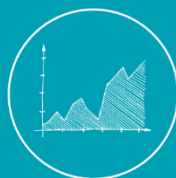
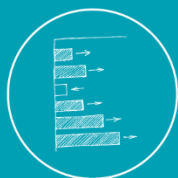




Aplicación de modelos de aprendizaje automático para el cálculo de horas de trabajo no remunerado en las Cuentas Satélites del Trabajo No Remunerado

2026



Autoridades:

Eva María Mera

Directora Ejecutiva

Marianita Granda

Subdirectora General

Andrea Molina Vera

Coordinadora General Técnica de Innovación en Métricas y Análisis de la Información

Cecilia Valdivia

Coordinadora General Técnica de Producción Estadística

Diana Barco

Directora de Estadísticas Económicas

Autores:

Henry Valdiviezo

Instituto Nacional de Estadística y Censos, Ecuador

Magaly Aguiar

Instituto Nacional de Estadística y Censos, Ecuador

Revisión:

Andrea Molina Vera

Coordinadora General Técnica de Innovación en Métricas y Análisis de la Información

Galo Egas

Director de Innovación en Métricas y Metodologías

Patricia Sánchez

Miembro de Equipo DINME

Sobre los Cuadernos de Trabajo

Los cuadernos de trabajo son documentos técnico-científicos orientados al desarrollo de investigaciones, tanto metodológicas como temáticas, que promueven la generación de conocimiento y el debate técnico.

Nota de responsabilidad

Las interpretaciones y opiniones expresadas en este documento pertenecen a los autores y no reflejan el punto de vista oficial del Instituto Nacional de Estadística y Censos (INEC). El INEC ha realizado una revisión del documento, no obstante, no garantiza la exactitud de los datos que figuran en el documento.

Citar como: [Valdiviezo, H]. & [Aguiar, M]. (2026). *Aplicación de modelos de aprendizaje automático para el cálculo de horas de trabajo no remunerado en las Cuentas Satélites del Trabajo No Remunerado*. Cuadernos de Trabajo, Instituto Nacional de Estadística y Censos (INEC), Quito, Ecuador.

Aplicación de modelos de aprendizaje automático para el cálculo de horas de trabajo no remunerado en las Cuentas Satélites del Trabajo No Remunerado

Henry Valdiviezo¹

Magaly Aguiar²

Resumen

En 2023, el módulo de uso del tiempo de la Encuesta Nacional de Empleo, Desempleo y Subempleo (ENEMDU) de Ecuador no levantó cinco de las quince preguntas habitualmente utilizadas para la construcción de las Cuentas Satélite de Trabajo No Remunerado de los Hogares (CSTNRH), comprometiendo la continuidad de la serie estadística. Ante esta limitación, el estudio propone la imputación de los valores faltantes mediante técnicas de aprendizaje automático (Machine Learning, ML). Se utilizan como base de entrenamiento los módulos de 2015 y 2017. El modelado se estructura en dos etapas: clasificación (participación en la actividad) mediante regresión logística con regularización y regresión (horas semanales) mediante XGBoost. Los modelos se estiman de forma independiente por sexo, dada la brecha de género existente en la distribución del trabajo no remunerado.

Las predicciones para las actividades domésticas regulares: tender las camas (UT31), compras diarias (UT47) y vigilancia del hogar (UT58) resultan consistentes con las tendencias históricas. Mientras, las actividades de participación comunitaria (UT100) y cuidado de personas con discapacidad (UT126) presentan mayor volatilidad, asociada a la baja incidencia poblacional de estas preguntas (clases altamente desbalanceadas). En general, los resultados confirman la persistencia de una carga horaria diferencial por sexo, con predominio de la mujer.

Se concluye que el ML constituye una herramienta válida para la imputación de información faltante en encuestas de hogares, especialmente para actividades de alta frecuencia, aunque se recomienda cautela en actividades de baja incidencia en tareas no remuneradas.

Palabras clave: Trabajo no remunerado, uso del tiempo, predicción ML, machine learning, investigación social, XGBoost.

Línea de Investigación: Sociodemográfica

Operación Estadística: Cuenta Satélite de Trabajo No Remunerado de los Hogares

¹ Investigador de la Dirección de Estadísticas Económicas del Instituto Nacional de Estadística y Censos. Dirección de correspondencia henry_valdiviezo@inec.gob.ec.

² Investigadora de la Dirección de Estadísticas Económicas del Instituto Nacional de Estadística y Censos. Dirección de correspondencia magaly_aguiar@inec.gob.ec.

1 Introducción

El trabajo no remunerado (TNR) constituye un pilar fundamental en el funcionamiento de la economía del hogar, no obstante, esta contribución no es valorada como parte de las Cuentas Nacionales y por lo tanto no se cuantifica en el Producto Interno Bruto (PIB). Frente a este vacío, se han desarrollado instrumentos analíticos como las Cuentas Satélite de Trabajo No Remunerado de los Hogares (CSTNRH), cuyo objetivo principal es visibilizar el valor económico que crea el trabajo no remunerado y poder compararlo frente al valor económico de bienes y servicios mercantiles medidos en el Producto Interno Bruto (PIB).

En Ecuador las CSTNRH son producidas por el Instituto Nacional de Estadística y Censos (INEC) y los resultados más recientes revelan que en 2019 el Valor Agregado Bruto (VAB) del TNR representó el 20,6% del PIB (INEC, 2025). A su vez, desde una perspectiva de género esta producción muestra una estructura similar al conjunto de países de América Latina; en donde, el trabajo no remunerado recae principalmente sobre las mujeres, quienes de manera global aportan con el 75% del valor económico total del trabajo no remunerado (INEC, 2025). En estos indicadores se ve reflejada la importancia del Trabajo No Remunerado respecto a la Economía Nacional, pero además desde su enfoque de género, la asignación inequitativa del esfuerzo de trabajo en el hogar es palpable, así se puede observar en la sobrecarga de tareas domésticas para la mujer y su consecuente participación limitada en el mercado laboral (Medeiros et al., 2007; PNUD-OIT, 2025, p. 26).

Para la medición de las CSTNRH, es fundamental obtener estadísticas del tiempo destinado al TNR, para lo cual, en Ecuador se viene levantando un módulo de uso del tiempo con 15 preguntas representativas y homologadas dentro de la Encuesta Nacional de Empleo, Desempleo y Subempleo (ENEMDU). Sin embargo, debido a factores técnicos y económicos, en el año 2023 el levantamiento del módulo presentó una reducción de cinco de las quince preguntas de uso del tiempo; esta ausencia de datos ha comprometido la integralidad de la medición y complica la continuidad de la serie del TNR.

Frente a esta problemática, resultó imperativo desarrollar un proceso de predicción de valores faltantes; para ello, se consideran como base, la existencia de 10 preguntas relacionadas al TNR, las que en conjunto con otras variables de control (características sociodemográficas) facilitan la modelización de la predicción de valores para las 5 preguntas no levantadas en 2023. Para esta predicción, el presente estudio propone recurrir a la potencia predictiva que el machine learning ofrece frente a los métodos tradicionales (Di Franco & Santurro, 2021).

Bajo este propósito, el estudio se organiza en cuatro apartados adicionales a esta introducción. En primer lugar, se presenta una revisión de literatura sobre el uso de técnicas de machine learning en problemas de clasificación y regresión vinculados al trabajo no remunerado. En segundo lugar, se desarrolla la sección de datos y metodología, en la que se describe la información utilizada, los modelos seleccionados y las métricas empleadas para su evaluación. Posteriormente, se expone la estrategia de modelado aplicada. Luego se presentan los resultados de predicción y finalmente

se discute su validez como insumo para la medición del trabajo no remunerado en las cuentas satélite.

2 Revisión de la literatura

En el siguiente apartado se da una contextualización del trabajo no remunerado de los hogares y como los modelos de aprendizaje automático aplicados a la investigación social pueden contribuir en la imputación de las variables ausentes en el módulo de uso del tiempo de la encuesta ENEMDU.

2.1 Contextualización del trabajo no remunerado de los hogares

El análisis del trabajo no remunerado de los hogares (TNRH) ha adquirido relevancia creciente en la literatura económica y social, consolidándose como una línea de investigación orientada a visibilizar actividades históricamente excluidas de los sistemas tradicionales de medición económica (Benería, 1999). Desde una perspectiva conceptual, la exclusión del trabajo no remunerado de las estadísticas económicas oficiales genera una subestimación del aporte real de los hogares a la economía y refuerza desigualdades estructurales, especialmente de género. Estudios como los de Folbre (2006) y Esquivel (2010) destacan que la distribución del trabajo no remunerado es marcadamente desigual, recayendo mayoritariamente en las mujeres, lo que limita su participación en el mercado laboral, afecta sus trayectorias profesionales y profundiza brechas económicas y sociales. En este sentido, el estudio del trabajo no remunerado se ha consolidado como una herramienta clave para el análisis de desigualdades, particularmente aquellas asociadas al género y a la distribución del cuidado en los hogares (Benería, 1999).

En el ámbito estadístico, la medición del trabajo no remunerado se ha fortalecido a través de las encuestas de uso del tiempo, las cuales permiten cuantificar el tiempo dedicado a distintas actividades productivas y no productivas dentro del hogar (Enríquez, s. f.). En Ecuador, el INEC ha desarrollado encuestas específicas de uso del tiempo que han servido como insumo principal para la elaboración de la Cuenta Satélite del Trabajo No Remunerado de los Hogares (INEC, 2025). Esta cuenta satélite amplía el marco del Sistema de Cuentas Nacionales al asignar un valor económico al tiempo destinado a actividades no remuneradas, utilizando metodologías consistentes con las recomendaciones internacionales de organismos como la CEPAL y las Naciones Unidas (Armas et al., 2009).

Las CSTNRH permiten visibilizar el aporte económico del trabajo no remunerado y analizar su evolución en el tiempo, así como su distribución según variables sociodemográficas como sexo, edad, nivel educativo y composición del hogar. No obstante, uno de los principales desafíos en su construcción es la disponibilidad y continuidad de la información, dado que las encuestas de uso del tiempo no se levantan de manera anual (INEC, 2025). Esta limitación genera vacíos temporales que dificultan el análisis longitudinal y la actualización periódica de las cuentas satélite, especialmente en contextos de cambios estructurales como los provocados por la pandemia de COVID-19.

En este contexto, resulta necesario complementar los enfoques tradicionales de medición con metodologías que permitan imputar información faltante y proyectar el comportamiento del trabajo no remunerado en años sin levantamiento de encuestas. Estos desafíos abren la puerta al uso de técnicas analíticas más avanzadas, como el aprendizaje automático, que permite capturar patrones complejos en los datos y fortalecer la producción de estadísticas oficiales sobre el TNRH, garantizando su comparabilidad, consistencia y utilidad para la formulación de políticas públicas.

2.2 Machine Learning en la investigación social

El aprendizaje automático ha demostrado buena capacidad predictiva y de generalización en el análisis de fenómenos sociales complejos. En particular, Wang et al. (2022) señalan que el ML preserva mejor la estructura interna de la encuesta frente a los métodos estadísticos tradicionales, lo que reduce la pérdida de información y mejora la consistencia de los indicadores derivados. Estas bondades, presentan al ML como una alternativa metodológica especialmente pertinente para contextos de información compleja tales como la imputación en encuestas de hogares.

En el ámbito del mercado laboral, los modelos de ML se vienen aplicando en la predicción y análisis de variables como la situación laboral, horas trabajadas, niveles salariales e ingresos. El INEC (2023), en el estudio *"Imputación de Información en Encuestas de Hogares a través de técnicas de Machine Learning: Análisis del caso ecuatoriano"*, imputó información cuantitativa del ingreso laboral e información dicotómica sobre tenencia del Registro Único de Actividades del Servicio de Rentas Internas en la ENEMDU, utilizando el algoritmo XGBoost como modelo de mejor desempeño frente a otras alternativas probadas. Esta aplicación directa sobre la ENEMDU se constituye como un referente metodológico muy cercano al presente estudio, que justifica la adopción de un enfoque similar para la imputación de variables de uso del tiempo presentes en la misma encuesta.

En la misma línea, Qutub et al. (2021) en el estudio *"Prediction of Employee Attrition Using Machine Learning and Ensemble Methods"*, analizan los factores que predicen la pérdida de empleos, aplicando modelos de Random Forests, Gradient Boosting y Logistic Regression sobre datos de recursos humanos, concluyendo que los métodos de ensamble (boosting) superan consistentemente a los modelos paramétricos simples en variables dicotómicas con datos desbalanceados. Este hallazgo es relevante dado que las actividades del trabajo no remunerado, de manera irregular, presentan datos desbalanceados.

Otros estudios enfocados en la predicción e imputación de salarios e ingresos mediante técnicas ML son *"A Novel Salary Prediction System Using Machine Learning Techniques"* (Asaduzzaman et al. 2024) y *"Predicting Salary of an Employee using Machine Learning"* (ManiKumar et al., 2024), en éstos, los autores desarrollan múltiples algoritmos de ML (árboles de decisión, redes neuronales y modelos de ensamble) evaluando su desempeño frente a modelos convencionales y obteniendo resultados para la estimación salarial que ratifican la robustez de estos enfoques en un contexto de relaciones no lineales en los predictores.

Por su parte, Rosati, (2021), en *"Métodos de Machine Learning como alternativa para la imputación de datos perdidos"*, realiza un ejercicio de imputación en la Encuesta

Permanente de Hogares de Argentina, comparando el desempeño de ML frente a métodos tradicionales de imputación múltiple como el Hot Deck, y concluye que los algoritmos de ML producen imputaciones más precisas y con menor sesgo en variables de ingreso con distribución asimétrica, características similares a las variables de horas de trabajo no remunerado que se modelan en el presente estudio.

En el ámbito específico del trabajo no remunerado, las aplicaciones son escasas; no obstante, Mulder et al. (2024) encuentran una aplicación del ML para la predicción de asignaciones de tiempo en tareas del hogar, en este estudio los autores proponen un modelo de ML que combina datos de rastreo de actividad física con variables socioeconómicas para predecir el tiempo destinado a distintas actividades, entre ellas el trabajo no remunerado. Sus resultados muestran que los modelos de ensamble logran predicciones razonables en actividades de alta frecuencia, pero presentan mayor error en actividades de baja frecuencia, este hallazgo anticipa el problema en actividades del trabajo no remunerado como la participación en actividades comunitarias y el cuidado de personas con discapacidad debido a su baja incidencia y frecuencia.

Finalmente, en un contexto más específico Tripathi (2024) en "Predicting Unpaid Care Work in India Using Random Forest: An Analysis of Socioeconomic and Demographic Factors" aplica el algoritmo Random Forest para predecir la participación en trabajo de cuidado no remunerado en India a partir de variables sociodemográficas; sus hallazgos muestran que las características del hogar, nivel educativo, situación laboral y edad son predictores relevantes y ratifica el mejor desempeño de las predicciones basadas en ML frente a métodos tradicionales. Estos hallazgos sugieren la pertinencia del uso de ML para la predicción y son consistentes con la selección de variables explicativas especificadas en los modelos de predicción del presente estudio.

En síntesis, la revisión de la literatura evidencia que el ML constituye una herramienta metodológica sólida para la predicción e imputación de variables en encuestas de hogares, especialmente cuando se trata de variables dicotómicas de participación y variables continuas. Los antecedentes revisados revelan estrategias de modelado basadas en la selección del mejor modelo. Estos estudios no advierten "a priori" la superioridad de un determinado algoritmo, por lo tanto, su elección depende de los resultados obtenidos y las características propias de la información de cada estudio. Asimismo, la literatura advierte sobre las limitaciones de estos modelos en actividades de baja frecuencia, aspecto que se considerará en la interpretación de los resultados en variables con estas limitaciones.

3 Datos y Metodología

En esta sección se describe las fuentes de información, el proceso de preparación de datos y la metodología empleada para estimar las variables de trabajo no remunerado que no fueron levantadas en la ENEMDU 2023.

3.1 Información base del trabajo no remunerado

Desde una perspectiva de género, las CSTNRH son cruciales para visibilizar el aporte económico de las personas en la realización de tareas no remuneradas dentro del hogar; para construir este tipo de estadísticas, se utilizan principalmente dos insumos:

- 1) Encuestas especializadas de uso del tiempo, para proveer información de tiempo destinado a tareas no remuneradas; y
- 2) Encuestas de empleo, para proveer información sobre salarios en actividades de mercado análogas a las realizadas en el hogar (INEC, 2025).

En Ecuador, la información para la medición de salarios proviene de la ENEMDU, mientras el tiempo del TNR proviene de la Encuesta de Uso del Tiempo (EUT); no obstante, esta última no cumple con el requisito de periodicidad; así en el país se han realizado únicamente dos encuestas completas de uso del tiempo (años 2007 y 2012), con 66 preguntas orientadas al trabajo no remunerado. Esta limitación surge principalmente por su elevado costo económico, frente a ello, para sostener la medición económica de las CSTNRH, con cierta regularidad se viene levantando, en la encuesta ENEMDU un módulo de 15 preguntas de uso del tiempo relacionadas con el TNR.

La ENEMDU es una encuesta dirigida a hogares, su variable de diseño es el desempleo y está orientada a mercado laboral, sin embargo, también constituye la fuente oficial de indicadores de pobreza. Una de las mayores ventajas de esta encuesta, es su continuidad, tiene una periodicidad mensual con acumulaciones trimestrales y anuales disponibles, pero con diferentes niveles de representatividad estadística de acuerdo con la frecuencia y tamaño de muestra levantada en cada ronda de investigación.

Aunque la encuesta está orientada a proporcionar indicadores del mercado laboral, en la práctica se ha posicionado como una encuesta de propósitos múltiples; pues en ella, se han incorporado preguntas y secciones de información ajenas a su objetivo central de estudio del mercado laboral, tales como, seguridad, consumo, uso del tiempo, entre otros.

Bajo esta situación, dentro de la encuesta ENEMDU se viene levantando el módulo de uso del tiempo en TNR, como un insumo necesario para actualizar las CSTNRH y cuyo levantamiento ocurre generalmente en septiembre de cada año. De todos modos, es necesario subrayar que, el levantamiento de este módulo de 15 preguntas del TNR es un recurso emergente frente a la imposibilidad de realizar levantamientos anuales de encuestas completas y específicas de uso del tiempo³, por lo tanto, su levantamiento es un mínimo esencial que permite actualizar el valor del trabajo no remunerado en el país a través de las CSTNRH.

Entre los años 2015 y 2023, se realizaron cuatro levantamientos del módulo de uso del tiempo. Con excepción del año 2023, se levantaron las mismas 15 preguntas representativas del TNR (ver tabla 1), garantizando un levantamiento bajo un mismo esquema metodológico, mediante informante directo, lo que permite conocer de manera efectiva las actividades y tiempo destinado por los individuos a tareas no remuneradas ejecutadas durante la última semana de referencia según la entrevista.

Tabla 1. Actividades levantadas en el módulo uso de tiempo

Actividades / Tareas	2015	2017	2019	2023
UT15: Cocinar y preparar alimentos	√	√	√	√
UT31: Tender las camas	√	√	√	X
UT33: Limpiar la casa	√	√	√	√

³ El número de preguntas acerca del TNR que se levantó en la última encuesta completa y específica de uso del tiempo (2012) comprende un total de 66 preguntas.

UT41: Lavar la ropa	✓	✓	✓	✓
UT46: Compras semanales	✓	✓	✓	✓
UT47: Compras diarias	✓	✓	✓	X
UT58: Vigilar la seguridad del hogar	✓	✓	✓	X
UT61: Dar de comer a un niño/a	✓	✓	✓	✓
UT63: Jugar con niño/a	✓	✓	✓	✓
UT77: Reparaciones de vivienda	✓	✓	✓	✓
UT98: Ayudar a otro hogar	✓	✓	✓	✓
UT100: Participar en actividades sociales	✓	✓	✓	X
UT124: Dar de comer a PCD	✓	✓	✓	✓
UT125: Bañar, asear o vestir a PCD	✓	✓	✓	✓
UT126: Practicar alguna terapia especial a PCD	✓	✓	✓	X

Nota: Los años 2015 y 2017 fueron levantados en septiembre y 2023 en noviembre como parte de la ENEMDU, mientras el año 2019 fue levantado en diciembre como parte de la Encuesta Multipropósito.

En la línea de construir series de datos homologados, el levantamiento de información del año 2019, presentó dos aspectos limitantes: el primero, responde a su levantamiento en una encuesta distinta a la ENEMDU, abiertamente de propósitos múltiples, pues el objetivo de la Encuesta Multipropósito (actualmente no vigente) fue proveer información en diversas temáticas para la medición de indicadores del Plan Nacional de Desarrollo de ese año; y el segundo, el módulo de uso del tiempo fue levantado en el mes de diciembre; un mes plenamente atípico frente al resto del año, debido a festividades navideñas y de fin de año, aspectos que indudablemente inciden en el comportamiento y actividades diarias realizadas por las personas. Frente a estas limitaciones, la información de las 15 preguntas del módulo de uso del tiempo del año 2019, se consideran referenciales y no forman parte del modelamiento en la presente investigación.

Para el año 2023, el levantamiento presentó una limitación de cobertura: de las 15 preguntas representativas asociadas al trabajo no remunerado, únicamente se levantaron 10, es decir, se redujo un tercio de la información necesaria para actualizar el TNR. Las preguntas no levantadas en 2023 se relacionan con las siguientes actividades:

- UT31: Tender las camas.
- UT47: Compras diarias.
- UT58: Vigilar la seguridad del hogar.
- UT100: Participación en actividades de servicio comunitario no remunerado.
- UT126: Practicar terapia a persona con discapacidad.

La omisión de estas preguntas compromete la medición integral del trabajo no remunerado en el año 2023 y dificulta el cálculo de la serie de datos de horas promedio destinadas al TNR, las cuales finalmente son utilizadas para la medición económica en las CSTNRH⁴. Por esta razón, la base de datos de 2023 y particularmente las 5 preguntas ausentes, son el objetivo de predicción de valores del presente estudio; mientras, las 10

⁴ Para conocer en detalle sobre la metodología para la medición de las Cuentas Satélites del Trabajo No Remunerado (CSTNRH) se puede consultar: <https://www.ecuadorencifras.gob.ec/cuenta-satelite-del-trabajo-no-remunerado/>

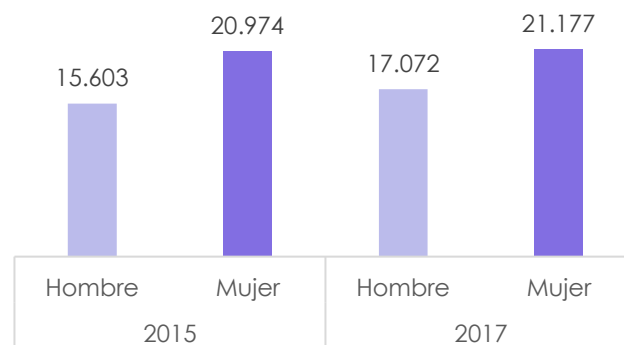
preguntas sí levantadas en 2015, 2017 y 2023 se constituyen como la información estadística clave a partir de la cual se busca predecir la información faltante.

Es así como, para implementar la predicción de valores ausentes, se utilizaron como insumos base, los módulos de uso del tiempo de la ENEMDU de los años 2015, 2017 y 2023. La estrategia de modelado se diseñó con un esquema temporal: los años 2015 y 2017 se emplearon para el entrenamiento y validación de los modelos, mientras el año 2023 se destinó al ejercicio de predicción/imputación.

3.2 Datos

En la figura 1 se observa que la muestra efectivamente levantada es superior a los 35 mil registros en cada año, esta magnitud facilita el modelamiento bajo el enfoque de machine learning, pues estos procesos demandan gran cantidad registros para la implementación de los algoritmos de aprendizaje.

Figura 1. Tamaño de muestra de personas que realizaron algún TNR (UT01 = 1) según sexo



Fuente: INEC - ENEMDU (2015 – 2017) - Módulo uso del tiempo

Estadísticas descriptivas del uso del tiempo en trabajo no remunerado

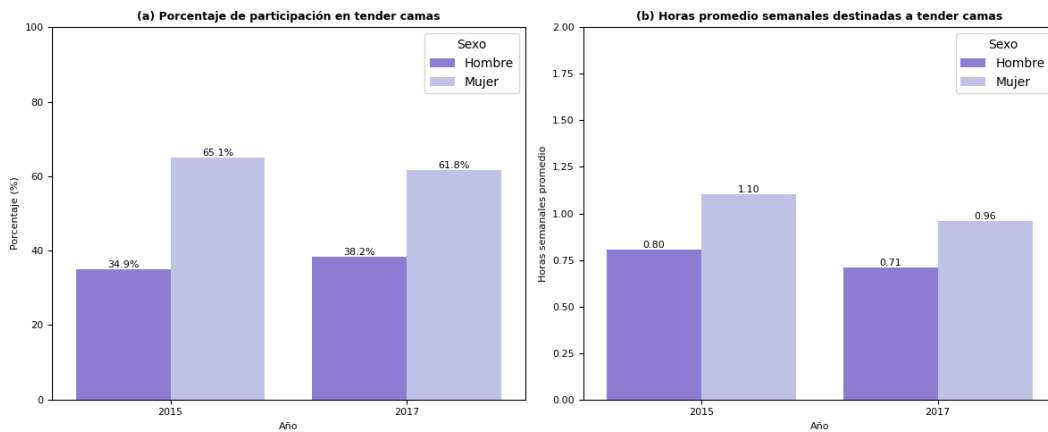
Los datos históricos de los módulos de uso del tiempo levantados en la ENEMDU son la fuente base para el entrenamiento de los modelos de aprendizaje automático para la predicción, por ello, conviene observar algunas estadísticas descriptivas de esta información base.

Estadísticas descriptivas de las variables a imputar

- **UT31:** Tender las camas

En la figura 2 se presenta la participación y el tiempo promedio semanal destinado a tender las camas (UT31) según sexo. Se observa que las mujeres registran mayor participación en los años analizados representando el 65,1% en 2015 y 61,8% en 2017. Por otro lado, la participación masculina aumenta ligeramente entre el 2015 y el 2017 de 34,9% a 38,2%. Respecto a las horas promedio, tanto para hombres y mujeres se reduce el tiempo promedio, sin embargo, para el 2017 la mujer continúa destinando más tiempo promedio semanal a tender las camas.

Figura 2. Porcentaje de participación y tiempo dedicado a tender camas (UT31) por sexo

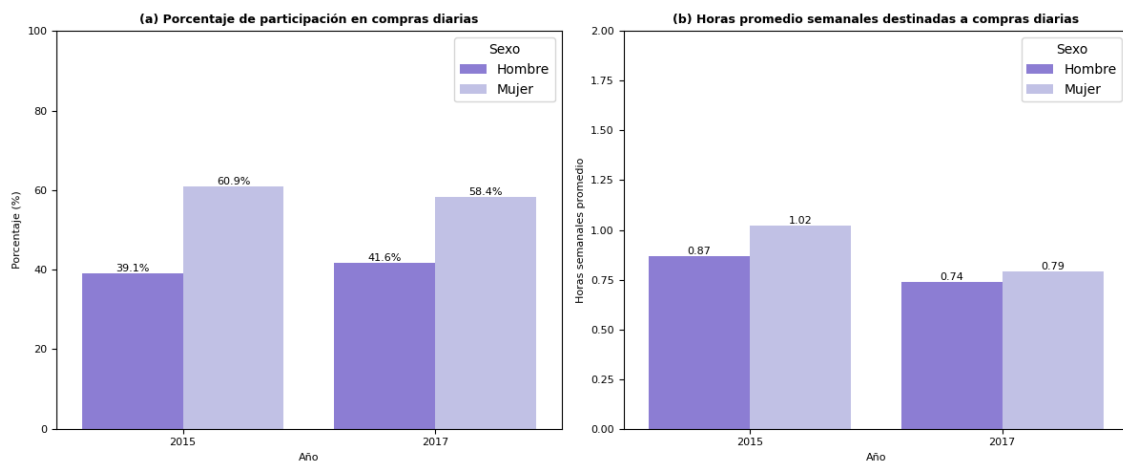


Fuente: INEC - ENEMDU (2015 – 2017) - Módulo uso del tiempo

- **UT47:** Compras diarias

En la figura 3 se presenta la participación y el tiempo promedio semanal destinado a las compras diarias (UT47) según sexo. Se observa que las mujeres registran mayor participación en los años analizados representando el 60,9% en 2015 y 58,4% en 2017. Por otro lado, la participación masculina aumenta ligeramente entre el 2015 y el 2017. Respecto a las horas promedio, tanto para hombres y mujeres el tiempo disminuye, sin embargo, la mujer continúa destinando más tiempo promedio semanal a esta actividad.

Figura 3. Porcentaje de participación y tiempo dedicado a las compras diarias (UT47) por sexo



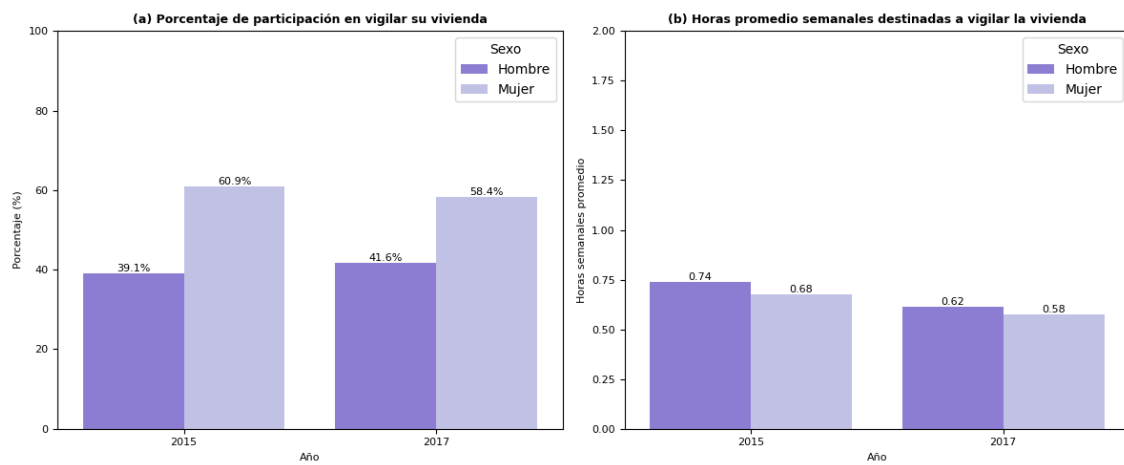
Fuente: INEC - ENEMDU (2015 – 2017) - Módulo uso del tiempo

- **UT58:** Vigilar la seguridad del hogar

En la figura 4 se presenta la participación y el tiempo promedio semanal destinado a vigilar la seguridad del hogar (UT58) según sexo. Se observa que los hombres registran un leve incremento en su participación en el año 2017 respecto al 2015. Sin embargo, las mujeres tienen la mayor participación en los dos años analizados. Respecto a las horas promedio, el tiempo disminuye tanto para hombres como para mujeres. En la actividad

de vigilar la seguridad de la vivienda el hombre dedica mayor tiempo promedio semanal con respecto a la mujer.

Figura 4. Porcentaje de participación y tiempo dedicado a vigilar la seguridad del hogar (UT58) por sexo

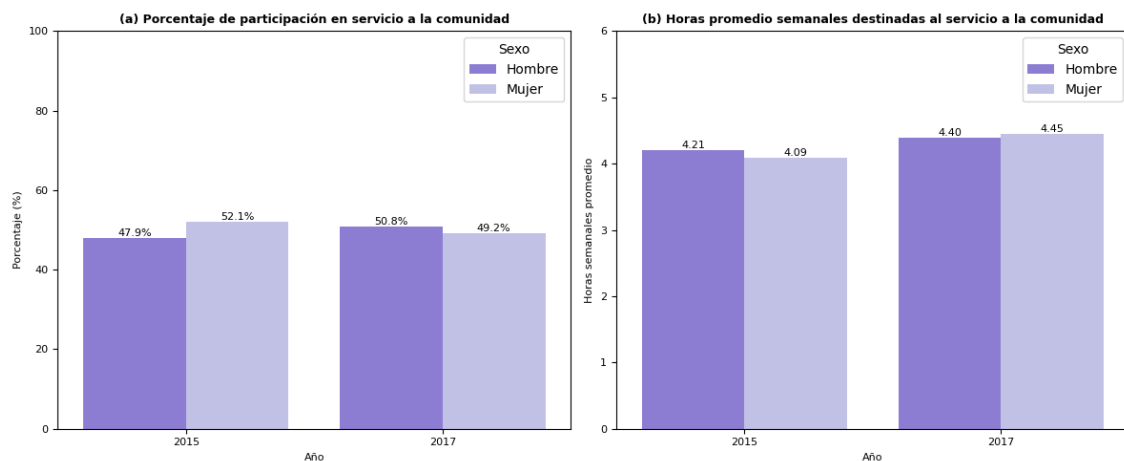


Fuente: INEC - ENEMDU (2015 – 2017) - Módulo uso del tiempo

- **UT100:** Participar en actividades sociales y comunitarias

La figura 5 presenta la participación y el tiempo promedio semanal destinado a realizar actividades destinadas al servicio a la comunidad (UT100) según sexo. La participación se invierte del año 2015 al 2017, dado que las mujeres reducen su participación, y los hombres tienen la mayor participación en el 2017. Respecto a las horas promedio, tanto para hombres y mujeres el tiempo promedio aumenta.

Figura 5. Porcentaje de participación y tiempo dedicado al servicio a la comunidad por sexo

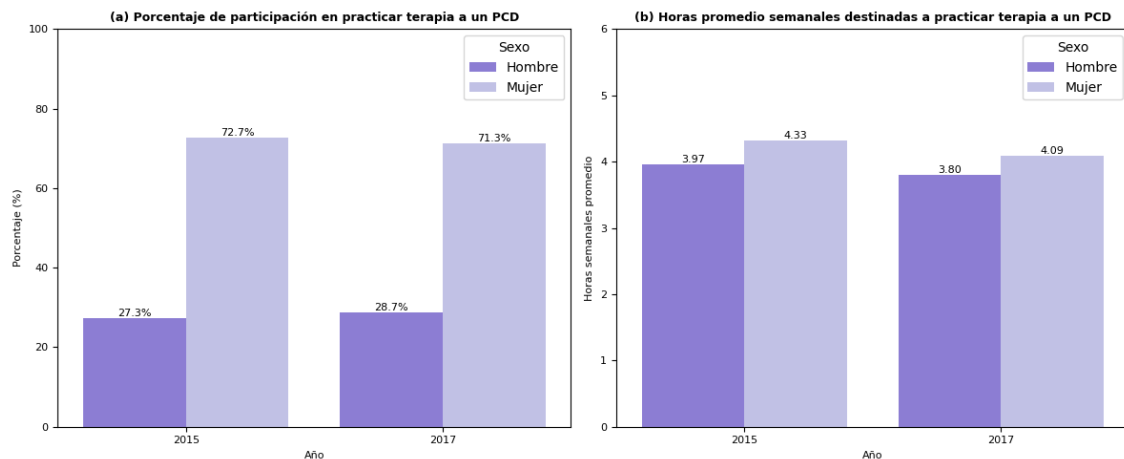


Fuente: INEC - ENEMDU (2015 – 2017) - Módulo uso del tiempo

- **UT126:** Practicar alguna terapia especial a personas con discapacidad (PCD)

La figura 6 presenta la participación y el tiempo promedio semanal destinado a practicar alguna terapia especial a un PCD (UT126) según sexo. En esta actividad la participación de la mujer es notable, dado que su participación engloba el 71,3% en el 2017. Se observa que los hombres incrementan su participación del 2015 al 2017. Respecto a las horas promedio, tanto para hombres y mujeres el tiempo promedio disminuye, siendo la mujer la que mayor tiempo dedica a esta actividad de cuidado.

Figura 6. Porcentaje de participación y tiempo dedicado a practicar terapia a un PCD por sexo



Fuente: INEC - ENEMDU (2015 – 2017) - Módulo uso del tiempo

Como se ha señalado previamente, estas cinco preguntas no fueron levantadas en la ENEMDU 2023, lo que genera vacíos de información para la medición del trabajo no remunerado. En este contexto, se ha visto necesario implementar un proceso de predicción que permita estimar los valores faltantes y garantizar la continuidad de la serie estadística utilizada en las CSTNRH. En la sección de anexos 7.1, se presentan las matrices de correlación que evidencian la relación entre las variables disponibles y las variables objetivo. Adicional, se observa la información proveniente de las diez preguntas de uso de tiempo levantadas en 2023, complementadas con variables sociodemográficas de control, los cuales son insumos suficientes para la construcción de modelos predictivos orientados a la estimación de valores faltantes. La metodología utilizada se describe en la siguiente sección.

3.3 Metodología

El aprendizaje automático es una subcategoría de la inteligencia artificial enfocada en construir sistemas que aprenden los patrones de los datos de entrenamiento y que posteriormente permiten hacer inferencias sobre nueva información (M. Chen, 2024). Estos sistemas buscan imitar la inteligencia humana ofreciendo un conjunto de métodos potentes de predicción de información en ambientes de información compleja; en este sentido, la potencia predictiva de las técnicas de machine learning se ha aplicado al ámbito de la investigación social (Di Franco & Santurro, 2021) destacando su fortaleza predictiva frente a métodos tradicionales.

Modelos de clasificación

El presente estudio propone reconstruir la información faltante de uso del tiempo, evaluando el desempeño de diferentes enfoques algorítmicos en condiciones reales de información con alta dimensionalidad (características socioeconómicas), desbalance de clases entre la proporción de personas que sí o no realizan TNR, y enfrentándose a cantidades atomizadas de tiempo destinado por las personas para realizar tareas domésticas y de cuidados en el hogar.

En función de la naturaleza de las variables (binarias y continuas) de las 5 preguntas ausentes en el módulo de uso del tiempo 2023, la presente investigación utiliza modelos “ML supervisados de clasificación basados en la regresión logística” para la predicción de variables dicotómicas (para identificar si la persona realizó o no TNR en cada pregunta faltante); y “Modelos ML supervisados” para la predicción del tiempo destinado a esas tareas del TNR.

Para la etapa de clasificación se optó por modelos de regresión logística debido a su carácter paramétrico, interpretabilidad y robustez en problemas binarios, lo cual permite establecer una línea base clara del desempeño predictivo. Posteriormente, se incorporaron variantes regularizadas mediante penalización L1 (Lasso) y penalización L2 (Ridge)⁵, con el objetivo de controlar el sobreajuste, especialmente considerando la alta dimensionalidad del conjunto de variables explicativas y la posible presencia de multicolinealidad entre predictores sociodemográficos.

Regresión logística

La regresión logística modela la probabilidad de ocurrencia de un evento binario a partir de una combinación lineal de variables predictoras. Esta probabilidad se transforma en intervalo [0,1] mediante la función logística o sigmoide, lo que permite interpretar la salida del modelo como la probabilidad de pertenecer a una de las dos clases (Hosmer & Lemeshow, 2012):

$$P(Y = 1|X) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 X_1 + \dots + \beta_k X_k)}}$$

Donde:

Y : es la variable de respuesta dicotómica

X_i : son las variables explicativas, y

β_i : son los coeficientes por estimar.

Dado que varias variables binarias del módulo de uso del tiempo presentan un marcado desbalance de clases, se empleó la estrategia de ajuste del umbral de decisión a partir de las probabilidades generadas por el modelo. (Domínguez, 2021).

Métricas de evaluación para modelos de clasificación

Las métricas que permiten evaluar los modelos de clasificación, en particular en el modelo logístico, se menciona a continuación:

- **Matriz de confusión:** La matriz de confusión es una tabla que resume el desempeño de un modelo de clasificación comparando las predicciones obtenidas del modelo frente a los valores reales (Campo, J, 2026).

⁵ . La regularización L1 (Least Absolute Shrinkage and Selection Operator, LASSO) incorpora una penalización basada en la suma de los valores absolutos de los coeficientes, favoreciendo la selección de variables. Por su parte, la regularización L2 (Ridge) utiliza una penalización basada en la suma de los coeficientes al cuadrado, contribuyendo a reducir la multicolinealidad y mejorar la estabilidad del modelo.

		Real	
		Positivo	Negativo
Predicción	Positivo	TP	FP
	Negativo	FN	TN

Donde:

TP (True positives): Casos donde el modelo predice correctamente la clase positiva⁶.

TN (True negatives): Casos donde el modelo predice correctamente la clase negativa⁷.

FP (False positives): Casos donde el modelo predice la clase positiva de manera incorrecta.

FN (False negatives): Casos donde el modelo predice la clase negativa de manera incorrecta.

- **Accuracy:** El accuracy mide la proporción de predicciones correctas entre el total de las observaciones (Campo, J, 2026). Su definición es:

$$Accuracy = \frac{\text{Número de predicciones correctas}}{\text{Total de predicciones}} = \frac{TP + TN}{TP + TN + FP + FN}$$

- **Precisión:** Esta medida indica el porcentaje de predicciones positivas que realmente son positivas (Campo, J, 2026).

$$Precisión = \frac{TP}{TP + FP}$$

- **Recall/Sensitivity:** La sensibilidad mide la capacidad del modelo para identificar correctamente los casos positivos (Fawcett, 2006).

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

- **F1-Score:** Esta métrica combina la precisión con la sensibilidad (recall) en un solo valor. Se define como la media armónica entre ambas métricas. Es una métrica importante de analizarla cuando las clases están desbalanceadas (Fawcett, 2006).

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

- **ROC Curve y AUC:** Esta métrica permite evaluar la capacidad en la que el modelo de clasificación distingue entre la clase positiva y negativa. La curva ROC representa la relación entre la Tasa de Verdaderos Positivos (TPR) y la Tasa de Falsos Positivos (FPR) para distintos umbrales de decisión (Campo, J, 2026).

La TPR y FPR se define como:

$$TPR = \frac{TP}{TP + FN}$$

⁶ Clase positiva (Clase 1): Es la categoría que representa la presencia del evento o fenómeno de interés

⁷ Clase negativa (Clase 0): Es la categoría que representa la ausencia del evento o fenómeno de interés

$$FPR = \frac{FP}{FP + TN}$$

Modelo de regresión

Para la estimación de las horas semanales dedicadas a las actividades del hogar se parte de un modelo de regresión lineal simple como línea base, debido a su simplicidad, interpretabilidad y capacidad para establecer una referencia inicial del comportamiento de la variable objetivo (Roustaei, 2024). Posteriormente, se incorporaron variantes regularizadas Ridge (L2) y Lasso (L1) con el propósito de mejorar la estabilidad del modelo frente a la multicolinealidad entre variables explicativas y reducir el riesgo de sobreajuste. Adicional se incluyó XGBoost como modelo no lineal basado en árboles de decisión, capaz de capturar relaciones complejas e interacciones entre variables que no pueden ser representadas adecuadamente por modelos lineales. Esta combinación de enfoques permitió comparar modelos paramétricos y no paramétricos, evaluando la interpretabilidad y capacidad predictiva.

Regresión lineal (base)

La regresión lineal es un tipo de algoritmo de aprendizaje automático supervisado que aprende de los conjuntos de datos etiquetados y asigna los puntos de datos con las funciones lineales más optimizadas, lo que permite realizar predicciones en nuevos conjuntos de datos. Supone que existe una relación lineal entre la entrada y la salida, lo que significa que la salida cambia a un ritmo constante a medida que cambia la entrada. Esta relación se representa mediante una línea recta (Regresión Lineal En El Aprendizaje Automático, s. f.).

Regresión lineal Ridge y Lasso

Las extensiones penalizadas introducen restricciones sobre los coeficientes para mejorar el desempeño predictivo y la robustez del modelo.

- Ridge (L2): introduce una penalización cuadrática sobre los coeficientes, lo que disminuye su varianza y permite manejar escenarios de alta correlación entre predictores (multicolinealidad), estabilizando las estimaciones (Cessie & Houwelingen, 1992).
- Lasso (L1): aplica una penalización absoluta que no solo reduce la magnitud de los coeficientes, sino que puede forzar algunos a cero, funcionando como un mecanismo de selección implícita de variables. Este rasgo resulta especialmente útil cuando se sospecha que no todas las covariables aportan de manera significativa al modelo (Tibshirani, 1996).

Extreme Gradient Boosting (XGBoost)

XGBoost es un algoritmo de ensamble que implementa la técnica de boosting con árboles de decisión como clasificadores base, contruidos de manera secuencial para corregir los errores de los modelos anteriores (T. Chen & Guestrin, 2016). A diferencia de enfoques lineales, este método optimiza de forma iterativa una función de pérdida mediante gradiente descendente, lo que le permite capturar relaciones no lineales y efectos de interacción entre predictores (García, 2024). En regresión, XGBoost busca predecir valores numéricos continuos minimizando las funciones de pérdida (p. ej., Raíz

del Error Cuadrático Medio - RMSE o Error Cuadrático Medio - MSE), a la vez que incorpora regularización para evitar el sobreajuste (XGBoost para regresión, s. f.).

Una de sus principales fortalezas radica en la inclusión de términos de regularización L1 y L2, que reducen el sobreajuste y estabilizan las estimaciones en escenarios de alta dimensionalidad. Adicionalmente, el algoritmo maneja de manera interna los valores faltantes, asignando rutas de división óptimas en cada nodo sin requerir imputación previa, lo cual es especialmente útil en encuestas sociales con patrones de no respuesta (Wei et al., 2019).

Asimismo, la combinación de técnicas como *shrinkage*⁸, *subsampling*⁹ de observaciones y predictores, y estrategias de paralelización, permiten que XGBoost sea computacionalmente eficiente y adecuado para trabajar con bases extensas. Estas características justifican su elección como complemento a los modelos lineales y penalizados, dado que contribuye a mejorar el ajuste global y reducir sesgos sistemáticos que se pueden producir en el proceso de predicción de variables del módulo de uso del tiempo 2023 (T. Chen & Guestrin, 2016).

Métricas de evaluación para modelos de regresión

Las métricas que se evalúan en un modelo de regresión son:

- **Error Cuadrático Medio (MSE):** es la medida promedio de los cuadrados de los errores, donde cada error corresponde a la diferencia entre el valor observado y la predicción del modelo (Madrigal, E, 2022).

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

- **Raíz del Error Cuadrático Medio (RMSE):** Es una métrica de evaluación utilizada ampliamente en aprendizaje automático para problemas de regresión. Representa la raíz cuadrada del error cuadrático medio (MSE), lo cual permite ver que la métrica esté expresada en las mismas unidades que la variable objetivo, facilitando la interpretación del rendimiento del modelo (Madrigal, E, 2022).

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2}$$

- **Error Absoluto Medio (MAE):** El MAE mide el error promedio sin considerar su dirección y penaliza de forma lineal la magnitud de las desviaciones, lo que lo hace más robusto frente a valores atípicos en comparación con MSE o RMSE (Madrigal, E, 2022).

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i|$$

⁸ Tasa de aprendizaje (Hilloulin & Umnakwe, 2024)

⁹ Técnica que reduce el tamaño de los datos o la información para optimizar el rendimiento en modelos como XGBoost (Gao et al., 2017)

- **Coefficiente de determinación (R^2):** En ML, representa la fracción de la varianza explicada por las predicciones. Se usa para evaluar el grado de ajuste de modelos como regresión lineal, Random Forest, Gradient Boosting y redes neuronales (Madrigal, E, 2022).

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

4 Estrategia de modelado

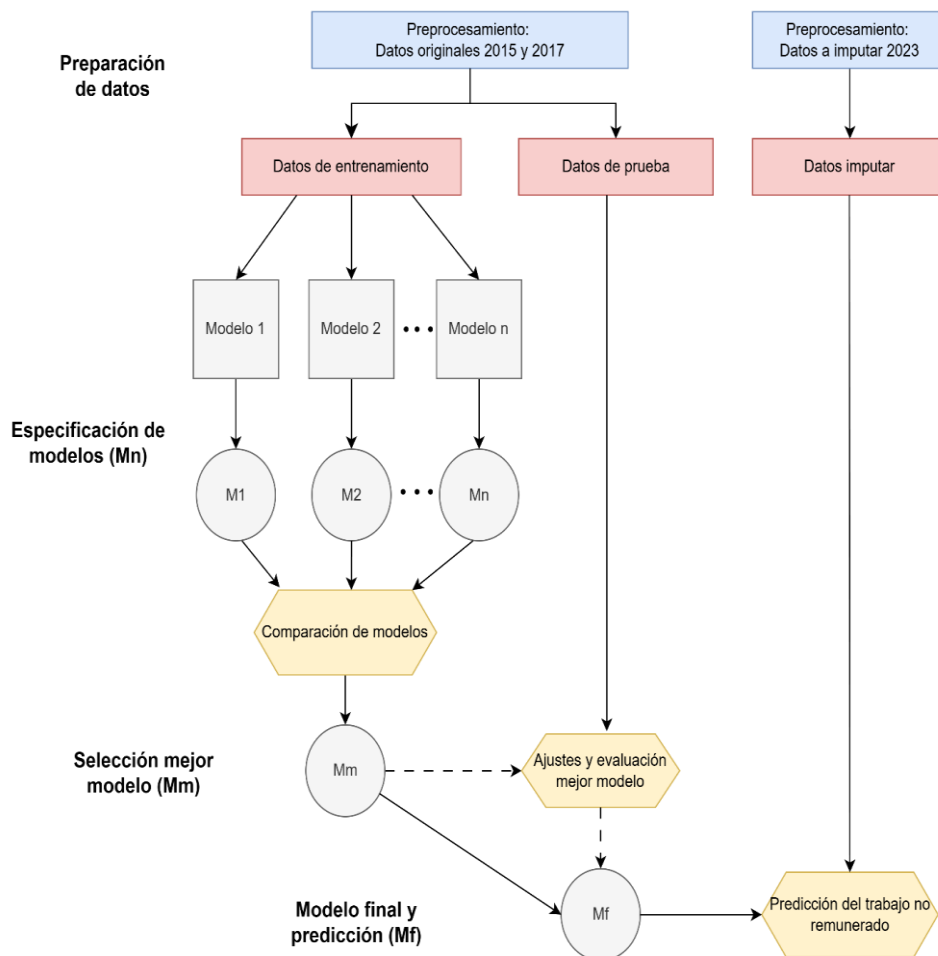
Este apartado constituye un componente central del documento metodológico, pues describe los flujos, técnicas, procedimientos y decisiones analíticas consideradas a la predicción de las variables de interés.

Flujo implementación del modelo

El proceso de modelado, ilustrado en la figura 7, consta de tres etapas:

- preparación de datos,
- especificación de modelos, y
- comparación y selección de modelos.

Figura 7. Flujo de implementación del modelo



Fuente: Elaboración propia

4.1 Preparación de datos

En la preparación de datos se incluyen varios pasos que determinan la calidad del modelado. El objetivo de la preparación de datos es reducir el tamaño de los datos e identificar las relaciones entre las variables y normalizarlas.

En esta etapa, en primer lugar, la base se restringe a individuos de 12 años y más, de acuerdo con la población objetivo de las CSTNRH. Luego, se considera únicamente los individuos que declaran realizar actividades de trabajo no remunerado ($ut01 = 1$). Estos criterios permiten que la población bajo análisis corresponda al grupo de estudio, evitando la inclusión de observaciones que no pertenecen al universo de interés.

Posteriormente, se fijan como años de referencia 2015 y 2017, utilizando aquellas variables que presentan un porcentaje superior al 10% de valores no perdidos, con el fin de garantizar la robustez de los predictores en la fase de modelado. Luego se revisan las variables categóricas para verificar la coherencia de sus categorías y frecuencias. Asimismo, las variables de respuesta binaria (UT01, UT31, UT47, UT58, UT100, UT126) se recodifican de modo que la clase mayoritaria quede representada como 1 y la minoritaria como 0, para facilitar la interpretación de métricas.

El conjunto de datos se divide en subconjuntos de entrenamiento y prueba, utilizando una proporción del 70% para entrenamiento (train) y 30% para prueba (test), asegurando aleatoriedad mediante una semilla fija¹⁰. Para el entrenamiento del modelo se emplea la función objetivo de regresión con error cuadrático y el método de construcción de árboles basado en histogramas, configuración que favorece la eficiencia computacional en conjuntos de datos de gran tamaño.

Es importante mencionar que el proceso de modelización se realiza de manera independiente para hombres y mujeres, considerando las diferencias estructurales observadas en la participación y distribución del tiempo destinado al trabajo no remunerado de los hogares. Esta estrategia permite capturar patrones específicos de comportamiento para cada grupo poblacional y mejorar la capacidad predictiva de los modelos.

En consecuencia, para cada una de las cinco preguntas objeto de imputación se estiman modelos diferenciados por sexo. Adicionalmente, debido a que cada pregunta requiere una predicción de participación (variable binaria) y una predicción de tiempo dedicado (variable continua), se desarrollan modelos de clasificación y regresión de forma independiente.

4.2 Especificación de modelos

El proceso de modelamiento se estructura diferenciando tipo de problema (clasificación y regresión) y sexo (hombres y mujeres), considerando la heterogeneidad en la participación y el tiempo dedicado a las actividades de trabajo no remunerado. En total, la estrategia metodológica contempla la estimación de veinte modelos

¹⁰ Es un valor que se utiliza para inicializar un generador de números aleatorios (GNA), un componente fundamental en muchas operaciones de IA (Nikunj Vaghasiya, 2024)

predictivos, correspondientes a cinco actividades, dos tipos de modelado (clasificación y regresión) y dos grupos poblacionales.

Variables dependientes

Las variables dependientes corresponden a indicadores de participación (clasificación) y tiempo dedicado (regresión) en actividades específicas del trabajo no remunerado. Estas se definen a partir de las preguntas del módulo de uso del tiempo, según se detalla a continuación:

a) Modelos de clasificación (participación en la actividad)

Variables binarias que indican si la persona realizó o no la actividad:

UT31: Tender la cama

UT47: Realizar compras diarias

UT58: Participar en la seguridad de la vivienda

UT100: Realizar servicio a la comunidad

UT126: Ayudar en terapias y atención a una persona con discapacidad (PCD)

Estas variables toman el valor de 1 si la persona realizó la actividad y 0 en caso contrario.

b) Modelos de regresión (tiempo semanal dedicado)

Para los individuos que reportan participación en la actividad, se modela el tiempo semanal dedicado, medido en horas promedio semanales:

UT31_semt: Tiempo semanal en tender la cama

UT47_semt: Tiempo semanal en realizar compras diarias

UT58_semt: Tiempo semanal en participar en la seguridad de la vivienda

UT100_semt: Tiempo semanal en realizar servicio a la comunidad

UT126_semt: Tiempo semanal en ayudar en terapias y atención a una persona con discapacidad

Variables independientes

Las variables independientes corresponden a un conjunto de características sociodemográficas, económicas y del hogar consideradas como potenciales factores explicativos de la participación y del tiempo dedicado al trabajo no remunerado. Si bien existe un núcleo común de variables utilizado en los modelos, la especificación final de cada uno se ajusta en función de la actividad analizada y de la disponibilidad de información pertinente.

De manera general, el conjunto de variables explicativas incluye:

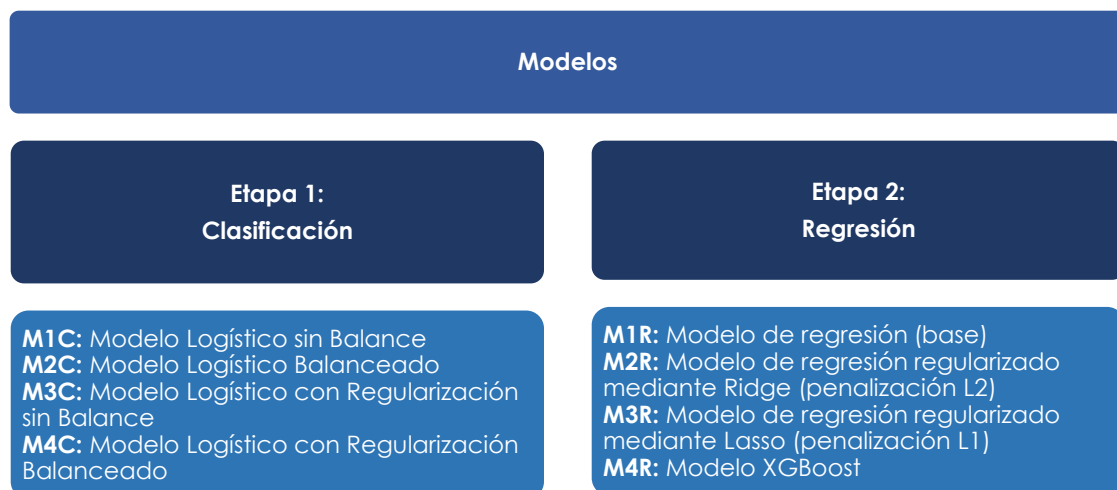
Tabla 2. Variables independientes

Variable	Descripción de variable
AREA	Área
P02	Sexo
P03	Edad
P04	Relación de Parentesco
P05A	Seguro Social - Alternativa 1

P06	Estado civil
P12A	Nivel de instrucción
P08	Como se considera
P15	Trabajó la semana anterior
INGPC	Ingreso Per cápita
POBREZA	Pobreza
CONDACTN	Condición de actividad Nueva
UT15	¿Cocina o prepara alimentos para el hogar?
UT33	¿Hace limpieza general de casa?
UT41	¿Lava ropa de miembros del hogar?
UT46	¿Realiza compras periódicas para el hogar?
UT61	¿Da de comer a niño/a menor de 12 años del hogar?
UT63	¿Juega, conversa con niño/a menor de 12 años del hogar?
UT77	¿Efectúa reparaciones en la vivienda del hogar?
UT98	¿Hace quehaceres domésticos gratuitos en otros hogares?
UT124	¿Da de comer a PCD del hogar?
UT125	¿Baña, asea, viste a PCD del hogar?
UT15_SEMT	Tiempo de cocinar o preparar alimentos para el hogar
UT33_SEMT	Tiempo de limpieza general de casa
UT41_SEMT	Tiempo en lavar ropa de miembros del hogar
UT46_SEMT	Tiempo en realizar compras periódicas para el hogar
UT61_SEMT	Tiempo en dar de comer a niño/a menor de 12 años del hogar
UT63_SEMT	Tiempo en jugar, conversar con niño/a menor de 12 años del hogar
UT77_SEMT	Tiempo en efectuar reparaciones en la vivienda del hogar
UT98_SEMT	Tiempo en quehaceres domésticos gratuitos en otros hogares
UT124_SEMT	Tiempo en dar de comer a PCD del hogar
UT125_SEMT	Tiempo en bañar, asear, vestir a PCD del hogar

Una vez definidas las variables dependientes e independientes, se especifica las etapas de los modelos implementados. El enfoque metodológico considera dos etapas, las cuales se observan en el siguiente esquema:

Figura 8. Esquema de modelado



4.3 Comparación y selección de modelos

Con el fin de comparar y seleccionar el mejor modelo según su desempeño predictivo, a continuación se muestran los resultados de las métricas de evaluación de los modelos de clasificación y regresión. Para simplificar la exposición del estudio, de forma

detallada se presentan las configuraciones de hiperparámetros¹¹ y métricas¹² correspondientes a la actividad **UT31: Tender la cama**, cuya selección responde a criterios expositivos y no implica representatividad respecto del conjunto de actividades modeladas.

Para el resto de variables (**UT47, UT58, UT100, UT126**) se expone de manera resumida el proceso de selección del mejor modelo al final de las respectivas secciones de clasificación y regresión, los resultados completos se pueden observar en el Anexo 8.2.

Modelos de clasificación (participación en la actividad UT31)

Para la predicción de la participación en actividades de trabajo no remunerado, se compara variantes del algoritmo de regresión logística según la naturaleza de la variable objetivo. Se estiman cuatro configuraciones del modelo:

- i. **(M1C):** Un modelo logístico base con desbalance de clases.
- ii. **(M2C):** Un modelo logístico con ponderación de clases.
- iii. **(M3C):** Un modelo logístico con regularización sin ponderación de clases incorporando una búsqueda del valor óptimo del parámetro de regularización C ¹³.
- iv. **(M4C):** Un modelo logístico con regularización con ponderación de clases incorporando una búsqueda del valor óptimo del parámetro de regularización C .

En los dos primeros casos se utiliza regularización L2 (**penalty= 'l2'**), el optimizador Newton-CG¹⁴ (**solver="newton-cg"**) y un máximo de 300 iteraciones (**max_iter = 300**), manteniendo un valor de **C= 1**. La diferencia entre ambos modelos radica en la ponderación de las clases: M1C se entrenó sin ponderación (**class_weight = None**), mientras que: M2C incorporó ponderación balanceada (**class_weight = 'balanced'**).

Para los modelos M3C y M4C se mantuvo la misma configuración, sin embargo, el parámetro de regularización C se optimizó mediante una búsqueda de malla (**Grid Search**) con validación cruzada.

Tabla 3. Hiperparámetros evaluados con grilla para los modelos M3C y M4C

Hiperparámetros	Valores evaluados
penalty	l2
solver	newton-cg
λ	np.logspace(-4, 4, 20)
C	$\frac{1}{\lambda}$
Cross validation(cv)	folds=5
max_iter	300

¹¹ Son variables de configuración que se establecen manualmente para gestionar el proceso de entrenamiento de un modelo de machine learning.

¹² Son medidas que permiten entender de forma objetiva qué tan bien está funcionando un modelo.

¹³ Hiperparámetro que controla la fuerza de la regularización, actuando como un balance entre la complejidad del modelo y su capacidad para cometer errores en los datos de entrenamiento.

¹⁴ Método que utiliza derivadas de segundo orden (Hessiana) para encontrar el óptimo en menos pasos.

El tuneo¹⁵ de los hiperparámetros junto con la grilla y la validación cruzada permitió controlar el sobreajuste debido a la temática analizada. En la tabla 4 se observa las métricas de evaluación en los datos de entrenamiento de los modelos según el sexo y tipo de modelo de clasificación.

Tabla 4. Métricas de evaluación en los datos de entrenamiento por modelo de clasificación según sexo – UT31

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,83	0,96	0,85	0,98	0,88	1,00
(M2C): Modelo Logístico Balanceado	0,89	0,98	0,81	0,87	0,74	0,78
(M3C): Modelo Logístico con Regularización sin Balance	0,83	0,96	0,85	0,98	0,88	1,00
(M4C): Modelo Logístico con Regularización Balanceado	0,89	0,98	0,81	0,87	0,74	0,78

En la tabla 5 se muestra las métricas de desempeño en los datos de prueba utilizadas para evaluar los modelos de clasificación. Se observa que los resultados obtenidos en el conjunto de entrenamiento son consistentes con los del conjunto de prueba, lo que muestra estabilidad y descarta sobreajuste de datos.

Tabla 5. Métricas de evaluación en los datos de test por modelo de clasificación según sexo – UT31

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,83	0,96	0,85	0,98	0,88	1,00
(M2C): Modelo Logístico Balanceado	0,89	0,98	0,81	0,87	0,74	0,78
(M3C): Modelo Logístico con Regularización sin Balance	0,83	0,96	0,85	0,98	0,88	1,00
(M4C): Modelo Logístico con Regularización Balanceado	0,88	0,98	0,80	0,87	0,74	0,78

El desempeño de los modelos se lo evalúa mediante métricas sensibles al desbalance de clases como el recall, F1-score y la precisión de la clase positiva, por lo tanto, se da prioridad a aquellas métricas que reflejan adecuadamente la capacidad del modelo para identificar la participación en actividades de trabajo no remunerado.

En esta evaluación se identifica que el modelo logístico balanceado y el modelo logístico con penalización L2 y ponderación de clase (class_weight = 'balanced') presentaron resultados similares. Para efectos de esta investigación, se considera el modelo con el mejor desempeño global tanto para hombres como para mujeres el **(M4C): Modelo logístico con regularización balanceado**.

En función de esta selección, la tabla 6 muestra las métricas analizadas para distintos umbrales de decisión, considerando un rango de 0,10 a 0,50 umbrales sobre las probabilidades predichas. Este análisis permite examinar el comportamiento del modelo

¹⁵ Búsqueda sistemática de la combinación de valores de un modelo de aprendizaje automático con el fin de maximizar el desempeño predictivo

frente a diferentes puntos de corte. El objetivo es mejorar la detección de individuos que realizan la actividad sin generar sobreestimación.

Tabla 6. Métricas de evaluación en los datos de test por modelo de clasificación según umbrales por sexo – UT31

Mejor modelo de clasificación: (M4C)	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
0,10	0,99	1,00	0,49	0,08	0,33	0,04
0,15	0,97	0,99	0,52	0,09	0,35	0,05
0,20	0,94	0,96	0,54	0,09	0,38	0,05
0,25	0,92	0,93	0,56	0,10	0,41	0,06
0,30	0,89	0,88	0,58	0,12	0,43	0,06
0,35	0,86	0,83	0,60	0,13	0,46	0,07
0,40	0,84	0,77	0,62	0,15	0,49	0,09
0,45	0,81	0,71	0,63	0,18	0,52	0,10
0,50	0,77	0,65	0,64	0,20	0,55	0,12

En ambos sexos, el mejor umbral es el de 0,30 dado que muestra mejor equilibrio entre el recall de la clase positiva y un F1-score aceptable, evitando la sobreestimación de la actividad. Por lo tanto, para ambos casos, **el umbral óptimo es 0,30** para la predicción final.

Selección del modelo de las preguntas UT47, UT58, UT100, UT126

Para estas variables se siguió el mismo procedimiento metodológico, donde se tunean los hiperparámetros mediante una búsqueda cruzada y la comparación de las distintas especificaciones evaluadas. La selección del modelo se realizó considerando conjuntamente las métricas de desempeño, la estabilidad entre conjuntos de entrenamiento y prueba y el objetivo de maximizar la capacidad predictiva para cada variable. En la tabla 7 se observa los modelos seleccionados y el umbral que maximiza el recall en clase 1 según cada pregunta.

Tabla 7. Selección del modelo y del umbral de decisión para las variables de clasificación según sexo

Pregunta	Sexo	Modelo seleccionado	Umbral
UT47	Hombre	M4C	0,30
	Mujer		
UT58	Hombre	M4C	0,35
	Mujer		
UT100	Hombre	M4C	0,10
	Mujer		
UT126	Hombre	M4C	0,20
	Mujer		

Modelos de regresión (tiempo semanal dedicado)

Para la estimación del tiempo semanal dedicado a las actividades de trabajo no remunerado, se emplean modelos de regresión, dado que la variable objetivo corresponde a una variable continua expresada en horas semanales. El análisis se realiza considerando únicamente a las personas que reportan participación en la actividad, con el fin de modelar de manera consistente la intensidad del tiempo dedicado.

Para el modelamiento se estiman cuatro modelos de regresión. Como referencia inicial se incluyó un **(M1R)**: Modelo de regresión lineal simple, el cual permitió evaluar la relación entre las variables explicativas y el tiempo semanal dedicado, sirviendo como punto de comparación frente a enfoques más robustos ante posibles problemas de multicolinealidad y sobreajuste. Los modelos estimados son variantes regularizadas del modelo de regresión lineal y un modelo XGBoost basado en aprendizaje automático, específicamente:

- i. **(M1R)**: Modelo de regresión lineal simple,
- ii. **(M2R)**: Modelo de regresión regularizado mediante Ridge (penalización L2) y
- iii. **(M3R)**: Modelo de regresión regularizado mediante Lasso (penalización L1)
- iv. **(M4R)**: Modelo XGBoost

En el caso del modelo **(M2R)** y **(M3R)** la selección del parámetro de regularización (α) se realiza mediante un procedimiento de búsqueda exhaustiva Grid Search con validación cruzada con **folds = 5**. Una vez identificado el valor óptimo de α , se estimó el modelo final de regresión Ridge y Lasso utilizando dicha configuración.

Tabla 8. Hiperparámetros evaluados con grilla para los modelos M2R y M3R

Hiperparámetros	Valores evaluados
estimator	Ridge, Lasso
λ	np.logspace(-4, 4, 10)
α	[0.0001, 0.001, 0.01, 0.1, 1, 20]
Cross validation(cv)	folds = 5
scoring	neg_mean_absolute_error (MAE)
return_train_score	True

Por otro lado, se estimó un modelo no lineal basado en XGBoost **(M4R)** para regresión. Este enfoque permite contrastar el desempeño de los modelos lineales frente a un método de ensamble basado en árboles de decisión, orientado a maximizar la capacidad predictiva.

Para el modelo XGBoost se especificó una función objetivo de regresión y se utilizó el método de construcción de árboles basado en histogramas. La selección del número óptimo de iteraciones se realiza mediante validación cruzada de cinco particiones utilizando la función **xgb.cv**¹⁶, aplicando un criterio de detención temprana. El modelo final se estima utilizando el número de iteraciones que minimiza el **RMSE promedio** en los conjuntos de validación.

Tabla 9. Hiperparámetros evaluados con grilla para los modelos M4R

Hiperparámetros	Valores evaluados
reg	squarederror
tree_method	hist
num_boost_round	100

¹⁶ La función xgb.cv() de la librería XGBoost implementa un procedimiento de validación cruzada para evaluar el desempeño del modelo en múltiples particiones de los datos. En este estudio se utilizó para determinar el número óptimo de iteraciones para minimizar el error cuadrático medio de la raíz promedio de los conjuntos de validación.

En la tabla 10, se puede observar las métricas obtenidas en el conjunto de entrenamiento de todos los modelos de regresión probados para la variable continua (ut31_semt).

Tabla 10. Métricas de evaluación en los datos de entrenamiento por modelo de regresión según sexo – UT31_semt

Modelo de regresión	R^2		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,07	0,07	0,48	0,89	0,69	0,94
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,07	0,07	0,48	0,89	0,69	0,94
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,06	0,07	0,48	0,89	0,69	0,95
(M4R): Modelo XGBoost	0,15	0,12	0,43	0,84	0,66	0,92

Tabla 11. Métricas de evaluación en la muestra test por modelo de regresión según sexo – UT31_semt

Modelo de regresión	R^2		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,04	0,07	0,47	0,99	0,69	0,99
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,04	0,07	0,47	0,99	0,68	0,99
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,04	0,07	0,47	0,99	0,68	0,99
(M4R): Modelo XGBoost	0,02	0,07	0,48	0,99	0,69	0,99

En la tabla 10 se puede verificar que las métricas son aproximadamente similares a las métricas de evaluación en los datos de prueba (tabla 11), lo que lleva a verificar estabilidad en los modelos y descartar sobreajuste en los datos.

Los resultados presentados en la tabla 11 muestran que los modelos de regresión lineal, Ridge y Lasso exhiben un desempeño prácticamente idéntico tanto para hombres como para mujeres. En el caso de los hombres, los tres modelos alcanzan un coeficiente de determinación (R^2) de 0,04, mientras que para las mujeres el valor es de 0,07. Estos resultados indican que las variables incluidas en el modelo explican únicamente entre el 4% y el 7 % de la variabilidad observada en el tiempo semanal dedicado a la actividad UT31.

En términos de error de predicción, los modelos lineales también presentan resultados muy similares. Para los hombres se registran valores de MSE entre 0,47 y 0,48 y RMSE entre 0,68 y 0,69, mientras que para las mujeres el MSE se sitúa en 0,99 y el RMSE en 0,99. La similitud de estas métricas muestra que la incorporación de regularización mediante Ridge (L2) y Lasso (L1) no genera mejoras sustanciales respecto al modelo de regresión lineal base.

Por otra parte, el modelo XGBoost tampoco evidencia mejoras significativas. En hombres, el R^2 disminuye ligeramente a 0,02 y los errores de predicción permanecen prácticamente inalterados (MSE = 0,48 y RMSE = 0,69). En mujeres, las métricas son idénticas a las obtenidas por los modelos lineales (R^2 = 0,07; MSE = 0,99; RMSE = 0,99).

En conjunto, los resultados muestran que las variables disponibles (predictores) poseen una capacidad limitada para explicar la variación en el tiempo dedicado a la actividad UT31. Asimismo, se observa que no existen mejoras relevantes al utilizar técnicas de regularización o modelos no lineales, lo cual indica que la información contenida en los predictores actuales resulta insuficiente para incrementar de manera sustancial la precisión de las estimaciones. Luego de analizar los indicadores de ajuste, para esta sección del estudio se selecciona el **(M4R): Modelo XGBoost** como la alternativa para la estimación del tiempo semanal asociado a la actividad UT31.

Si bien XGBoost constituye un algoritmo robusto, estable y altamente escalable, capaz de capturar relaciones no lineales y complejas entre las variables, en este estudio dichas ventajas no se tradujeron en mejoras sustanciales del desempeño predictivo. Esto muestra que la información contenida en los predictores disponibles presenta una capacidad explicativa limitada, por lo que el uso de modelos de mayor complejidad no aporta ganancias relevantes en precisión.

De hecho, las predicciones obtenidas mediante XGBoost y los modelos regularizados (Ridge y Lasso) resultan prácticamente equivalentes, evidenciando una convergencia hacia valores estimados similares (horas promedio estimadas). En consecuencia, pese a la selección del **XGBoost**, es necesario señalar que cualquiera de estos enfoques (**Ridge y Lasso**) constituye una alternativa metodológicamente válida para la estimación de las variables analizadas.

Selección del modelo de las preguntas UT47, UT58, UT100, UT126

En estas variables se aplicó el mismo procedimiento metodológico de búsqueda exhaustiva y validación cruzada con el tuneo de hiperparámetros antes señalado (tabla 9). La selección del modelo se realizó considerando también las métricas del R^2 , MSE y RMSE, así como la consistencia del desempeño entre las muestras de entrenamiento y prueba. Adicionalmente, cuando dos o más modelos presentaron resultados similares, la elección se fundamentó en criterios metodológicos relacionados con la interpretabilidad, estabilidad y capacidad de generalización.

Tabla 12. Selección de modelos de regresión para la estimación del tiempo semanal según variable y sexo.

Pregunta	Sexo	Modelo seleccionado	RMSE
UT47_semt	Hombre	M4R	0,75
	Mujer		0,79
UT58_semt	Hombre	M4R	0,82
	Mujer		0,62
UT100_semt	Hombre	M3R	2,74
	Mujer		2,64
UT126_semt	Hombre	M3R	3,24
	Mujer		4,30

En la tabla 12 se presenta los modelos de regresión finalmente seleccionados para cada variable, junto con los valores de RMSE obtenidos en los datos de prueba, utilizados como principal criterio para la selección; el detalle del resto de métricas de evaluación se presenta en anexos (8.2).

5 Resultados

Una vez identificadas las especificaciones con mejor desempeño según las métricas de evaluación consideradas, estas se adoptaron para la etapa de **predicción**. Los resultados obtenidos para las variables binarias de participación en tareas domésticas, de cuidado y comunitarias no remuneradas y de horas promedio semanal destinadas a cada una de estas tareas se muestran a continuación.

5.1 Participación en la actividad del trabajo no remunerado

La tabla 13 presenta la participación en trabajo no remunerado desagregada por sexo, considerando como referencia los valores observados en 2015 y 2017 y la estimación proyectada para 2023. En términos generales, los resultados proyectados se mantienen dentro de los rangos históricos observados, sin mostrar cambios estructurales significativos en la tendencia.

Tabla 13. Comparación entre valores históricos (2017, 2019) y predicción 2023, sobre la incidencia de personas que realizan TNR en 5 preguntas, desagregadas por sexo

Pregunta	2015		2017		2023	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
Realizó TNRH en UT31: Tender las camas	69%	95%	73%	96%	86%	95%
Realizó TNRH en UT47: Realizar compras diarias	49%	57%	46%	52%	63%	64%
Realizó TNRH en UT58: Participar en la seguridad de la vivienda	72%	58%	76%	63%	71%	68%
Realizó TNRH en UT100: Realizar servicio a la comunidad	2%	2%	2%	2%	18%	15%
Realizó TNRH en UT126: Ayudar en terapias y atención a una persona con discapacidad (PCD)	0,80%	1,50%	0,80%	1,60%	1,40%	1,20%

Fuente: INEC - ENEMDU - Módulo de uso de tiempo (2015-2017 y 2023 Imputada)

Para **UT31: Tender la cama**, el análisis de valores históricos frente al obtenido para el 2023 se muestra en mujeres un comportamiento estable, mientras que en los hombres se observa un incremento de participación. Sin embargo, se encuentra dentro de una tendencia marcada por los históricos.

Para **UT47: Realizar compras diarias**, si bien se observa una ligera disminución entre los años 2015 y 2017, para el año 2023 se evidencia una subida en la participación en ambos sexos. Dicho comportamiento puede interpretarse como ajuste en la población de acuerdo con los patrones observados.

Para **UT58: Participar en la seguridad de la vivienda**, en los hombres la incidencia proyectada se mantiene alineada con los niveles históricos mostrando estabilidad en el tiempo. En mujeres, se identifica una tendencia creciente, por lo tanto, se observa un patrón sostenido en la predicción del 2023.

Para **UT100: Realizar servicio a la comunidad**, esta actividad presenta una incidencia históricamente muy baja y estable, lo que implica una limitada variabilidad en los datos de entrenamiento. La estimación del 2023 muestra un incremento significativo en la participación, que no sigue el patrón observado en los años históricos.

Dicho comportamiento de la estimación muestra limitaciones en la capacidad predictiva del modelo para variables de baja prevalencia.

Para **UT126: Ayudar en terapias y atención a una persona con discapacidad (PCD)**, la baja frecuencia de ocurrencia limita la capacidad del modelo para capturar patrones robustos, lo que genera inestabilidad en las predicciones. Los resultados deben interpretarse con sensibilidad dado que poseen limitaciones.

En el caso de las actividades **UT100 y UT126**, se identifican particularidades asociadas a su baja prevalencia histórica. Estas características limitan la capacidad del modelo para capturar patrones estables, generando estimaciones más sensibles a variaciones en las probabilidades predichas. En consecuencia, los resultados para estas variables se interpretan considerando las restricciones poblacionales de los históricos.

La tabla 14 presenta la participación en trabajo no remunerado para las preguntas faltantes de la ENEMDU 2023, considerando como referencia los valores observados en 2015 y 2017 y la estimación proyectada para 2023.

Tabla 14. Comparación entre valores históricos (2015, 2017) y predicción 2023, sobre la participación según sexo en 5 tareas no remuneradas.

Pregunta	2015		2017		2023	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
UT31: Tender la cama	35%	65%	38%	62%	40%	60%
UT47: Realizar compras diarias	40%	60%	42%	58%	41%	59%
UT58: Participar en la seguridad de la vivienda	48%	52%	49%	51%	43%	57%
UT100: Realizar servicio a la comunidad	48%	52%	51%	49%	47%	53%
UT126: Ayudar en terapias y atención a una persona con discapacidad (PCD)	27%	73%	29%	71%	42%	58%

Fuente: INEC - ENEMDU - Módulo de uso de tiempo (2015-2017 y 2023 Imputada)

El caso de la UT126 se requiere un análisis particular, dado que en ambos sexos los resultados muestran mayor volatilidad y diferencias marcadas respecto a los años previos. La baja proporción de casos afirmativos en la UT126 evidencia un fuerte desbalance en la distribución de la variable, lo que reduce la cantidad de observaciones para entrenamiento y prueba, y en consecuencia, limita la potencia predictiva del modelo en la fase de estimación.

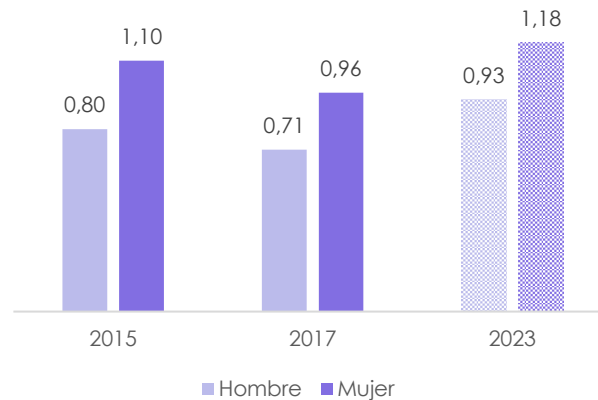
5.2 Tiempo semanal en las actividades del trabajo no remunerado

El objetivo final del estudio es predecir las horas promedio semanales dedicadas al trabajo no remunerado por sexo en las actividades TNR con datos ausentes. Los valores proyectados para 2023, en términos generales se mantienen cercanos a los observados en los años de referencia (2015 y 2017), lo que sugiere consistencia en la imputación realizada y estabilidad en los patrones de tiempo dedicados a estas actividades. Sin embargo, al analizar cada pregunta este patrón debe tratarse con precaución, especialmente para las preguntas de baja incidencia en TNR, como se muestra a continuación.

UT31: Tender la cama, en la figura 9, se observa una ligera variación en las horas promedio dedicadas por ambos sexos. Mientras que, en los años observados los

hombres reportaban entre 0,71 y 0,80 horas semanales y las mujeres entre 0,96 y 1,10 horas, la estimación para 2023 muestra un incremento moderado en ambos casos. Sin embargo, esta variación mantiene la brecha de género observada históricamente, donde las mujeres dedican mayor tiempo a esta actividad doméstica.

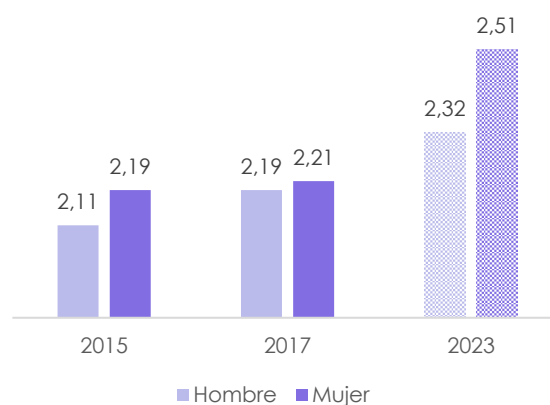
Figura 9. Horas promedio de trabajo no remunerado por sexo: valores observados (2015-2017) y predicción 2023. UT31 "Tender la cama"



Fuente: INEC - ENEMDU - Módulo de uso de tiempo – ut31 (2015-2017 y 2023 Imputada)

Para la actividad **UT47: Realizar compras diarias**, en la figura 10 se observa una evolución gradual en el tiempo de dedicación a esta actividad. En el caso de los hombres se observa un crecimiento entre años; sin embargo, para el año 2023, la predicción arroja un crecimiento de aproximadamente 6% respecto al año anterior observado (2017). De igual manera para las mujeres, se muestra en crecimiento del 14% respecto al año anterior observado (2017). En términos generales, los resultados evidencian que la estimación mantiene los patrones de comportamiento diferenciados históricamente. No obstante, las cifras muestran un incremento relativamente mayor en las horas promedio semanales estimadas, por lo que, su implementación en la construcción de las Cuentas Satélite de Trabajo No Remunerado de los Hogares dependerá del análisis final.

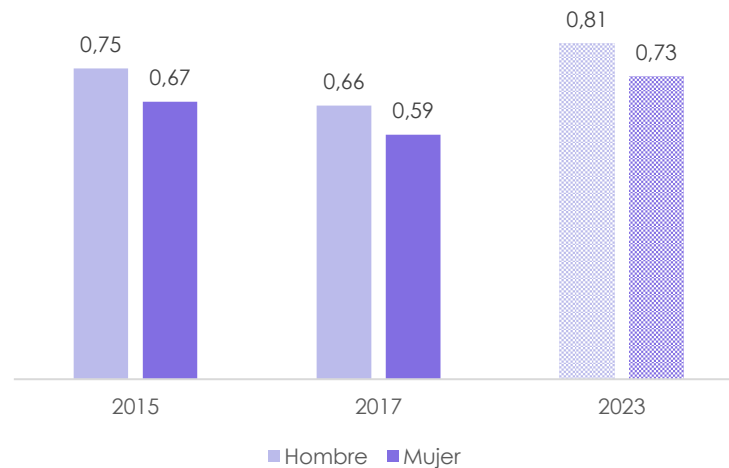
Figura 10. Horas promedio de trabajo no remunerado por sexo: valores observados (2015-2017) y predicción 2023. UT47 "Realizar compras diarias"



Fuente: INEC - ENEMDU - Módulo de uso de tiempo – ut47 (2015-2017 y 2023 Imputada)

En **UT58: Participar en la seguridad de la vivienda**, la figura 11 muestra que la predicción mantiene niveles similares a los registrados en los años de referencia, lo que indica que el modelo en general capta los patrones. Sin embargo, es importante mencionar que las variaciones son evaluadas en etapas posteriores de la construcción de las Cuentas Satélite de Trabajo No Remunerado de los Hogares, con el fin de verificar que no generen distorsiones en la serie.

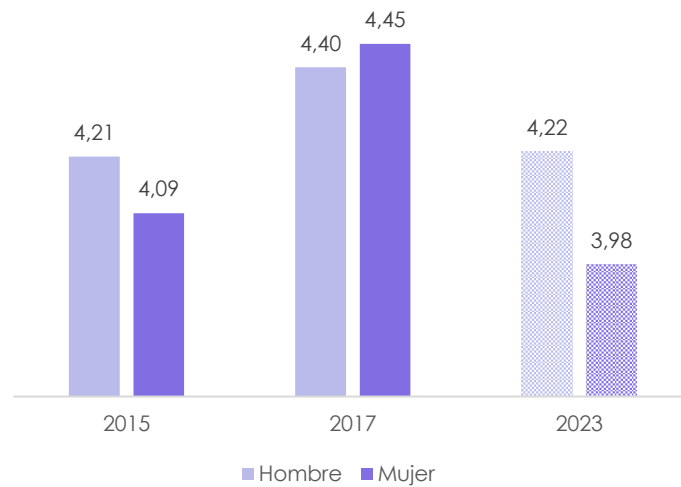
Figura 11. Horas promedio de trabajo no remunerado por sexo: valores observados (2015-2017) y predicción 2023. UT58 "Participar en la seguridad de la vivienda"



Fuente: INEC - ENEMDU - Módulo de uso de tiempo – ut58 (2015-2017 y 2023 Imputada)

Respecto a **UT100: Realizar servicio a la comunidad**, en la figura 12 se observa que las horas promedio estimadas presentan variaciones moderadas entre los años analizados. Entre 2015 y 2017 se observa un incremento en el tiempo promedio dedicado a esta actividad tanto en hombres como en mujeres, pasando de 4,21 a 4,40 horas en el caso de los hombres y de 4,09 a 4,45 horas promedio en mujeres. Para 2023, las estimaciones muestran una ligera reducción en ambos grupos. Estas variaciones se mantienen dentro de rangos coherentes con los patrones observados; sin embargo, por los porcentajes de participación de la misma, el proceso de construcción de las CSTNRH define su aplicación.

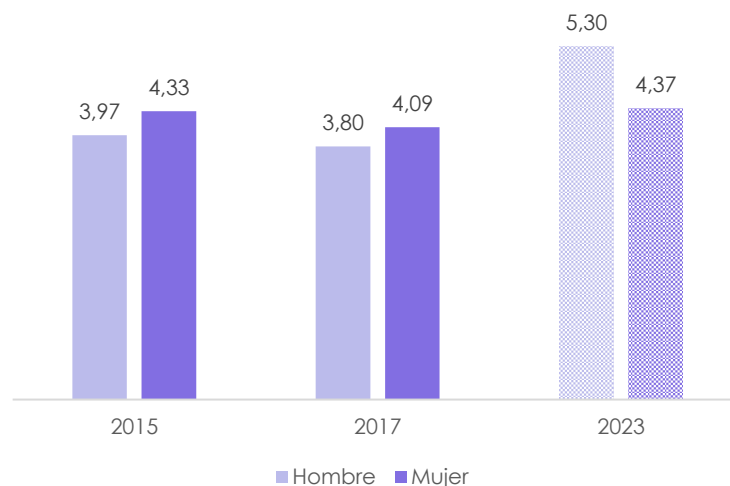
Figura 12. Horas promedio de trabajo no remunerado por sexo: valores observados (2015-2017) y predicción 2023. UT100 "Realizar servicio a la comunidad"



Fuente: INEC - ENEMDU - Módulo de uso de tiempo – ut100 (2015-2017 y 2023 Imputada)

Finalmente, en **UT126: Ayudar en terapias y atención a una persona con discapacidad**, en la figura 13 se identifican variaciones marcadas en comparación con las otras actividades predichas. Este comportamiento puede estar asociado a la baja incidencia de esta actividad en la población, lo que genera mayor sensibilidad en los promedios estimados, por lo tanto, su utilización final de la construcción de las CSTNRH podría no ser consistente.

Figura 13. Horas promedio de trabajo no remunerado por sexo: valores observados (2015-2017) y predicción 2023. UT126 "Ayudar en terapias y atención a una persona con discapacidad"



Fuente: INEC - ENEMDU - Módulo de uso de tiempo – ut126 (2015-2017 y 2023 Imputada)

En conjunto, los resultados muestran que las estimaciones preservan los patrones estructurales del trabajo no remunerado de los hogares observados en los años de referencia, particularmente en lo relacionado con la distribución por sexo y la magnitud relativa de las horas promedio dedicadas a cada actividad. No obstante, algunas actividades presentan variaciones, por lo tanto, cambian la consistencia con la

tendencia histórica observada. En este contexto, únicamente las estimaciones que mantienen coherencia con la evolución histórica tanto en tiempo como en participación se considerarán como insumo para la construcción de las CSTNRH, mientras que las que evidencian volatilidad se excluyen de dicho proceso.

6 Conclusiones

El presente estudio ha demostrado la viabilidad y pertinencia de aplicar modelos de aprendizaje automático para la imputación de información faltante en el módulo de uso del tiempo de la ENEMDU 2023, específicamente para las variables relacionadas con el trabajo no remunerado de los hogares. A partir de los resultados obtenidos, se derivan las siguientes conclusiones.

Los resultados de la predicción muestran un desempeño diferenciado según el tipo de actividad. Para las actividades domésticas regulares: **UT31 Tender las camas; UT46: Compras a diario; y UT58: Vigilar la seguridad del hogar**; las imputaciones generadas presentan una consistencia sólida con las tendencias históricas observadas en los años de referencia (2015 y 2017). Las tasas de participación y las horas semanales promedio estimadas se mantienen dentro de rangos esperados, lo que respalda su utilización para la continuidad de la serie estadística de las CSTNRH. Este hallazgo muestra que, para actividades con frecuencias relativamente altas y patrones estables en el tiempo, los modelos ML constituyen una herramienta confiable

La predicción de resultados en las preguntas: **UT100 Participación en actividades de servicio comunitario; y UT126 Practicar terapia a persona con discapacidad**, muestra un comportamiento muy volátil y diferente entre hombres y mujeres, reflejando inestabilidad en el tiempo, por lo que su aplicación no debería utilizarse bajo el método de predicción aquí propuesto. Otra especificación y supuestos más generalizados podrían adecuarse más en contextos de preguntas con bajas tasas de participación y de datos altamente desbalanceados.

Para **UT31 “Tender la cama”** el umbral de decisión afecta la precisión de la predicción, sin embargo, esta ayuda a mejorar significativamente el recall de la clase 1. El recall dada la estructura de los datos, se constituye como la métrica de evaluación de mayor relevancia. Las métricas de clasificación con un umbral seleccionado, tanto en hombres como en mujeres (0,30) da como resultado una precisión inferior al umbral por defecto (0,50), dado que el incremento de falsos positivos minimiza la precisión, sin embargo, mejora la detección de verdaderos positivos (recall – clase 1).

El desempeño de los modelos mostró diferencias entre hombres y mujeres por los patrones inmersos que existen en el trabajo no remunerado de los hogares, sin embargo, el análisis se enfocó en obtener mayor recall pese a sacrificar las demás métricas. Por lo tanto, para hombres se obtuvo un recall de 89% para la clase 1 y en mujeres un 88% para la clase 1 (realizan la actividad de tender la cama como trabajo no remunerado del hogar).

Por otro lado, en la variable continua específicamente en el tiempo promedio dedicado a tender la cama, la selección del modelo se enfocó en minimizar el error de predicción en la estimación de horas, por lo tanto, el modelo seleccionado para hombre tuvo un RMSE = 0,69 horas y para mujeres RMSE = 0,99 horas.

En conjunto, los hallazgos muestran que algoritmos flexibles como la regresión logística balanceada con regularización y XGBoost presentan un buen desempeño predictivo, en variables dicotómicas, por lo tanto, estos modelos alcanzaron mejores resultados dada la naturaleza de la problemática.

Los resultados imputados demuestran que existe una persistente brecha de género en el trabajo no remunerado, junto con una tendencia leve al alza en la participación masculina y la configuración de un patrón más complejo en la distribución del tiempo destinado a estas actividades, probablemente condicionado por el efecto de la pandemia de COVID-19.

6.1 Líneas futuras de investigación

Los resultados obtenidos evidencian que la metodología propuesta constituye una alternativa viable para estimar actividades de trabajo no remunerado cuando existe información faltante. Sin embargo, para futuras investigaciones se recomienda profundizar el análisis de aquellas actividades que presentaron mayor inestabilidad predictiva, particularmente el servicio comunitario (UT100) y el cuidado de personas con discapacidad (UT126) y no fueron utilizadas con insumo para las CSTNRH.

En este sentido, resulta pertinente evaluar la sensibilidad de los modelos ante la exploración de diferentes configuraciones metodológicas e incorporar técnicas alternativas de aprendizaje automático, y balanceo de clases. Todo lo anterior mencionado permitirá mejorar la capacidad predictiva en actividades caracterizadas por una baja frecuencia de ocurrencia y una alta variabilidad en la población de estudio.

Asimismo, futuras investigaciones podrían extender el análisis a nuevos períodos de observación, incorporando variables contextuales y socioeconómicas adicionales que contribuyan a explicar de mejor manera los patrones del trabajo no remunerado en el Ecuador.

7 Referencias

- Armas, A., Contreras, J., & Vásquez, A. (2009). *La economía del cuidado, el trabajo no remunerado y remunerado en Ecuador*. Comisión de Transición Instituto Nacional de Estadísticas y Censos Fondo de Desarrollo de las Naciones Unidas para la Mujer. <https://repositorio.iaen.edu.ec/handle/24000/4332>
- Asaduzzaman, A., Uddin, M. R., Woldeyes, Y., & Sibai, F. N. (2024). A Novel Salary Prediction System Using Machine Learning Techniques. 2024 Joint International Conference on Digital Arts, Media and Technology with ECTI Northern Section Conference on Electrical, Electronics, Computer and Telecommunications Engineering (ECTI DAMT & NCON), Chiang-mai, Thailand, 2024, pp. 38-43, DOI: 10.1109/ECTIDAMTNCN60518.2024.10480058.
- Benería, L. (1999). El debate inconcluso sobre el trabajo no remunerado. *Revista Internacional del Trabajo*, 118(3), 321-346. <https://doi.org/10.1111/j.1564-913X.1999.tb00136.x>
- Campo Yepes, John Jairo. (2026). Evaluación de métricas en modelos predictivos de clasificación en machine learning [EAN Universidad]. <https://repository.universidadean.edu.co/server/api/core/bitstreams/1579de58-ccb3-4c27-adab-54455d2c503f/content>
- Cessie, S. L., & Houwelingen, J. C. V. (1992). Ridge Estimators in Logistic Regression. *Applied Statistics*, 41(1), 191. <https://doi.org/10.2307/2347628>
- Chen, M. (2024). Desmitificando el aprendizaje automático: Una guía completa. <https://www.oracle.com/latam/artificial-intelligence/machine-learning/what-is-machine-learning/>
- 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*, 785-794. <https://doi.org/10.1145/2939672.2939785>
- Di Franco, G., & Santurro, M. (2021). Machine learning, artificial neural networks and social research. *Quality & Quantity*, 55(3), 1007-1025. <https://doi.org/10.1007/s11135-020-01037-y>
- Domínguez, M. (2021). *Regresión Logística y Técnicas de Aprendizaje. Aplicaciones*. Universidad Zaragoza.

- Enríquez, C. R. (s. f.). El trabajo de cuidado no remunerado en Argentina: Un análisis desde la evidencia del Módulo de Trabajo no Remunerado.
- Esquivel, V. (2010). Trabajadores del cuidado en la Argentina. En el cruce entre el orden laboral y los servicios de cuidado. *Revista Internacional Del Trabajo*, 129(4), 529-547. <https://doi.org/10.1111/j.1564-9148.2010.00099.x>
- Esteban Madrigal. (2022, octubre 28). Conoce las métricas de precisión más comunes para Modelos de Regresión. Conoce las métricas de precisión más comunes para Modelos de Regresión. <https://www.growupcr.com/post/metricas-precision>
- Fawcett, T. (2006). An introduction to ROC analysis. *Pattern Recognition Letters*, 27(8), 861-874. <https://doi.org/10.1016/j.patrec.2005.10.010>
- Folbre, N. (2006). Measuring Care: Gender, Empowerment, and the Care Economy. *Journal of Human Development*, 7(2), 183-199. <https://doi.org/10.1080/14649880600768512>
- Gao, X., Fan, S., Li, X., Guo, Z., Zhang, H., Peng, Y., & Diao, X. (2017). An improved XGBoost based on weighted column subsampling for object classification. 2017 4th International Conference on Systems and Informatics (ICSAI), 1557-1562. <https://doi.org/10.1109/ICSAI.2017.8248532>
- García, M. (2024). Machine Learning avanzado con algoritmos híbridos. Universitat D'Alacant.
- Hilloulin, B., & Umunnakwe, R. (2024). Machine learning-aided prediction of shrinkage in modern concrete: Focus on mix proportions and SCMs. *Journal of Building Engineering*, 98, 111410. <https://doi.org/10.1016/j.jobbe.2024.111410>
- Hosmer, D., & Lemeshow, S. (2012). Applied Logistic Regression. En *Applied Logistic Regression* (2da edición, pp. 1-56). A Wiley-Interscience Publication.
- INEC. (2023, mayo). Imputación de información en encuestas de hogares a través de técnicas de Machine Learning: Análisis del caso ecuatoriano. https://www.ecuadorencifras.gob.ec/documentos/web-inec/Bibliotecas/Libros/cuadernos_trabajo/Imputacion_ENEMDU_ML.pdf
- INEC. (2025). Principales Resultados de Estadísticas de Cuentas Satélite de Trabajo No Remunerado de los Hogares. <https://www.ecuadorencifras.gob.ec/cuenta-satelite-del-trabajo-no-remunerado/>

- ManiKumar, A. N. V., Manohar, T. A., Akhil, D., Manideep, C., & Rao, A. V. S. (2024). Predicting Salary of an Employee using Machine Learning. *International Journal of Scientific Development and Research*, 9(4), 137-140.
- Mulder, J., Hocuk, S., Kilic, T., Zezza, A., & Kumar, P. (2024). Making Time Count: A Machine Learning Approach to Predict Time Use in Low-Income Countries from Physical Activity Tracking Data. *Policy Research Working Paper Series*, Policy Research Working Paper Series, Article 10835. <https://ideas.repec.org//p/wbk/wbrwps/10835.html>
- Nikunj Vaghasiya. (2024, noviembre 22). Comprender las semillas en la IA: la clave para la reproducibilidad y la creatividad. [Medium]. Comprender Las Semillas En La IA: La Clave Para La Reproducibilidad y La Creatividad. <https://medium.com/@nikunj.vaghasiya2050/understanding-seeds-in-ai-the-key-to-reproducibility-and-creativity-edcfd3bf649c>
- Qutub, A., Al-Mehmadi, A., Al-Hssan, M., & Alghamdi, H. S. (2021). Prediction of Employee Attrition Using Machine Learning and Ensemble Methods. *International Journal of Machine Learning and Computing*, 11(2), 110-114. <https://doi.org/10.18178/ijmlc.2021.11.2.1022>
- Regresión lineal en el aprendizaje automático. (s. f.). [Educativa]. Geeks for Geeks. Recuperado 4 de febrero de 2026, de <https://www.geeksforgeeks.org/>
- Rosati, G. (2021). Métodos de Machine Learning como alternativa para la imputación de datos perdidos. Un ejercicio en base a la Encuesta Permanente de Hogares. *Estudios del trabajo*, (61), 122-145.
- Roustaei, N. (2024). Application and interpretation of linear-regression analysis. *Medical Hypothesis, Discovery & Innovation Ophthalmology Journal*, 13(3), 151-159. <https://doi.org/10.51329/mehdiophthal1506>
- Tibshirani, R. (1996). Regression Shrinkage and Selection Via the Lasso. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 58(1), 267-288. <https://doi.org/10.1111/j.2517-6161.1996.tb02080.x>
- Tripathi, S. (2024). Predicting Unpaid Care Work in India Using Random Forest: An Analysis of Socioeconomic and Demographic Factors. *Social Policy and Society*, 1-17. <https://doi.org/10.1017/S1474746424000447>

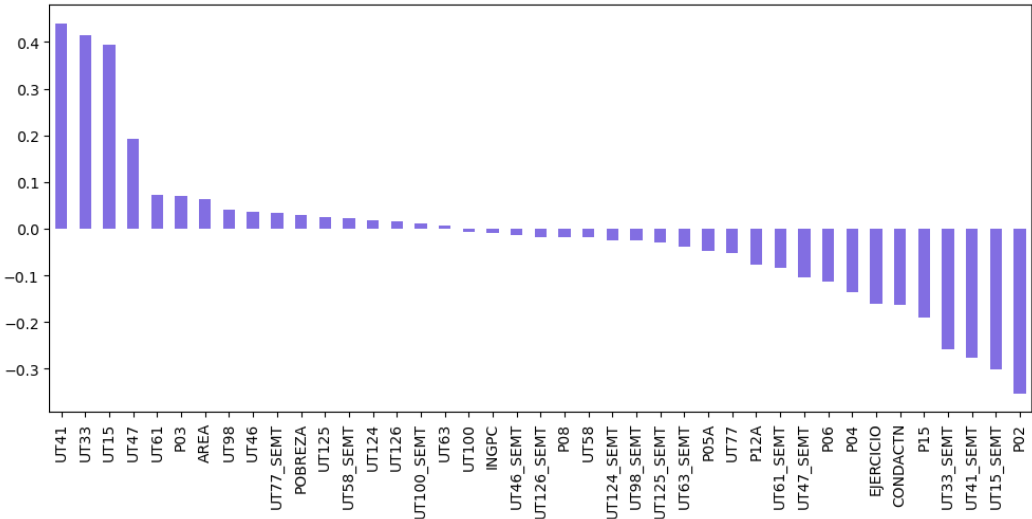
- Wang, Z., Akande, O., Poulos, J., & Li, F. (2022). Are deep learning models superior for missing data imputation in large surveys? Evidence from an empirical comparison. arXiv. <https://doi.org/10.48550/arXiv.2103.09316>
- Wei, J., Chu, X., Sun, X., Xu, K., Deng, H., Chen, J., Wei, Z., & Lei, M. (2019). Machine learning in materials science. *InfoMat*, 1 (3), 338-358. <https://doi.org/10.1002/inf2.12028>
- XGBoost para regresión. (s. f.). [Educativa]. Geeks for Geeks. Recuperado <https://www.geeksforgeeks.org/machine-learning/xgboost-for-regression/>

8 Anexos

8.1 Matrices de correlación

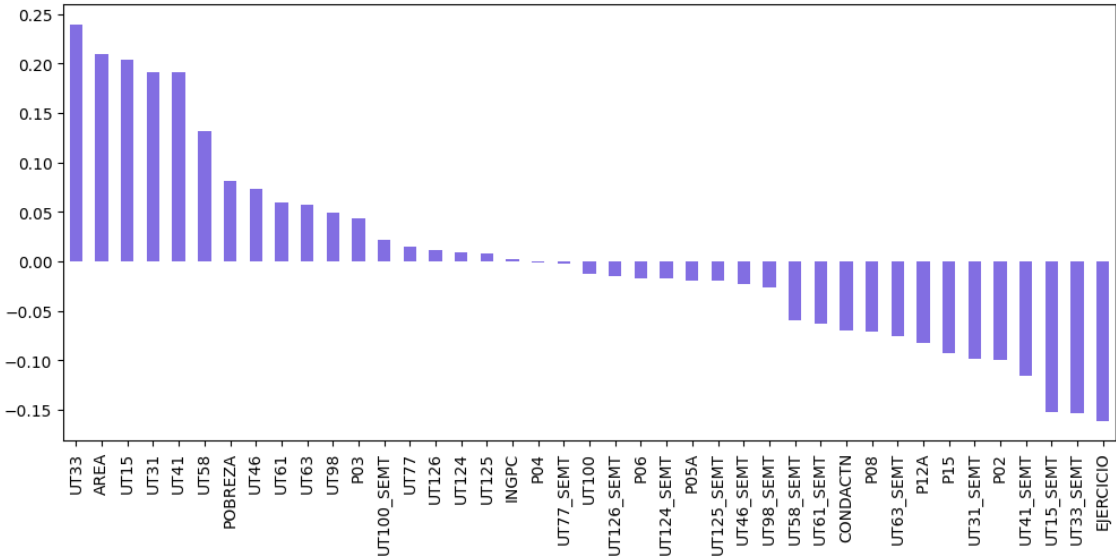
- **UT31:** Tender las camas

Figura 14. Correlaciones entre variable dependiente e independientes – UT31



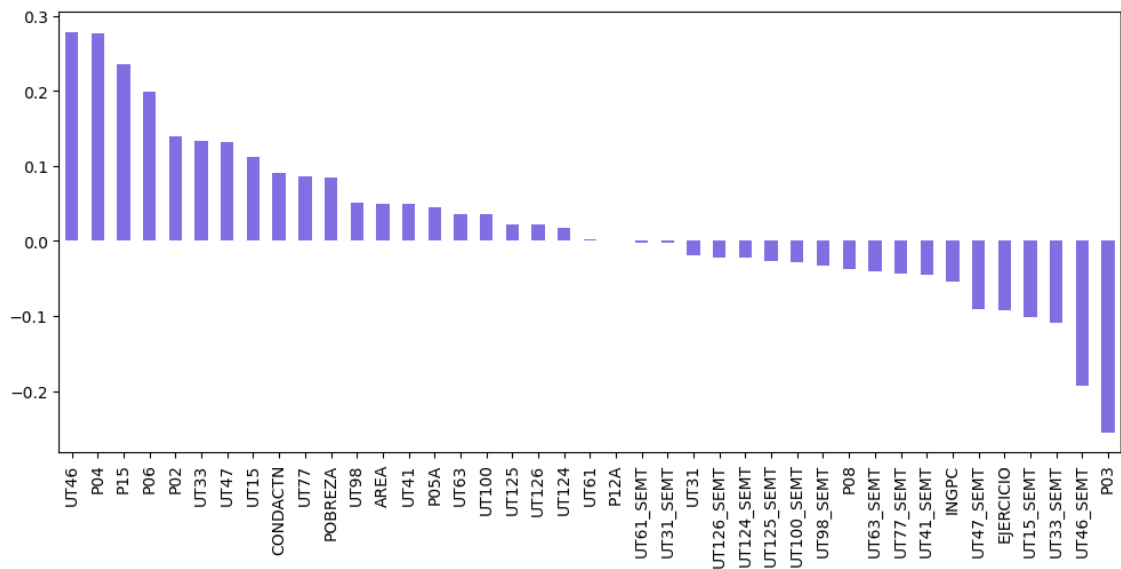
- **UT47:** Compras diarias

Figura 15. Correlaciones entre variable dependiente e independientes – UT47



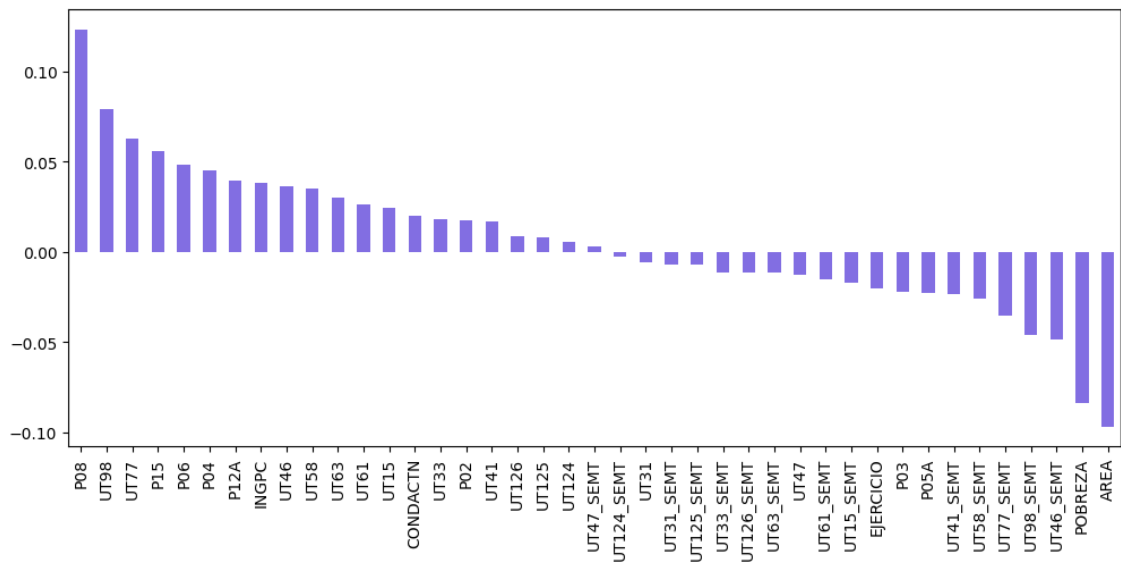
- **UT58:** Vigilar la seguridad del hogar

Figura 16. Correlaciones entre variable dependiente e independientes – UT58



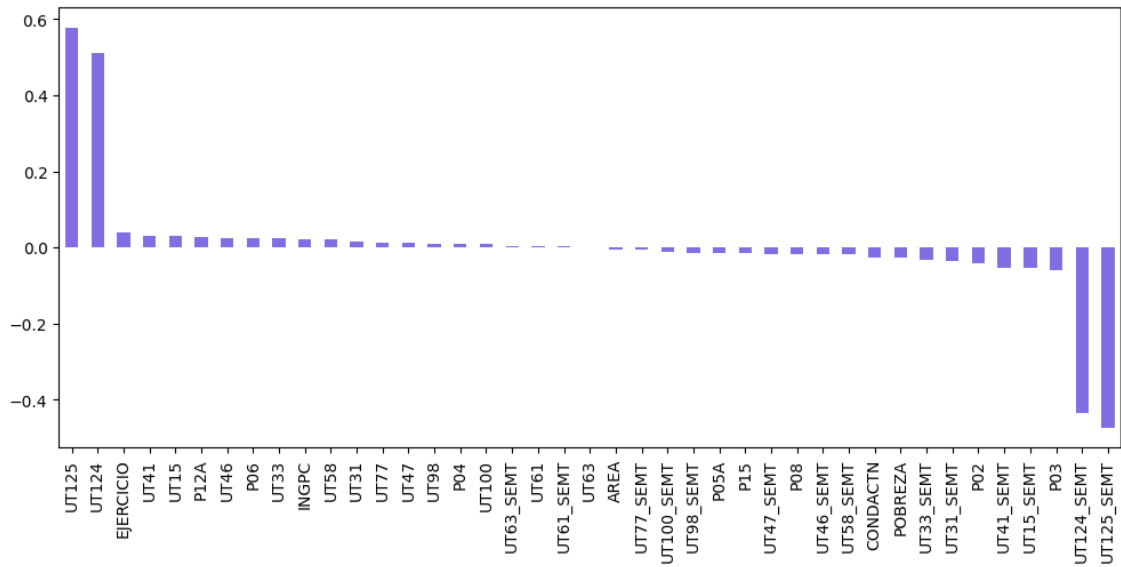
- **UT100:** Participar en actividades sociales

Figura 17. Correlaciones entre variable dependiente e independientes – UT100



- **UT126:** Practicar alguna terapia especial a PCD

Figura 18. Correlaciones entre variable dependiente e independientes – UT126



8.2 Métricas de evaluación del modelo

- **UT47:** Compras diarias

Modelos de clasificación

Tabla 13. Métricas de evaluación en los datos de entrenamiento por modelo de clasificación según sexo - UT47

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,62	0,65	0,60	0,69	0,57	0,74
(M2C): Modelo Logístico Balanceado	0,60	0,67	0,62	0,67	0,64	0,67
(M3C): Modelo Logístico con Regularización sin Balance	0,62	0,65	0,60	0,69	0,57	0,74
(M4C): Modelo Logístico con Regularización Balanceado	0,60	0,67	0,62	0,67	0,64	0,67

Tabla 14. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según sexo – UT47

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,62	0,65	0,59	0,69	0,56	0,73
(M2C): Modelo Logístico Balanceado	0,61	0,67	0,62	0,67	0,63	0,66
(M3C): Modelo Logístico con Regularización sin Balance	0,62	0,65	0,59	0,69	0,56	0,73
(M4C): Modelo Logístico con Regularización Balanceado	0,60	0,67	0,62	0,67	0,63	0,66

Tabla 15. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según umbrales por sexo – UT47

Modelo de clasificación: Modelo Logístico con Regularización Balanceado	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
0,10	1,00	1,00	0,69	0,62	0,52	0,45
0,15	1,00	1,00	0,69	0,62	0,52	0,45
0,20	1,00	1,00	0,69	0,62	0,53	0,45
0,25	0,98	0,99	0,69	0,63	0,53	0,46
0,30	0,95	0,95	0,70	0,63	0,55	0,48
0,35	0,89	0,88	0,70	0,64	0,57	0,51
0,40	0,82	0,79	0,69	0,64	0,59	0,54
0,45	0,73	0,70	0,67	0,63	0,62	0,57
0,50	0,62	0,61	0,64	0,60	0,65	0,60

Modelos de regresión

Tabla 16. Métricas de evaluación en los datos de entrenamiento por modelo de regresión según sexo - UT47_semt

Modelo de regresión	R ²		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,05	0,09	0,61	0,67	0,78	0,82
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,05	0,09	0,61	0,67	0,78	0,82
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,05	0,09	0,61	0,67	0,78	0,82
(M4R): Modelo XGBoost	-0,15	0,20	0,74	0,59	0,86	0,77

Tabla 17. Métricas de evaluación en los datos de prueba (test) por modelo de regresión según sexo – UT47_semt

Modelo de regresión	R^2		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,06	0,08	0,45	0,65	0,67	0,80
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,06	0,08	0,45	0,65	0,67	0,80
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,06	0,08	0,45	0,64	0,67	0,80
(M4R): Modelo XGBoost	-0,18	0,10	0,57	0,63	0,75	0,79

Nota: Si bien se seleccionó XGBoost como modelo de referencia, el modelo regularizado Lasso (l1) obtuvo métricas de desempeño comparables y valores estimados que convergen hacia resultados similares respecto a las horas promedio estimadas. En futuras actualizaciones se considerará su reevaluación como alternativa del modelo.

- **UT58:** Vigilar la seguridad del hogar

Modelos de clasificación

Tabla 18. Métricas de evaluación en los datos de entrenamiento por modelo de regresión según sexo - UT58

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,84	0,72	0,87	0,78	0,90	0,84
(M2C): Modelo Logístico Balanceado	0,90	0,76	0,82	0,74	0,76	0,71
(M3C): Modelo Logístico con Regularización sin Balance	0,83	0,72	0,87	0,78	0,92	0,85
(M4C): Modelo Logístico con Regularización Balanceado	0,90	0,76	0,82	0,74	0,76	0,73

Tabla 19. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según sexo - UT58

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,83	0,72	0,86	0,78	0,90	0,84
(M2C): Modelo Logístico Balanceado	0,90	0,76	0,82	0,74	0,75	0,72
(M3C): Modelo Logístico con Regularización sin Balance	0,83	0,72	0,87	0,78	0,91	0,85
(M4C): Modelo Logístico con Regularización Balanceado	0,89	0,76	0,82	0,74	0,76	0,73

Tabla 20. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según umbrales por sexo – UT58

Modelo de clasificación: Modelo Logístico con Regularización Balanceado	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
0,10	1,00	1,00	0,43	0,57	0,28	0,39
0,15	0,98	1,00	0,47	0,57	0,31	0,40
0,20	0,95	0,99	0,50	0,58	0,34	0,41
0,25	0,91	0,97	0,53	0,60	0,38	0,43
0,30	0,87	0,94	0,56	0,61	0,41	0,45
0,35	0,84	0,88	0,58	0,62	0,45	0,48
0,40	0,81	0,80	0,60	0,63	0,48	0,52
0,45	0,78	0,71	0,62	0,63	0,51	0,56
0,50	0,75	0,64	0,62	0,62	0,53	0,60

Modelos de regresión

Tabla 21. Métricas de evaluación en los datos de entrenamiento por modelo de regresión según sexo – UT58_semt

Modelo de regresión	R ²		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,03	0,03	0,47	0,42	0,69	0,65
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,03	0,03	0,47	0,42	0,69	0,65
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,03	0,03	0,47	0,42	0,69	0,65
(M4R): Modelo XGBoost	0,09	0,04	0,44	0,42	0,66	0,65

Tabla 22. Métricas de evaluación en los datos de prueba (test) por modelo de regresión según sexo – UT58_semt

Modelo de regresión	R ²		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,02	0,02	0,66	0,38	0,81	0,62
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,02	0,02	0,66	0,38	0,81	0,62
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,02	0,02	0,66	0,38	0,81	0,62
(M4R): Modelo XGBoost	0,01	0,01	0,66	0,39	0,82	0,62

Nota: Si bien se seleccionó XGBoost como modelo de referencia, el modelo regularizado Lasso (l1) obtuvo métricas de desempeño comparables y valores estimados que convergen hacia resultados similares respecto a las horas promedio estimadas. En futuras actualizaciones se considerará su reevaluación como alternativa del modelo.

- **UT100:** Participar en actividades sociales

Modelos de clasificación

Tabla 23. Métricas de evaluación en los datos de entrenamiento por modelo de clasificación según sexo – UT100

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,29	1,00	0,01	0,00	0,00	0,00
(M2C): Modelo Logístico Balanceado	0,07	0,06	0,12	0,11	0,74	0,76
(M3C): Modelo Logístico con Regularización sin Balance	-	-	-	-	-	-
(M4C): Modelo Logístico con Regularización Balanceado	0,07	0,06	0,12	0,11	0,74	0,76

Nota: Las métricas del modelo M3C no se reportan debido a que el modelo no generó predicciones para la clase positiva, como consecuencia del desbalance de clases, impidiendo una evaluación representativa mediante las métricas de clasificación consideradas.

Tabla 24. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según sexo – UT100

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	1,00	1,00	0,01	0,02	0,01	0,01
(M2C): Modelo Logístico Balanceado	0,07	0,05	0,12	0,10	0,77	0,68
(M3C): Modelo Logístico con Regularización sin Balance	-	-	-	-	-	-
(M4C): Modelo Logístico con Regularización Balanceado	0,07	0,05	0,12	0,09	0,77	0,68

Nota: Las métricas del modelo M3C no se reportan debido a que el modelo no generó predicciones para la clase positiva, como consecuencia del desbalance de clases, impidiendo una evaluación representativa mediante las métricas de clasificación consideradas.

Tabla 25. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según umbrales por sexo – UT100

Modelo de clasificación: Modelo Logístico con Regularización Balanceado	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
0,10	0,98	0,98	0,98	0,98	0,98	0,99
0,15	0,96	0,96	0,97	0,97	0,99	0,99
0,20	0,94	0,94	0,96	0,96	0,99	0,99
0,25	0,92	0,93	0,95	0,96	0,99	0,99
0,30	0,90	0,91	0,94	0,95	0,99	0,99
0,35	0,88	0,88	0,93	0,93	0,99	0,99
0,40	0,85	0,86	0,92	0,92	0,99	0,99
0,45	0,82	0,83	0,90	0,90	0,99	0,99
0,50	0,78	0,80	0,87	0,88	0,99	0,99

Modelos de regresión

Tabla 26. Métricas de evaluación en los datos de entrenamiento por modelo de regresión según sexo – UT100_semt

Modelo de regresión	R^2		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,29	0,23	5,63	4,92	2,37	2,22
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,29	0,23	5,64	4,93	2,37	2,22
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,15	0,21	6,74	5,05	2,60	2,25
(M4R): Modelo XGBoost	-1,76	0,22	21,73	4,94	4,66	2,22

Tabla 27. Métricas de evaluación en los datos de prueba (test) por modelo de regresión según sexo – UT100_semt

Modelo de regresión	R^2		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,06	0,05	7,74	7,55	2,78	2,75
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,06	0,06	7,71	7,46	2,78	2,73
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,09	0,12	7,51	6,99	2,74	2,64
(M4R): Modelo XGBoost	-1,95	0,05	24,27	7,56	4,93	2,75

- **UT126:** Practicar alguna terapia especial a PCD

Modelos de clasificación

Tabla 28. Métricas de evaluación en los datos de entrenamiento por modelo de clasificación según sexo – UT126

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,62	0,63	0,40	0,57	0,29	0,52
(M2C): Modelo Logístico Balanceado	0,06	0,29	0,11	0,43	0,73	0,83
(M3C): Modelo Logístico con Regularización sin Balance	0,60	0,65	0,27	0,53	0,17	0,45

(M4C): Modelo Logístico con Regularización Balanceado	0,21	0,49	0,31	0,59	0,60	0,75
--	------	------	------	------	------	------

Tabla 29. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según sexo – UT126

Modelo de clasificación	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1C): Modelo Logístico sin Balance	0,66	0,60	0,39	0,49	0,28	0,41
(M2C): Modelo Logístico Balanceado	0,06	0,28	0,11	0,42	0,67	0,79
(M3C): Modelo Logístico con Regularización sin Balance	0,71	0,60	0,29	0,45	0,18	0,36
(M4C): Modelo Logístico con Regularización Balanceado	0,27	0,46	0,38	0,55	0,63	0,69

Tabla 30. Métricas de evaluación en los datos de prueba (test) por modelo de clasificación según umbrales por sexo – UT126

Modelo de clasificación: Modelo Logístico con Regularización Balanceado	Recall (clase 1)		F1-Score		Precisión	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
0,10	1,00	1,00	1,00	0,99	0,99	0,99
0,15	1,00	1,00	1,00	0,99	0,99	0,99
0,20	1,00	0,99	1,00	0,99	0,99	0,99
0,25	1,00	0,99	1,00	0,99	1,00	0,99
0,30	1,00	0,99	1,00	0,99	1,00	0,99
0,35	1,00	0,99	1,00	0,99	1,00	0,99
0,40	0,99	0,99	1,00	0,99	1,00	0,99
0,45	0,99	0,99	0,99	0,99	1,00	1,00
0,50	0,99	0,99	0,99	0,99	1,00	1,00

Modelos de regresión

Tabla 31. Métricas de evaluación en los datos de entrenamiento por modelo de regresión según sexo – UT126_semt

Modelo de regresión	R^2		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	0,40	0,16	8,82	13,03	2,97	3,61
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	0,39	0,16	8,98	13,04	3,00	3,61
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,21	0,00	11,56	15,58	3,40	3,95
(M4R): Modelo XGBoost	-0,87	-0,86	27,41	28,94	5,24	5,38

Tabla 32. Métricas de evaluación en los datos de prueba (test) por modelo de regresión según sexo – UT126_semt

Modelo de regresión	R^2		MSE		RMSE	
	Hombre	Mujer	Hombre	Mujer	Hombre	Mujer
(M1R): Modelo de regresión (base)	-0,11	-0,02	12,57	18,80	3,55	4,34
(M2R): Modelo de regresión regularizado mediante Ridge (L2)	-0,07	-0,01	12,16	18,66	3,49	4,32
(M3R): Modelo de regresión regularizado mediante Lasso (L1)	0,07	-0,00	10,51	18,48	3,24	4,30
(M4R): Modelo XGBoost	-0,75	-0,79	19,85	33,07	4,46	5,75



@ecuadorencifras



@ecuadorencifras



@InecEcuador



t.me/equadorencifras



INEC/Ecuador



INECEcuador

Administración Central (Quito)
Juan Larrea N15-36 y José Riofrio,
Teléfonos: (02) 2544 326 - 2544 561 Fax: (02) 2509 836
Código postal: 170410
correo-e: inec@inec.gob.ec

www.ecuadorencifras.gob.ec