



Modelo explicable de aprendizaje supervisado y no supervisado para caracterizar la cobertura de programas sociales en territorios vulnerables del Perú

Explainable model of supervised and unsupervised learning to characterize the coverage of social programs in vulnerable territories of Peru

Marin-Rodriguez, William J.^{1*}

Vellón-Flores, Viviana I.¹

Ausejo-Sánchez, José L.¹

Solano-Armas, Timoteo¹

Maguiña-Poma, Yolanda E.¹

Samanamud-Malca, Sixto¹

¹Universidad Nacional José Faustino Sánchez Carrión, Huacho, Perú

Recibido: 03 Ene. 2026 | **Aceptado:** 26 Jun. 2026 | **Publicado:** 20 Jul. 2026

Autor de correspondencia*: wmarin@unjfsc.edu.pe

Cómo citar este artículo: Marin-Rodriguez, W. J., Vellón-Flores, V. I., Ausejo-Sánchez, J. L., Solano-Armas, T., Maguiña-Poma, Y. E. & Samanamud-Malca, S. (2026). Modelo explicable de aprendizaje supervisado y no supervisado para caracterizar la cobertura de programas sociales en territorios vulnerables del Perú. *Revista Científica de Sistemas e Informática*, 6(2), e1266. <https://doi.org/10.51252/rcsi.v6i2.1266>

RESUMEN

Este estudio desarrolló y evaluó un enfoque explicable de aprendizaje supervisado y no supervisado para caracterizar la cobertura de programas sociales e identificar territorios históricamente vulnerables del Perú. Se integraron 35 952 observaciones distrito-mes, correspondientes a 1 893 distritos y 19 periodos, con información del MIDIS y mapas de pobreza. A partir de 13 indicadores se compararon cuatro algoritmos de agrupamiento y tres clasificadores mediante validación cruzada y prueba independiente. K-means con dos conglomerados fue seleccionado como solución principal, con una silueta de 0,540, e identificó un perfil de cobertura y vulnerabilidad histórica. En la clasificación, Elastic Net basado únicamente en indicadores de cobertura alcanzó un F1-score de 0,655 y, en la prueba independiente, una sensibilidad de 0,829, un F1-score de 0,647 y un ROC-AUC de 0,871, sin diferencias significativas frente a los demás modelos. La explicabilidad destacó la cobertura de Juntos, los usuarios de FONCODES y el número de programas presentes. Los resultados muestran que estos registros pueden apoyar la priorización territorial, complementados con información actualizada y revisión institucional.

Palabras clave: agrupamiento; focalización social; inteligencia artificial; pobreza histórica; regresión logística

ABSTRACT

This study developed and evaluated an explainable supervised and unsupervised learning approach to characterize social program coverage and identify historically vulnerable territories in Peru. A total of 35,952 district-month observations from 1,893 districts and 19 periods were integrated using MIDIS administrative data and poverty maps. Based on 13 coverage indicators, four clustering algorithms and three classifiers were compared through cross-validation and an independent test set. K-means with two clusters was selected as the main solution, achieving a silhouette score of 0.540 and identifying a profile associated with greater coverage and historical vulnerability. For classification, Elastic Net using only coverage indicators achieved an F1-score of 0.655 and, in the independent test set, a sensitivity of 0.829, an F1-score of 0.647, and a ROC-AUC of 0.871, with no significant differences compared with the other models. Explainability analysis highlighted Juntos coverage, estimated FONCODES users, and the number of programs present. The findings show that administrative records can support territorial prioritization when complemented with updated information and institutional review.

Keywords: clustering; social targeting; artificial intelligence; historical poverty; logistic regression



1. INTRODUCCIÓN

La pobreza continúa siendo uno de los principales desafíos para el desarrollo territorial del Perú, pese a la reducción observada durante los últimos años. Según el Instituto Nacional de Estadística e Informática (INEI, 2026), en 2025 la pobreza monetaria afectó al 25,7 % de la población nacional; sin embargo, su incidencia alcanzó el 35,5 % en el área rural y superó el 40 % en departamentos como Cajamarca y Loreto. Estas diferencias evidencian que la mejora de los indicadores agregados no ha eliminado la concentración territorial de las privaciones y que los promedios nacionales pueden ocultar realidades distritales marcadamente heterogéneas. En este escenario, identificar los territorios con mayores condiciones de vulnerabilidad resulta prioritario para orientar la intervención pública y distribuir los recursos de protección social con mayor pertinencia territorial. Además de su dimensión monetaria, la pobreza comprende privaciones relacionadas con el ingreso, la alimentación, los servicios básicos y las condiciones de la vivienda, por lo que su análisis requiere una perspectiva multidimensional (García et al., 2024).

Los programas sociales constituyen instrumentos clave para proteger a la población en situación de pobreza y ampliar su acceso a servicios básicos y transferencias. Sin embargo, sus resultados dependen de la focalización, las características de la población atendida y las condiciones estructurales de cada territorio (Borga & D'Ambrosio, 2021). En el Perú, factores como la educación, la infraestructura, el empleo y las restricciones institucionales generan efectos diferenciados entre zonas urbanas y rurales (Hinojosa Pérez et al., 2024; Maco, 2025; Paucara-Charca et al., 2024). Por ello, resulta necesario analizar de manera conjunta la presencia de los programas sociales y su configuración territorial.

La disponibilidad de datos administrativos abiertos ofrece una oportunidad para avanzar hacia este tipo de análisis. El Ministerio de Desarrollo e Inclusión Social publica información distrital sobre la cobertura de Cuna Más, Juntos, Foncodes, Pensión 65, Wasi Mikuna, Contigo y PAIS, expresada mediante el número de usuarios, hogares, beneficiarios, atenciones, instituciones, establecimientos y proyectos registrados (MIDIS, 2026). Sin embargo, la diversidad de unidades de medida, las diferencias poblacionales entre distritos y la coexistencia de varios programas en un mismo territorio dificultan la interpretación directa de estos registros. En consecuencia, se requieren procedimientos que normalicen los indicadores y examinen simultáneamente sus relaciones, de modo que la información administrativa pueda transformarse en perfiles territoriales comprensibles y útiles para la toma de decisiones.

En este contexto, el aprendizaje automático ha generado nuevas posibilidades para analizar grandes volúmenes de información social y reconocer patrones que difícilmente podrían identificarse mediante procedimientos descriptivos convencionales. En Togo, la combinación de datos de telefonía móvil y modelos de aprendizaje automático redujo entre 4 % y 21 % los errores de exclusión respecto de las alternativas de focalización geográfica consideradas para distribuir asistencia humanitaria (Aiken et al., 2022). En Nigeria, los mapas de pobreza de alta resolución contruidos mediante información geoespacial y aprendizaje automático disminuyeron los errores de inclusión y exclusión en la asignación territorial de transferencias sociales (Smythe & Blumenstock, 2022). Asimismo, la integración de datos administrativos, imágenes, conectividad y otras fuentes espaciales permitió mejorar la identificación de hogares vulnerables en contextos urbanos con disponibilidad limitada de información socioeconómica (Jung et al., 2026). Estos

avances evidencian el potencial de los modelos computacionales para complementar los instrumentos tradicionales de diagnóstico territorial.

El aprendizaje no supervisado y el supervisado ofrecen capacidades complementarias para analizar la cobertura social. Mientras los algoritmos de agrupamiento permiten identificar territorios con patrones semejantes sin categorías previas, los modelos de clasificación evalúan su capacidad para discriminar condiciones de vulnerabilidad. Sin embargo, en el ámbito social no basta con alcanzar un buen rendimiento predictivo; también es necesario explicar qué variables influyen en las estimaciones y cómo se relacionan con los resultados. La literatura advierte que la transparencia y la interpretabilidad aún son limitadas en varios estudios computacionales sobre pobreza (Hall et al., 2022). Por ello, la inteligencia artificial aplicada a políticas públicas requiere modelos comprensibles y auditables que respalden, sin reemplazar, el criterio de los responsables de la toma de decisiones (Lisboa et al., 2023; Papadakis et al., 2024).

A pesar de estos avances, aún es limitada la caracterización integrada de la cobertura de múltiples programas sociales a escala distrital mediante técnicas de aprendizaje supervisado, no supervisado y explicabilidad. En el Perú, los estudios se han concentrado principalmente en programas específicos, restricciones de focalización o clasificación de beneficiarios, sin analizar conjuntamente los patrones territoriales de los servicios del MIDIS. Por ello, el objetivo del estudio fue desarrollar y evaluar un enfoque analítico explicable para caracterizar la cobertura de los programas sociales e identificar territorios históricamente vulnerables. Para este propósito, se integraron registros administrativos distritales, se compararon algoritmos de agrupamiento y clasificación y se aplicaron métodos de explicabilidad. El enfoque propuesto aporta una metodología reproducible para transformar datos públicos heterogéneos en evidencia útil que apoye el diagnóstico y la planificación territorial de la protección social.

2. MATERIALES Y MÉTODOS

2.1 Diseño y alcance del estudio

La investigación fue de tipo aplicada, con enfoque cuantitativo, diseño no experimental, retrospectivo y ecológico, y alcance exploratorio-analítico. La unidad de análisis fue el distrito peruano, identificado mediante el código de Ubicación Geográfica de seis dígitos (UBIGEO).

La base inicial presentó una estructura longitudinal de tipo distrito-mes. Para el análisis mediante aprendizaje automático, los registros mensuales se agregaron a nivel distrital mediante el promedio temporal de las variables de cobertura, obteniéndose una sola observación por distrito.

2.2 Procedimiento metodológico

El estudio integró la preparación de las fuentes oficiales, la construcción de indicadores territoriales y el análisis de los patrones de cobertura social. Posteriormente, se aplicaron técnicas de aprendizaje no supervisado y supervisado para caracterizar los territorios y clasificar su vulnerabilidad histórica. El análisis se complementó con procedimientos de validación estadística, explicabilidad y evaluación de los errores de clasificación. Como se observa en la Figura 1, estas actividades se articularon en un flujo secuencial que permitió mantener la trazabilidad de los datos desde su obtención hasta la interpretación de los modelos.

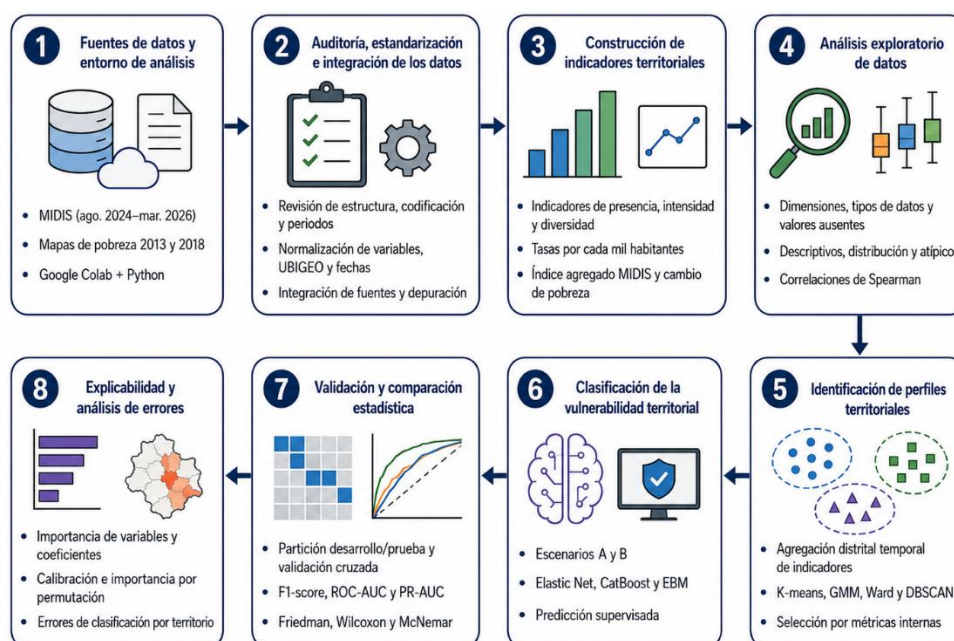


Figura 1. Flujo metodológico del estudio

Fuentes de datos y entorno de análisis

Se utilizaron dos fuentes oficiales de acceso público. La primera estuvo conformada por 19 archivos mensuales en formato CSV publicados por el Ministerio de Desarrollo e Inclusión Social, con registros distritales de cobertura comprendidos entre agosto de 2024 y marzo de 2026 (MIDIS, 2026). El corte de noviembre de 2024 fue excluido debido a una inconsistencia identificada en el campo de fecha. La información correspondió a los programas Cuna Más, Juntos, FONCODES, Pensión 65, Wasi Mikuna, Contigo y PAIS, e incluyó registros de usuarios, hogares, beneficiarios, atenciones, instituciones, tambos y proyectos, según la estructura operativa de cada programa.

La segunda fuente estuvo constituida por los Mapas de Pobreza Provincial y Distrital de 2013 y 2018, disponibles en formato Excel (INEI, 2015, 2020). De estos archivos se extrajeron el UBIGEO, departamento, provincia, distrito, población estimada, límites inferior y superior de pobreza, grupo robusto y posición en el ranking distrital.

El procesamiento se realizó en Google Colaboratory mediante Python. Para la manipulación e integración de los datos se emplearon pandas y NumPy; para la visualización, Matplotlib; para el aprendizaje automático, scikit-learn, CatBoost e InterpretML; y para el análisis estadístico, SciPy y Statsmodels.

Auditoría, estandarización e integración de los datos

Los archivos fueron auditados según su estructura, codificación, delimitador, dimensiones, columnas disponibles, valores ausentes y periodos registrados. Para garantizar su lectura, se evaluaron distintas codificaciones y separadores, seleccionándose automáticamente la combinación que permitió recuperar correctamente la estructura tabular (Van den Broeck et al., 2005).

Los nombres de las variables se normalizaron mediante la conversión a minúsculas, eliminación de tildes y caracteres especiales y sustitución de espacios por guiones bajos. Además, se implementó un mapeo de equivalencias para reconocer cambios de denominación entre periodos,

como la transición de Qali Warma a Wasi Mikuna y las variantes utilizadas para el Programa Nacional PAIS. El UBIGEO fue depurado, conservando únicamente caracteres numéricos y completándolo hasta alcanzar seis dígitos. Las fechas de corte se convirtieron desde el formato AAAAMMDD y las variables de cobertura fueron transformadas a formato numérico (Harron et al., 2017).

Después de consolidar los archivos mensuales, se verificó la unicidad de la combinación periodo-UBIGEO. Como regla de depuración, ante la presencia de registros repetidos se conservó la última observación disponible para cada distrito y periodo. En los mapas de pobreza se utilizaron las hojas “Cdr01” para 2013 y “Anexo1” para 2018. Se eliminaron encabezados, códigos no válidos, registros sin identificación territorial y observaciones que no cumplieron criterios lógicos de población y porcentaje de pobreza.

La pobreza promedio distrital se calculó como el punto medio del intervalo reportado:

$$Pobreza\ promedio = \frac{Límite\ inferior\ de\ pobreza + Límite\ superior\ de\ pobreza}{2}$$

Las fuentes fueron integradas mediante el UBIGEO, manteniendo como estructura principal el panel mensual del MIDIS. Los valores ausentes de pobreza no fueron imputados. En las variables de cobertura, los valores faltantes se reemplazaron por cero como criterio operativo y se conservó un indicador del número de ausencias originales para mantener su trazabilidad. No se aplicó interpolación temporal.

Construcción de indicadores territoriales

A partir de las variables originales se construyeron indicadores de presencia, intensidad y diversidad de la cobertura social. Para cada programa se generó una variable binaria con valor uno cuando al menos una de sus prestaciones presentó un registro mayor que cero y valor cero en caso contrario. Asimismo, se calcularon el número de programas presentes y el número de variables de cobertura con valores positivos en cada distrito y periodo (Amarante et al., 2026).

Con el propósito de controlar las diferencias poblacionales entre distritos, diez variables de cobertura fueron expresadas por cada mil habitantes:

$$Tasa\ de\ cobertura_j = \frac{Cobertura\ del\ programa_j}{Población\ del\ distrito} \times 1000$$

La normalización se aplicó al cuidado diurno y acompañamiento a familias de Cuna Más; hogares afiliados y abonados de Juntos; usuarios estimados de FONCODES; usuarios de Pensión 65; niñas y niños atendidos por Wasi Mikuna; usuarios de Contigo; y atenciones y beneficiarios del Programa Nacional PAIS.

Las diez tasas fueron sumadas para construir un índice agregado de intensidad de cobertura MIDIS por cada mil habitantes. Este indicador representó la magnitud conjunta de los registros administrativos y no el número de beneficiarios únicos, debido a que una persona podía participar en más de un programa o modalidad de atención (Ferreira et al., 2022). Como variable contextual se calculó la variación de la pobreza distrital:

$$Cambio\ de\ pobreza = Pobreza\ promedio\ en\ 2018 - Pobreza\ promedio\ en\ 2013$$

Para la clasificación supervisada se definió operacionalmente como territorio de alta vulnerabilidad histórica aquel distrito cuya pobreza promedio en 2018 fue igual o superior al

percentil 75 de la distribución distrital. Los demás distritos fueron asignados a la categoría de menor vulnerabilidad relativa. Las variables de pobreza, ranking y grupos robustos se conservaron únicamente para la definición e interpretación del resultado y no fueron incorporadas como predictores.

Análisis exploratorio de datos

El análisis exploratorio incluyó la revisión de las dimensiones, tipos de datos, valores ausentes, duplicados, cobertura temporal y disponibilidad de información por distrito (Behrens, 1997). Para las variables numéricas se calcularon medidas de tendencia central, dispersión y posición, incluyendo media, mediana, desviación estándar, valores mínimo y máximo, y los percentiles 1, 5, 25, 75, 95 y 99.

Los valores potencialmente atípicos fueron identificados mediante el rango intercuartílico, sin ser eliminados en esta etapa. También se examinaron la distribución del número de programas presentes, el índice agregado de intensidad de cobertura, la evolución mensual de las principales prestaciones y las diferencias territoriales entre departamentos y distritos.

Las asociaciones entre los indicadores de cobertura se evaluaron mediante el coeficiente de correlación de Spearman, debido a la asimetría de sus distribuciones. Asimismo, se analizó su relación bivariada con la pobreza promedio distrital de 2018. Estas asociaciones se interpretaron únicamente con fines descriptivos y no como relaciones causales (Schober et al., 2018).

Identificación de perfiles territoriales

Para el análisis no supervisado, los registros mensuales fueron agregados a nivel distrital mediante el promedio temporal de 13 indicadores: diez tasas de cobertura por cada mil habitantes, el número de programas presentes, el número de variables de cobertura con valores positivos y el índice agregado de intensidad de cobertura MIDIS. Los valores ausentes fueron reemplazados por cero; posteriormente, las variables se transformaron mediante $\log_{1p}$, se winsorizaron entre los percentiles 1 y 99 y se escalaron con RobustScaler.

Se compararon cuatro algoritmos de agrupamiento. K-means se evaluó entre dos y diez conglomerados, con 100 inicializaciones y semilla aleatoria 42. El modelo de mezclas gaussianas se examinó entre dos y diez componentes, con matriz de covarianza completa, 20 inicializaciones y la misma semilla. El agrupamiento jerárquico aglomerativo se evaluó en igual intervalo mediante enlace de Ward. Para DBSCAN, los valores de ϵ se obtuvieron de los percentiles 60 a 95 de las distancias al vecino más cercano, y $\min_samples$ tomó los valores 5, 13, 16 y 19 (Jain, 2010).

La calidad interna de las soluciones se evaluó mediante el coeficiente de silueta y los índices Davies-Bouldin y Calinski-Harabasz. Para K-means también se analizaron la inercia y el método del codo; para el modelo de mezclas gaussianas, los criterios AIC y BIC; y para DBSCAN, la proporción de observaciones clasificadas como ruido. En este último algoritmo, las métricas se calcularon únicamente con las observaciones no consideradas ruido y se conservaron las soluciones que generaron entre dos y diez conglomerados, con un máximo de 30 % de ruido (Arbelaitz et al., 2013).

La comparación final incluyó las soluciones de K-means determinadas por el método del codo y la mejor silueta, el modelo de mezclas gaussianas con menor BIC, el agrupamiento aglomerativo con mejor silueta y la mejor configuración válida de DBSCAN. Para cada candidato se calculó el puntaje:

$$P_C = R_S + R_{DB} + R_{CH} + P_R$$

donde R_S y R_{CH} representaron los rangos descendentes de Silhouette y Calinski-Harabasz, mientras que R_{DB} correspondió al rango ascendente de Davies-Bouldin. La penalización por ruido (P_R) fue cero hasta el 10 %, un punto cuando superó el 10 % y dos puntos cuando excedió el 20 %. La solución con el menor puntaje fue seleccionada como principal.

De manera complementaria, se examinó una solución K-means de cuatro conglomerados para obtener perfiles territoriales más detallados. El análisis de componentes principales con tres dimensiones se empleó únicamente para la representación gráfica. Las variables de pobreza no participaron en la formación de los conglomerados y se utilizaron posteriormente para caracterizar los perfiles identificados.

Clasificación de la vulnerabilidad territorial

La clasificación supervisada utilizó como variable objetivo la vulnerabilidad territorial alta de 2018 y como predictores los 13 indicadores de cobertura social. Antes del entrenamiento se realizó una auditoría para excluir variables asociadas con la definición del objetivo, como la pobreza, el ranking distrital, los grupos robustos y las variables derivadas de vulnerabilidad. Estas se conservaron únicamente para la interpretación posterior de los resultados.

Se evaluaron dos escenarios. El **escenario A** utilizó exclusivamente los 13 indicadores de cobertura del MIDIS. El **escenario B** incorporó, además, la pertenencia a cuatro conglomerados obtenidos mediante K-means, representada mediante tres variables binarias y un conglomerado de referencia. Para evitar fuga de información, el preprocesamiento y K-means se ajustaron únicamente con los datos de entrenamiento de cada partición (Kapoor & Narayanan, 2023).

En ambos escenarios se compararon tres clasificadores:

Regresión logística Elastic Net: solucionador SAGA, proporción L1 de 0,5, $C=1$, ponderación balanceada de clases, máximo de 5 000 iteraciones y semilla aleatoria 42 (Zou & Hastie, 2005).

CatBoost: función de pérdida Logloss, métrica AUC, 500 iteraciones, tasa de aprendizaje de 0,03, profundidad máxima de 4, regularización L2 de 5, balance automático de clases y semilla 42.

Explainable Boosting Machine: cinco términos de interacción, cuatro bolsas externas, 256 intervalos máximos, tasa de aprendizaje de 0,03, máximo de 2 500 rondas y procesamiento paralelo. El desbalance de clases se compensó mediante pesos muestrales.

Las probabilidades estimadas se transformaron en clases utilizando un umbral de 0,50.

Validación y comparación estadística

La base distrital se dividió de manera estratificada en un conjunto de desarrollo del 75 % y un conjunto de prueba independiente del 25 %, utilizando una semilla aleatoria de 42. Este último permaneció aislado durante el entrenamiento, la validación cruzada y la selección de los modelos.

En el conjunto de desarrollo se aplicó una validación cruzada estratificada de diez particiones, con barajado aleatorio y los mismos folds para todos los modelos y escenarios. En cada partición, la imputación con cero, la transformación log1p, la winsorización, el escalamiento robusto y, en el escenario B, el ajuste de K-means se realizaron exclusivamente con los datos de entrenamiento. Los parámetros obtenidos se aplicaron posteriormente a la partición de validación correspondiente (Varma & Simon, 2006).

El desempeño se evaluó mediante exactitud, exactitud balanceada, precisión, sensibilidad, especificidad, F1-score, ROC-AUC, PR-AUC, matriz de confusión y tiempo de entrenamiento. Debido al desbalance de clases y al interés por identificar los territorios vulnerables, el F1-score se estableció como métrica principal. Para cada combinación de modelo y escenario se calcularon la media, desviación estándar, error estándar e intervalo de confianza del 95 % obtenidos en los diez folds.

La distribución del F1-score se examinó mediante la prueba de Shapiro-Wilk y gráficos cuantil-cuantil. La comparación global se realizó con la prueba de Friedman y el tamaño del efecto se estimó mediante el coeficiente W de Kendall. Como análisis pareado complementario se aplicó la prueba de rangos con signo de Wilcoxon, con corrección de Holm para comparaciones múltiples y un nivel de significancia de 0,05 (García et al., 2010).

El modelo final se seleccionó mediante el menor rango promedio de F1-score en los folds compartidos. Como criterios complementarios se consideraron los rangos de PR-AUC, ROC-AUC y exactitud balanceada. Posteriormente, cada modelo se reentrenó con todo el conjunto de desarrollo y se evaluó una sola vez en la prueba independiente.

Las predicciones de la prueba independiente se compararon mediante la prueba exacta de McNemar y un bootstrap pareado de 1 000 repeticiones para estimar las diferencias de F1-score y sus intervalos de confianza del 95 %. Los valores de significancia de ambas comparaciones fueron ajustados mediante el método de Holm.

Explicabilidad y análisis de errores

La explicabilidad se desarrolló mediante procedimientos globales y específicos para cada modelo. En el conjunto de desarrollo se examinaron las correlaciones entre predictores y el factor de inflación de la varianza, con el propósito de identificar multicolinealidad y orientar la interpretación de los resultados.

Para la regresión logística Elastic Net se obtuvieron los coeficientes estandarizados, las razones de momios (*odds ratios*) y los efectos marginales aproximados sobre la probabilidad de vulnerabilidad. Su estabilidad se evaluó mediante 200 remuestreos *bootstrap*, a partir de los cuales se calcularon intervalos de confianza del 95 % y la proporción de repeticiones en las que cada coeficiente conservó su signo.

Para CatBoost y Explainable Boosting Machine se extrajeron las medidas de importancia nativa. Además, en todos los modelos se calculó la importancia por permutación con diez repeticiones por predictor, utilizando PR-AUC, ROC-AUC y F1-score como criterios de evaluación (Marcinkevičs & Vogt, 2023).

La calibración de las probabilidades se examinó mediante el Brier score y curvas construidas con diez intervalos definidos por cuantiles (Van Calster et al., 2019). Finalmente, las predicciones se clasificaron en verdaderos positivos, verdaderos negativos, falsos positivos y falsos negativos. Los errores se analizaron globalmente y según departamento y perfil territorial; estas variables se utilizaron únicamente con fines interpretativos y no participaron en el entrenamiento de los modelos.

3. RESULTADOS Y DISCUSIÓN

3.1 Conformación de la base territorial

La integración de los registros administrativos del MIDIS con los mapas distritales de pobreza generó una base de 35 952 observaciones distrito-mes, correspondientes a 1 893 distritos y 19 periodos, sin duplicados en la combinación periodo-UBIGEO. Para el análisis supervisado se conservaron 1 872 distritos con información válida de pobreza en 2018. La pobreza distrital promedio fue de 43,95 % en 2013 y de 34,16 % en 2018, mientras que el cambio promedio entre los distritos con información disponible en ambos años fue de -9,31 puntos porcentuales. La alta vulnerabilidad histórica se definió operacionalmente mediante el percentil 75 de la pobreza distrital de 2018, equivalente a 46,80 %, como se sintetiza en la Tabla 1.

Tabla 1. Características generales de la base territorial

Indicador	Resultado
Registros distrito-mes	35 952
Distritos analizados	1 893
Periodos disponibles	19
Distritos con pobreza 2018	1 872
Distritos sin pobreza 2018	21
Distritos sin pobreza 2013	90
Pobreza distrital promedio 2013	43,95 %
Pobreza distrital promedio 2018	34,16 %
Cambio promedio 2013-2018	-9,31 puntos porcentuales
Punto de corte de vulnerabilidad	46,7974 %
Distritos vulnerables	468
Distritos de menor vulnerabilidad relativa	1 404

3.2 Cobertura de los programas sociales

La presencia registrada de los programas sociales mostró diferencias considerables entre los distritos y periodos analizados. Pensión 65 y Juntos presentaron la mayor extensión, con registros positivos en más del 99 % de las observaciones, mientras que PAIS y FONCODES mostraron una presencia más concentrada, como se detalla en la Tabla 2.

Tabla 2. Presencia de los programas sociales en la base analítica

Programa	Registros con presencia	Porcentaje de registros	Distritos con presencia
Pensión 65	35 923	99,92 %	1 891
Juntos	35 789	99,55 %	1 888
Wasi Mikuna	35 127	97,71 %	1 891
Contigo	34 928	97,15 %	1 848
Cuna Más	28 642	79,67 %	1 545
FONCODES	14 950	41,58 %	918
PAIS	7 679	21,36 %	407

En el ámbito departamental, Huancavelica presentó el mayor índice promedio de intensidad de cobertura, con 710,96 unidades por cada mil habitantes, seguida de Ayacucho, Apurímac y Huánuco, con 683,56, 653,10 y 594,46, respectivamente. Los valores más bajos se registraron en el Callao, Ica y Lambayeque, con 54,68, 144,76 y 179,07 unidades por cada mil habitantes.

La evolución mensual del índice también presentó variaciones marcadas, con un mínimo de 319,13 unidades por cada mil habitantes en febrero de 2025 y un máximo de 544,34 en abril del mismo año. Estas fluctuaciones se interpretaron con cautela, debido a que podían estar relacionadas con

diferencias en la actualización, acumulación o reporte de los registros administrativos, como se observa en la Figura 2.

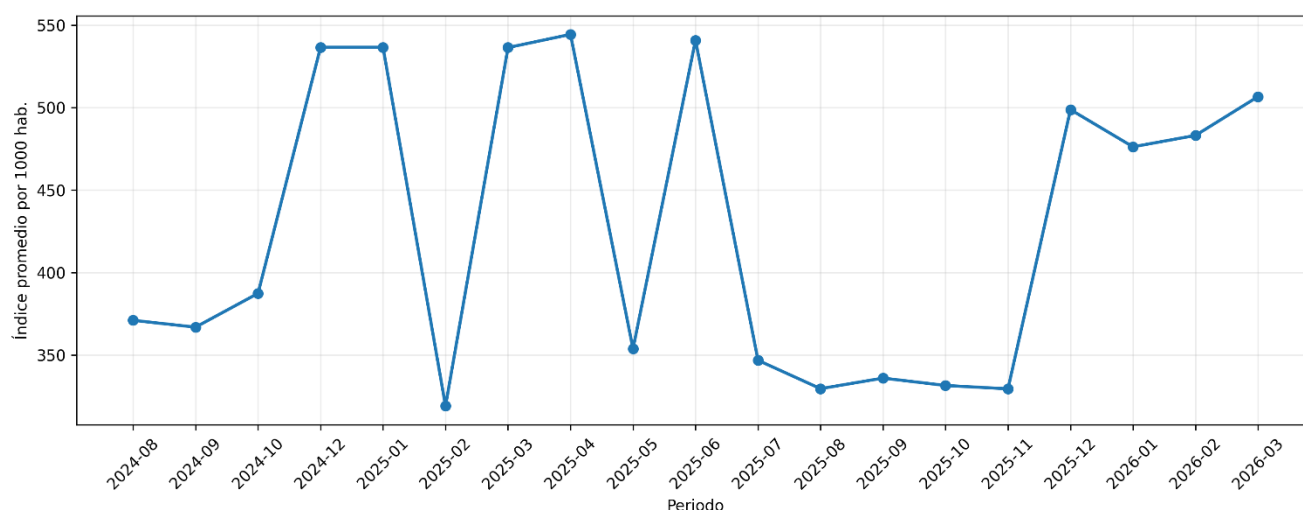


Figura 2. Evolución mensual del índice de intensidad de cobertura MIDIS

3.3 Cobertura social y vulnerabilidad territorial

Los distritos clasificados con alta vulnerabilidad histórica presentaron mayor intensidad y diversidad de cobertura social. El índice agregado de intensidad de cobertura MIDIS alcanzó 621,10 unidades por cada mil habitantes, frente a 369,76 en los distritos de menor vulnerabilidad relativa. Asimismo, el número promedio de programas presentes fue de 6,05 y 5,15, respectivamente.

Como se muestra en la Tabla 3, la pobreza distrital de 2018 presentó asociaciones positivas con la mayoría de los indicadores de cobertura. Los coeficientes más elevados correspondieron a los hogares abonados y afiliados de Juntos, con valores de 0,672 y 0,664, seguidos por el índice agregado de intensidad de cobertura MIDIS, con 0,599. En contraste, el cuidado diurno de Cuna Más presentó la única asociación negativa, con un coeficiente de $-0,261$. Estos resultados reflejaron principalmente la concentración de los programas sociales en territorios con mayores condiciones históricas de vulnerabilidad y no un efecto causal de la cobertura sobre la pobreza.

Tabla 3. Correlaciones entre los indicadores de cobertura y la pobreza distrital de 2018

Variable	Correlación de Spearman
Hogares abonados por Juntos	0,672
Hogares afiliados a Juntos	0,664
Índice agregado de cobertura MIDIS	0,599
Acompañamiento a familias de Cuna Más	0,565
Número de programas presentes	0,549
Número de variables de cobertura presentes	0,480
Usuarios de Contigo	0,469
Usuarios de Pensión 65	0,448
Beneficiarios de PAIS	0,301
Usuarios estimados de FONCODES	0,286
Niñas y niños atendidos por Wasi Mikuna	0,256
Cuidado diurno de Cuna Más	$-0,261$

3.4 Perfiles territoriales de cobertura social

La comparación de los algoritmos mostró que DBSCAN alcanzó el coeficiente de silueta más alto, con 0,606; sin embargo, clasificó como ruido al 27,05 % de los distritos. En cambio, K-means con

dos conglomerados obtuvo un coeficiente de silueta de 0,540, un índice Davies-Bouldin de 0,696 y un índice Calinski-Harabasz de 1 946,77. Aunque su silueta fue menor que la de DBSCAN, esta solución asignó la totalidad de los distritos a un perfil y presentó un mejor equilibrio entre separación, compactación e interpretabilidad. Por ello, K-means con dos conglomerados fue seleccionado como la solución principal del análisis no supervisado.

Los tres primeros componentes principales explicaron conjuntamente el 85,13 % de la variabilidad de los indicadores de cobertura: 56,14 % en el primer componente, 20,97 % en el segundo y 8,01 % en el tercero. La distribución de los distritos en este espacio tridimensional evidenció una diferenciación consistente entre los dos perfiles territoriales identificados, como se muestra en la Figura 3.

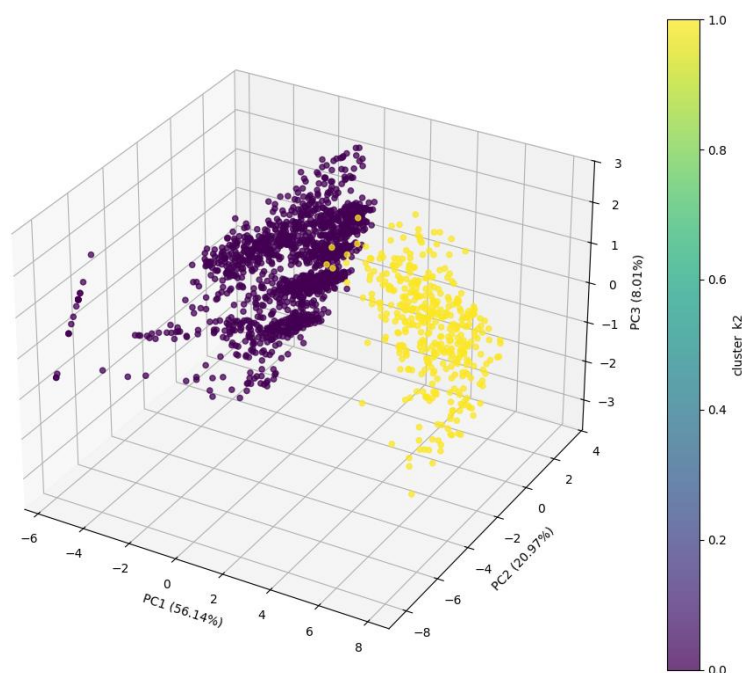


Figura 1. Distribución tridimensional de los perfiles territoriales obtenidos mediante K-means (k=2)

La solución principal identificó un perfil de menor intensidad y diversidad de cobertura, conformado por 1 497 distritos, y otro de mayor intensidad y diversidad, integrado por 396. Como se presenta en la Tabla 4, el segundo perfil registró un índice de intensidad de cobertura MIDIS superior al doble del observado en el primero, además de una mayor cantidad promedio de programas y variables de cobertura con valores positivos. Asimismo, estuvo integrado por distritos menos poblados y con mayores niveles históricos de pobreza.

Tabla 4. Características de los perfiles territoriales obtenidos mediante K-means

Indicador	Perfil 1: menor vulnerabilidad y cobertura	Perfil 2: mayor vulnerabilidad y cobertura
Distritos	1 497	396
Porcentaje de distritos	79,08 %	20,92 %
Pobreza promedio 2018	31,38 %	44,53 %
Pobreza promedio 2013	40,35 %	57,04 %
Cambio de pobreza	-8,46 puntos	-12,40 puntos
Población promedio 2020	20 155	7 254
Índice de cobertura MIDIS	347,88	729,89
Programas presentes	5,05	6,59
Variables de cobertura presentes	7,46	11,28

La diferencia de 13,16 puntos porcentuales en la pobreza promedio de 2018 entre ambos perfiles resultó especialmente relevante, debido a que esta variable no participó en la formación de los conglomerados. En consecuencia, los patrones administrativos de cobertura permitieron distinguir territorios asociados con distintos niveles de vulnerabilidad histórica. No obstante, la mayor cobertura observada en el segundo perfil no debe interpretarse como un efecto causal de los programas sobre la pobreza, sino como una expresión de su focalización en distritos históricamente más desfavorecidos.

De manera complementaria, la solución de K-means con cuatro conglomerados obtuvo un coeficiente de silueta de 0,362 y permitió reconocer patrones territoriales más específicos. Sin embargo, debido a su menor separación interna, se conservó únicamente como segmentación complementaria para evaluar su aporte en uno de los escenarios de clasificación supervisada.

3.5 Desempeño de los modelos supervisados

La vulnerabilidad territorial se clasificó en dos escenarios: el escenario A empleó los 13 indicadores de cobertura social, mientras que el escenario B incorporó adicionalmente tres variables binarias derivadas de la segmentación K-means con cuatro conglomerados. Como se observa en la Tabla 5, la regresión logística Elastic Net del escenario A alcanzó el mayor F1-score promedio en la validación cruzada, con 0,655, y el menor rango promedio, con 2,30. Las demás combinaciones obtuvieron valores de F1-score entre 0,645 y 0,654, lo que evidenció diferencias numéricas reducidas.

La prueba de Friedman no identificó diferencias globales estadísticamente significativas entre los modelos y escenarios, $\chi^2(5) = 10,310$; $p = 0,0669$, con un coeficiente W de Kendall de 0,206. Las comparaciones pareadas mediante Wilcoxon con corrección de Holm tampoco resultaron significativas. En consecuencia, Elastic Net del escenario A fue seleccionado por presentar el mejor rango promedio y el mayor F1-score en la validación cruzada, sin que ello implicara una superioridad estadística frente a los demás modelos.

Tabla 5. Desempeño de los modelos en la validación cruzada

Escenario y modelo	F1-score	Exactitud balanceada	Sensibilidad	ROC-AUC	PR-AUC	Rango promedio
Elastic Net, escenario A	0,655	0,796	0,818	0,878	0,666	2,30
Elastic Net, escenario B	0,654	0,795	0,821	0,877	0,666	2,80
EBM, escenario A	0,652	0,792	0,807	0,881	0,707	3,40
CatBoost, escenario B	0,653	0,785	0,752	0,874	0,707	3,80
EBM, escenario B	0,645	0,787	0,801	0,879	0,702	4,25
CatBoost, escenario A	0,649	0,782	0,752	0,875	0,718	4,45

En la prueba independiente, cuyos resultados se presentan en la Tabla 6, Elastic Net del escenario B obtuvo el mayor F1-score numérico, con 0,656, mientras que CatBoost del escenario A alcanzó el mayor PR-AUC, con 0,683. El modelo seleccionado, Elastic Net del escenario A, mantuvo un desempeño competitivo, con F1-score de 0,647, sensibilidad de 0,829, exactitud balanceada de 0,792 y ROC-AUC de 0,871. La incorporación de los conglomerados no generó mejoras consistentes

entre los tres clasificadores, por lo que su utilidad se concentró principalmente en la caracterización territorial. Las pruebas de McNemar y el bootstrap pareado, corregidas mediante Holm, tampoco identificaron diferencias significativas entre las predicciones de los modelos.

Tabla 6. Desempeño de los modelos en la prueba independiente

Modelo	Exactitud	Exactitud balanceada	Precisión	Sensibilidad	F1-score	ROC-AUC	PR-AUC
Elastic Net, escenario B	0,778	0,801	0,535	0,846	0,656	0,871	0,638
EBM, escenario A	0,782	0,795	0,542	0,821	0,653	0,862	0,643
EBM, escenario B	0,782	0,792	0,543	0,812	0,651	0,857	0,640
Elastic Net, escenario A	0,774	0,792	0,530	0,829	0,647	0,871	0,639
CatBoost, escenario B	0,786	0,783	0,552	0,778	0,645	0,864	0,674
CatBoost, escenario A	0,780	0,776	0,542	0,769	0,636	0,867	0,683

3.6 Explicabilidad, calibración y análisis de errores

La explicabilidad y el análisis de errores se realizaron sobre la regresión logística Elastic Net del escenario A, seleccionada previamente mediante validación cruzada. Como se observa en la Figura 4, el modelo clasificó correctamente a 97 de los 117 distritos vulnerables del conjunto de prueba independiente y dejó sin detectar a 20. Entre los 351 distritos de menor vulnerabilidad relativa, identificó correctamente a 265 y generó 86 falsas alertas. La sensibilidad de 0,829 evidenció que el modelo reconoció aproximadamente el 83 % de los territorios vulnerables; sin embargo, la precisión de 0,530 indicó que parte de los distritos señalados requeriría una revisión adicional.

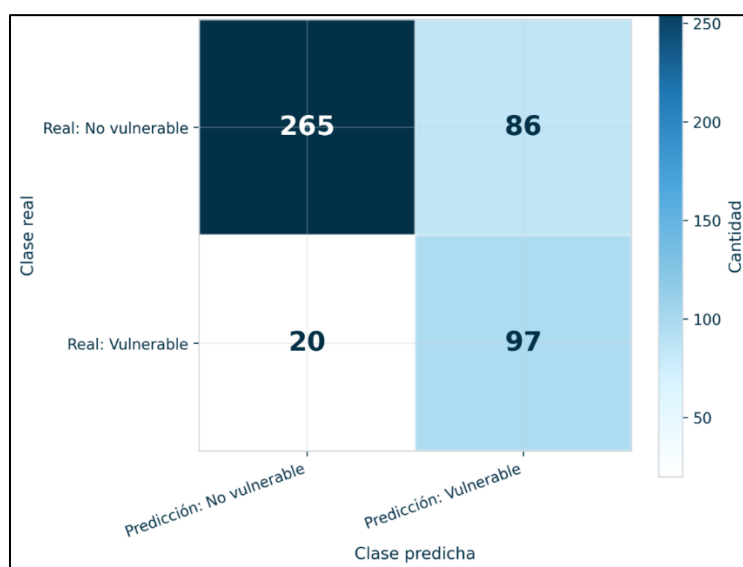


Figura 2. Matriz de confusión del modelo Elastic Net en el conjunto de prueba independiente

La contribución global de los predictores se examinó mediante importancia por permutación, utilizando el PR-AUC como métrica de referencia. En la Figura 5 se observa que los hogares abonados por Juntos presentaron la mayor importancia, con una disminución media de 0,152 al permutar sus valores. Les siguieron los usuarios estimados de FONCODES, con 0,142, y el número de programas presentes, con 0,102. Estos indicadores constituyeron las principales señales administrativas empleadas por el modelo para discriminar territorios históricamente vulnerables.

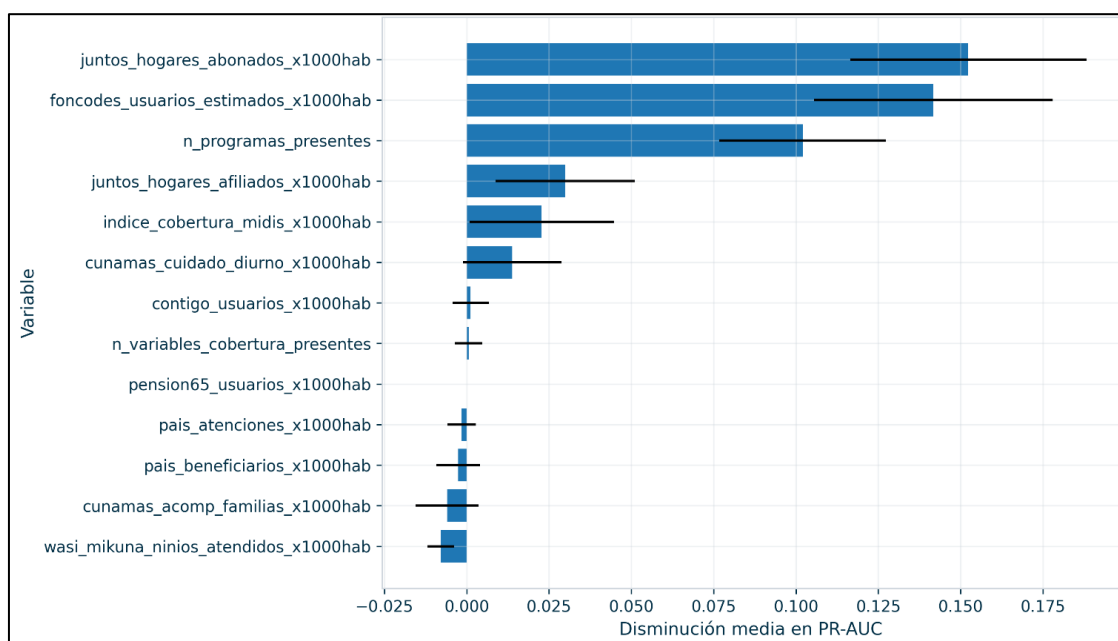


Figura 3. Importancia por permutación de las variables del modelo Elastic Net respecto del PR-AUC

El análisis de estabilidad mediante 200 remuestreos bootstrap mostró signos positivos consistentes para los hogares abonados y afiliados de Juntos, los usuarios estimados de FONCODES y el número de programas presentes. En contraste, el cuidado diurno de Cuna Más presentó un signo negativo estable. Estas relaciones expresaron asociaciones internas del modelo y no efectos causales de los programas sociales sobre la pobreza.

También se identificó una elevada multicolinealidad entre variables operativas de un mismo programa. Los hogares afiliados y abonados por Juntos presentaron factores de inflación de la varianza de 375,05 y 372,76, respectivamente, mientras que las atenciones y los beneficiarios de PAIS alcanzaron valores de 218,83 y 212,18. Aunque la regularización Elastic Net redujo la inestabilidad asociada con estas dependencias, los coeficientes individuales se interpretaron conjuntamente con los resultados del bootstrap y la importancia por permutación.

En cuanto a la calibración, CatBoost obtuvo los menores Brier scores en ambos escenarios, con 0,143, seguido por Elastic Net, con 0,152, y EBM, con valores cercanos a 0,154. Por tanto, CatBoost produjo probabilidades ligeramente mejor calibradas, aunque no alcanzó el menor rango promedio ni el mayor F1-score durante la validación cruzada. En conjunto, Elastic Net del escenario A mantuvo un equilibrio adecuado entre capacidad discriminativa, sensibilidad e interpretabilidad, sin presentar una superioridad estadísticamente significativa frente a las demás configuraciones evaluadas.

3.7 Discusión

Los resultados mostraron que los registros administrativos de cobertura contienen información útil para diferenciar territorios con distintos niveles históricos de vulnerabilidad. Los distritos con mayor pobreza presentaron una mayor intensidad y diversidad de intervenciones sociales. Este patrón debe interpretarse como una expresión de la focalización territorial y no como un efecto de los programas sobre la pobreza. Altındağ et al. (2021) demostraron que los datos administrativos pueden apoyar la focalización social cuando se someten a procesos adecuados de diseño y validación.

El agrupamiento mediante K-means identificó dos perfiles diferenciados, pese a que la pobreza no intervino en la formación de los conglomerados. El perfil con mayor cobertura presentó también mayores niveles históricos de pobreza y una población promedio menor. Esto indica que los patrones operativos de los programas permitieron recuperar parcialmente la estructura territorial de la vulnerabilidad. No obstante, estos perfiles no sustituyen una medición actual y multidimensional, pues los territorios priorizados pueden variar según el criterio de bienestar empleado (Henderson & Follett, 2022).

La regresión logística Elastic Net basada exclusivamente en indicadores de cobertura fue seleccionada por su desempeño en la validación cruzada y mantuvo resultados competitivos en la prueba independiente. La ausencia de diferencias significativas frente a CatBoost y EBM mostró que una mayor complejidad algorítmica no garantizó mejores resultados. Li et al. (2022) también encontraron que incorporar más variables o modelos complejos no mejora necesariamente la clasificación de pobreza. Además, la preferencia por una alternativa parsimoniosa coincide con Asri et al. (2022), quienes destacaron la utilidad operativa de criterios de focalización más simples y verificables.

La elevada sensibilidad del modelo evidenció una buena capacidad para identificar distritos vulnerables, aunque su precisión moderada produjo una cantidad relevante de falsas alertas. Este comportamiento puede resultar útil cuando se busca reducir errores de exclusión, siempre que las predicciones sean revisadas institucionalmente. Poulin et al. (2022) señalaron que los métodos de focalización basados en aprendizaje automático deben valorarse no solo por su desempeño, sino también por sus costos, transparencia y consecuencias operativas. Por ello, el modelo debe utilizarse como apoyo para la priorización territorial y no como mecanismo automático de asignación de beneficios.

La explicabilidad identificó como señales principales la cobertura de Juntos, los usuarios estimados de FONCODES y la diversidad de programas. Sin embargo, la elevada multicolinealidad impidió interpretar estas variables como efectos independientes o causales. Aunque CatBoost presentó una calibración ligeramente mejor, Elastic Net ofreció un equilibrio adecuado entre rendimiento, sensibilidad e interpretabilidad. En concordancia con Amarasinghe et al. (2023), las explicaciones aplicadas a políticas públicas deben responder a necesidades concretas de decisión y complementarse con información actualizada y criterio especializado.

CONCLUSIONES

La integración de los registros administrativos del MIDIS con los mapas distritales de pobreza permitió identificar perfiles territoriales asociados con la vulnerabilidad histórica y desarrollar un modelo explicable para su clasificación. K-means distinguió dos perfiles diferenciados, mientras que Elastic Net basado únicamente en indicadores de cobertura fue seleccionado por su equilibrio entre desempeño, sensibilidad e interpretabilidad, sin mostrar diferencias significativas frente a modelos más complejos. La incorporación de los conglomerados no mejoró de forma consistente la predicción, por lo que su principal aporte fue descriptivo. En consecuencia, el modelo puede apoyar la priorización y revisión territorial, siempre que se complemente con información socioeconómica actualizada y no se utilice como mecanismo automático de asignación de beneficios.

FINANCIAMIENTO

Los autores no recibieron ningún patrocinio para llevar a cabo este estudio-artículo.

CONFLICTO DE INTERESES

Los autores declaran que no existe ningún tipo de conflicto de interés relacionado con la materia del trabajo.

CONTRIBUCIÓN DE AUTORES

Conceptualización: Marin-Rodriguez, W. J.; Vellón-Flores, V. I. Curación de datos: Marin-Rodriguez, W. J.; Ausejo-Sánchez, J. L.; Solano-Armas, T. Análisis formal: Marin-Rodriguez, W. J.; Vellón-Flores, V. I.; Ausejo-Sánchez, J. L. Investigación: Marin-Rodriguez, W. J.; Vellón-Flores, V. I.; Solano-Armas, T.; Maguiña-Poma, Y. E. Metodología: Marin-Rodriguez, W. J.; Vellón-Flores, V. I.; Samanamud-Malca, S. Administración del proyecto: Marin-Rodriguez, W. J.; Samanamud-Malca, S. Software: Marin-Rodriguez, W. J.; Ausejo-Sánchez, J. L. Supervisión: Maguiña-Poma, Y. E.; Samanamud-Malca, S. Validación: Vellón-Flores, V. I.; Solano-Armas, T.; Maguiña-Poma, Y. E. Visualización: Marin-Rodriguez, W. J.; Ausejo-Sánchez, J. L. Redacción – borrador original: Marin-Rodriguez, W. J.; Vellón-Flores, V. I. Redacción – revisión y edición: Marin-Rodriguez, W. J.; Vellón-Flores, V. I.; Ausejo-Sánchez, J. L.; Solano-Armas, T.; Maguiña-Poma, Y. E.; Samanamud-Malca, S.

REFERENCIAS

- Aiken, E., Bellue, S., Karlan, D., Udry, C., & Blumenstock, J. E. (2022). Machine learning and phone data can improve targeting of humanitarian aid. *Nature*, 603(7903), 864–870. <https://doi.org/10.1038/s41586-022-04484-9>
- Altındağ, O., O’Connell, S. D., Şaşmaz, A., Balcioğlu, Z., Cadoni, P., Jerneck, M., & Foong, A. K. (2021). Targeting humanitarian aid using administrative data: Model design and validation. *Journal of Development Economics*, 148, 102564. <https://doi.org/10.1016/j.jdeveco.2020.102564>
- Amarante, V., Rossel, C., & Sánchez Laguardia, G. (2026). Social Protection Coverage in Latin America: Aggregate and Individual-level Transformations. *Journal of Human Development and Capabilities*, 1–27. <https://doi.org/10.1080/19452829.2025.2610195>
- Amarasinghe, K., Rodolfa, K. T., Lamba, H., & Ghani, R. (2023). Explainable machine learning for public policy: Use cases, gaps, and research directions. *Data & Policy*, 5, e5. <https://doi.org/10.1017/dap.2023.2>
- Arbelaitz, O., Gurrutxaga, I., Muguerza, J., Pérez, J. M., & Perona, I. (2013). An extensive comparative study of cluster validity indices. *Pattern Recognition*, 46(1), 243–256. <https://doi.org/10.1016/j.patcog.2012.07.021>
- Asri, V., Michaelowa, K., Panda, S., & Paul, S. B. (2022). The pursuit of simplicity: Can simplifying eligibility criteria improve social pension targeting? *Journal of Economic Behavior & Organization*, 200, 820–846. <https://doi.org/10.1016/j.jebo.2022.06.003>
- Behrens, J. T. (1997). Principles and procedures of exploratory data analysis. *Psychological Methods*, 2(2), 131–160. <https://doi.org/10.1037/1082-989X.2.2.131>
- Borga, L. G., & D’Ambrosio, C. (2021). Social protection and multidimensional poverty: Lessons from Ethiopia, India and Peru. *World Development*, 147, 105634. <https://doi.org/10.1016/j.worlddev.2021.105634>
- Ferreira, L. Z., Utazi, C. E., Huicho, L., Nilsen, K., Hartwig, F. P., Tatem, A. J., & Barros, A. J. D. (2022). Geographic inequalities in health intervention coverage – mapping the composite coverage index in Peru using geospatial modelling. *BMC Public Health*, 22(1), 2104.

<https://doi.org/10.1186/s12889-022-14371-7>

- García, S., Fernández, A., Luengo, J., & Herrera, F. (2010). Advanced nonparametric tests for multiple comparisons in the design of experiments in computational intelligence and data mining: Experimental analysis of power. *Information Sciences*, 180(10), 2044–2064. <https://doi.org/10.1016/j.ins.2009.12.010>
- García, V. H. B., Martínez, F. de M. G., Llamas Félix, B. I., & Esparza, R. M. V. (2024). Medición de la pobreza multidimensional en México mediante un análisis bibliométrico y de ecuaciones estructurales. *Iberoamerican Journal of Science Measurement and Communication*, 4(2), 1–23. <https://doi.org/10.47909/ijsmc.1354>
- Hall, O., Ohlsson, M., & Rögnvaldsson, T. (2022). A review of explainable AI in the satellite data, deep machine learning, and human poverty domain. *Patterns*, 3(10), 100600. <https://doi.org/10.1016/j.patter.2022.100600>
- Harron, K., Dibben, C., Boyd, J., Hjern, A., Azimaee, M., Barreto, M. L., & Goldstein, H. (2017). Challenges in administrative data linkage for research. *Big Data & Society*, 4(2), 205395171774567. <https://doi.org/10.1177/2053951717745678>
- Henderson, H., & Follett, L. (2022). Targeting social safety net programs on human capabilities. *World Development*, 151, 105741. <https://doi.org/10.1016/j.worlddev.2021.105741>
- Hinojosa Pérez, J. A., Avalos, H. R. B., Salazar, I. Y. V., & Carrasco Mamani, S. C. (2024). Social Programs and Socioeconomic Variables: Their Impact on Peruvian Regional Poverty (2013–2022). *Economies*, 12(8), 197. <https://doi.org/10.3390/economies12080197>
- INEI. (2015). *Mapa de Pobreza Provincial y Distrital 2013*. Gob.pe. <https://www.gob.pe/institucion/inei/informes-publicaciones/3204925-mapa-de-pobreza-provincial-y-distrital-2013>
- INEI. (2020). *Mapa de Pobreza Provincial y Distrital 2018*. Gob.pe. <https://www.gob.pe/institucion/inei/informes-publicaciones/3204872-mapa-de-pobreza-provincial-y-distrital-2018>
- INEI. (2026). *INEI: Pobreza monetaria alcanzó al 25,7% de la población en el Perú durante el año 2025*. Gob.pe. <https://www.gob.pe/institucion/inei/noticias/1387366-inei-pobreza-monetaria-alcanzo-al-25-7-de-la-poblacion-en-el-peru-durante-el-ano-2025>
- Jain, A. K. (2010). Data clustering: 50 years beyond K-means. *Pattern Recognition Letters*, 31(8), 651–666. <https://doi.org/10.1016/j.patrec.2009.09.011>
- Jung, W., Benotsmane, R., Stoeffler, Q., Kim, A. H., Ghadimi, S., Hosseini, M., Ntarlagiannis, D., Ammari, T., Lu, Y., & Steiner, J. (2026). Contextualized poverty targeting with multimodal spatial data and machine learning in Brazzaville, Congo. *Cities*, 170, 106429. <https://doi.org/10.1016/j.cities.2025.106429>
- Kapoor, S., & Narayanan, A. (2023). Leakage and the reproducibility crisis in machine-learning-based science. *Patterns*, 4(9), 100804. <https://doi.org/10.1016/j.patter.2023.100804>
- Li, Q., Yu, S., Échevin, D., & Fan, M. (2022). Is poverty predictable with machine learning? A study of DHS data from Kyrgyzstan. *Socio-Economic Planning Sciences*, 81, 101195. <https://doi.org/10.1016/j.seps.2021.101195>
- Lisboa, P. J. G., Saralajew, S., Vellido, A., Fernández-Domenech, R., & Villmann, T. (2023). The coming of age of interpretable and explainable machine learning models. *Neurocomputing*, 535, 25–39. <https://doi.org/10.1016/j.neucom.2023.02.040>
- Maco, V. (2025). Lock-in from within: Challenges to expanding cash transfer programs in Peru. *Journal of International and Comparative Social Policy*, 41(3), 335–349.

<https://doi.org/10.1017/ics.2025.10066>

- Marcinkevičs, R., & Vogt, J. E. (2023). Interpretable and explainable machine learning: A methods-centric overview with concrete examples. *WIREs Data Mining and Knowledge Discovery*, 13(3). <https://doi.org/10.1002/widm.1493>
- MIDIS. (2026). *Cobertura de los Programas Sociales adscritos al MIDIS- [Ministerio de Desarrollo e Inclusión Social - MIDIS]*. Gob.pe. <https://www.datosabiertos.gob.pe/dataset/cobertura-de-los-programas-sociales-adscritos-al-midis-ministerio-de-desarrollo-e-inclusion>
- Papadakis, T., Christou, I. T., Ipektsidis, C., Soldatos, J., & Amicone, A. (2024). Explainable and transparent artificial intelligence for public policymaking. *Data & Policy*, 6, e10. <https://doi.org/10.1017/dap.2024.3>
- Paucara-Charca, R., Almanza Huamán, L. M., Szczepansky Grobas, D., Landa Almanza, A. I., & Muñiz Laura, K. A. (2024). La eficacia de los efectos políticos en el desarrollo alcance y resultados de las políticas regionales orientadas para mitigar el impacto de la explotación laboral en menores de edad en especial en el sector minero de La Pampa Madre de Dios 2018. *Revista Amazónica De Ciencias Sociales*, 2(2), 12–21. <https://doi.org/10.55873/racs.v2i2.262>
- Poulin, C., Trimmer, J., Press-Williams, J., Yachori, B., Khush, R., Peletz, R., & Delaire, C. (2022). Performance of a novel machine learning-based proxy means test in comparison to other methods for targeting pro-poor water subsidies in Ghana. *Development Engineering*, 7, 100098. <https://doi.org/10.1016/j.deveng.2022.100098>
- Schober, P., Boer, C., & Schwarte, L. A. (2018). Correlation Coefficients: Appropriate Use and Interpretation. *Anesthesia & Analgesia*, 126(5), 1763–1768. <https://doi.org/10.1213/ANE.0000000000002864>
- Smythe, I. S., & Blumenstock, J. E. (2022). Geographic microtargeting of social assistance with high-resolution poverty maps. *Proceedings of the National Academy of Sciences*, 119(32). <https://doi.org/10.1073/pnas.2120025119>
- Van Calster, B., McLernon, D. J., van Smeden, M., Wynants, L., & Steyerberg, E. W. (2019). Calibration: the Achilles heel of predictive analytics. *BMC Medicine*, 17(1), 230. <https://doi.org/10.1186/s12916-019-1466-7>
- Van den Broeck, J., Argeseanu Cunningham, S., Eeckels, R., & Herbst, K. (2005). Data Cleaning: Detecting, Diagnosing, and Editing Data Abnormalities. *PLoS Medicine*, 2(10), e267. <https://doi.org/10.1371/journal.pmed.0020267>
- Varma, S., & Simon, R. (2006). Bias in error estimation when using cross-validation for model selection. *BMC Bioinformatics*, 7(1), 91. <https://doi.org/10.1186/1471-2105-7-91>
- Zou, H., & Hastie, T. (2005). Regularization and Variable Selection Via the Elastic Net. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 67(2), 301–320. <https://doi.org/10.1111/j.1467-9868.2005.00503.x>