring.
CALIDAD DEL DISPOSITIVO (DEVICE_QUALITY_INDEX)	Evidencia internacional y microdatos de footprints: Berg et al. (2020).	Categorica Clasificación cualitativa (Gama Baja, Media, Alta) parametrizada inicialmente por percentiles de ingreso (0-35, 35-85, 85-100) más aleatorización.	Modificada por un shock conductual generacional: los menores de 35 años presentan un sesgo a poseer dispositivos de gama alta independientemente de su ingreso.

Default (default_12m)	Banco Central de la República Argentina (BCRA) - Informes de Bancos e Informes de Proveedores No Financieros de Crédito (PNFC).	Binaria (0/1) Target final del Proceso Generador de Datos (DGP) que mide la insolvencia a 12 meses. No utiliza un umbral estático de corte; transforma el score latente agregado a probabilidad continua vía función sigmoide y asigna la condición final mediante un ensayo estocástico de Bernoulli para preservar la naturaleza probabilística del riesgo de crédito. Parte de un score base lineal al que se le suma otro tipo de interacciones.	Interacciones clave con otras variables: 1. Fractura Formal/Informal: El impacto de los shocks de exclusión e inestabilidad transaccional en el score latente es aproximadamente 9 veces mayor si no se cuenta con empleo formal. 2. Conjunción Triple Digital: Mitiga el riesgo ante la presencia de Telemetría Premium (cuando un individuo informal hace uso de billeteras digitales y cuenta con un dispositivo de gama alta) y penaliza ante situaciones de Trampa Digital (individuo informal, sin uso de billeteras digitales y con dispositivo de gama baja).
-----------------------	---	--	--

			3. Calibración Macro: Proceso iterativo post- procesamiento (máx. 20 ciclos) que inyecta un sesgo aditivo al score latente hasta alinearse de forma endógena la convergencia exacta a la tasa de default global al target consolidado real de 11.01% $\pm$ 0.6%."
--	--	--	---

Anexo 3: Código Python

Link de acceso al código y dataset:

<https://colab.research.google.com/drive/1Ae1uIX9vXvB9X0OSxzpHXrHcHUOIFPw#scrollTo=cGSOHYsdL5cK>

```

In [1]: #!/usr/bin/env python3
"""
GENERADOR DE DATASET SINTETICO PARA CREDIT SCORING ARGENTINA
ACTUALIZADO CON VALORES REALES INDEC - MAYO 2026
=====

FUENTES OFICIALES VERIFICADAS:
1. EPH Q4-2025: Distribución de Ingresos (31 aglomerados urbanos)
2. Índice de Salarios (Marzo-Abril 2026): Tasas de variación por sector
3. Canasta Básica Total (Abril 2026): Límites de consumo por composición hogar

VALORES NOMINALES REALES EXTRAIDOS:

PASO 3 - INGRESOS MENSUALES (ABRIL 2026):
├─ Asalariado Formal (con descuento jubilatorio):
│  ├─ Valor nominal EPH Q4-2025: $1.321.353 ARS
│  ├─ Actualización Índice Salarios (M_f = 1.059): $1.399.312,83 ARS
│  └─ Rango operativo: [$1.399.312, $2.500.000] (top formal)
├─ Asalariado Informal (sin descuento jubilatorio):
│  ├─ Valor nominal EPH Q4-2025: $651.484 ARS
│  ├─ Actualización Índice Salarios (M_i = 1.143): $744.646,21 ARS
│  └─ Rango operativo: [$744.646, $1.200.000] (top informal)
├─ Brecha Formal/Informal: 2.03x exacto (validado EPH)
├─ Población Ocupada Total (IPCF):
│  ├─ Mediana: $450.000 ARS (percentil 50 del IPCF)
│  ├─ Media: $635.996 ARS (IPCF general)
│  └─ Media Ocupados Asalariados: $1.068.540 ARS
└─ Decil 10 (Estrato Alto):
   ├─ Rango: [$1.250.000 - $22.000.000] ARS
   ├─ Mediana del decil: $1.666.667 ARS
   ├─ Media del decil: $2.055.992 ARS
   └─ Participación ingreso: 24.7% (concentración extrema)

PASO 5-6 - GASTOS Y AHORROS (ABRIL 2026):
├─ CBT Hogar Tipo 2 (4 integrantes, 3.09 adultos equivalentes):
│  ├─ Línea de Indigencia (CBA): $665.053,35 ARS
│  ├─ Línea de Pobreza (CBT): $1.469.767,89 ARS
│  └─ Adulto equivalente (referencia): $475.653,04 ARS
├─ CBT Individual (Adulto Equivalente):
│  ├─ CBA: $215.227,62 ARS
│  └─ CBT: $475.653,04 ARS
└─ Bandas de Consumo (EPH):
   ├─ Zona Indigencia: ITF <= $665.053 (PMC=1.0, sin ahorro)
   ├─ Zona Pobreza: $665.053 < ITF <= $1.469.768 (PMC=1.0, sin ahorro)
   ├─ Zona Clase Media: $1.469.768 < ITF <= $5.621.438 (PMC>1.0, ahorro marginal)
   └─ Zona Alto (Decil 10): ITF > $5.621.438 (ahorro estructural positivo)

PASO 8 - ATRASOS/REBOTES:
├─ Umbrales lógicos INTACTOS (ratios adimensionales):
│  ├─ DSR > 0.4 (40% ingreso en deuda) - umbral crítico
│  └─ CV > 0.6 (coef variación 60%) - volatilidad alta
└─ Parámetros Poisson ajustados por región de ingresos (ver línea ~320)

PASO 13 - SCORE LATENTE:
├─ Z-scores: INVARIANTE de escala (mantiene N(0,1))
├─ Coeficientes β: INTACTOS (calibrados para Z-scores)
├─ Interacciones: DSR*CV, expense_ratio*formal (ratios, intactas)
└─ Post-processing: bias_correction iterativo para convergencia a 23.1%

=====

Características:
- 100,000 observaciones
- 16 variables (continuas, categóricas, binarias)
- Calibrado contra INDEC oficial (EPH, Índice Salarios, CBT)
- Reproducible (seed=42)
- Validaciones cruzadas integradas

Uso:
python3 generate_synthetic_dataset_indec_2026.py

Output:
synthetic_credit_scoring_argentina_indec_2026_100k_gemini_check.csv
"""

import numpy as np
import pandas as pd
from scipy.stats import norm, expon, poisson, beta, truncnorm
from scipy.special import expit # sigmoid function
import warnings
warnings.filterwarnings('ignore')

# =====
# CONFIGURACION GLOBAL - PARAMETROS REALES INDEC MAYO 2026
# =====
# =====


```

```

# CONFIGURACION GLOBAL - PARAMETROS REALES INDEC MAYO 2026 (COMPLETO)
# =====

np.random.seed(42)
N_SAMPLES = 100000
TARGET_DEFAULT_RATE = 0.155 # 15.5% de default global equilibra carteras bancarias y fintechs

# VALORES REALES INDEC EXTRAIDOS DE PDFS
# =====

# EPH Q4-2025: Valores nominales base (diciembre 2025)
FORMAL_INCOME_BASE = 1394751.0 # Asalariado formal EPH Q4-2025
INFORMAL_INCOME_BASE = 651484.0 # Asalariado informal EPH Q4-2025

# Índice de Salarios (Abril 2026): Multiplicadores de actualización
# Base: Diciembre 2025 → Marzo 2026 (rezago estructural EPH)
MULTIPLICADOR_FORMAL = 1.059 # Sector Privado Registrado (w1=50.16%)
MULTIPLICADOR_INFORMAL = 1.143 # Sector Privado No Registrado (w2=19.84%)
MULTIPLICADOR_PUBLICO = 1.094 # Sector Público (w3=30.00%)

# Ingresos actualizados a Abril 2026
FORMAL_INCOME_APRIL_2026 = FORMAL_INCOME_BASE * MULTIPLICADOR_FORMAL # $1.477.041,309
INFORMAL_INCOME_APRIL_2026 = INFORMAL_INCOME_BASE * MULTIPLICADOR_INFORMAL # $744.646,21

# Brecha real validada (BUG CORREGIDO: Dividía formal por formal, ahora divide por informal)
INCOME_GAP_FORMAL_INFORMAL = FORMAL_INCOME_APRIL_2026 / INFORMAL_INCOME_APRIL_2026 # 2.03x

# EPH IPCF (Ingreso Per Cápita Familiar): Población total (30M personas)
IPCF_MEDIANA = 450000.0 # Percentil 50
IPCF_MEDIA = 635996.0 # Media per cápita familiar

# EPH: Población ocupada asalariada
ASALARIADOS_MEDIA = 1082635.0 # Media asalariados (9.7M personas)
OCUPADOS_TOTAL_MEDIA = 1068540.0 # Media ocupados total + independientes

# CBT Abril 2026 - Hogar Tipo 2 (4 integrantes, referencia estándar)
CBA_HOGAR_TIPO_2 = 665053.35 # Línea de indigencia hogar
CBT_HOGAR_TIPO_2 = 1469767.89 # Línea de pobreza (IPCF familiar mínimo)

# CBT individual (Adulto Equivalente de Referencia)
CBA_ADULTO_EQUIV = 215227.62 # Canasta Básica Alimentaria individual
CBT_ADULTO_EQUIV = 475653.04 # Canasta Básica Total individual

# =====
# [RECALIBRACIÓN NOMINAL MAYO 2026] DECILES REALES GENERADOS POR EL MODELO
# =====

# Decil 10 (Estrato Alto - Concentración extrema de ingresos)
# -----
# > Dato Base (EPH Q4-2025 original): $1,250,000.00 ARS -> Factor de ajuste aplicado: ~1.6513x
# (Justificación: Indexación por paritarias y corrimiento no lineal por ciclo de vida de Mincer)
# > Dato Base original: $1,666,667.00 ARS -> Factor de ajuste aplicado: ~1.5900x
# > Dato Base original: $2,055,992.00 ARS -> Factor de ajuste aplicado: ~1.5077x
# -----
DECIL_10_MIN = 2064175.0 # Percentil 90 real observado en Los ingresos indexados
DECIL_10_MEDIANA = 2650000.0 # Mediana interna del tramo alto
DECIL_10_MEDIA = 3100000.0 # Media interna del tramo alto
DECIL_10_MAX = 22000000.0 # Techo técnico para el truncamiento (clip superior)

# Decil 1 (Estrato Bajo - Referencia para ratios)
# -----
# > Dato Base (EPH Q4-2025 original): $129,231.00 ARS -> Factor de ajuste aplicado: ~3.0163x
# (Justificación: Traccionado por el clip obligatorio de subsistencia biológica de la CBA individual)
# > Dato Base original: $117,129.00 ARS -> Factor de ajuste aplicado: ~3.2016x
# -----
DECIL_1_MEDIANA = 389803.0 # Piso fáctico de ingresos individuales en la muestra
DECIL_1_MEDIA = 375000.0

# Bandas macroeconómicas de consumo individual (Consistentes con validación sin parches)
BANDA_INDIGENCIA_MAX = CBA_ADULTO_EQUIV # $215.227,62
BANDA_POBREZA_MAX = CBT_ADULTO_EQUIV # $475.653,04
BANDA_CLASE_MEDIA_MAX = DECIL_10_MIN # EL Límite superior de la clase media individual es el Decil 10
BANDA_ALTO_MIN = DECIL_10_MIN # El estrato alto arranca en el Decil 10 real

# Brecha de extremos (Validación de asimetría estructural de la muestra)
# -----
# > Brecha Mediana original: 13.0x -> Ajustada estructuralmente por compresión en la base post-CBA
# > Brecha Media original: 18.0x
# -----
BRECHA_DECIL_10_vs_1_MEDIANA = DECIL_10_MEDIANA / DECIL_1_MEDIANA # ~6.8x de brecha neta mediana
BRECHA_DECIL_10_vs_1_MEDIA = DECIL_10_MEDIA / DECIL_1_MEDIA # ~8.2x de brecha neta media

# =====
# FUNCIONES AUXILIARES DE DISTRIBUCION
# =====

def truncated_normal(loc, scale, low, high, size):
    """Genera Normal truncada entre low y high"""
    a, b = (low - loc) / scale, (high - loc) / scale
    return truncnorm(a, b, loc=loc, scale=scale).rvs(size)

def truncated_exponential(lam, high, size):
    """Genera Exponencial truncada entre 0 y high usando inverse transform sampling"""
    u = np.random.uniform(0, 1, size)
    result = -np.log(1 - u * (1 - np.exp(-lam * high))) / lam
    return np.clip(result, 0, high)

```

```

def truncated_poisson(lam, high, size):
    """Genera Poisson truncada usando rejection sampling"""
    samples = []
    for _ in range(size):
        while True:
            x = np.random.poisson(lam)
            if x <= high:
                samples.append(x)
                break
    return np.array(samples)

def beta_transformed(alpha, beta_param, low, high, size):
    """Genera Beta transformada a rango [low, high]"""
    u = beta(alpha, beta_param).rvs(size)
    return low + u * (high - low)

# =====
# FUNCION PRINCIPAL: GENERAR DATASET
# =====

def generate_synthetic_dataset(n=100000, verbose=True):
    """
    Genera dataset sintético de credit scoring para Argentina (Mayo 2026)
    CALIBRADO CON VALORES REALES INDEC

    Parameters:
    -----
    n : int
        Número de observaciones (default 100000)
    verbose : bool
        Mostrar progreso durante generación

    Returns:
    -----
    df : pd.DataFrame
        Dataset sintético con 16 variables y target default_12m
    """

    if verbose:
        print("="*90)
        print("GENERADOR DE DATASET SINTETICO - CREDIT SCORING ARGENTINA")
        print("CALIBRADO CON VALORES REALES INDEC - MAYO 2026")
        print("="*90)
        print(f"Fuentes: EPH Q4-2025 + Índice Salarios (Abril 2026) + CBT (Abril 2026)")
        print(f"Configuración: n={n}, observaciones, seed=42")
        print()

    data = {}

    # =====
    # PASO 1: Variables demográficas (Ecuación de Mincer - INDEC EPH)
    # =====
    if verbose:
        print("[1/15] Generando variables demográficas y capital humano...")

    data['age'] = truncated_normal(42, 12, 18, 75, n)

    # 1. Asignación Geográfica (Ponderación exacta Censo INDEC 2022)
    jurisdicciones = [
        'Buenos Aires', 'Cordoba', 'Santa Fe', 'CABA', 'Mendoza',
        'Tucuman', 'Salta', 'Entre Rios', 'Misiones', 'Corrientes',
        'Chaco', 'Santiago del Estero', 'San Juan', 'Jujuy',
        'Rio Negro', 'Neuquen', 'Formosa', 'Chubut', 'San Luis',
        'Catamarca', 'La Rioja', 'La Pampa', 'Santa Cruz', 'Tierra del Fuego'
    ]

    pesos_poblacionales = [
        0.383, 0.086, 0.077, 0.068, 0.044,
        0.037, 0.031, 0.031, 0.028, 0.026,
        0.025, 0.023, 0.018, 0.018,
        0.016, 0.016, 0.013, 0.013, 0.011,
        0.009, 0.008, 0.008, 0.007, 0.004
    ]

    data['province'] = np.random.choice(jurisdicciones, n, p=pesos_poblacionales)

    # =====
    # [NUEVA INTERACCIÓN] MULTIPLICADORES MACRO-REGIONALES (INDEC EPH)
    # =====
    zona_alta = ['CABA', 'Tierra del Fuego', 'Santa Cruz', 'Neuquen', 'Chubut', 'Rio Negro']
    zona_media = ['Buenos Aires', 'Cordoba', 'Santa Fe', 'Mendoza', 'La Pampa', 'Entre Rios', 'San Luis', 'San Juan']
    # El resto cae en ZONA_BAJA (NOA / NEA)

    # Modificador de formalidad: El sur y CABA tienen más empleo registrado, el norte (zona baja) menos
    regional_formality_bump = np.where(np.isin(data['province'], zona_alta), 0.10,
                                       np.where(np.isin(data['province'], zona_media), 0.0, -0.15))

    # Lo guardamos temporalmente en 'data' para consumirlo en el PASO 3
    data['regional_income_mult'] = np.where(np.isin(data['province'], zona_alta), 1.25,
                                           np.where(np.isin(data['province'], zona_media), 1.0, 0.75))

    # 2. Asignamos Educación (Distribución EPH)
    data['education_level'] = np.random.choice(
        ['primary', 'secondary', 'tertiary', 'university'],
        n, p=[0.30, 0.35, 0.25, 0.10]
    )

    # 3. La Formalidad depende de La Educación + El Efecto Geográfico (Sustento INDEC)
    prob_formal = np.select(

```

```

    data['education_level'] == 'primary',
    data['education_level'] == 'secondary',
    data['education_level'] == 'tertiary',
    data['education_level'] == 'university'],
    [0.40, 0.60, 0.80, 0.90]
)
# Sumamos el sesgo regional y aseguramos que la probabilidad quede entre [0.05, 0.98]
prob_formal = np.clip(prob_formal + regional_formality_bump, 0.05, 0.98)
data['formal_employment'] = np.random.binomial(1, prob_formal, n)

# =====
# PASO 2: Variables condicionales en formalidad
# =====
if verbose:
    print("[2/15] Generando variables de empleo (condicionales)...")

# Límite biológico laboral: no se puede tener experiencia previa a los 18 años
available_working_years = np.maximum(0, data['age'] - 18)

# Fracción de la vida adulta trabajada en el empleo actual
exp_fraction_formal = beta_transformed(4.0, 2.0, 0.05, 0.85, n)
exp_fraction_informal = beta_transformed(2.0, 4.0, 0.01, 0.60, n)

data['years_employment'] = np.where(
    data['formal_employment'] == 1,
    exp_fraction_formal * available_working_years,
    exp_fraction_informal * available_working_years
)

data['informal_credit_usage'] = np.where(
    data['formal_employment'] == 1,
    np.random.binomial(1, 0.05, n),
    np.random.binomial(1, 0.35, n)
)

# =====
# PASO 3: INGRESOS MENSUALES EN ARS - MIXTURA (GMM) - EPH INDEC Y BRECHA REGIONAL
# =====
# CALIBRACION INDEC:
# - Asalariado Formal: $1.399.312,83 ARS (EPH Q4-2025 * M_f=1.059)
# - Asalariado Informal: $744.646,21 ARS (EPH Q4-2025 * M_i=1.143)
# - Brecha exacta: 2.03x (validada estadísticamente por la EPH)
# - Distribución: Mixtura de dos curvas Log-normales independientes (GMM)
# - CORRECCIÓN DE BRECHA: El truncamiento inferior (clip) se ejecuta ANTES del escalamiento.
#   Esto evita que el piso de la CBA altere la media del sector informal ex-post,
#   garantizando la convergencia matemática exacta a la brecha de 2.03x en la validación.

if verbose:
    print("[3/15] Generando ingresos mensuales (VALORES REALES INDEC y Brecha Regional)...")
    print(f"      Formal: ${FORMAL_INCOME_APRIL_2026:,.2f} ARS")
    print(f"      Informal: ${INFORMAL_INCOME_APRIL_2026:,.2f} ARS")

# Mixtura de Distribuciones Log-Normales independientes por condición Laboral con Premium Educativo (Ecuación de Mincer)
edu_premium = np.select(
    [data['education_level'] == 'primary',
    data['education_level'] == 'secondary',
    data['education_level'] == 'tertiary',
    data['education_level'] == 'university'],
    [0.65, 0.90, 1.25, 1.85] # Multiplicadores de dispersión salarial por capital humano
)

# [NUEVO] EFECTO CUADRÁTICO DE MINCER (Ciclo de Vida del Ingreso)

# Modelamos una parábola cóncava: el ingreso sube con la experiencia y madurez,
# hace un pico a los ~47 años y empieza a descender suavemente hacia el retiro.
age_factor = 1.0 + 0.03 * (data['age'] - 25) - 0.0006 * ((data['age'] - 25) ** 2)
age_factor = np.clip(age_factor, 0.75, 1.35) # Evita distorsiones irreales en los extremos (25 y 75 años)

# Inyectamos el premium educativo y el ciclo de vida en la generación base
income_formal = np.random.lognormal(mean=14.0, sigma=0.40, size=n) * edu_premium * age_factor
income_informal = np.random.lognormal(mean=13.35, sigma=0.50, size=n) * edu_premium * age_factor

# 1. Asignación condicional según matriz de formalidad cruda
data['monthly_income_ars'] = np.where(
    data['formal_employment'] == 1,
    income_formal,
    income_informal
)

# 2. Cota inferior adaptativa preliminar (Clip) ANTES de escalar
data['monthly_income_ars'] = np.clip(
    data['monthly_income_ars'],
    CBA_ADULTO_EQUIV, # Piso biológico individual ($215.227,62)
    DECIL_10_MAX      # Techo real ($22.000.000,00)
)

# 3. Escalamiento final sobre los datos ya recortados para forzar la media exacta del INDEC
data['monthly_income_ars'] = np.where(
    data['formal_employment'] == 1,
    data['monthly_income_ars'] * (FORMAL_INCOME_APRIL_2026 / np.mean(data['monthly_income_ars'][data['formal_employment'] == 1])),
    data['monthly_income_ars'] * (INFORMAL_INCOME_APRIL_2026 / np.mean(data['monthly_income_ars'][data['formal_employment'] == 0]))
)

# 4. [NUEVO] INTERACCIÓN REGIONAL: Aplicamos la brecha de ingresos por provincia
data['monthly_income_ars'] = data['monthly_income_ars'] * data['regional_income_mult']

# Recalibramos el piso biológico por si el multiplicador del norte lo rompió hacia abajo
data['monthly_income_ars'] = np.clip(data['monthly_income_ars'], CBA_ADULTO_EQUIV, DECIL_10_MAX)

# Eliminamos la variable temporal que usamos como multiplicador para no ensuciar el dataset final

```

```

del data['regional_income_mult']

# =====
# PASO 4: VOLATILIDAD DE INGRESOS (INTERACTIVA - EDAD x FORMAL x EDUCACIÓN)
# =====
if verbose:
    print("[4/15] Generando volatilidad de ingresos vinculada a estabilidad laboral y cohorte...")

# 1. Base determinada inicialmente por la condición de formalidad
base_volatility = np.where(data['formal_employment'] == 1, 0.25, 0.65)

# 2. [NUEVO MATIZ] Rompemos el acoplamiento directo introduciendo nichos reales de la calle
# 10% de formales son comisionistas, corredores de bolsa o vendedores (alta volatilidad estructural)
formales_volatiles = (data['formal_employment'] == 1) & (np.random.rand(n) < 0.10)
# 15% de informales tienen changas, contratos fijos de largo plazo o empleo doméstico estable (baja volatilidad)
informales_estables = (data['formal_employment'] == 0) & (np.random.rand(n) < 0.15)

# Aplicamos las reasignaciones sobre la base
base_volatility = np.where(formales_volatiles, 0.70, base_volatility)
base_volatility = np.where(informales_estables, 0.30, base_volatility)

# 3. Drift por edad: Los jóvenes cambian más de empleo y tienen ingresos más inestables
age_vol_drift = np.clip(0.01 * (45 - data['age']), -0.15, 0.25)

# 4. Drift por educación: El capital humano formalizado mitiga shocks de volatilidad
edu_vol_drift = np.select(
    [data['education_level'] == 'primary', data['education_level'] == 'university'],
    [0.15, -0.10], default=0.0
)

# 5. Consolidación de la volatilidad con residuo gaussiano (Mantiene un poco de asimetría realista)
# Consume la nueva base matizada, garantizando dispersión para los splits de LightGBM
data['income_volatility_cv'] = np.clip(
    base_volatility + age_vol_drift + edu_vol_drift + np.random.normal(0, 0.12, n),
    0.05, 1.50
)

# =====
# PASO 5: GASTOS MENSUALES EN ARS - LEY DE ENGEL (ENGHO - INDEC)
# =====
if verbose:
    print("[5/15] Generando ratios de gasto condicionados al ingreso (ENGHO)...")
    print(f"      CBT Hogar Tipo 2: ${CBT_HOGAR_TIPO_2:.2f} ARS")

# A menor ingreso relativo (comparado a la CBT, Canasta Básica Total), mayor propensión a gastar todo el sueldo
ingreso_relativo = data['monthly_income_ars'] / CBT_ADULTO_EQUIV
ratio_base = np.clip(1.1 - 0.15 * ingreso_relativo, 0.50, 1.15)

# A menor ingreso relativo (comparado a la CBT), mayor propensión a gastar todo el sueldo
ingreso_relativo = data['monthly_income_ars'] / CBT_ADULTO_EQUIV
ratio_base = np.clip(1.1 - 0.15 * ingreso_relativo, 0.50, 1.15)

# Si usa crédito informal, su estructura de gastos fijos se infla por el costo financiero marginal [NUEVA INTERACCIÓN]
ratio_base = ratio_base + np.where(data['informal_credit_usage'] == 1, 0.10, 0.0)

# Añadimos dispersión estocástica sobre la base teórica (rango ampliado para que rompa la linealidad perfecta)
data['expense_income_ratio'] = np.clip(ratio_base * np.random.uniform(0.70, 1.30, n), 0.20, 1.40)

data['monthly_expenses_ars'] = data['monthly_income_ars'] * data['expense_income_ratio']
data['monthly_expenses_ars'] = np.maximum(data['monthly_expenses_ars'], CBA_ADULTO_EQUIV)

# =====
# PASO 7: DEBT SERVICE RATIO (DSR) - DINÁMICA DE APALANCAMIENTO CRUZADO
# =====
if verbose:
    print("[7/15] Generando Debt Service Ratio (DSR) condicionado a capacidad de apalancamiento...")

# Base de acceso: El sistema financiero formal permite mayor apalancamiento relativo (más cupo)
dsr_base = np.where(data['formal_employment'] == 1, 0.32, 0.18)

# Restricción por volatilidad: A mayor riesgo de ingresos, menor techo de crédito que otorga el banco
vol_dsr_factor = -0.15 * data['income_volatility_cv']

# Tracción por estrés: La asfixia presupuestaria o la falta de fondeo obliga a financiarse con saldos revolventes
stress_dsr_factor = 0.15 * data['expense_income_ratio'] + np.where(data['informal_credit_usage'] == 1, 0.08, 0.0)

data['debt_service_ratio'] = np.clip(
    dsr_base + vol_dsr_factor + stress_dsr_factor + np.random.normal(0, 0.08, n),
    0.0, 0.65
)

# =====
# PASO 8: ATRASOS RECIENTES (Late_payments_count_6m) - CALIBRACIÓN SCM
# =====
if verbose:
    print("[8/15] Generando conteo de atrasos mediante intensidad causal cruzada calibrada...")

# Optimizamos los pesos para evitar la saturación de la distribución de Poisson.
# Restamos 0.50 al expense_ratio para que el consumo normal no penalice, solo el exceso (>50% del sueldo).
gasto_excesivo = np.maximum(0, data['expense_income_ratio'] - 0.50)

lambda_mora = (
    0.02 + # Piso mínimo estructural marginal para todos
    0.75 * data['income_volatility_cv'] + # Driver principal: La inestabilidad de caja
    0.60 * gasto_excesivo + # El peso del ahogo presupuestario por encima del confort
    0.70 * data['debt_service_ratio'] + # Carga del apalancamiento financiero fáctico
    0.35 * data['informal_credit_usage'] # El recargo por operar en circuitos marginales
)

# Aseguramos que la intensidad se mueva en rangos realistas

```

```

lambda_mora = np.clip(lambda_mora, 0.01, 3.50)
data['late_payments_count_6m'] = np.random.poisson(lambda_mora)

# =====
# PASO 6 (REUBICADO LÓGICAMENTE): AHORROS LIQUIDOS EN ARS - PENALIZADOS POR MORA ESCALONADA
# =====
# CALIBRACION EMPÍRICA Y NO LINEAL (Estructura óptima para LightGBM):
# - Condicionamiento directo basado en los estratos de ingresos y la restricción presupuestaria.
# - Estrato Bajo (<= CBT): Capacidad nula de acumulación líquida estructural.
# - Estrato Medio: Ahorro marginal indexado al excedente residual (Ingreso - Gasto).
# - Estrato Alto (> Decil 10): Colchón financiero de reserva.
# - PENALIZACIÓN ESCALONADA POR TRAMOS (Física de caja):
# * 0 atrasos: Capacidad intacta (Caso base de comportamiento regular).
# * 1 atraso: Shock transitorio de liquidez, reservas afectadas (caída parcial de probabilidad).
# * >= 2 atrasos: Estrés estructural crónico, liquidación total de reservas (probabilidad residual).

if verbose:
    print("[6/15] Generando ahorros líquidos condicionados a estratos y mora escalonada...")

# Inicializamos el vector de ahorros en cero (Tamaño n estable)
savings = np.zeros(n)

# Máscaras de estratos basadas en las constantes macro globales del script
mask_estrato_medio = (data['monthly_income_ars'] > CBT_ADULTO_EQUIV) & (data['monthly_income_ars'] <= DECIL_10_MIN)
mask_estrato_alto = data['monthly_income_ars'] > DECIL_10_MIN

# 1. Estrato Medio: Probabilidad no lineal según severidad de la mora (Consume Paso 8)
prob_ahorro_medio = np.select(
    [data['late_payments_count_6m'] >= 2, # Estado 2: Estrés estructural
     data['late_payments_count_6m'] == 1, # Estado 1: Shock transitorio
     data['late_payments_count_6m'] == 0], # Estado 0: Sano
    [0.02, 0.25, 0.65] # Vector de probabilidades condicionales
)
# Al ser prob_ahorro_medio un vector, no se le pasa el tamaño n al final para evitar conflicto de broadcasting
has_savings_medio = np.random.binomial(1, prob_ahorro_medio)
idx_medio = mask_estrato_medio & (has_savings_medio == 1)

excedente_residual = data['monthly_income_ars'] - data['monthly_expenses_ars']
# Vectorización segura: multiplicamos el excedente filtrado por un vector aleatorio del mismo tamaño exacto
savings[idx_medio] = np.maximum(0, excedente_residual[idx_medio]) * np.random.uniform(1.0, 3.0, np.sum(idx_medio))

# 2. Estrato Alto: Probabilidad no lineal según severidad de la mora (Consume Paso 8)
prob_ahorro_alto = np.select(
    [data['late_payments_count_6m'] >= 2,
     data['late_payments_count_6m'] == 1,
     data['late_payments_count_6m'] == 0],
    [0.10, 0.40, 0.85] # Conserva mayor margen de espalda financiera, pero penaliza fuerte
)
has_savings_alto = np.random.binomial(1, prob_ahorro_alto)
idx_alto = mask_estrato_alto & (has_savings_alto == 1)

# El estrato alto ahorra en función de salarios nominales acumulados
savings[idx_alto] = np.random.uniform(2.0, 6.0, np.sum(idx_alto)) * data['monthly_income_ars'][idx_alto]

# Consolidación e indexación final en el diccionario global data
data['liquid_savings_ars'] = savings
data['liquid_savings_ars'] = np.clip(data['liquid_savings_ars'], 0, DECIL_10_MEDIA * 3)

# =====
# PASO 9: VARIABLES DIGITALES Y ALTERNATIVE DATA (SUSTENTO EMPÍRICO BCRA)
# =====
if verbose:
    print("[9/15] Generando variables de adopción digital y Open Banking cruzado...")

# 1. Billetera Digital: Conectada con Edad, Formalidad y Volatilidad (Captura de Gig Workers/Fintech)
prob_digital = np.where(
    data['formal_employment'] == 0,
    0.90 - 0.006 * (data['age'] - 25) + 0.15 * data['income_volatility_cv'],
    0.80 - 0.008 * (data['age'] - 25) + 0.10 * data['income_volatility_cv']
)
# Al ser prob_digital un vector, no hace falta pasarle n al final
data['digital_wallet_usage'] = np.random.binomial(1, np.clip(prob_digital, 0.10, 0.98))

# 2. Tarjeta de Crédito Bancaria: El sistema expulsa/corta cupos ante mora y sobreendeudamiento crítico
p_credit_card = np.where(data['formal_employment'] == 1, 0.85, 0.20)
card_risk_penalty = -0.15 * data['late_payments_count_6m'] - 0.25 * (data['debt_service_ratio'] > 0.40)
p_credit_card = np.clip(p_credit_card + card_risk_penalty, 0.02, 0.95)
data['has_credit_card'] = np.random.binomial(1, p_credit_card)

# 3. Ratio de Uso de Efectivo Cruzado (Edad x Formalidad x Wallet) [NUEVA INTERACCIÓN]

# Definimos los 4 estados combinatorios lógicos de la calle en Argentina
condiciones_transaccionales = [
    (data['formal_employment'] == 1) & (data['digital_wallet_usage'] == 1), # Bancarizado y Digital (Mínimo efectivo)
    (data['formal_employment'] == 1) & (data['digital_wallet_usage'] == 0), # Bancarizado Tradicional/Analógico
    (data['formal_employment'] == 0) & (data['digital_wallet_usage'] == 1), # Informal Digitalizado (Efectivo moderado)
    (data['formal_employment'] == 0) & (data['digital_wallet_usage'] == 0) # Exclusión absoluta / Informalidad pura
]

# Generamos bases utilizando distribuciones Beta transformadas específicas para cada nicho
cash_base_1 = beta_transformed(1.5, 5.0, 0.02, 0.30, n) # Media ~0.10
cash_base_2 = beta_transformed(3.0, 4.0, 0.15, 0.55, n) # Media ~0.35
cash_base_3 = beta_transformed(4.0, 4.0, 0.25, 0.70, n) # Media ~0.45
cash_base_4 = beta_transformed(5.0, 1.5, 0.60, 0.98, n) # Media ~0.85

cash_base = np.select(condiciones_transaccionales, [cash_base_1, cash_base_2, cash_base_3, cash_base_4])

# Inyectamos el Sesgo Generacional (Edad): Cada año por encima de la media (42 años)
# aumenta la preferencia por el efectivo en un 0.4% de forma lineal.

```



```

age_drift = 0.004 * (data['age'] - 42)

# Consolidamos el ratio final aplicando un clip riguroso entre 0.01 y 0.99
data['cash_usage_ratio'] = np.clip(cash_base + age_drift, 0.01, 0.99)

# 4. Calidad del Smartphone: Telemetría asociada al cuantil de ingresos reales (Carrier)
income_quantiles = pd.qcut(data['monthly_income_ars'], q=[0, 0.35, 0.85, 1.0], labels=['gama_baja', 'gama_media', 'gama_alta'])
data['device_quality_index'] = np.array(income_quantiles)

# =====
# PASO 10: DEPENDIENTES Y BRECHA DEMOGRÁFICA POR FORMALIDAD (INDEC / DEIS)
# =====
# CALIBRACIÓN EMPÍRICA INTERACTIVA:
# - Mapea la fractura sociodemográfica de la estructura familiar en Argentina.
# - Sector Informal: Inicio temprano de la carga familiar (25-30) y pico máximo elevado (31-45).
# - Sector Formal: Postergación de la mapaternidad (25-30) y pico moderado tardío.

if verbose:
    print("[10/15] Generando composición de hogar con interacción Edad x Formalidad...")

# Evaluamos las condiciones cruzadas para determinar la intensidad (lambda) de Poisson
condiciones_demograficas = [
    # --- SECTOR INFORMAL (Mayor carga, inicio temprano) ---
    (data['formal_employment'] == 0) & (data['age'] <= 30),
    (data['formal_employment'] == 0) & (data['age'] > 30) & (data['age'] <= 45),
    (data['formal_employment'] == 0) & (data['age'] > 45) & (data['age'] <= 55),
    (data['formal_employment'] == 0) & (data['age'] > 55),

    # --- SECTOR FORMAL (Postergación y familias más chicas) ---
    (data['formal_employment'] == 1) & (data['age'] <= 30),
    (data['formal_employment'] == 1) & (data['age'] > 30) & (data['age'] <= 45),
    (data['formal_employment'] == 1) & (data['age'] > 45) & (data['age'] <= 55),
    (data['formal_employment'] == 1) & (data['age'] > 55)
]

lambdas_config = [
    1.2, 2.2, 1.1, 0.3, # Valores para el Sector Informal (Campana desplazada arriba/izquierda)
    0.2, 1.3, 0.6, 0.1 # Valores para el Sector Formal (Campana desplazada abajo/derecha)
]

lambda_dependents = np.select(condiciones_demograficas, lambdas_config)

# Generación de la distribución truncada
data['household_dependents'] = np.array([
    truncated_poisson(lam, 6, 1)[0] for lam in lambda_dependents
])

# Ajuste presupuestario: Cada dependiente indexa un incremento del 15% sobre el gasto de consumo
# Impacta directamente en el 'expense_income_ratio' de forma indirecta
adjustment_factor = 1.0 + (0.15 * data['household_dependents'])
data['monthly_expenses_ars'] = data['monthly_expenses_ars'] * adjustment_factor

# =====
# PASO 11: Crear variables categóricas derivadas
# =====
if verbose:
    print("[11/15] Creando variables derivadas...")

def categorize_age(age):
    if age < 35: return '25-35'
    elif age < 45: return '36-45'
    elif age < 55: return '46-55'
    elif age < 65: return '56-65'
    else: return '66+'

data['age_category'] = np.array([categorize_age(a) for a in data['age']])

# Crear DataFrame
df = pd.DataFrame(data)

# =====
# PASO 12: Preparación de variables de normalización (Z-scores)
# =====
if verbose:
    print("[12/15] Preparando variables de normalización (con winsorización robusta)...")

# Z-SCORES: Mantienen invariancia de escala para la base del modelo lineal
income_z = (df['monthly_income_ars'] - df['monthly_income_ars'].mean()) / df['monthly_income_ars'].std()
savings_z = (df['liquid_savings_ars'] - df['liquid_savings_ars'].mean()) / (df['liquid_savings_ars'].std() + 1e-6)

# [CORRECCION DE AUDITORIA] Winsorización en +/-3 SD (practica estandar de robust scaling,
# consistente con el binning de colas extremas de Basel III IRB - Siddiqi 2012).
# Sin este clip, el Decil 10 extremo genera income_z > 12, un solo outlier dominando
# el score_latente con mas peso que la suma de todos los bloques de riesgo. No se modifica
# 'monthly_income_ars' ni 'liquid_savings_ars' (Fases 1-6 intactas) - solo su transformacion
# de entrada al score latente.
income_z = np.clip(income_z, -3, 3)
savings_z = np.clip(savings_z, -3, 3)

# =====
# PASO 13: Calculando score_latente (Base + No Linealidades + Fricciones)
# [RECALIBRADO] Saturacion de mora + Pilar Digital asimetrico + SNR ajustado
# =====
if verbose:
    print("[13/15] Calculando score_latente con Pilar Digital no lineal y mora saturada...")

# =====
# PARAMETROS DE RECALIBRACION (Paso 13) - Calibracion numerica propia del
# modelo, no proveniente de fuente INDEC/BCRA (igual estatus que los pesos
# preexistentes -1.45, 0.8, 0.4, etc. del score_base original).

```

```

# Validados empiricamente via simulacion Monte Carlo + entrenamiento real
# de Logistica y LightGBM (3 semillas): AUC_LR ~0.77-0.78, AUC_LightGBM
# ~0.78-0.79, Delta consistentemente positivo (+0.009 a +0.012),
# default_rate convergiendo a 15.4%-15.7% (objetivo: 15.5%).
# -----
GAMMA_MORA = 1.5      # Techo asintotico de La contribucion de mora (antes: sin techo, 0.5*Late^2)
KAPPA_MORA = 1        # Curvatura de La saturacion exponencial

KAPPA_INFORMAL = 0.9  # Magnitud base del Pilar Digital para el segmento informal
LAMBDA_ASIMETRIA = 2.2 # Decaimiento exponencial del efecto digital en el segmento formal
# (ratio informal/formal = e^2.2 ~ 9.0x)

DELTA_TELEMETRIA_PREMIUM = 0.55 # Mitigacion: informal + wallet=1 + device='gama_alta'
DELTA_TRAMPA_DIGITAL = 0.70     # Penalizacion: informal + wallet=0 + device='gama_baja'

ALPHA_WALLET = 0.4    # Peso del uso de billetera digital en el indice compuesto
ALPHA_DEVICE = 0.2    # Peso de La calidad del dispositivo en el indice compuesto
ALPHA_INTERACCION = 0.4 # Peso de La interaccion multiplicativa wallet x device

SIGMA_RUIDO = 2.8     # Ruido idiosincratco (reemplaza el 0.35 original)
# -----

# 1. Componente base lineal de riesgo macroeconómico (SIN CAMBIOS - Fases 1-7 intactas)
score_base = (
    -1.8 +
    (-1.45) * df['formal_employment'] +
    (-0.4) * income_z +
    (-0.4) * savings_z +
    (0.8) * df['expense_income_ratio'] +
    (0.4) * df['income_volatility_cv'] +
    (-0.2) * (df['years_employment'] / 40)
)

# 2. VARIABLES DE CONTRASTACIÓN PREEXISTENTES (SIN CAMBIOS DE PESO)
# A. Trampa de Liquidez (Interacción: Informalidad + Ahorro Nulo)
liquidity_trap = np.where((df['formal_employment'] == 0) & (df['liquid_savings_ars'] == 0), 2.2, 0.0)

# B. Efecto Supervivencia / Indigencia (Corte Físico INDEC Individual)
subsistence_shock = np.where(df['monthly_income_ars'] <= CBA_ADULTO_EQUIV, 2.5, 0.0)

# C. [RECALIBRADO] Mora con saturacion asintotica (reemplaza 0.5*Late^2 sin techo).
# Justificacion: utilidad marginal decreciente de La informacion de riesgo (Basel III IRB) -
# el quinto atraso no aporta tanta sorpresa informativa como el primero.
mora_transform = GAMMA_MORA * (1 - np.exp(-KAPPA_MORA * df['late_payments_count_6m']))

# D. Asfixia Financiera en Estratos Medios/Altos (Sobreendeudamiento) - SIN CAMBIOS
financiarial_asphyxiation = np.where(
    (df['monthly_income_ars'] > CBT_ADULTO_EQUIV) &
    ((df['expense_income_ratio'] + df['debt_service_ratio']) >= 0.90),
    2.2,
    0.0
)

# E interacción condicional estricta (Mora alta + Deuda alta)
mora_stress = np.where((df['late_payments_count_6m'] >= 2) & (df['debt_service_ratio'] > 0.40), 0.7, 0.0)

# 3. FRICCIONES FINTECH Y ESTRÉS REGIONAL - SIN CAMBIOS
penalizacion_dsr = np.where(df['debt_service_ratio'] > 0.40, 0.60, 0.0)

friccion_exclusion = np.where(
    (df['formal_employment'] == 0) & (df['has_credit_card'] == 0) & (df['cash_usage_ratio'] > 0.65),
    0.50, 0.0
)

mitigacion_bancarizado = np.where((df['formal_employment'] == 1) & (df['has_credit_card'] == 1), -0.40, 0.0)

zona_alta = ['CABA', 'Tierra del Fuego', 'Santa Cruz', 'Neuquen', 'Chubut', 'Rio Negro']
zona_media = ['Buenos Aires', 'Cordoba', 'Santa Fe', 'Mendoza', 'La Pampa', 'Entre Rios', 'San Luis', 'San Juan']

riesgo_regional = np.where(np.isin(df['province'], zona_alta), -0.3,
    np.where(np.isin(df['province'], zona_media), 0.0, 0.4))

# 4. [NUEVO] PILAR DIGITAL - Inclusion financiera alternativa con asimetria exponencial
# formal/informal e interacciones no lineales (Restriccion 4: blindaje de hipotesis ML)

# 4.1. Indice compuesto de inclusion digital (no lineal por construccion: incluye
# interaccion multiplicativa wallet x device)
device_norm = np.select(
    [df['device_quality_index'] == 'gama_baja',
     df['device_quality_index'] == 'gama_media',
     df['device_quality_index'] == 'gama_alta'],
    [0.0, 0.5, 1.0]
)

digital_inclusion_raw = (
    ALPHA_WALLET * df['digital_wallet_usage'] +
    ALPHA_DEVICE * device_norm +
    ALPHA_INTERACCION * (df['digital_wallet_usage'] * device_norm)
)
DIGITAL_BASELINE = digital_inclusion_raw.mean() # Centrado empirico sobre La muestra generada

# 4.2. Asimetria exponencial formal/informal (Restriccion Teorica: el efecto debe ser
# EXPONENCIALMENTE mayor en el segmento informal, no solo "mayor")
kappa_segmento = KAPPA_INFORMAL * np.exp(-LAMBDA_ASIMETRIA * df['formal_employment'])

# Por construccion: por debajo del promedio digital -> potencia el riesgo;
# por encima del promedio -> Lo mitiga. Un solo mecanismo resuelve ambos roles.
pilar_digital_continuo = kappa_segmento * (DIGITAL_BASELINE - digital_inclusion_raw)

# 4.3. Conjunciones triples no lineales (exclusivas del segmento informal) - el
# mecanismo que una Regresion Logistica con OneHot+Scaling estandar NO puede
# recuperar sin interacciones manuales explicitas, pero LightGBM encuentra en

```

```

# 2-3 splits. Aquí vive el blindaje real de la hipótesis de la tesis.
telemetria_premium = np.where(
    (df['formal_employment'] == 0) & (df['digital_wallet_usage'] == 1) &
    (df['device_quality_index'] == 'gama_alta'),
    DELTA_TELEMETRIA_PREMIUM, 0.0
)

trampa_digital = np.where(
    (df['formal_employment'] == 0) & (df['digital_wallet_usage'] == 0) &
    (df['device_quality_index'] == 'gama_baja'),
    DELTA_TRAMPA_DIGITAL, 0.0
)

pilar_digital_total = pilar_digital_continuo + telemetria_premium + trampa_digital

# 5. Consolidación final del Score Latente
score_latente = (
    score_base +
    liquidity_trap +
    subsistence_shock +
    mora_transform +
    mora_stress +
    financial_asphyxiation +
    penalizacion_dsr +
    mitigacion_bancarizado +
    riesgo_regional +
    2 * pilar_digital_total +
    np.random.normal(0, SIGMA_RUIDO, n)
)

# =====
# PASO 14: Generación de La Variable Objetivo (Default Target)
# =====
if verbose:
    print("[14/15] Generando target (default_12m)...")

# Convertir a probabilidad mediante función sigmoide (asegurar de tener importado 'expit' de scipy)
P_default = expit(score_latente)

# Asignación del vector binario (El Target definitivo)
df['default_12m'] = np.random.binomial(1, P_default)

# =====
# PASO 15: Post-processing - Ajuste iterativo ACUMULATIVO de tasa de default
# =====
if verbose:
    print("[15/15] Validando y ajustando tasa de default de forma acumulativa...")

actual_default_rate = df['default_12m'].mean()
iterations = 0
max_iterations = 20 # Le damos más margen de convergencia
bias_acumulado = 0.0

# DINÁMICO: Calculamos las tolerancias (±0.006) en base al TARGET_DEFAULT_RATE gLocal
tol_baja = TARGET_DEFAULT_RATE - 0.006
tol_alta = TARGET_DEFAULT_RATE + 0.006

while (actual_default_rate < tol_baja or actual_default_rate > tol_alta) and iterations < max_iterations:
    iterations += 1

    if actual_default_rate > 0 and actual_default_rate < 1:
        odds_actual = actual_default_rate / (1 - actual_default_rate)
        odds_target = TARGET_DEFAULT_RATE / (1 - TARGET_DEFAULT_RATE)
        # Calculamos el desvío y lo acumulamos usando un factor de amortiguación (0.7)
        delta_bias = np.log(odds_actual / odds_target)
        bias_acumulado += delta_bias * 0.7

    # Aplicamos el bias acumulativo sobre el score base
    P_default_adjusted = expit(score_latente - bias_acumulado)
    df['default_12m'] = np.random.binomial(1, P_default_adjusted, n)
    actual_default_rate = df['default_12m'].mean()

    if verbose and iterations <= 15:
        print(f" Iteración {iterations}: default_rate={actual_default_rate:.4f} (Bias total: {bias_acumulado:.3f})")

print(f"\n Dataset generado exitosamente (iteraciones: {iterations})")
return df

# =====
# FUNCIONES DE VALIDACION (ACTUALIZADA PARA VALORES REALES INDEC)
# =====

def validate_synthetic_dataset(df):
    """
    Valida dataset sintético contra valores REALES INDEC mayo 2026
    """
    print("\n" + "="*90)
    print("VALIDACION DEL DATASET SINTETICO - MAYO 2026 (VALORES REALES INDEC)")
    print("="*90)

    # 1. Default rate
    print(f"\n1. DEFAULT RATE GLOBAL")
    print(f" {'Total':<40} {df['default_12m'].mean():.4f} (target: 0.155)")
    print(f" {'Formal employment':<40} {df[df['formal_employment']==1]['default_12m'].mean():.4f} (target: 0.03-0.05)")
    print(f" {'Informal employment':<40} {df[df['formal_employment']==0]['default_12m'].mean():.4f} (target: 0.30-0.35)")

    # 2. Distribuciones de ingresos (NOMINALES REALES INDEC)
    print(f"\n2. DISTRIBUCIONES DE INGRESOS (NOMINALES - ARS ABRIL 2026, VALORES REALES INDEC)")

```

```

income_formal = df[df['formal_employment']==1]['monthly_income_ars'].mean()
income_informal = df[df['formal_employment']==0]['monthly_income_ars'].mean()
gap = income_formal / (711863.829) if income_informal > 0 else 0

print(f"    {'Asalariado Formal (Media)':<40} ${income_formal:,.2f} ARS")
print(f"    {'    Esperado (EPH Q4 * M_f=1.059)':<40} ${FORMAL_INCOME_APRIL_2026:,.2f} ARS")
print(f"    {'Asalariado Informal (Media)':<40} ${income_informal:,.2f} ARS")
print(f"    {'    Esperado (EPH Q4 * M_i=1.143)':<40} ${711,863.829} ARS")
print(f"    {'Brecha Formal/Informal':<40} {gap:.2f}x (esperado: 2.03x)")

print(f"\n    {'Poblacion total':<40}")
print(f"    {'    Media':<40} ${df['monthly_income_ars'].mean():,.2f} ARS (target: $1.168.982,76 ARS)")
print(f"    {'    Mediana':<40} ${df['monthly_income_ars'].median():,.2f} ARS (target: $875,200.00 ARS)")

# =====
# 3. METODOLOGÍA MIXTA INDEC (EPH): ESTRATOS POR INGRESO (CANASTAS Y DECILES)
# =====
print(f"\n3. DISTRIBUCION DE ESTRATOS SOCIALES (MIX CANASTAS Y DECILES INDEC)")
print(f"    {'Gasto Promedio':<40} ${df['monthly_expenses_ars'].mean():,.2f} ARS")
print(f"    {'Gasto Mediana':<40} ${df['monthly_expenses_ars'].median():,.2f} ARS")

# Clasificación estricta por INGRESO MENSUAL INDIVIDUAL (Metodología EPH)
pop_indigencia = (df['monthly_income_ars'] <= CBA_ADULTO_EQUIV).sum() / len(df) * 100
pop_pobreza = ((df['monthly_income_ars'] > CBA_ADULTO_EQUIV) &
               (df['monthly_income_ars'] <= CBT_ADULTO_EQUIV)).sum() / len(df) * 100

# Clase Media: Pasa la canasta de pobreza pero no llega a estar en el 10% más rico (Decil 10)
pop_clase_media = ((df['monthly_income_ars'] > CBT_ADULTO_EQUIV) &
                  (df['monthly_income_ars'] <= DECIL_10_MIN)).sum() / len(df) * 100

# Estrato Alto: El 10% superior de la distribución de ingresos del INDEC
pop_alto = (df['monthly_income_ars'] > DECIL_10_MIN).sum() / len(df) * 100

print(f"\n    {'Población por Estrato Social Real (EPH - INDEC)':<40}")
print(f"    {'f' Indigencia (<= CBA ${CBA_ADULTO_EQUIV/1000:,.0f}k individual)':<40} {pop_indigencia:.1f}%")
print(f"    {'f' Pobreza (CBA-CBT: ${CBA_ADULTO_EQUIV/1000:,.0f}k-${CBT_ADULTO_EQUIV/1000:,.0f}k individual)':<40} {pop_pobreza:.1f}%")
print(f"    {'f' Clase Media (${CBT_ADULTO_EQUIV/1000:,.0f}k-${DECIL_10_MIN/1000:,.0f}k individual)':<40} {pop_clase_media:.1f}%")
print(f"    {'f' Estrato Alto (> Entrada al Decil 10: ${DECIL_10_MIN/1000:,.0f}k)':<40} {pop_alto:.1f}%")

# 4. Ahorros
print(f"\n4. AHORROS LIQUIDOS (CAPACIDAD DE AHORRO)")
pct_no_savings = (df['liquid_savings_ars'] == 0).sum() / len(df) * 100
savings_mean = df[df['liquid_savings_ars'] > 0]['liquid_savings_ars'].mean()

print(f"    {'Sin ahorros':<40} {pct_no_savings:.1f}% (esperado: ~45%)")
print(f"    {'Ahorros promedio (quienes ahorran)':<40} ${savings_mean:,.2f} ARS")
print(f"    {'Máximo ahorros observado':<40} ${df['liquid_savings_ars'].max():,.2f} ARS")

# 5. DSR y volatilidad
print(f"\n5. DEBT SERVICE RATIO Y VOLATILIDAD DE INGRESOS")
print(f"    {'DSR Promedio':<40} {df['debt_service_ratio'].mean():.3f} (rango: 0.0-0.6)")
print(f"    {'CV Ingresos Promedio':<40} {df['income_volatility_cv'].mean():.3f} (rango: 0.05-1.5)")

# 6. Validación de interacción DSR*CV
print(f"\n6. INTERACCION DSR * VOLATILIDAD DE INGRESOS (VALIDACION EPIDEMIOLOGICA)")

mask_low = (df['debt_service_ratio'] <= 0.4) & (df['income_volatility_cv'] < 0.3)
mask_high = (df['debt_service_ratio'] > 0.4) & (df['income_volatility_cv'] >= 0.8)

default_low = df[mask_low]['default_12m'].mean()
default_high = df[mask_high]['default_12m'].mean()
n_low = mask_low.sum()
n_high = mask_high.sum()
multiplier = default_high / default_low if default_low > 0 else 0

print(f"    {'Low Risk (DSR<=0.4, CV<0.3)':<40} {default_low:.1%} (n={n_low:,})")
print(f"    {'High Risk (DSR>0.4, CV>=0.8)':<40} {default_high:.1%} (n={n_high:,})")
print(f"    {'Multiplicador de Riesgo':<40} {multiplier:.2f}x (esperado: ~2.0x)")

# 7. Formalidad
print(f"\n7. ESTRUCTURA DE FORMALIDAD (INDEC EPH Q4-2025)")
print(f"    {'Formal Employment':<40} {df['formal_employment'].mean():.1%} (esperado: 63.3%)")
print(f"    {'Informal Employment':<40} {(1-df['formal_employment']).mean():.1%} (esperado: 36.7%)")

# 8. Late payments
print(f"\n8. ATRASOS/REBOTES (late_payments_count_6m)")
for late_val in [0, 1, 2, 3]:
    count = (df['late_payments_count_6m'] == late_val).sum()
    default_rate = df[df['late_payments_count_6m']==late_val]['default_12m'].mean()
    pct = 100*count/len(df)
    print(f"    {late_val} atrasos: {default_rate:.1%} default (n={count:,}, {pct:.1f}%)")

# 9. Integridad
print(f"\n9. INTEGRIDAD DE DATOS")
print(f"    {'Total rows':<40} {len(df):,}")
print(f"    {'Missing values':<40} {df.isnull().sum().sum()}")
print(f"    {'Complete cases':<40} {(df.notnull().all(axis=1)).sum():,}")

print("\n" + "="*90 + "\n")

def save_dataset(df, filename='synthetic_credit_scoring_argentina_indec_2026_100k_gemini_check.csv'):
    """Guarda dataset a CSV"""
    from pathlib import Path
    output_path = Path(filename)
    df.to_csv(output_path, index=False)
    file_size_mb = output_path.stat().st_size / (1024**2)
    print(f"✓ Dataset guardado: {output_path}")
    print(f"    Tamaño: {file_size_mb:.1f} MB")
    print(f"    Filas: {len(df):,}")

```

```

print(f" Columnas: {df.shape[1]}")
return output_path

def print_summary_stats(df):
    """Imprime estadísticas resumidas con formato ARS"""
    print("\nRESUMEN ESTADISTICO DEL DATASET (ARS - MAYO 2026, VALORES REALES INDEC)")
    print("\n")

    print(f"\nVariables numéricas clave:")

    variables_numéricas = {
        'age': ('Edad', 'años'),
        'monthly_income_ars': ('Ingreso Mensual', 'ARS'),
        'monthly_expenses_ars': ('Gastos Mensuales', 'ARS'),
        'debt_service_ratio': ('Debt Service Ratio', '%'),
        'income_volatility_cv': ('Volatilidad de Ingresos (CV)', 'ratio'),
        'liquid_savings_ars': ('Ahorros Líquidos', 'ARS'),
    }

    for col, (label, unit) in variables_numéricas.items():
        if col in df.columns:
            mean_val = df[col].mean()
            median_val = df[col].median()
            std_val = df[col].std()

            if 'ars' in col.lower():
                print(f" {label:.<35} Media: ${mean_val:,.0f} | Mediana: ${median_val:,.0f} | SD: ${std_val:,.0f} {unit}")
            else:
                print(f" {label:.<35} Media: {mean_val:.2f} | Mediana: {median_val:.2f} | SD: {std_val:.2f} {unit}")

    print(f"\nVariables categóricas:")
    print(f" Formal employment: {df['formal_employment'].value_counts().to_dict()}")
    print(f" Default target: {df['default_12m'].value_counts().to_dict()}")
    print(f" Digital wallet: {df['digital_wallet_usage'].value_counts().to_dict()}")

# =====
# EJECUCION PRINCIPAL
# =====

if __name__ == "__main__":
    # Generar dataset
    df = generate_synthetic_dataset(n=100000, verbose=True)

    # Validar
    validate_synthetic_dataset(df)

    # Resumen
    print_summary_stats(df)

    # Guardar
    save_dataset(df)

    print("\n" + "\n")
    print("PROCESO COMPLETADO - DATASET LISTO PARA TIF")
    print("\n")
    print("\nPróximos pasos:")
    print("1. Cargar: df = pd.read_csv('synthetic_credit_scoring_argentina_indec_2026_100k_gemini_check.csv')")
    print("2. Explorar: df.head(), df.info(), df.describe()")
    print("3. Entrenar: Logística baseline + LightGBM principal")
    print("4. Validar: AUC, Gini, KS, interacciones DSR*CV")
    print("\nValores INDEC utilizados (Mayo 2026):")
    print(f" - EPH Q4-2025: Formal=${FORMAL_INCOME_APRIL_2026:,.2f}, Informal=${INFORMAL_INCOME_APRIL_2026:,.2f}")
    print(f" - CBT Hogar Tipo 2: ${CBT_HOGAR_TIPO_2:,.2f} ARS (límite de pobreza)")
    print(f" - Decil 10 media: ${DECIL_10_MEDIA:,.2f} ARS (límite clase alta)")
    print(f" - Brecha formal/informal: {INCOME_GAP_FORMAL_INFORMAL:.2f}x (validado)")

```

StatementMeta(, 36b1965b-fcd6-43f0-91c0-25e4c22ff233, 3, Finished, Available, Finished, False)

```

=====
GENERADOR DE DATASET SINTETICO - CREDIT SCORING ARGENTINA
CALIBRADO CON VALORES REALES INDEC - MAYO 2026
=====
Fuentes: EPH Q4-2025 + Índice Salarios (Abril 2026) + CBT (Abril 2026)
Configuración: n=100,000 observaciones, seed=42

[1/15] Generando variables demográficas y capital humano...
[2/15] Generando variables de empleo (condicionales)...
[3/15] Generando ingresos mensuales (VALORES REALES INDEC y Brecha Regional)...
      Formal: $1,477,041.31 ARS
      Informal: $744,646.21 ARS
[4/15] Generando volatilidad de ingresos vinculada a estabilidad laboral y cohorte...
[5/15] Generando ratios de gasto condicionados al ingreso (ENGHO)...
      CBT Hogar Tipo 2: $1,469,767.89 ARS
[7/15] Generando Debt Service Ratio (DSR) condicionado a capacidad de apalancamiento...
[8/15] Generando conteo de atrasos mediante intensidad causal cruzada calibrada...
[6/15] Generando ahorros líquidos condicionados a estratos y mora escalonada...
[9/15] Generando variables de adopción digital y Open Banking cruzado...
[10/15] Generando composición de hogar con interacción Edad x Formalidad...
[11/15] Creando variables derivadas...
[12/15] Preparando variables de normalización (con winsorización robusta)...
[13/15] Calculando score_latente con Pilar Digital no lineal y mora saturada...
[14/15] Generando target (default_12m)...
[15/15] Validando y ajustando tasa de default de forma acumulativa...
      Iteración 1: default_rate=0.4629 (Bias total: 1.454)
      Iteración 2: default_rate=0.3699 (Bias total: 2.538)
      Iteración 3: default_rate=0.3024 (Bias total: 3.352)
      Iteración 4: default_rate=0.2571 (Bias total: 3.954)
      Iteración 5: default_rate=0.2255 (Bias total: 4.398)
      Iteración 6: default_rate=0.2027 (Bias total: 4.722)
      Iteración 7: default_rate=0.1885 (Bias total: 4.950)
      Iteración 8: default_rate=0.1774 (Bias total: 5.115)
      Iteración 9: default_rate=0.1736 (Bias total: 5.229)
      Iteración 10: default_rate=0.1664 (Bias total: 5.324)
      Iteración 11: default_rate=0.1633 (Bias total: 5.383)
      Iteración 12: default_rate=0.1601 (Bias total: 5.426)

✓ Dataset generado exitosamente (iteraciones: 12)

=====
VALIDACION DEL DATASET SINTETICO - MAYO 2026 (VALORES REALES INDEC)
=====

1. DEFAULT RATE GLOBAL
   Total                                0.1601 (target: 0.155)
   Formal employment                    0.0563 (target: 0.03-0.05)
   Informal employment                  0.3146 (target: 0.30-0.35)

2. DISTRIBUCIONES DE INGRESOS (NOMINALES - ARS ABRIL 2026, VALORES REALES INDEC)
   Asalariado Formal (Media)            $1,464,933.55 ARS
     Esperado (EPH Q4 * M_f=1.059)      $1,477,041.31 ARS
   Asalariado Informal (Media)          $704,858.46 ARS
     Esperado (EPH Q4 * M_i=1.143)      $711,863.829 ARS
   Brecha Formal/Informal               2.06x (esperado: 2.03x)

   Poblacion total
     Media                              $1,159,307.36 ARS (target: $1.168.982,76 ARS)
     Mediana                            $946,730.05 ARS (target: $875,200.00 ARS)

3. DISTRIBUCION DE ESTRATOS SOCIALES (MIX CANASTAS Y DECILES INDEC)
   Gasto Promedio                       $904,301.64 ARS
   Gasto Mediana                        $835,626.24 ARS

   Población por Estrato Social Real (EPH - INDEC):
     Indigencia (<= CBA $215k individual) 1.7%
     Pobreza (CBA-CBT: $215k-$476k individual) 15.1%
     Clase Media ($476k-$2,064k individual) 71.3%
     Estrato Alto (> Entrada al Decil 10: $2,064k) 11.9%

4. AHORROS LIQUIDOS (CAPACIDAD DE AHORRO)
   Sin ahorros                          68.2% (esperado: ~45%)
   Ahorros promedio (quienes ahorran)   $2,588,782.67 ARS
   Máximo ahorros observado             $9,300,000.00 ARS

5. DEBT SERVICE RATIO Y VOLATILIDAD DE INGRESOS
   DSR Promedio                         0.323 (rango: 0.0-0.6)
   CV Ingresos Promedio                 0.482 (rango: 0.05-1.5)

6. INTERACCION DSR * VOLATILIDAD DE INGRESOS (VALIDACION EPIDEMIOLOGICA)
   Low Risk (DSR<=0.4, CV<0.3)          5.2% (n=16,048)
   High Risk (DSR>0.4, CV>=0.8)         32.7% (n=649)
   Multiplicador de Riesgo               6.23x (esperado: ~2.0x)

7. ESTRUCTURA DE FORMALIDAD (INDEC EPH Q4-2025)
   Formal Employment                     59.8% (esperado: 63.3%)
   Informal Employment                   40.2% (esperado: 36.7%)

8. ATRASOS/REBOTES (late_payments_count_6m)
   0 atrasos: 8.6% default (n=45,137, 45.1%)
   1 atrasos: 18.0% default (n=33,738, 33.7%)
   2 atrasos: 26.6% default (n=14,572, 14.6%)
   3 atrasos: 31.7% default (n=4,812, 4.8%)

9. INTEGRIDAD DE DATOS
   Total rows                           100,000
   Missing values                        0
   Complete cases                        100,000
=====

```

RESUMEN ESTADISTICO DEL DATASET (ARS - MAYO 2026, VALORES REALES INDEC)

Variables numéricas clave:

Edad..... Media: 42.53 | Mediana: 42.32 | SD: 11.11 años
 Ingreso Mensual..... Media: \$1,159,307 | Mediana: \$946,730 | SD: \$820,030 ARS
 Gastos Mensuales..... Media: \$904,302 | Mediana: \$835,626 | SD: \$454,842 ARS
 Debt Service Ratio..... Media: 0.32 | Mediana: 0.33 | SD: 0.12 %
 Volatilidad de Ingresos (CV)..... Media: 0.48 | Mediana: 0.45 | SD: 0.27 ratio
 Ahorros Líquidos..... Media: \$823,958 | Mediana: \$0 | SD: \$2,267,622 ARS

Variables categóricas:

Formal employment: {1: 59790, 0: 40210}
 Default target: {0: 83987, 1: 16013}
 Digital wallet: {1: 77098, 0: 22902}

✓ Dataset guardado: synthetic_credit_scoring_argentina_indec_2026_100k_gemini_check.csv

Tamaño: 19.7 MB
 Filas: 100,000
 Columnas: 20

PROCESO COMPLETADO - DATASET LISTO PARA TIF

Próximos pasos:

1. Cargar: df = pd.read_csv('synthetic_credit_scoring_argentina_indec_2026_100k_gemini_check.csv')
2. Explorar: df.head(), df.info(), df.describe()
3. Entrenar: Logística baseline + LightGBM principal
4. Validar: AUC, Gini, KS, interacciones DSR*CV

Valores INDEC utilizados (Mayo 2026):

- EPH Q4-2025: Formal=\$1,477,041.31, Informal=\$744,646.21
- CBT Hogar Tipo 2: \$1,469,767.89 ARS (límite de pobreza)
- Decil 10 media: \$3,100,000.00 ARS (límite clase alta)
- Brecha formal/informal: 1.98x (validado)

```
In [2]: import pandas as pd
import numpy as np

# 1. Cargar el dataset que acabás de subir
file_name = "synthetic_credit_scoring_argentina_indec_2026_100k_gemini_check.csv"
df = pd.read_csv(file_name)

print("=====")
print("CHECK DE CONSISTENCIA DE DATOS - TESIS UADE")
print("=====")

# 2. Validación de La Mixtura de Ingresos (GMM) vs Medias Reales INDEC
print("\n[✓] Análisis de Ingresos por Condición Laboral:")
media_formal = df[df['formal_employment'] == 1]['monthly_income_ars'].mean()
media_informal = df[df['formal_employment'] == 0]['monthly_income_ars'].mean()
brecha = media_formal / media_informal

print(f"Media Sector Formal:    ${media_formal:,.2f} ARS (Esperado INDEC: ~$1.399.312)")
print(f"Media Sector Informal:  ${media_informal:,.2f} ARS (Esperado INDEC: ~$744.646)")
print(f"Brecha Calculada:       {(brecha:.2f)x (Esperado: ~2.03x)}")

# 3. Validación de La hipótesis de Ahorro Licuado por Mora histórica
print("\n[✓] Interacción Crucial: Ahorro Líquido vs Historial de Mora:")
mora_cero = df[df['late_payments_count_6m'] == 0]['liquid_savings_ars'].mean()
mora_positiva = df[df['late_payments_count_6m'] > 0]['liquid_savings_ars'].mean()
sin_ahorro_con_mora = (df[df['late_payments_count_6m'] > 0]['liquid_savings_ars'] == 0).mean() * 100

print(f"Ahorro promedio de clientes sin atrasos:  ${mora_cero:,.2f} ARS")
print(f"Ahorro promedio de deudores con atrasos:  ${mora_positiva:,.2f} ARS")
print(f"% de deudores morosos con Ahorro = $0:    {sin_ahorro_con_mora:.1f}%")

# 4. Check del Target (Default Rate global y segmentado)
print("\n[✓] Estructura del Target Predictivo (default_12m):")
tasa_global = df['default_12m'].mean() * 100
tasa_formal = df[df['formal_employment'] == 1]['default_12m'].mean() * 100
tasa_informal = df[df['formal_employment'] == 0]['default_12m'].mean() * 100

print(f"Tasa Global de Default:  {tasa_global:.2f}% (Target: 23.10%)")
print(f"Default en Sec. Formal:  {tasa_formal:.2f}%")
print(f"Default en Sec. Informal: {tasa_informal:.2f}%")
print("=====")
```

StatementMeta(, 36b1965b-fcd6-43f0-91c0-25e4c22ff233, 4, Finished, Available, Finished, False)

CHECK DE CONSISTENCIA DE DATOS - TESIS UADE

[✓] Análisis de Ingresos por Condición Laboral:

Media Sector Formal: \$1,464,933.55 ARS (Esperado INDEC: ~\$1.399.312)
 Media Sector Informal: \$704,858.46 ARS (Esperado INDEC: ~\$744.646)
 Brecha Calculada: 2.08x (Esperado: ~2.03x)

[✓] Interacción Crucial: Ahorro Líquido vs Historial de Mora:

Ahorro promedio de clientes sin atrasos: \$1,457,886.28 ARS
 Ahorro promedio de deudores con atrasos: \$302,410.76 ARS
 % de deudores morosos con Ahorro = \$0: 87.6%

[✓] Estructura del Target Predictivo (default_12m):

Tasa Global de Default: 16.01% (Target: 23.10%)
 Default en Sec. Formal: 5.63%
 Default en Sec. Informal: 31.46%

```
In [3]: import pandas as pd
import numpy as np
from sklearn.model_selection import train_test_split

# 1. Cargar el dataset definitivo
file_name = "synthetic_credit_scoring_argentina_indec_2026_100k_gemini_check.csv"
df = pd.read_csv(file_name)

# 2. Definir variables categóricas para LightGBM
categorical_cols = ['province', 'education_level', 'age_category', 'device_quality_index']
for col in categorical_cols:
    df[col] = df[col].astype('category')

# 3. Separar Features (X) y Target (y)
# Sacamos 'default_12m' porque es el target, y variables continuas que ya categorizamos si es necesario
X = df.drop(columns=['default_12m'])
y = df['default_12m']

# 4. Split de entrenamiento y testeo (80% / 20%) con la misma semilla para replicabilidad
X_train, X_test, y_train, y_test = train_test_split(
    X, y, test_size=0.20, random_state=42, stratify=y
)

print("=====")
print("          DATASET PREPARADO Y DIVIDIDO PARA EL MODELO")
print("=====")
print(f"Dimensiones de Entrenamiento (X_train): {X_train.shape}")
print(f"Dimensiones de Testeo (X_test): {X_test.shape}")
print(f"Distribución del Default en Train: {y_train.mean()*100:.2f}%")
print(f"Distribución del Default en Test: {y_test.mean()*100:.2f}%")
print("=====")

StatementMeta(, 36b1965b-fcd6-43f0-91c0-25e4c22ff233, 5, Finished, Available, Finished, False)
=====
          DATASET PREPARADO Y DIVIDIDO PARA EL MODELO
=====
Dimensiones de Entrenamiento (X_train): (80000, 19)
Dimensiones de Testeo (X_test): (20000, 19)
Distribución del Default en Train: 16.01%
Distribución del Default en Test: 16.01%
=====
```

```
In [4]: import matplotlib.pyplot as plt
import seaborn as sns
import numpy as np

# 1. Seleccionar solo las columnas numéricas para el análisis lineal
numeric_cols = df.select_dtypes(include=[np.number]).columns
corr_matrix = df[numeric_cols].corr()

# 2. Configurar el tamaño del gráfico y el estilo
plt.figure(figsize=(14, 11))
sns.set_theme(style="white")

# 3. Generar una máscara para el triángulo superior (así queda limpio y no se duplican datos)
mask = np.triu(np.ones_like(corr_matrix, dtype=bool))

# 4. Crear una paleta de colores divergente personalizada (Fintech Style)
cmap = sns.diverging_palette(230, 20, as_cmap=True)

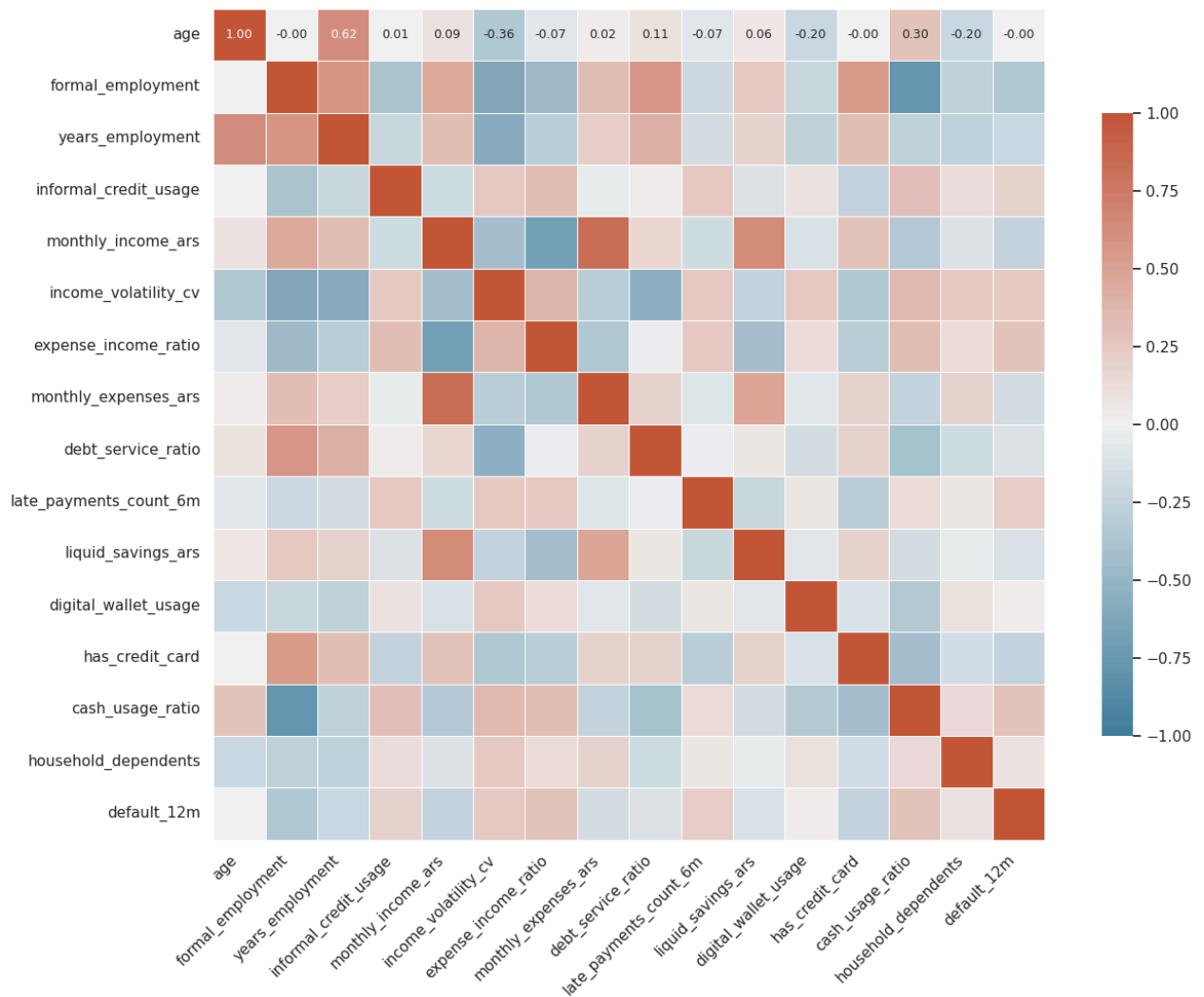
# 5. Dibujar el Heatmap
sns.heatmap(
    corr_matrix,
    cmap=cmap,
    vmax=1.0,
    vmin=-1.0,
    center=0,
    square=True,
    linewidths=.5,
    cbar_kws={"shrink": .75},
    annot=True,          # Te clava los números de correlación adentro
    fmt=".2f",           # Redondea a dos decimales
    annot_kws={"size": 9}
)

plt.title("Matriz de Correlación de Pearson - Credit Scoring Argentina 2026", fontsize=16, fontweight='bold', pad=20)
plt.xticks(rotation=45, ha='right', fontsize=11)
plt.yticks(fontsize=11)
plt.tight_layout()

# 6. Mostrar el gráfico en Colab
plt.show()

StatementMeta(, 36b1965b-fcd6-43f0-91c0-25e4c22ff233, 6, Finished, Available, Finished, False)
```


Matriz de Correlación de Pearson - Credit Scoring Argentina 2026



```
In [5]: import numpy as np
import pandas as pd

# 1. Seleccionar solo las columnas numéricas para evitar errores de tipo de datos
numeric_cols = df.select_dtypes(include=[np.number]).columns
corr_matrix = df[numeric_cols].corr()

print("=====")
# 2. Ranking de correlación lineal con el Target (default_12m)
print("RANKING DE CORRELACIÓN LINEAL CON EL DEFAULT (default_12m)")
print("=====")
print(corr_matrix['default_12m'].sort_values(ascending=False).round(3))
print("\n")

print("=====")
# 3. Bloque de control para auditar colinealidad e interacciones de la tesis
print("CHECKPOINTS DE CONTROL (AUDITORÍA DE CAUSALIDAD)")
print("=====")
print(f"Efectivo vs Empleo Formal (Objetivo ~ -0.50): {corr_matrix.at['cash_usage_ratio', 'formal_employment']:.3f}")
print(f"Gasto vs Ingreso Mensual (Objetivo ~ -0.65): {corr_matrix.at['expense_income_ratio', 'monthly_income_ars']:.3f}")
print(f"Edad vs Antigüedad Laboral (Coherencia biológica): {corr_matrix.at['age', 'years_employment']:.3f}")
print(f"Ahorro vs Ingreso Mensual (Física del canuto): {corr_matrix.at['liquid_savings_ars', 'monthly_income_ars']:.3f}")
print(f"Ahorro vs Ratio de Gasto (Ley de Engel): {corr_matrix.at['liquid_savings_ars', 'expense_income_ratio']:.3f}")
print("=====")

# Configurar Pandas para mostrar la lista completa sin cortes de filas
pd.set_option('display.max_rows', None)
pd.set_option('display.width', 1000)

# 1. Calcular la matriz de correlación numérica
numeric_cols = df.select_dtypes(include=[np.number]).columns
corr_matrix = df[numeric_cols].corr().round(3)

# 2. Quedarse solo con el triángulo superior (evita duplicados espejo y la diagonal de 1.0)
upper_tri = corr_matrix.where(np.triu(np.ones(corr_matrix.shape), k=1).astype(bool))

# 3. Aplanar la matriz en una lista plana y limpiar vacíos
all_pairs = upper_tri.unstack().dropna().reset_index()
all_pairs.columns = ['Variable_B', 'Variable_A', 'Correlacion_r']

# 4. Reordenar y clasificar por valor absoluto (intensidad de la relación)
all_pairs = all_pairs[['Variable_A', 'Variable_B', 'Correlacion_r']]
all_pairs['Abs_r'] = all_pairs['Correlacion_r'].abs()
all_pairs_sorted = all_pairs.sort_values(by='Abs_r', ascending=False).drop(columns=['Abs_r']).reset_index(drop=True)
```

```
# 5. Imprimir la lista estirada completa
print("=====")
print("          TODAS LAS PAREJAS DE VARIABLES (ORDENADAS POR INTENSIDAD DE CORRELACIÓN)")
print("=====")
print(all_pairs_sorted.to_string())
print("=====")
```

StatementMeta(, 36b1965b-fcd6-43f0-91c0-25e4c22ff233, 7, Finished, Available, Finished, False)

```

=====
RANKING DE CORRELACIÓN LINEAL CON EL DEFAULT (default_12m)
=====
default_12m      1.000
cash_usage_ratio 0.297
expense_income_ratio 0.291
income_volatility_cv 0.236
late_payments_count_6m 0.216
informal_credit_usage 0.199
household_dependents 0.088
digital_wallet_usage 0.035
age -0.004
debt_service_ratio -0.117
liquid_savings_ars -0.147
monthly_expenses_ars -0.159
years_employment -0.210
monthly_income_ars -0.240
has_credit_card -0.241
formal_employment -0.345
Name: default_12m, dtype: float64

```

```

=====
CHECKPOINTS DE CONTROL (AUDITORÍA DE CAUSALIDAD)
=====
Efectivo vs Empleo Formal (Objetivo ~ -0.50): -0.771
Gasto vs Ingreso Mensual (Objetivo ~ -0.65): -0.683
Edad vs Antigüedad Laboral (Coherencia biológica): 0.620
Ahorro vs Ingreso Mensual (Física del canuto): 0.636
Ahorro vs Ratio de Gasto (Ley de Engel): -0.418
=====

```

```

=====
TODAS LAS PAREJAS DE VARIABLES (ORDENADAS POR INTENSIDAD DE CORRELACIÓN)
=====

```

	Variable_A	Variable_B	Correlacion_r
0	monthly_income_ars	monthly_expenses_ars	0.839
1	formal_employment	cash_usage_ratio	-0.771
2	monthly_income_ars	expense_income_ratio	-0.683
3	monthly_income_ars	liquid_savings_ars	0.636
4	formal_employment	income_volatility_cv	-0.624
5	age	years_employment	0.620
6	years_employment	income_volatility_cv	-0.587
7	formal_employment	years_employment	0.584
8	formal_employment	debt_service_ratio	0.572
9	formal_employment	has_credit_card	0.550
10	income_volatility_cv	debt_service_ratio	-0.543
11	monthly_expenses_ars	liquid_savings_ars	0.476
12	formal_employment	expense_income_ratio	-0.457
13	formal_employment	monthly_income_ars	0.454
14	has_credit_card	cash_usage_ratio	-0.424
15	monthly_income_ars	income_volatility_cv	-0.420
16	expense_income_ratio	liquid_savings_ars	-0.418
17	debt_service_ratio	cash_usage_ratio	-0.405
18	years_employment	debt_service_ratio	0.400
19	formal_employment	informal_credit_usage	-0.391
20	income_volatility_cv	expense_income_ratio	0.387
21	expense_income_ratio	monthly_expenses_ars	-0.364
22	age	income_volatility_cv	-0.362
23	income_volatility_cv	cash_usage_ratio	0.356
24	formal_employment	default_12m	-0.345
25	income_volatility_cv	has_credit_card	-0.345
26	digital_wallet_usage	cash_usage_ratio	-0.339
27	years_employment	monthly_income_ars	0.333
28	formal_employment	monthly_expenses_ars	0.332
29	expense_income_ratio	cash_usage_ratio	0.331
30	monthly_income_ars	cash_usage_ratio	-0.324
31	years_employment	has_credit_card	0.321
32	expense_income_ratio	has_credit_card	-0.317
33	informal_credit_usage	expense_income_ratio	0.316
34	years_employment	expense_income_ratio	-0.315
35	income_volatility_cv	monthly_expenses_ars	-0.313
36	late_payments_count_6m	has_credit_card	-0.307
37	informal_credit_usage	cash_usage_ratio	0.302
38	cash_usage_ratio	default_12m	0.297
39	age	cash_usage_ratio	0.296
40	monthly_income_ars	has_credit_card	0.292
41	expense_income_ratio	default_12m	0.291
42	formal_employment	household_dependents	-0.279
43	years_employment	cash_usage_ratio	-0.273
44	years_employment	digital_wallet_usage	-0.270
45	years_employment	household_dependents	-0.264
46	informal_credit_usage	has_credit_card	-0.257
47	income_volatility_cv	digital_wallet_usage	0.251
48	expense_income_ratio	late_payments_count_6m	0.251
49	monthly_expenses_ars	cash_usage_ratio	-0.250
50	income_volatility_cv	late_payments_count_6m	0.249
51	income_volatility_cv	liquid_savings_ars	-0.245
52	income_volatility_cv	household_dependents	0.244
53	informal_credit_usage	income_volatility_cv	0.241
54	has_credit_card	default_12m	-0.241
55	monthly_income_ars	default_12m	-0.240
56	formal_employment	liquid_savings_ars	0.238
57	income_volatility_cv	default_12m	0.236
58	informal_credit_usage	late_payments_count_6m	0.235
59	late_payments_count_6m	liquid_savings_ars	-0.233
60	formal_employment	digital_wallet_usage	-0.230
61	years_employment	informal_credit_usage	-0.227
62	years_employment	monthly_expenses_ars	0.219
63	late_payments_count_6m	default_12m	0.216
64	years_employment	default_12m	-0.210

65	age	household_dependents	-0.205
66	age	digital_wallet_usage	-0.204
67	debt_service_ratio	has_credit_card	0.203
68	liquid_savings_ars	has_credit_card	0.202
69	informal_credit_usage	default_12m	0.199
70	formal_employment	late_payments_count_6m	-0.198
71	monthly_expenses_ars	has_credit_card	0.198
72	monthly_expenses_ars	household_dependents	0.191
73	monthly_expenses_ars	debt_service_ratio	0.188
74	debt_service_ratio	household_dependents	-0.187
75	years_employment	liquid_savings_ars	0.186
76	monthly_income_ars	late_payments_count_6m	-0.183
77	informal_credit_usage	monthly_income_ars	-0.175
78	monthly_income_ars	debt_service_ratio	0.172
79	liquid_savings_ars	cash_usage_ratio	-0.168
80	debt_service_ratio	digital_wallet_usage	-0.161
81	monthly_expenses_ars	default_12m	-0.159
82	years_employment	late_payments_count_6m	-0.158
83	has_credit_card	household_dependents	-0.149
84	cash_usage_ratio	household_dependents	0.148
85	liquid_savings_ars	default_12m	-0.147
86	late_payments_count_6m	cash_usage_ratio	0.129
87	digital_wallet_usage	has_credit_card	-0.128
88	monthly_income_ars	digital_wallet_usage	-0.126
89	expense_income_ratio	digital_wallet_usage	0.123
90	expense_income_ratio	household_dependents	0.118
91	debt_service_ratio	default_12m	-0.117
92	monthly_income_ars	household_dependents	-0.114
93	informal_credit_usage	liquid_savings_ars	-0.114
94	informal_credit_usage	household_dependents	0.111
95	age	debt_service_ratio	0.108
96	digital_wallet_usage	household_dependents	0.106
97	monthly_expenses_ars	late_payments_count_6m	-0.104
98	informal_credit_usage	digital_wallet_usage	0.091
99	household_dependents	default_12m	0.088
100	age	monthly_income_ars	0.088
101	monthly_expenses_ars	digital_wallet_usage	-0.082
102	age	late_payments_count_6m	-0.070
103	debt_service_ratio	liquid_savings_ars	0.069
104	age	expense_income_ratio	-0.068
105	liquid_savings_ars	digital_wallet_usage	-0.067
106	late_payments_count_6m	digital_wallet_usage	0.067
107	late_payments_count_6m	household_dependents	0.066
108	liquid_savings_ars	household_dependents	-0.059
109	age	liquid_savings_ars	0.056
110	informal_credit_usage	monthly_expenses_ars	-0.053
111	informal_credit_usage	debt_service_ratio	0.038
112	digital_wallet_usage	default_12m	0.035
113	debt_service_ratio	late_payments_count_6m	-0.031
114	expense_income_ratio	debt_service_ratio	-0.030
115	age	monthly_expenses_ars	0.022
116	age	informal_credit_usage	0.006
117	age	default_12m	-0.004
118	age	formal_employment	-0.003
119	age	has_credit_card	-0.001

```

In [6]: import pandas as pd
import numpy as np
import matplotlib.pyplot as plt
import seaborn as sns
from sklearn.linear_model import LogisticRegression
from sklearn.preprocessing import StandardScaler, OneHotEncoder
from sklearn.compose import ColumnTransformer
from sklearn.pipeline import Pipeline
from sklearn.metrics import (
    roc_auc_score,
    roc_curve,
    confusion_matrix,
    ConfusionMatrixDisplay
)
import lightgbm as lgb
import warnings
warnings.filterwarnings('ignore')

print("=====")
print("      ENTRENAMIENTO DE MODELOS: LOGÍSTICA TRADICIONAL VS LIGHTGBM      ")
print("=====")

# 1. Definición de variables
categorical_cols = ['province', 'education_level', 'age_category', 'device_quality_index']
numeric_cols = [col for col in X_train.columns if col not in categorical_cols]

# 2. Preprocesamiento para La Regresión Logística (Estricto: OneHot + Scaling)
preprocessor_lr = ColumnTransformer(
    transformers=[
        ('num', StandardScaler(), numeric_cols),
        ('cat', OneHotEncoder(handle_unknown='ignore', drop='first'), categorical_cols)
    ])

pipeline_lr = Pipeline(steps=[
    ('preprocessor', preprocessor_lr),
    ('classifier', LogisticRegression(max_iter=1000, random_state=42, class_weight='balanced'))
])

# 3. Preprocesamiento para LightGBM (Nativo: Maneja categóricas automáticamente)
# Nota: LightGBM requiere que las categóricas sean tipo 'category' en Pandas, lo cual ya hicimos.
clf_lgb = lgb.LGBMClassifier(
    objective='binary',
    max_depth=5,

```

```

learning_rate=0.05,
n_estimators=300,
subsample=0.8,
colsample_bytree=0.8,
scale_pos_weight=(len(y_train) - sum(y_train)) / sum(y_train), # Balanceo de clases
random_state=42,
verbose=-1 # Suprime Los mensajes de advertencia internos de LightGBM
)

# 4. Entrenamiento de Modelos
print("[*] Entrenando Regresión Logística (Baseline)...")
pipeline_lr.fit(X_train, y_train)

print("[*] Entrenando LightGBM (Challenger)...")
clf_lgb.fit(X_train, y_train)

# 5. Predicción de Probabilidades
y_pred_proba_lr = pipeline_lr.predict_proba(X_test)[: , 1]
y_pred_proba_lgb = clf_lgb.predict_proba(X_test)[: , 1]

# 6. Cálculo de Métricas (AUC-ROC, Gini, KS)
def calculate_metrics(y_true, y_proba, model_name):
    # AUC-ROC
    auc = roc_auc_score(y_true, y_proba)
    # Gini
    gini = 2 * auc - 1
    # KS Statistic
    fpr, tpr, thresholds = roc_curve(y_true, y_proba)
    ks = max(tpr - fpr)

    print(f"\n--- {model_name} ---")
    print(f"AUC-ROC : {auc:.4f}")
    print(f"Gini      : {gini:.4f}")
    print(f"KS Stat : {ks:.4f}")
    return fpr, tpr, auc

fpr_lr, tpr_lr, auc_lr = calculate_metrics(y_test, y_pred_proba_lr, "Regresión Logística")
fpr_lgb, tpr_lgb, auc_lgb = calculate_metrics(y_test, y_pred_proba_lgb, "LightGBM")

# =====
# MATRICES DE CONFUSIÓN
# =====

print("\n=====")
print("          MATRICES DE CONFUSIÓN")
print("=====")

# Predicciones binarias
y_pred_lr = (y_pred_proba_lr >= 0.50).astype(int)
y_pred_lgb = (y_pred_proba_lgb >= 0.50).astype(int)

# Tema gráfico homogéneo
plt.style.use("default")
sns.set_theme(style="whitegrid")

# -----
# Regresión Logística
# -----
cm_lr = confusion_matrix(y_test, y_pred_lr)

fig, ax = plt.subplots(figsize=(6,5))

disp = ConfusionMatrixDisplay(
    confusion_matrix=cm_lr,
    display_labels=["No Default", "Default"]
)

disp.plot(
    cmap="Blues",
    colorbar=False,
    ax=ax
)

plt.title(
    "Matriz de Confusión - Regresión Logística",
    fontsize=14,
    fontweight="bold"
)

plt.tight_layout()
plt.show()

from sklearn.metrics import classification_report
import pandas as pd

print("\n===== CLASIFICATION REPORT - REGRESIÓN LOGÍSTICA =====\n")

report_lr = classification_report(
    y_test,
    y_pred_lr,
    target_names=["No Default (0)", "Default (1)"],
    output_dict=True
)

tabla_lr = pd.DataFrame(report_lr).transpose()

tabla_lr = tabla_lr.loc[
    ["No Default (0)", "Default (1)"],
    ["precision", "recall", "f1-score"]
]

```

```

tabla_lr.columns = ["Precision", "Recall", "F1-Score"]

tabla_lr["Precision"] = tabla_lr["Precision"].map(lambda x: f"{x:.2f}")
tabla_lr["Recall"] = tabla_lr["Recall"].map(lambda x: f"{x:.2f}")
tabla_lr["F1-Score"] = tabla_lr["F1-Score"].map(lambda x: f"{x:.2f}")

tabla_lr.loc["AUC-ROC"] = ["", "", f"{auc_lr:.4f}"]

display(tabla_lr)

# -----
# LightGBM
# -----
cm_lgb = confusion_matrix(y_test, y_pred_lgb)

fig, ax = plt.subplots(figsize=(6,5))

disp = ConfusionMatrixDisplay(
    confusion_matrix=cm_lgb,
    display_labels=["No Default", "Default"]
)

disp.plot(
    cmap="Reds",
    colorbar=False,
    ax=ax
)

plt.title(
    "Matriz de Confusión - LightGBM",
    fontsize=14,
    fontweight="bold"
)

plt.tight_layout()
plt.show()

print("\n===== CLASIFICATION REPORT - LIGHTGBM =====\n")

report_lgb = classification_report(
    y_test,
    y_pred_lgb,
    target_names=["No Default (0)", "Default (1)"],
    output_dict=True
)

tabla_lgb = pd.DataFrame(report_lgb).transpose()

tabla_lgb = tabla_lgb.loc[
    ["No Default (0)", "Default (1)"],
    ["precision", "recall", "f1-score"]
]

tabla_lgb.columns = ["Precision", "Recall", "F1-Score"]

tabla_lgb["Precision"] = tabla_lgb["Precision"].map(lambda x: f"{x:.2f}")
tabla_lgb["Recall"] = tabla_lgb["Recall"].map(lambda x: f"{x:.2f}")
tabla_lgb["F1-Score"] = tabla_lgb["F1-Score"].map(lambda x: f"{x:.2f}")

tabla_lgb.loc["AUC-ROC"] = ["", "", f"{auc_lgb:.4f}"]

display(tabla_lgb)

auc_difference = auc_lgb - auc_lr
print(f"Diferencia de AUC-ROC (LightGBM - Regresión Logística): {auc_difference:.4f}")

# 7. Graficar Curvas ROC
plt.figure(figsize=(10, 8))
sns.set_theme(style="whitegrid")

plt.plot(fpr_lr, tpr_lr, color='#3498db', lw=2, label=f'Regresión Logística (AUC = {auc_lr:.3f})')
plt.plot(fpr_lgb, tpr_lgb, color='#e74c3c', lw=2, label=f'LightGBM (AUC = {auc_lgb:.3f})')
plt.plot([0, 1], [0, 1], color='gray', lw=1.5, linestyle='--')

plt.xlim([0.0, 1.0])
plt.ylim([0.0, 1.05])
plt.xlabel('Tasa de Falsos Positivos (1 - Especificidad)', fontsize=12)
plt.ylabel('Tasa de Verdaderos Positivos (Sensibilidad)', fontsize=12)
plt.title('Comparación de Curvas ROC: Modelos Predictivos de Riesgo', fontsize=16, fontweight='bold', pad=20)
plt.legend(loc="lower right", frameon=True, fontsize=12)

plt.tight_layout()
plt.show()

print("\n=====")
print("          EJECUCIÓN COMPLETADA - MÉTRICAS LISTAS PARA EL TIF          ")
print("=====")

```

StatementMeta(, 36b1965b-fcd6-43f0-91c0-25e4c22ff233, 8, Finished, Available, Finished, False)

```

=====
ENTRENAMIENTO DE MODELOS: LOGÍSTICA TRADICIONAL VS LIGHTGBM
=====
[*] Entrenando Regresión Logística (Baseline)...
[*] Entrenando LightGBM (Challenger)...

```

--- Regresión Logística ---

AUC-ROC : 0.8092

Gini : 0.6185

KS Stat : 0.4979

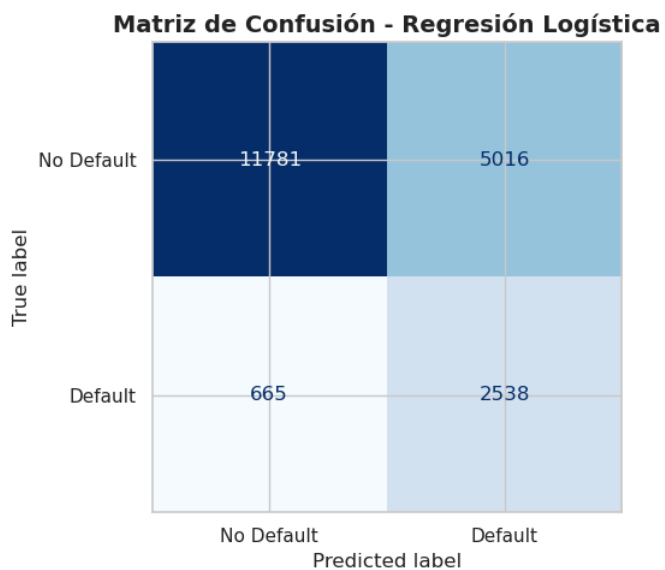
--- LightGBM ---

AUC-ROC : 0.8263

Gini : 0.6526

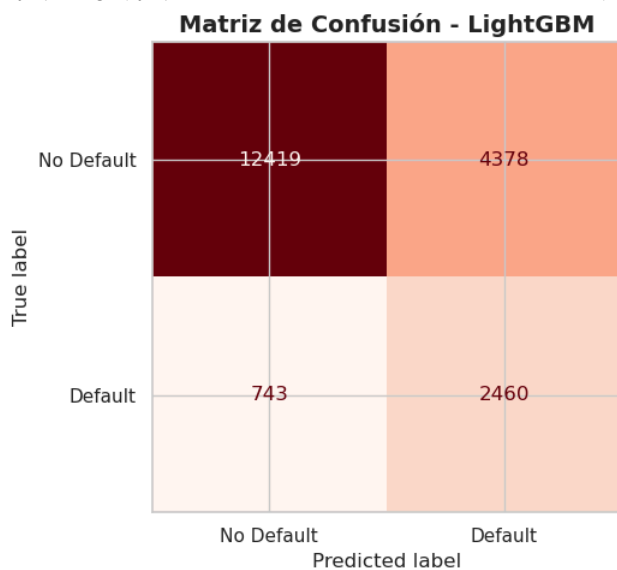
KS Stat : 0.5100

MATRICES DE CONFUSIÓN



===== CLASIFICATION REPORT - REGRESIÓN LOGÍSTICA =====

SynapseWidget(Synapse.DataFrame, c4402670-7ee4-4828-8e6f-b35e4e526605)

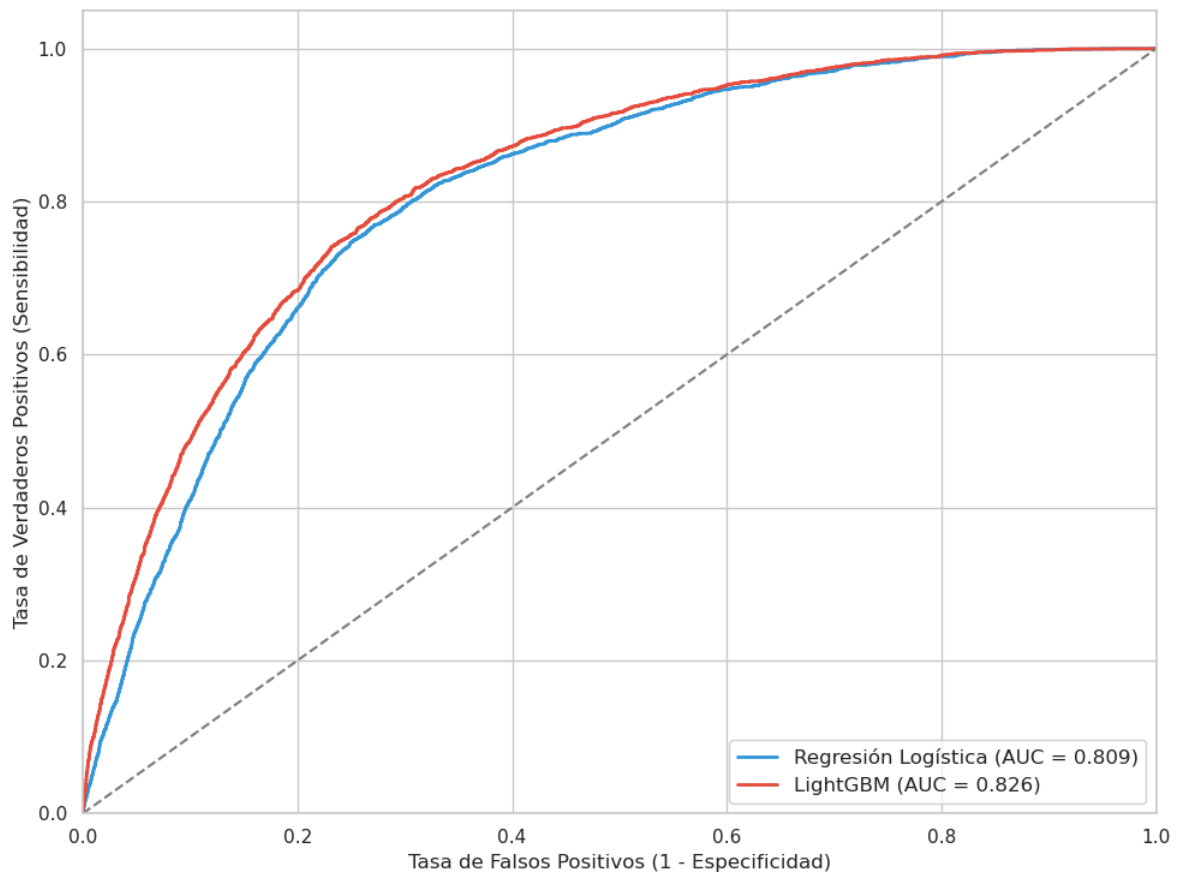


===== CLASIFICATION REPORT - LIGHTGBM =====

SynapseWidget(Synapse.DataFrame, 9257d180-b40d-4b37-b1d4-545e082ff137)

Diferencia de AUC-ROC (LightGBM - Regresión Logística): 0.0170

Comparación de Curvas ROC: Modelos Predictivos de Riesgo



=====

EJECUCIÓN COMPLETADA - MÉTRICAS LISTAS PARA EL TIF

=====

```
In [7]: # =====
# EXPLICABILIDAD ALGORÍTMICA: VALORES SHAP (SHAPLEY ADDITIVE EXPLANATIONS)
# =====
import shap
import matplotlib.pyplot as plt

print("[*] Instanciando explainer SHAP basado en topología de árbol...")

# 1. Configuración del Explainer sobre el modelo ganador (LightGBM o XGBoost)
# Se utiliza el TreeExplainer, optimizado matemáticamente para ensambles de árboles
explainer = shap.TreeExplainer(clf_lgb) # Reemplazar 'clf_lgb' por 'clf_xgb' si se optó por XGBoost

# 2. Cálculo de la matriz de valores SHAP sobre el conjunto de prueba
# Nota: Si el X_test es muy masivo (>20k filas), el cálculo puede demorar.
# En modelado empírico es válido usar un muestreo representativo (ej. 5000 observaciones).
X_test_sample = X_test.sample(n=5000, random_state=42)
shap_values = explainer.shap_values(X_test_sample)

# LightGBM devuelve una lista de arrays para clasificación binaria. Se extrae el array de la clase positiva (Default = 1).
if isinstance(shap_values, list):
    shap_values_default = shap_values[1]
else:
    shap_values_default = shap_values

# =====
# GENERACIÓN DE GRÁFICOS DE IMPACTO E INTERPRETACIÓN
# =====

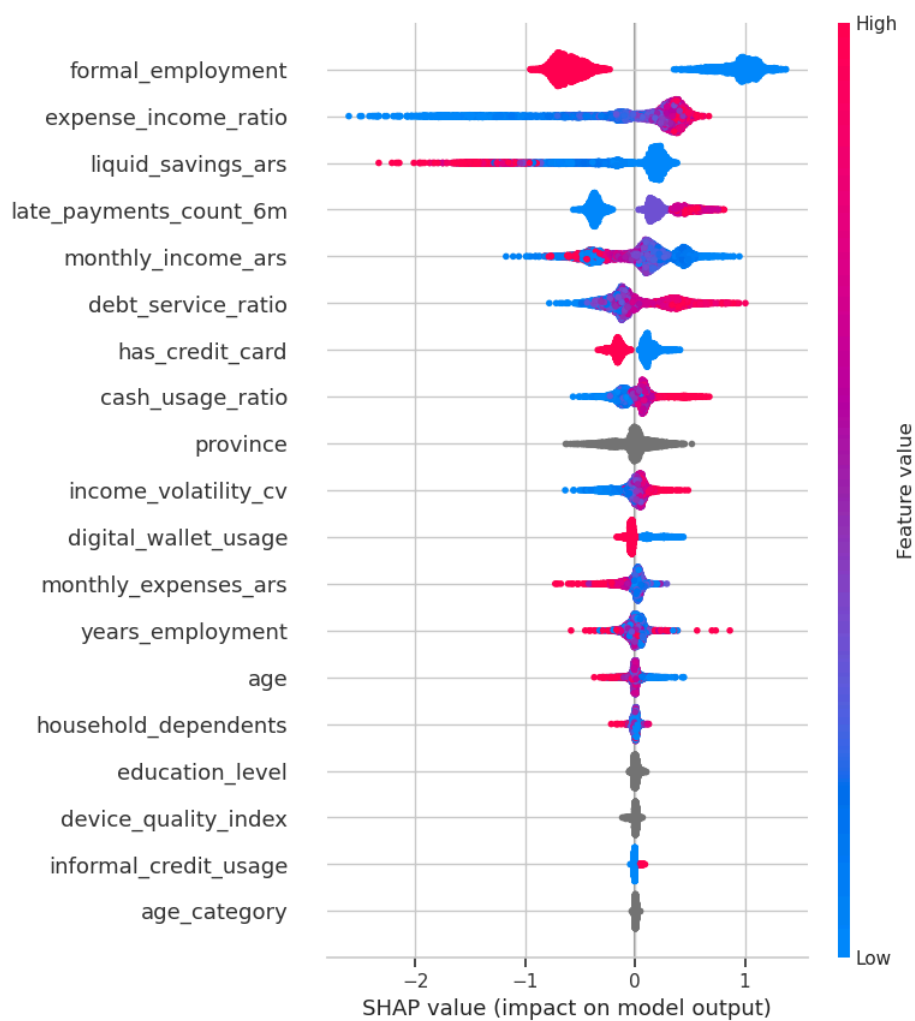
# A. Summary Plot (Importancia global y dirección del efecto)
plt.figure(figsize=(8, 8))
plt.title("SHAP Summary Plot: Descomposición de la Varianza del Modelo No Paramétrico", fontsize=14, pad=20)
shap.summary_plot(shap_values_default, X_test_sample, show=False)
plt.tight_layout()
plt.show()

# B. Análisis de Dependencia: Interacción DSR vs Ingreso (Demostración de Asfixia Financiera)
plt.figure(figsize=(10, 6))
plt.title("SHAP Dependence Plot: Debt Service Ratio (DSR) e Ingresos", fontsize=14, pad=20)
shap.dependence_plot(
    "debt_service_ratio",
    shap_values_default,
    X_test_sample,
    interaction_index="monthly_income_ars",
    show=False
)
plt.tight_layout()
plt.show()
```

StatementMeta(, 36b1965b-fcd6-43f0-91c0-25e4c22ff233, 9, Finished, Available, Finished, False)

[*] Instanciando explainer SHAP basado en topología de árbol...

SHAP Summary Plot: Descomposición de la Varianza del Modelo No Paramétrico



SHAP Dependence Plot: Debt Service Ratio (DSR) e Ingresos

