



Modelos predictivos de satisfacción en los programas de tratamiento de adicciones (2022-2024)

Autor: Adrián Fernández Figueroa

Tutor: Luis Ángel Calvo Pascual

MADRID | JUNIO 2026

RESUMEN

Este Trabajo de Fin de Grado tiene como objetivo desarrollar modelos predictivos capaces de clasificar el nivel de satisfacción de los usuarios de los Centros de Atención a las Drogodependencias (CAD) e identificar los factores que más influyen en dicha satisfacción.

Para ello, se ha utilizado un conjunto de datos procedente de las encuestas de satisfacción realizadas durante los años 2022, 2023 y 2024. Tras un proceso de preparación, limpieza y transformación, se han desarrollado distintos modelos supervisados de Machine Learning, incluyendo entre ellos, regresión logística, árboles de decisión, Random Forest, XGBoost y redes neuronales artificiales.

Los resultados obtenidos muestran que los modelos basados en ensamblados de árboles presentan el mejor rendimiento predictivo. En concreto, el modelo Random Forest optimizado fue seleccionado como mejor modelo final al obtener el mejor equilibrio entre capacidad predictiva, estabilidad e interpretabilidad.

Asimismo, los análisis de interpretabilidad han permitido identificar diversos factores clave asociados con la satisfacción de los usuarios, como el trato recibido por parte del personal, la coordinación entre profesionales, la rapidez de respuesta ante situaciones urgentes, la disponibilidad de opciones terapéuticas y la ausencia de experiencias de discriminación.

Estos factores han destacado de forma consistente en los distintos modelos predictivos y técnicas aplicadas, lo que sugiere que constituyen elementos muy relevantes en la percepción global de la calidad del servicio por parte de los usuarios.

Los resultados evidencian el potencial que poseen las técnicas de Business Analytics y de Machine Learning para apoyar la toma de decisiones y contribuir a la mejora de la calidad de los servicios ofrecidos por los Centros de Atención a las Drogodependencias (CAD).

Palabras Clave: Business Analytics, Machine Learning, satisfacción de usuarios, analítica predictiva, Random Forest, interpretabilidad de modelos, gestión pública.

ABSTRACT

This Bachelor's Degree Final Project aims to develop predictive models capable of classifying the level of satisfaction of users of the Drug Addiction Care Centres (CAD) and identifying the factors that most influence this satisfaction.

To this end, a dataset obtained from satisfaction surveys conducted during the years 2022, 2023, and 2024 was used. After a process of preparation, cleaning, and transformation, different supervised Machine Learning models were developed, including logistic regression, decision trees, Random Forest, XGBoost, and artificial neural networks.

The results obtained show that tree-based ensemble models deliver the best predictive performance. Specifically, the optimised Random Forest model was selected as the best final model, as it achieved the best balance between predictive capacity, stability and interpretability.

Furthermore, the interpretability analyses made it possible to identify several key factors associated with user satisfaction, such as the treatment received from staff, coordination among professionals, response speed in urgent situations, the availability of therapeutic options and the absence of experiences of discrimination.

These factors have consistently stood out across the different predictive models and techniques applied, suggesting that they constitute highly relevant elements in users' overall perception of service quality.

The results demonstrate the potential of Business Analytics and Machine Learning techniques to support decision-making and contribute to improving the quality of the services offered by Drug Addiction Care Centres (CAD).

Keywords: Business Analytics, Machine Learning, user satisfaction, predictive analytics, Random Forest, model interpretability, public management.

ÍNDICE

INTRODUCCIÓN	7
1.1 CONTEXTO	7
1.2 OBJETIVO DEL TFG	9
1.3 METODOLOGÍA Y ANTECEDENTES	10
1.4 ESTRUCTURA DEL TFG	13
 CAPÍTULO I: LIMPIEZA Y PREPARACIÓN DE LOS DATOS	14
2.1 INTEGRACIÓN Y ESTRUCTURAS DE LAS BASES DE DATOS	14
2.2 LIMPIEZA Y PREPROCESAMIENTO	17
2.3 CONSTRUCCIÓN DE VARIABLES OBJETIVO	23
2.4 ANÁLISIS PRELIMINAR DE RELEVANCIA DE VARIABLES	25
 CAPÍTULO II: MODELOS PREDICTIVOS SUPERVISADOS	31
3.1 PLANTEAMIENTO DEL PROBLEMA	31
3.2 MODELOS PREDICTIVOS BASE	33
3.3 MODELO AVANZADOS DE MACHINE LEARNING	37
3.4 OPTIMIZACIÓN Y VALIDACIÓN DE MODELOS	41
3.5 INTERPRETABILIDAD Y EXPLICACIÓN DE RESULTADOS	48
3.6 COMPARACIÓN INTERANUAL DE MODELOS	55
3.7 FAST INTERPRETABLE GREEDY-TREE SUMS (FIGS)	59
 CAPÍTULO III: DISCUSIÓN DE RESULTADOS	63
4.1 PRINCIPALES HALLAZGOS DEL ESTUDIO	63
4.2 FACTORES DETERMINANTES DE LA SATISFACCIÓN DE LOS USUARIOS	65
4.3 IMPLICACIONES PARA LA GESTIÓN DE LOS CAD	66
4.4 LIMITACIONES DEL ESTUDIO	67
 CONCLUSIONES Y RECOMENDACIONES	69
 FUTURAS LÍNEAS DE INVESTIGACIÓN	71

ÍNDICE DE ILUSTRACIONES

IMAGEN 1. VARIABLES CON MAYOR PORCENTAJE DE VALORES FALTANTES	17
IMAGEN 2. PORCENTAJE DE VALORES FALTANTES POR AÑO	18
IMAGEN 3. DISTRIBUCIÓN DE FRECUENCIAS DE LA VARIABLE OBJETIVO (SATISFACCIÓN GENERAL CON EL SERVICIO DEL CAD)	24
IMAGEN 4. DISTRIBUCIÓN FINAL DE LA VARIABLE OBJETIVO	25
IMAGEN 5. VARIABLES CATEGÓRICAS MÁS ASOCIADAS CON EL NIVEL DE SATISFACCIÓN	26
IMAGEN 6. VARIABLES NUMÉRICAS MÁS ASOCIADAS CON EL NIVEL DE SATISFACCIÓN	27
IMAGEN 7. TIEMPO TOTAL EN TRATAMIENTO SEGÚN EL NIVEL DE SATISFACCIÓN	28
IMAGEN 8. MATRIZ DE CORRELACIONES DE PEARSON CON LAS 5 VARIABLES MÁS CORRELACIONADAS	29
IMAGEN 9. VARIABLES CON MAYOR RELEVANCIA CON RESPECTO A LA VARIABLE OBJETIVO	30
IMAGEN 10. DISTRIBUCIÓN DE CLASES ANTES Y DESPUÉS DE APLICAR SMOTE SOBRE EL SUBCONJUNTO DE ENTRENAMIENTO (TRAIN SET)	32
IMAGEN 11. MATRIZ DE CONFUSIÓN DEL MODELO DE REGRESIÓN LOGÍSTICA MULTINOMIAL	33
IMAGEN 12. VARIABLES CON MAYOR IMPORTANCIA DEL MODELO DE REGRESIÓN LOGÍSTICA	34
IMAGEN 13. MATRIZ DE CONFUSIÓN DEL MODELO DE ÁRBOL DE DECISIÓN	35
IMAGEN 14. NIVELES DEL ÁRBOL DE DECISIÓN E INTERPRETACIÓN DE REGLAS PRINCIPALES	36
IMAGEN 15. COMPARACIÓN DEL RENDIMIENTO DE LOS MODELOS PREDICTIVOS BASE	36
IMAGEN 16. MATRIZ DE CONFUSIÓN DEL MODELO DE RANDOM FOREST	37
IMAGEN 17. VARIABLES MÁS IMPORTANTES EN EL MODELO DE RANDOM FOREST	38
IMAGEN 18. MATRIZ DE CONFUSIÓN DEL MODELO XGBOOST	39
IMAGEN 19. VARIABLES MÁS IMPORTANTES EN EL MODELO XGBOOST	39
IMAGEN 20. MATRIZ DE CONFUSIÓN DEL MODELO DE REDES NEURONALES (MLP)	40
IMAGEN 21. COMPARACIÓN DE LAS MÉTRICAS ENTRE MODELOS AVANZADOS	41
IMAGEN 22. COMPARACIÓN DE MODELOS MEDIANTE VALIDACIÓN CRUZADA	42

IMAGEN 23. MATRIZ DE CONFUSIÓN DEL MODELO RANDOM FOREST OPTIMIZADO	43
IMAGEN 24. COMPARACIÓN DEL F1-SCORE ANTES Y DESPUÉS DE LA OPTIMIZACIÓN	44
IMAGEN 25. COMPARACIÓN GLOBAL DEL RENDIMIENTO PREDICTIVO DE LOS MODELOS EVALUADOS MEDIANTE F1-SCORE	45
IMAGEN 26. COMPARACIÓN GLOBAL DE LOS MODELOS PREDICTIVOS MEDIANTE ROC-AUC	46
IMAGEN 27. CURVAS ROC DE LOS MODELOS RANDOM FOREST Y XGBOOST OPTIMIZADOS	47
IMAGEN 28. VARIABLES MÁS RELEVANTES DEL MODELO RANDOM FOREST OPTIMIZADO	49
IMAGEN 29. VARIABLES MÁS RELEVANTES DEL MODELO XGBOOST OPTIMIZADO	50
IMAGEN 30. RANKING CONJUNTO DE VARIABLES MÁS RELEVANTES DE AMBOS MODELOS	50
IMAGEN 31. SHAP GRÁFICO RESUMEN PARA LA CLASE DE ALTA SATISFACCIÓN	52
IMAGEN 32. INFLUENCIA DE LAS OPCIONES TERAPÉUTICAS SUFICIENTES	53
IMAGEN 33. INFLUENCIA DE LA COORDINACIÓN ENTRE PROFESIONALES	53
IMAGEN 34. INFLUENCIA DE EXPERIENCIAS DISCRIMINATORIAS EN EL CAD	53
IMAGEN 35. EXPLICACIÓN INDIVIDUAL DE UNA PREDICCIÓN MEDIANTE SHAP	54
IMAGEN 36. REGLA DE ASOCIACIÓN IDENTIFICADA MÁS RELEVANTE DEL ANÁLISIS	55
IMAGEN 37. RENDIMIENTO PREDICTIVO POR AÑO Y POR MODELO	56
IMAGEN 38. COMPARACIÓN INTERANUAL DEL F1-SCORE DE LOS MODELOS RANDOM FOREST Y XGBOOST	57
IMAGEN 39. COMPARACIÓN INTERANUAL DEL ROC-AUC DE LOS MODELOS RANDOM FOREST Y XGBOOST	57
IMAGEN 40. VARIABLES COMUNES ENTRE 2022, 2023 Y 2024 EN RANDOM FOREST	58
IMAGEN 41. VARIABLES COMUNES ENTRE 2023 Y 2024 EN XGBOOST	58
IMAGEN 42. MATRIZ DE CONFUSIÓN DEL MODELO FIGS	60
IMAGEN 43. VARIABLES MÁS RELEVANTES IDENTIFICADAS POR FIGS	61
IMAGEN 44. COMPARACIÓN DE RENDIMIENTO ENTRE FIGS, RANDOM FOREST Y XGBOOST	61

INTRODUCCIÓN

1.1 Contexto

La evaluación de la **calidad** de los servicios sanitarios y sociosanitarios se ha convertido en una de las principales preocupaciones de las administraciones públicas (Sitzia & Wood, 1997; Mira & Aranaz, 2000; Saturno et al., 1990) y organismos encargados de prestar este tipo de servicios.

En una situación en la que los **recursos** disponibles son limitados y las necesidades de las personas son cada vez más complejas (Donabedian, 1988; Simsekler et al., 2021), resulta imprescindible conocer el grado de satisfacción de las personas y sus factores más relevantes para identificar las áreas de mejora y orientar el camino hacia una mejor toma de decisiones basada en evidencia.

En este contexto, la satisfacción de los usuarios es uno de los indicadores más relevantes, ya que permite valorar no solo la eficacia de los tratamientos, sino también la **accesibilidad**, atención recibida, capacidad de respuesta y, sobre todo, factores relacionales, que influyen en la percepción de la calidad de los servicios (Donabedian, 1988; Parasuraman et al., 1988; Mira Solves et al., 1998).

Esta evaluación es especialmente relevante en los Centros de Atención a las **Drogodependencias** (CAD), donde la intervención de especialistas suele durar largos periodos de tiempo (Delegación del Gobierno para el Plan Nacional sobre Drogas, 2018; Ministerio de Sanidad, 2022; OEDA, 2023) y requiere distintos perfiles.

En ellos, la experiencia del usuario está condicionada por múltiples variables más allá del tratamiento clínico, como los socioeconómicos, organizativos y **relacionales**.

Tradicionalmente, se han empleado análisis descriptivos que permiten identificar tendencias generales o agrupar pacientes en distintos grupos de usuarios, aunque estas aproximaciones presentan **limitaciones** para explicar qué factores determinan realmente el nivel de satisfacción o anticipar comportamientos futuros.

En los últimos años, el desarrollo de nuevas técnicas de **Machine Learning** ha abierto nuevas posibilidades al permitir identificar patrones complejos en grandes volúmenes de datos, y ser capaces de generar modelos predictivos con altos niveles de precisión (Hastie et al., 2009; Breiman, 2001).

Asimismo, su aplicación en el contexto sanitario y sociosanitario ha crecido notablemente abarcando usos como el apoyo al diagnóstico, la estimación de riesgos clínicos, la personalización de tratamientos y la optimización de **recursos** (Topol, 2019; Obermeyer & Emanuel, 2016).

Estas técnicas pueden emplearse perfectamente para analizar la experiencia de los usuarios y detectar aquellos factores que poseen una mayor influencia sobre el grado de **satisfacción** de los pacientes.

En la analítica avanzada, la combinación de los **modelos predictivos** con los interpretativos resulta especialmente valiosa para transformar resultados en recomendaciones útiles para los gestores.

Por ello, este TFG combina técnicas de Business Analytics y Machine Learning para analizar la satisfacción de los usuarios de los CAD a partir de microdatos de las encuestas de 2022, 2023 y 2024, incorporando técnicas de interpretación como el análisis de importancia de variables, los valores **SHAP** y las reglas de asociación.

En definitiva, el trabajo busca demostrar que dichas herramientas de análisis de datos y de interpretación de modelos, pueden ayudar a los responsables de **gestión** de los centros a identificar los factores más influyentes en la satisfacción de los pacientes. De esta manera, se generaría información útil que posteriormente puedan orientar para la toma de decisiones y así, realizar un proceso de mejora continua en los Centros de Atención a las Drogodependencias (CAD).

1.2 Objetivo del TFG

El objetivo principal de este Trabajo de Fin de Grado es desarrollar **modelos predictivos** capaces de clasificar el nivel de satisfacción de los usuarios de los CAD, y evaluar e interpretar los resultados obtenidos a partir de esos modelos.

Para ello, se va a utilizar técnicas de **Machine Learning** y aplicarlas a los datos procedentes de las encuestas de satisfacción correspondientes a los años 2022, 2023 y 2024.

Además de la construcción de **modelos predictivos**, otro de los objetivos es determinar cuáles son los factores que poseen una mayor influencia sobre la satisfacción de los usuarios, para contribuir a la mejora de la calidad de los servicios del centro.

Para alcanzar estos objetivos generales, se plantean los siguientes objetivos específicos:

- Identificar los factores que presentan una mayor influencia sobre la **satisfacción** de los usuarios de los CAD.
- Construir un modelo predictivo capaz de clasificar el nivel de satisfacción de los **usuarios** a partir de la información recogida en las encuestas.
- Comparar el **rendimiento** de diferentes enfoques de aprendizaje automático para determinar cuáles ofrecen una mayor capacidad predictiva (Fawcett, 2006; Powers, 2011).
- Analizar la **estabilidad temporal** de los factores asociados a la satisfacción entre los años 2022 y 2024.
- Obtener conocimiento útil para apoyar la toma de **decisiones** y proponer actuaciones orientadas a mejorar la calidad de los servicios prestados.

En definitiva, el trabajo intenta demostrar el potencial de las técnicas de **Business Analytics** para transformar datos en outputs muy valiosos que permiten ayudar a los responsables de los centros a la toma de decisiones, mediante la aportación de evidencias objetivas.

1.3 Metodología y Antecedentes

1.3.1 Metodología

Este TFG se enmarca en el ámbito del Business Analytics, combinando técnicas estadísticas, análisis de datos y modelización predictiva para extraer, de grandes volúmenes de datos, información útil para la **gestión** y planificación estratégica de los centros.

Se han utilizado **microdatos** de encuestas de satisfacción de los CAD de 2022, 2023 y 2024, que incluyen numerosas dimensiones, entre ellas, la experiencia de los usuarios, la valoración de servicios y diferentes factores relacionales y de atención recibida.

El desarrollo se realizó en Jupyter Notebook (v 6.4.8) dentro del entorno de Anaconda con **Python 3.9.12**. Las librerías principales fueron Pandas y NumPy para el tratamiento de datos; Scikit-Learn, **XGBoost** e Imbalanced-Learn para la modelización predictiva; Matplotlib y Seaborn para la visualización; y SHAP e IModels para la interpretabilidad avanzada.

En primer lugar, se ha realizado un proceso de depuración, tratamiento y preparación de los datos, incluyendo el tratamiento de valores faltantes, transformación de variables, y la construcción de la **variable objetivo**.

En segundo lugar, se ha llevado a cabo una exploración inicial de los datos para comprender las características de los usuarios e identificar las variables más asociadas con la **satisfacción**.

Posteriormente, se han desarrollado y comparado distintos modelos **supervisados** de clasificación. Entre ellos, la regresión logística, el árbol de decisión, el Random Forest, el XGBoost y las redes neuronales.

Una vez entrenados, los modelos se han evaluado mediante el Accuracy, F1-score y ROC-AUC, junto con procedimientos de **validación cruzada** y optimización de

hiperparámetros (Stone, 1974; Bergstra & Bengio, 2012; Arlot & Celisse, 2010) para mejorar los modelos obtenidos.

Finalmente, se han empleado técnicas de interpretabilidad para comprender qué factores eran los más influyentes en los mejores modelos. Entre ellas, se encuentran el análisis de importancia de variables, **SHAP** y las reglas de asociación.

En definitiva, esta metodología combina la capacidad predictiva de los modelos **Machine Learning** con un enfoque más interpretativo orientado a la gestión de estos servicios, facilitando la identificación de los factores clave asociados con la satisfacción de los usuarios de los CAD.

1.3.2 Antecedentes

En los últimos años, el incremento en la disponibilidad de datos y el avance de las tecnologías **analíticas** han impulsado al desarrollo de nuevas técnicas capaces de extraer conclusiones a partir de grandes volúmenes de datos.

En este contexto, surge la analítica predictiva, rama del **Business Analytics**, que utiliza datos históricos para predecir comportamientos futuros (Davenport & Harris, 2007; Provost & Fawcett, 2013) y estimar la probabilidad de que ocurran determinados eventos.

La analítica predictiva se apoya cada vez más en el **Machine Learning**, al permitir identificar patrones complejos y generar predicciones sin la necesidad de establecer previamente alguna regla de decisión concreta.

Por su eficacia y buen funcionamiento, estas técnicas se han convertido en tan comunes y ampliamente utilizadas en ámbitos como las finanzas, el marketing, la logística y la **gestión sanitaria**.

Dentro del **Machine Learning**, el aprendizaje supervisado trabaja con datos etiquetados para predecir una variable concreta (Vapnik, 1995; James et al., 2021), mientras que el

aprendizaje no supervisado intenta descubrir agrupaciones, relaciones ocultas e incluso estructuras dentro de los datos.

Este trabajo se centra en el **aprendizaje supervisado**, dado que el objetivo es predecir el nivel de satisfacción de los usuarios a partir de las encuestas, encuadrándose en un problema de clasificación multiclase, donde cada usuario es asignado a uno de los niveles de satisfacción disponibles.

La aplicación de los modelos predictivos en el ámbito sanitario ha incrementado notablemente. Diversos estudios han querido demostrar que estas herramientas son muy útiles para apoyar en los diagnósticos, empleando modelos de **regresión logística**, árboles de decisión y redes neuronales (Acion et al., 2017; Chung & Teo, 2023), además de, estimar riesgos clínicos, mejorar la calidad de los servicios e incluso optimizar la asignación de recursos.

También, ha crecido el interés por analizar la experiencia y satisfacción de los pacientes (Zhang et al., 2021; Simsekler et al., 2021), aplicando **Random Forest** junto con técnicas de interpretabilidad para identificar los factores clave de la satisfacción del paciente en las etapas de registro y de consulta hospitalaria.

Otros estudios proponen marcos de **Machine Learning** más interpretables (Liu et al., 2021), que no solo predicen la satisfacción, sino que razonan los factores más influyentes de forma comprensible para los gestores sanitarios.

En esta línea, han surgido modelos altamente interpretables, como el Fast Interpretable Greedy-Tree Sums (**FIGS**), que combinan la capacidad predictiva de los árboles de decisión con reglas de clasificación transparentes. y fácilmente comprensibles (Tan et al., 2025; Rudin, 2019).

Así, la combinación del Machine Learning con herramientas de **interpretabilidad** resulta especialmente interesante para las organizaciones sanitarias (Simsekler et al., 2021; Lundberg & Lee, 2017), ya que permite predecir los niveles de satisfacción con precisión, y a la vez, comprender qué factores explican dichas predicciones.

A su vez, resulta esencial para las organizaciones utilizar esos resultados con el fin de diseñar estrategias para mejorar los **recursos** y servicios de los centros.

Por tanto, este trabajo combina analítica predictiva, aprendizaje automático e interpretación de resultados, aplicando técnicas de análisis de datos avanzadas a un problema de especial relevancia social.

1.4 Estructura del TFG

Este Trabajo de Fin de Grado se estructura en tres **capítulos** principales, además de las secciones de conclusiones, recomendaciones, futuras líneas de investigación y bibliografía.

El Capítulo I recoge la preparación y limpieza de los datos. Se han llevado a cabo las siguientes tareas: integración y **homogeneización** de las bases de 2022, 2023 y 2024, **limpieza** y transformación de variables, además de un análisis exploratorio previo.

El Capítulo II recoge el desarrollo empírico del trabajo. Se construyen los distintos modelos predictivos, el proceso de **validación cruzada** y la **optimización** de hiperparámetros, junto con la aplicación de técnicas de interpretabilidad.

El Capítulo III está dedicado a la discusión de resultados, analizando los principales hallazgos, interpretando los factores más influyentes, y planteando las **implicaciones** para la gestión de los Centros de Atención a las Drogodependencias (CAD).

Por último, el trabajo finaliza con las conclusiones, una serie de recomendaciones derivadas de los resultados y, diversas propuestas para investigaciones futuras.

CAPÍTULO I: LIMPIEZA Y PREPARACIÓN DE LOS DATOS

2.1 Integración y estructura de las bases de datos

2.1.1 Fuentes de información utilizadas

Para el desarrollo del trabajo se han empleado seis ficheros CSV. Tres con las respuestas de los **usuarios** de los programas de adicciones de 2022, 2023 y 2024, y otros tres con el diseño del cuestionario de cada año, actuando de diccionario de variables.

Los cuestionarios de 2023 y 2024 presentan prácticamente el mismo diseño, mientras que el de 2022, requiere de una mejor **armonización** y **homogeneización**.

Además, al no disponer de un identificador longitudinal que permita seguir a los mismos **usuarios** entre años, se plantea el análisis como una comparación entre tres muestras independientes.

2.1.2 Particularidades estructurales de la base de datos 2022

La base de 2022 presenta diferencias estructurales relevantes respecto a las bases de 2023 y 2024. La **nomenclatura** de variables sigue un formato basado en códigos “Q” (“Q01”, “Q02”, “Q13” o “Q15_Q1501”) frente a la nomenclatura de los años posteriores (“CAD”, “SEXO”, “EDAD”, “P13” o “P15_1”, respectivamente).

Además, la base de 2022 incluye la variable “COEF_POND_CAD” (coeficiente de ponderación por CAD) como primer campo, mientras que en 2023 aparece “ID”, y en 2024 “REGISTRO”.

También, las preguntas de **respuesta múltiple** que en 2023 y 2024 aparecen desagregadas en columnas, en 2022 aparecen recogidas en una sola columna. Por ello, la base de 2022 requiere transformar su estructura y homogeneizar nombres y escalas de variables.

2.1.3 Homogeneización de nombres de variables

Identificadas las diferencias estructurales, ha sido necesario un proceso de **homogeneización** de algunos nombres de variables, ya que, aunque los cuestionarios de

2023 y 2024 son muy similares, existían algunas columnas con una denominación diferentes en ambos años.

La base de 2022 ha requerido una **armonización** más amplia, mediante un diccionario de correspondencias con la nomenclatura de 2023 y 2024. De esta manera, variables como Q01, Q02, Q03, Q13 o Q63, han sido renombradas como CAD, SEXO, EDAD, Satisfacción general con el CAD y Recomendación del CAD, respectivamente

En el caso no disponer de una equivalencia exacta, como la Q40 de 2022, que indica el lugar de recogida de **metadona**, se le asignó un nombre específico como es P39_LUGAR_METADONA_2022, para no distorsionar las variables equivalentes de 2023 y 2024, que recoge la información sobre las alternativas a la metadona.

2.1.4 Homogeneización de respuestas y escalas

Además de los nombres, ha sido necesario unificar las **categorías** de respuestas, ya que algunas respuestas equivalentes aparecían formuladas de forma diferente entre años. Por ejemplo, en 2022 la categoría de EDAD denominada “menor 25” ha sido transformada en “Hasta 24 años”, y las respuestas de “NS/NS” en “No sabe/No contesta”.

También, se han normalizado **escalas** de valoración que poseían una formulación ligeramente diferente, ya que en algunas variables del 2022 aparecían categorías como “Muy bueno general” “Bueno general” o “Mala” y, por tanto, han sido transformadas a las categorías comunes utilizadas en los años 2023 y 2024, como “Muy bueno”, “Bueno” o “Mal”, según corresponda.

Finalmente, las preguntas de **respuesta múltiple**, que en 2022 aparecían en una sola columna, han sido transformadas en variables binarias donde “1” indica que la opción fue seleccionada y “0” que no fue seleccionada, garantizando la compatibilidad con las bases de 2023 y 2024.

2.1.5 Incorporación de la variable temporal

Una vez normalizadas las bases y variables, se ha incorporado una variable denominada “año”, que identifica el año de cada registro, fundamental para poder estudiar posteriormente las **diferencias** entre 2022, 2023 y 2024.

Adicionalmente, se han creado variables binarias para 2023 y 2024, y poder evaluar si pertenecer a un año u a otro aporta capacidad predictiva sobre la **satisfacción** o la recomendación del servicio.

2.1.6 Unión de la base de datos

Antes de la unión final de las tres bases, se han alineado las **columnas** de todas ellas, ya que no todas las variables están disponibles en todos los años. Por tanto, se han conservado esas variables donde la pregunta no estuviera disponible, asignando valores en blanco.

Posteriormente, las tres bases de datos se han **concatenado** verticalmente en una única base de datos, con todas las observaciones y variables temporales añadidas. Además, se reordenaron las columnas para situar al inicio las variables identificadoras y sociodemográficas, facilitando el análisis posterior.

2.1.7 Comprobaciones iniciales de estructura y consistencia

Tras la integración, se han realizado unas **comprobaciones** iniciales de estructura. En primer lugar, el número total de observaciones y variables del conjunto final (1133 observaciones y 178 variables), seguido de la distribución por año, detección de registros duplicados y la tipología de las variables.

Finalmente, la base armonizada se ha guardado en un CSV, como punto de partida para los siguientes pasos del análisis.

2.2 Limpieza y preprocesamiento

2.2.1 Análisis inicial de calidad de los datos

El conjunto final está integrado por un total de 1133 observaciones y 178 variables, incluyendo; sociodemográficas, características clínicas, valoraciones de satisfacción, percepción de calidad, uso de recursos y preguntas de **respuesta múltiple**.

Además, se ha revisado también la **tipología** de variables presentes en la base de datos, y se dispone de:

- Variables categóricas nominales
- Variables ordinales de escala tipo Likert
- Variables binarias derivadas de preguntas de respuesta múltiple
- Variables numéricas discretas
- Variables de texto libre

El análisis de valores ausentes ha mostrado que la ausencia de información no se distribuye de forma homogénea entre las variables del dataset. Las variables “CAD” y “Sustancia” (P7), presentan un 100% de **missing values**, lo que supone una limitación de la información disponible.

Además de estas variables completamente vacías, se han identificado otras con porcentajes muy elevados de **valores faltantes**:

Variable	Descripción	Porcentaje de valores perdidos
P7	Sustancia	100.00%
CAD	CAD	100.00%
P39_LUGAR_METADONA_OTROS_2022	Lugar de recogida de metadona en 2022: texto abierto	99.91%
P47	Valoración de piso de apoyo a la reinserción	95.94%
P50	Valoración de comunidad terapéutica	94.35%
P44	Valoración de piso de apoyo al tratamiento	93.73%
P39_LUGAR_METADONA_2022	Lugar de recogida de metadona en 2022	93.20%
P41_DISTANCIA_METADONA_2022	Valoración de distancia al lugar de recogida de metadona	93.20%
P53	Valoración de centro de patología dual	92.67%
P39A	Tratamiento actual con metadona	92.23%

Imagen 1. Variables con mayor porcentaje de valores faltantes. Fuente: elaboración propia mediante Python

Estos resultados son coherentes con la estructura de los cuestionarios. Muchas variables son condicionales, lo que son únicamente aplicables a determinados usuarios según las respuestas previas. Por ejemplo, las relacionadas con **metadona**, pisos de apoyo o comunidades terapéuticas solo las responden quienes han tenido contacto con esos recursos.

Por tanto, una elevada proporción de **valores faltantes** en dichas variables no indica una mala calidad de los datos, sino como algo lógico tal y como está diseñado el cuestionario.

Por otro lado, se observan diferencias en el porcentaje medio de valores **ausentes** entre años, como muestra a siguiente imagen.

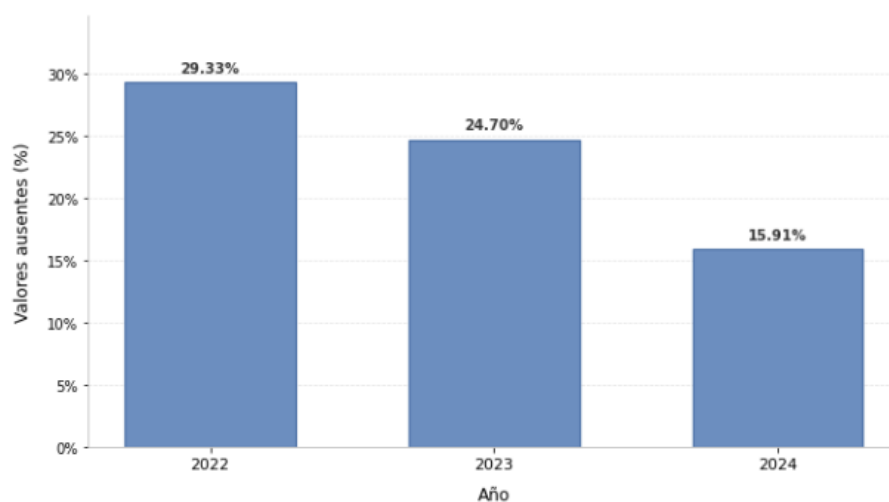


Imagen 2. Porcentaje de valores faltantes por año. Fuente: elaboración propia mediante Python

Podemos observar una pequeña mejora progresiva en la **completitud** del cuestionario a lo largo de los años, aunque puede deberse también a cambios estructurales o en la forma en la que se recogen los datos.

2.2.2 Tratamiento de los valores faltantes

Tras el análisis inicial, se ha procedido al **tratamiento** de los **valores faltantes** para minimizar la pérdida de información y garantizar la robustez de los modelos.

Primero, las variables “Sustancia” y “CAD”, con un 100% de valores nulos han sido eliminadas del conjunto final para la **modelización**.

También, se han identificado otras variables con porcentajes muy elevados de **valores faltantes**, pero en lugar, de eliminar automáticamente todas estas variables con altos porcentajes de ausencia, se ha realizado una evaluación de las mismas debido a su posible utilidad en fases posteriores del estudio, ya que resulta especialmente relevante en investigaciones relacionadas con el **ámbito sanitario**, qué determinadas variables pueden tener un interés superior pese a su baja frecuencia de respuestas.

2.2.3 Detección y tratamiento de valores atípicos (outliers)

Es especialmente relevante controlar los valores atípicos, ya que puede haber algunos algoritmos que pueden verse afectados por ellos. Para identificarlos, se ha empleado el criterio del **rango intercuartílico (IQR)**, considerando outliers aquellos valores que queden fuera del siguiente intervalo:

$$[Q1 - 1,5 \cdot IQR, Q3 + 1,5 \cdot IQR]$$

Donde:

- Q1 representa el primer cuartil,
- Q3 el tercer cuartil,
- e IQR el rango intercuartílico

Además del criterio, se llevaron a cabo un par de representaciones gráficas mediante **boxplots** para poder visualizar la distribución de estas variables numéricas.

Se han detectado registros extremadamente elevados en las variables de **duración** del tratamiento (P9 y P9_B, tiempo total en meses y en años), con valores de 999 meses o 1.000 años, cifras incompatibles e interpretados como códigos internos de no respuesta.

Por tanto, se ha establecido un **umbral** máximo de 480 meses y 40 años, reemplazando los valores superiores a estos umbrales como valores ausentes.

Y tras la eliminación de valores imposibles, el análisis mediante los *boxplots* ha mostrado tanto la existencia de *outliers* como de valores algo extremos pero que tiene cabida dentro

del **rango intercuartílico** y que se corresponden con trayectorias prolongadas dentro del tratamiento. Algo posiblemente normal dentro de los procesos de adicción y seguimiento terapéutico, por lo que estos valores algo extremos han sido considerados válidos y por tanto se mantienen dentro de la muestra.

Así, solo se han eliminado los registros claramente incompatibles, depurando además los valores fuera del **rango intercuartílico** una vez apartado dichos valores imposibles.

2.2.4 Creación de variables derivadas

Se han generado nuevas variables derivadas para facilitar el posterior análisis predictivo. En concreto, se ha creado una variable de **duración** total del tratamiento en meses unificada, integrando la información expresada en meses y en años.

2.2.5 Recodificación de variables ordinarias y binarias

2.2.5.1 Recodificación de escalas ordinales tipo Likert

Las variables ordinales de **satisfacción** y valoración se han transformado en escalas numéricas ordinales del 1 al 5, donde 1, representa la valoración más negativa posible y 5 la más positiva.

Y las respuestas correspondientes a “No sabe/No contesta”, han sido tratadas como **valores ausentes** para evitar sesgos en los análisis posteriores.

Por tanto, esta transformación ha servido para mantener el orden natural de las respuestas categóricas, mejorar la **interpretabilidad**, y facilitar el cálculo de medias, correlaciones y la aplicación de modelos predictivos, garantizando la consistencia estructural del conjunto integrado.

2.2.5.2 Recodificación de variables binarias

Las variables binarias se han transformado a formato numérico mediante **codificación dicotómica**. Las respuestas afirmativas como “Sí”, se han codificado como 1, mientras que las negativas como “No”, en 0. Y, los registros sin respuesta han sido considerados

como valores ausentes. Esto, permite el uso de estas variables en algoritmos de Machine Learning que requieren de entrada numérica.

2.2.6 Tratamiento de preguntas con respuesta múltiple

Las preguntas de respuesta múltiple identificadas son; P10 (como conoció el CAD), P11 (motivos del tratamiento), P12 (personas de apoyo al inicio), P35 (motivos para no participar en terapias grupales), P56 (razones para no recomendar el servicio de orientación laboral) y P62 (motivos percibidos de **discriminación**).

En 2023 y 2024, estas variables ya venían desagregadas, cada una con una opción de respuesta. En cambio, en 2022 presentaban una codificación diferente, ya que la opción seleccionada por los **usuarios** en el cuestionario aparecía marcada como “Sí”.

Por tanto, el tratamiento aplicado ha sido la **recodificación** de las opciones, las cuales, las seleccionadas por el usuario van a aparecer como 1 y las no seleccionadas como 0.

2.2.7 Imputación de valores perdidos (Multiple Imputation)

Posteriormente, se ha procedido a la imputación de valores faltantes. Se trata de una fase que se ha realizado al final del preprocesamiento para evitar que los valores que se van a imputar alterasen la distribución original de las variables, sobre todo, en el análisis de los **outliers** o en la construcción de las variables derivadas.

Antes de la imputación, el dataset contenía 49.280 **valores faltantes**, concentrados en variables condicionales, preguntas de texto libre y variables de servicio no aplicables a todos los usuarios.

Dada la diferente naturaleza de las variables, se han aplicado dos estrategias de imputación. Para las variables categóricas, se ha utilizado la imputación por moda mientras que para variables numéricas, se ha empleado la técnica de **Multiple Imputation** (van Buuren & Groothuis-Oudshoorn, 2011).

Esta técnica se basa en la estimación de los valores faltantes mediante la información disponible en el resto de las variables, repitiendo el proceso de forma iterativa hasta completar el conjunto de datos, concretamente, basada en **Random Forest**.

Tras el proceso, el conjunto final ha quedado compuesto por 1133 observaciones, 180 variables y 0 **valores faltantes**, frente a los 49.280 que se tenían anteriormente.

2.2.8 Identificación final de tipos de variables

Se ha verificado que el dataset ha quedado compuesto por 141 variables numéricas (indicadores anuales, binarias de respuesta múltiple, variables de **duración** del tratamiento, ordinales recodificadas a numéricas y variables derivadas) y 39 categóricas (sociodemográficas, texto libre, y preguntas abiertas) como SEXO, EDAD, ESTUDIOS o SITUACION.

La verificación ha confirmado la ausencia de **valores faltantes** y la correcta clasificación de las variables. Aunque, dado que aún hay variables categóricas y algunos algoritmos de Machine Learning requieren de variables numéricas, será necesario aplicar técnicas de codificación adicionales.

2.2.9 Construcción del dataset analítico final

Finalmente, tras integrar, limpiar, transformar, recodificar e imputar los datos, el **dataset** final quedó compuesto por 1.133 observaciones y 180 variables con distribución equilibrada por año (372 en 2022, 379 en 2023 y 382 en 2024), sin valores faltantes ni registros duplicados.

Este dataset se ha guardado en formato CSV y posee la base analítica definitiva del trabajo para el desarrollo de modelos supervisados y las técnicas de **interpretabilidad**.

2.3 Construcción de variables objetivo

2.3.1 Fundamentación teórica de la variable objetivo

La construcción de la **variable objetivo** es uno de los pasos metodológicos más importantes, dado que estamos definiendo la variable que será modelizada mediante técnicas estadísticas y de machine Learning.

Se ha elegido la variable “**Satisfacción** general con el servicio del CAD” (P13) como indicador principal, dado que, recoge la valoración global de los encuestados sobre el servicio recibido, la percepción del centro y la atención recibida.

La elección de esta variable se basa en distintos motivos metodológicos. En primer lugar, se trata de una pregunta relacionada directamente con el objetivo de este estudio, ya que se orienta a identificar los factores asociados a una mejor **percepción** de los programas del centro. En segundo lugar, el conjunto integrado presenta una elevada calidad del dato y una **distribución** más equilibrada que otras variables candidatas a variable **objetivo** del cuestionario como la “Recomendación del CAD” (P63). Y finalmente, dada su naturaleza, vamos a poder recodificarla para poder realizar diferentes estrategias de modelización que necesitan de una adaptación de la naturaleza del tipo de variable.

Desde una perspectiva analítica, la satisfacción general es un indicador muy relevante en la evaluación de servicios sociosanitarios, ya que existen diversos estudios que asocian la buena adherencia al servicio y la **confianza** en el centro con mejores resultados clínicos y mayor estabilidad personal.

2.3.2 Selección de la variable objetivo

El análisis descriptivo inicial ha mostrado una clara distribución concentrada en las categorías superiores de la escala, aunque más equilibrada que otras variables potencialmente usables como **variable objetivo**.

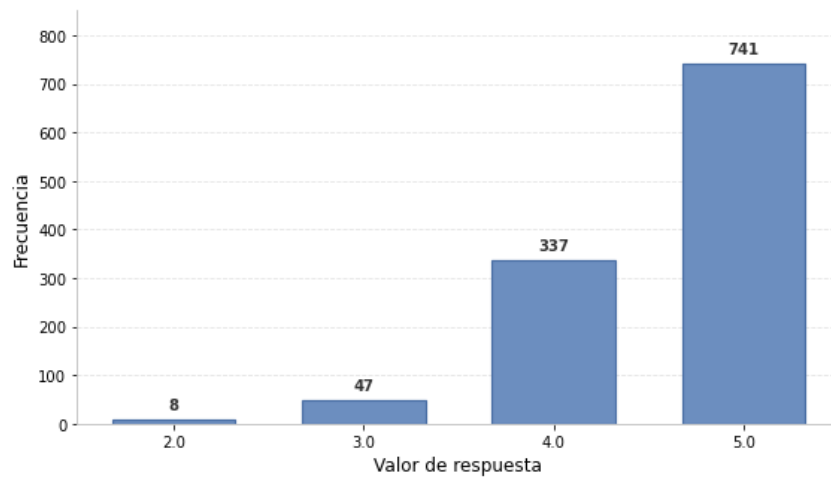


Imagen 3. Distribución de frecuencias de la variable objetivo (P13 - Satisfacción general con el servicio del CAD). Fuente: elaboración propia mediante Python

La gran mayoría de las observaciones se concentra en las categorías 4 y 5, lo que representan niveles altos en cuanto a la satisfacción general. Sin embargo, las categorías bajas e intermedias tienen una presencia un tanto suficiente para poder realizar una **clasificación** más estable y metodológicamente correcta y robusta.

Por este motivo, se ha considerado esta variable como la alternativa más adecuada para el análisis principal del estudio.

2.3.3 Recodificación y construcción de la variable objetivo

La variable objetivo final se ha reconstruido mediante la **recodificación** agrupada de la variable original “Satisfacción general con el servicio del CAD” (P13). Para mejorar el **equilibrio** entre las clases y la estabilidad de los modelos, las categorías han sido reorganizadas en los siguientes niveles:

- Satisfacción baja: valores 1, 2 y 3
- Satisfacción media: valor 4
- Satisfacción alta: 5

Esta **recodificación** reduce la distribución asimétrica existente, sobre todo en las categorías poco frecuentes, manteniendo la interpretación de las variables y obteniendo una distribución considerablemente más equilibrada.

La variable objetivo resultante se ha denominado “*target_satisfaccion*”, constituyéndose la variable principal para los modelos predictivos de los siguientes capítulos.

2.3.4 Análisis descriptivo de la variable objetivo

La siguiente imagen muestra la distribución final de la variable objetivo. Muestra un claro dominio de la **satisfacción** alta, seguida de un volumen intermedio de registros de media satisfacción y un porcentaje menor todavía el de satisfacción baja.

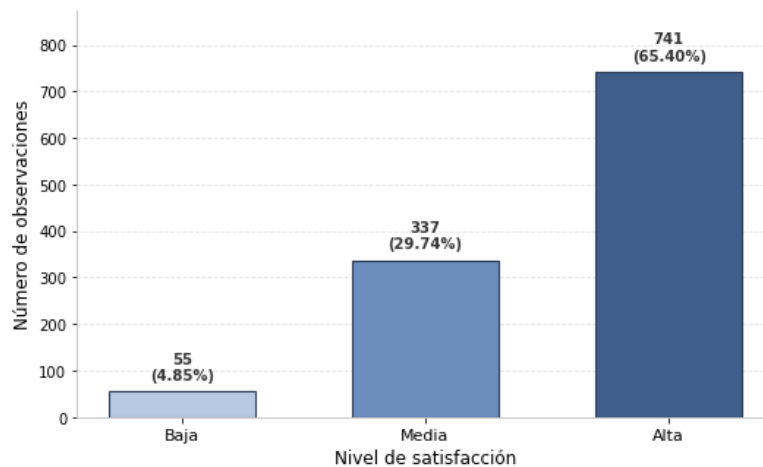


Imagen 4. Distribución final de la variable objetivo. Fuente: elaboración propia mediante Python

2.3.5 Evaluación del equilibrio de clases

La **distribución** final presenta una mayor frecuencia, aunque moderada de la categoría alta, aunque manteniendo una representación suficiente de las categorías media y baja.

Desde el punto de vista metodológico, esta distribución de frecuencias resulta más apropiada para modelos de **aprendizaje supervisado**, ya que reduce el sesgo hacia la clase mayoritaria y favorece la estabilidad de los modelos.

2.4 Análisis preliminar de relevancia de variables

2.4.1 Selección inicial de variables predictoras

Se ha realizado una selección preliminar de variables candidatas para la construcción de los **modelos predictivos**.

Para ello, se han excluido las variables **identificadoras**, las derivadas **redundantes** de la variable objetivo y las de **texto libre**, al no poder incorporarse a los modelos sin un preprocesamiento adicional.

Adicionalmente, se ha eliminado la variable original (“Satisfacción general con el CAD” - P13) que se ha usado para construir la variable objetivo (*target_satisfaccion*) y el conjunto final de variables predictoras ha quedado compuesto por variables sociodemográficas, clínicas, calidad asistencial y percepción del servicio, todo relacionado con los tratamientos recibidos por cada uno de los usuarios.

2.4.2 Asociación entre variables categóricas y variable objetivo

Para evaluar la relación entre las variables categóricas y la objetivo, se ha aplicado el contraste de **Chi-cuadrado** de independencia, junto con el estadístico **V de Cramér** que muestra el tamaño del efecto.

Variable	Descripción	Chi-cuadrado	p-valor	V de Cramer
P61	Experiencia de discriminación en el CAD	121.22	0.0000	0.231
SITUACION	Situación laboral	22.94	0.0613	0.101
P39	Tratamiento con metadona	10.86	0.0044	0.098
P40	Información sobre alternativas a metadona	19.47	0.0006	0.093
SEXO	Sexo	9.72	0.0078	0.093
EDAD	Grupo de edad	14.20	0.5836	0.079
P6	Tiempo en el tratamiento actual	14.04	0.0807	0.079
P17	Responsabilidades de cuidado	6.76	0.0341	0.077
P8	Primera vez en tratamiento en un CAD	6.48	0.0392	0.076
ESTUDIOS	Nivel de estudios	8.50	0.5799	0.061
P18	Dificultad por responsabilidades de cuidado	3.34	0.5019	0.038
P38A	Motivo de no participación familiar	2.80	0.9464	0.035

Imagen 5. Variables categóricas más asociadas con el nivel de satisfacción. Fuente: elaboración propia mediante Python

Los resultados muestran asociaciones estadísticamente significativas entre diversas variables y el nivel de satisfacción. Entre las variables con **mayor intensidad**, destaca

principalmente P61, relacionada con la experiencia de discriminación en el CAD, con el mayor valor de V de Cramér (0,231).

También, muestran asociaciones relevantes, en menor medida, el tratamiento con la metadona (P39), la información sobre **alternativas terapéuticas** a la metadona (P40), sexo, edad, situación laboral y variables relacionadas con el acceso al servicio.

En general, los **tamaños de efecto** medidos por la V de Cramér, son bajos o moderados, algo habitual en este tipo de estudios sociales y sanitarios con variables tanto de comportamiento como subjetivas.

2.4.3 Relación entre variables numéricas y variable objetivo

Posteriormente, se han analizado las diferencias entre grupos de satisfacción para las variables numéricas mediante la técnica de ANOVA.

Variable	Descripción	Estadístico F	p-valor
P31_2	Personal general: coordinación entre profesionales	132.05	< 0.0001
P31_1	Personal general: trato recibido	119.69	< 0.0001
P31_4	Personal general: rapidez ante urgencias	114.71	< 0.0001
P21_2	Agilidad para resolver problemas en recepción	91.92	< 0.0001
P63	Recomendación del CAD	82.88	< 0.0001
P21_1	Trato recibido en recepción	78.28	< 0.0001
P25_2	Trabajador/a social: duración de consulta	74.46	< 0.0001
P28_2	Enfermero/a: duración de consulta	71.06	< 0.0001
P30_2	Terapeuta ocupacional: duración de consulta	67.53	< 0.0001
P33_3	Libertad para elegir ofertas terapéuticas	67.36	< 0.0001
P16	Agilidad en la primera cita	65.21	< 0.0001
P26_3	Psicólogo/a: conocimientos y ayuda	64.14	< 0.0001
P26_2	Psicólogo/a: duración de consulta	63.81	< 0.0001
P36	Valoración de terapia de grupo	63.01	< 0.0001
P59	Opciones terapéuticas suficientes	62.38	< 0.0001

Imagen 6. Variables numéricas más asociadas con el nivel de satisfacción. Fuente: elaboración propia mediante Python

Los resultados muestran diferencias estadísticamente significativa entre los distintos niveles de **satisfacción** con respecto a aquellas variables relacionadas con la percepción de la atención recibida.

Destacan, con los mayores valores del **estadístico “F”**, P31_2 (coordinación entre profesionales), P31_1 (trato del personal general), P31_4 (rapidez ante urgencias), P21_2 (agilidad para resolver problemas en recepción) y P63 (recomendación del CAD), lo que indica diferencias muy relevantes entre los distintos niveles de satisfacción.

Adicionalmente, un *boxplot* de la variable “tiempo_total_meses” segmentado por niveles de satisfacción, muestra como existe una mayor dispersión en el grupo de baja satisfacción, además de tratamientos con mayor **duración**.

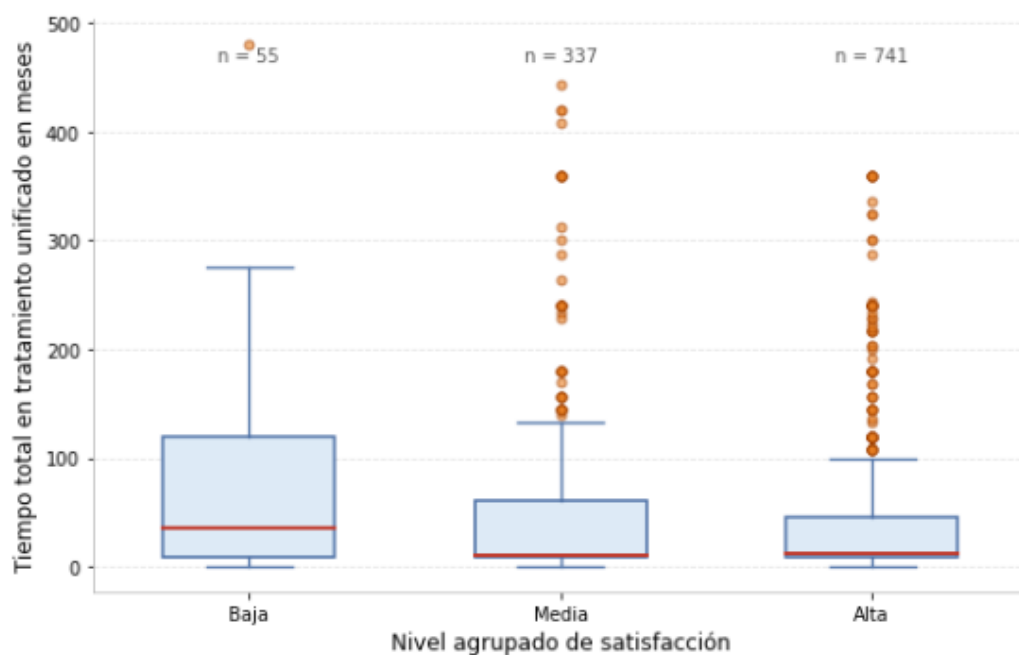


Imagen 7. Tiempo total en tratamiento según el nivel de satisfacción. Fuente: elaboración propia mediante Python

Estos resultados sugieren que ciertas características clínicas o percepciones de los **recursos** pueden estar asociadas con diferentes niveles de satisfacción.

2.4.4 Matriz de correlaciones entre variables numéricas

Para identificar las posibles relaciones lineales entre las variables numéricas y la objetivo, se ha realizado un análisis de correlaciones mediante el coeficiente de **Pearson**.

Dado el elevado número de variables, se han representado únicamente aquellas que mostraban una mayor **asociación** con la variable objetivo.

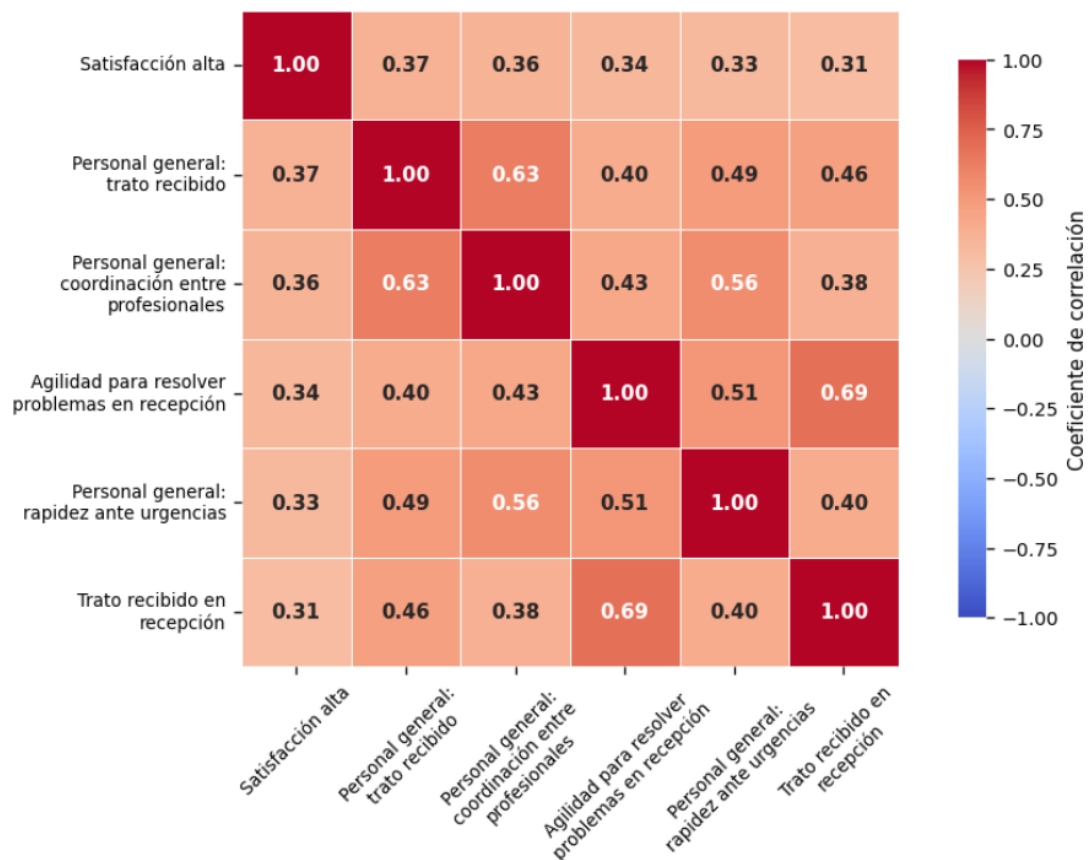


Imagen 8. Matriz de correlaciones entre la satisfacción y las variables numéricas más relacionadas. Fuente: elaboración propia mediante Python

Los resultados muestran **correlaciones** moderadas y positivas entre la satisfacción y diversas variables relacionadas con la calidad percibida del servicio.

Destacan el trato del personal general ($\rho = 0,37$), la coordinación entre profesionales ($\rho = 0,36$), la agilidad para resolver problemas en recepción ($\rho = 0,34$), la rapidez de respuesta ante situaciones urgentes ($\rho = 0,33$) y el trato recibido en recepción ($\rho = 0,31$).

También, se puede observar **correlaciones** relativamente fuertes entre algunas de estas variables.

Por ejemplo, la **relación** entre la agilidad para resolver problemas en recepción y el trato recibido en la recepción alcanza un valor de 0,69, mientras que la correlación entre el

trato recibido por el personal general y la coordinación entre profesionales alcanza un valor de 0,63. Dichos resultados sugieren que estas dimensiones de la **calidad asistencial** del servicio tienden a estar relacionadas entre sí.

En conjunto, las **correlaciones** son moderadas, indicando que la satisfacción no depende únicamente de un solo factor aislado, sino de la combinación múltiples factores relacionados con la atención recibida.

2.4.5 Selección preliminar de variables relevantes

Finalmente, se ha realizado una **selección preliminar** de las variables potencialmente más relevantes según los resultados del análisis de asociación, ANOVA y correlaciones.

Variable	Descripción	Variable	Descripción
P61	Experiencia de discriminación en el CAD	P31_2	Personal general: coordinación entre profesionales
SITUACION	Situación laboral	P31_1	Personal general: trato recibido
P39	Tratamiento con metadona	P31_4	Personal general: rapidez ante urgencias
P40	Información sobre alternativas a metadona	P21_2	Agilidad para resolver problemas en recepción
SEXO	Sexo	P63	Recomendación del CAD
EDAD	Grupo de edad	P21_1	Trato recibido en recepción
P6	Tiempo en el tratamiento actual	P25_2	Trabajador/a social: duración de consulta
P17	Responsabilidades de cuidado	P28_2	Enfermero/a: duración de consulta
P8	Primera vez en tratamiento en un CAD	P30_2	Terapeuta ocupacional: duración de consulta
ESTUDIOS	Nivel de estudios	P33_3	Libertad para elegir ofertas terapéuticas

Imagen 9. Variables con mayor relevancia con respecto a la variable objetivo (categóricas y numéricas, respectivamente). Fuente: elaboración propia mediante Python

Entre ellas, las relacionadas con la percepción de la atención recibida, la coordinación entre los profesionales, el trato del personal, la rapidez de respuesta, la recomendación del CAD y aspectos subjetivos como la percepción de discriminación en el centro.

Por tanto, este análisis preliminar ha permitido identificar un conjunto sólido de variables **candidatas** para el capítulo de modelización y una pequeña evaluación de su importancia relativa.

CAPÍTULO II: MODELOS PREDICTIVOS SUPERVISADOS

3.1 Planteamiento del problema

3.1.1 Definición de la variable objetivo

En este capítulo, se plantea un problema de **clasificación supervisado**, orientado a predecir el nivel de satisfacción general de las personas encuestadas en los programas de tratamientos de adicciones.

La variable objetivo es “Satisfacción general con el servicio del CAD”, que recoge la valoración global de todos los servicios del centro. A partir de ella, se ha construido una **clasificación agrupada** de la satisfacción en tres niveles: baja, media y alta satisfacción.

La baja satisfacción agrupa a aquellos usuarios cuya respuesta fue entre los valores 1 y 3, la media corresponde al valor 4 y la alta satisfacción recoge el valor máximo, que es el 5. Esta agrupación conserva la **naturaleza** de la variable, al mismo tiempo que reúne las categorías más minoritarias y facilita la aplicación de los modelos de clasificación.

La distribución final muestra una asimetría hacia las categorías más altas, siendo la categoría baja la menos predominante. Pese al **desbalanceo**, esta formulación resulta más adecuada para la modelización que otras formulaciones binarias exploradas (por ejemplo, agrupar 1-3 como “Insatisfecho” y 4-5 como “Satisfecho”).

3.1.2 Preparación de datos para modelización

Una vez definida la **variable objetivo**, se ha procedido a la verificación del dataset construido anteriormente con respecto al conjunto de variables predictoras que va a ser utilizado para el entrenamiento de los modelos supervisados. En él, ya no se encuentran variables identificativas, de campos de texto libre, variables no disponibles o comparables entre años, o derivadas redundantes.

Posterior a la verificación, se han transformado las variables categóricas mediante técnicas de codificación **one-hot encoding**, lo que permite obtener un conjunto de

variables completamente numéricas aptas para los algoritmos de Machine Learning previstos a utilizar.

Finalmente, el conjunto de datos se ha dividido en entrenamiento y prueba (80% - 20%), aplicando **estratificación** sobre la variable objetivo para mantener la distribución original de las categorías en ambos subconjuntos.

3.1.3 Tratamiento del desbalanceo de clases

La variable objetivo sigue presentando una distribución moderadamente desequilibrada entre las distintas clases, sobre todo, por el menor número de registros de pacientes con baja satisfacción, lo que puede afectar **negativamente** a los modelos, favoreciendo la **predicción** de clases mayoritarias.

Para corregirlo, se ha aplicado la técnica **SMOTE** (*Synthetic Minority Over-sampling Technique*) sobre el subconjunto de entrenamiento, que genera nuevas observaciones de las clases minoritarias mediante la interpolación de registros similares, permitiendo equilibrar la distribución sin necesidad de replicar las observaciones originales directamente.

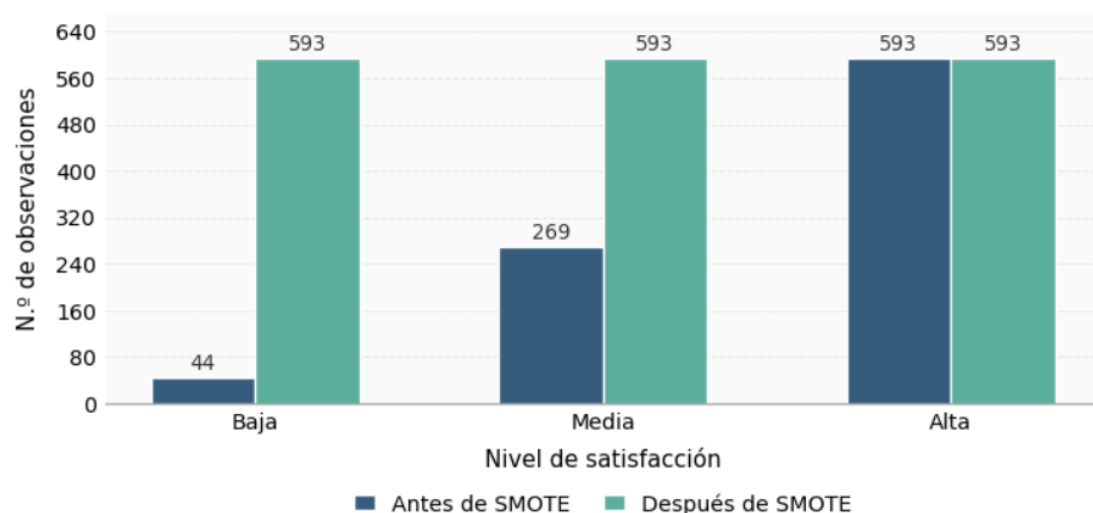


Imagen 10. Distribución de clases antes y después de aplicar SMOTE sobre el subconjunto de entrenamiento (train set). Fuente: elaboración propia mediante Python

Tras aplicar **SMOTE**, el subconjunto de entrenamiento ha pasado a tener una distribución completamente equilibrada y balanceada entre las 3 categorías disponibles; baja, media y alta satisfacción.

3.2 Modelos predictivos base

3.2.1 Regresión logística multinomial

3.2.1.1 Entrenamiento del modelo

En primer lugar, se ha desarrollado un modelo de **regresión logística multinomial** como primera aproximación al problema de clasificación categórica planteado.

Se trata de uno de los modelos supervisados más usados en problemas de clasificación por su simplicidad, su **interpretabilidad**, y su capacidad de estimación de probabilidades.

El modelo ha sido entrenado utilizando el subconjunto de entrenamiento balanceado mediante **SMOTE**, para reducir el desbalanceo existente y así, mejorar la capacidad predictiva sobre las clases minoritarias.

3.2.1.2 Evaluación del rendimiento

La evaluación sobre el subconjunto de test muestra un Accuracy del 65,64% y un F1-score de 0,6557, lo que refleja un rendimiento bastante correcto para un problema de clasificación multiclase con un elevado desbalanceo.

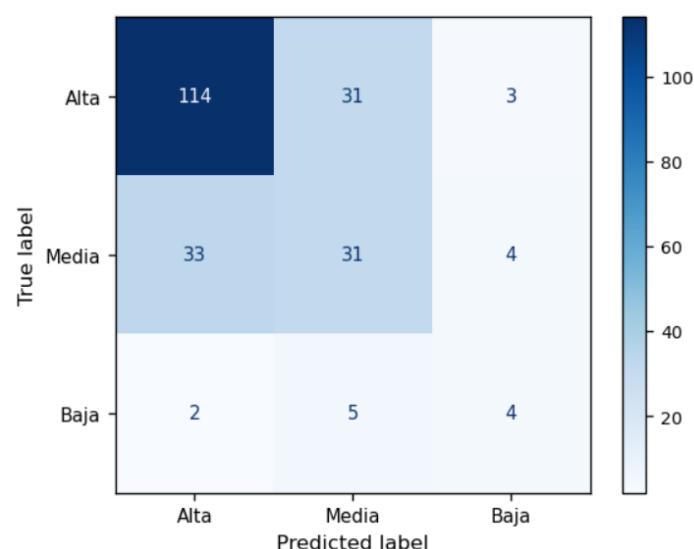


Imagen 11. Matriz de confusión del modelo de regresión logística multinomial. Fuente: elaboración propia mediante Python

La **matriz de confusión** muestra que el modelo posee una elevada capacidad de clasificación para la alta satisfacción, mientras que las categorías baja y media presentan una mayor dificultad de separación.

Por otro lado, pese al número reducido de observaciones de baja **satisfacción**, el modelo ha conseguido identificar correctamente 4 de 11 casos, algo especialmente relevante.

3.2.1.3 Interpretación de resultados

Posteriormente, se han analizado las variables con mayor peso mediante el análisis de los coeficientes absolutos de la **regresión logística**. Entre ellas podemos observar:

Variable	Descripción	Importancia media absoluta
P26_3	Psicólogo/a: conocimientos y ayuda	0.9898
P32_7	Suficiencia de otro personal sanitario	0.9568
P23_2	Médico/a: duración de consulta	0.8935
P34	Participación en terapia de grupo	0.8918
P19_2	Comodidad de las instalaciones	0.8525
P39A	Tratamiento actual con metadona	0.8061
P16	Agilidad en la primera cita	0.7395
P32_6	Suficiencia de personal administrativo	0.7289
P63	Recomendación del CAD	0.6994
P14	Facilidad/dificultad de acceso al CAD	0.6409
P61_No, nunca	P61_No, nunca	0.6271
P31_2	Personal general: coordinación entre profesionales	0.6144
P23_4	Médico/a: respeto a la opinión	0.6036
P31_1	Personal general: trato recibido	0.5900
P26_1	Psicólogo/a: tiempo hasta la cita	0.5807

Imagen 12. Variables con mayor importancia del modelo de Regresión Logística. Fuente: elaboración propia mediante Python

Destacan diferentes **indicadores de calidad** de la asistencia, la atención de los profesionales, la percepción subjetiva del tratamiento y la experiencia general del usuario.

3.2.2 Árboles de decisión

3.2.2.1 Entrenamiento del modelo

A continuación, se ha desarrollado un **árbol de decisión**, una técnica de aprendizaje supervisado muy usada en problemas de clasificación por su capacidad de identificar relaciones no lineales junto con interacciones complejas entre las variables predictoras.

3.2.2.2 Evaluación del rendimiento

El árbol de decisión ha obtenido un Accuracy del 63,44% y un F1-score de 0,6424 sobre el test. Aunque el rendimiento es ligeramente inferior a la regresión logística, el modelo mantiene una **capacidad** de clasificación **razonable**.

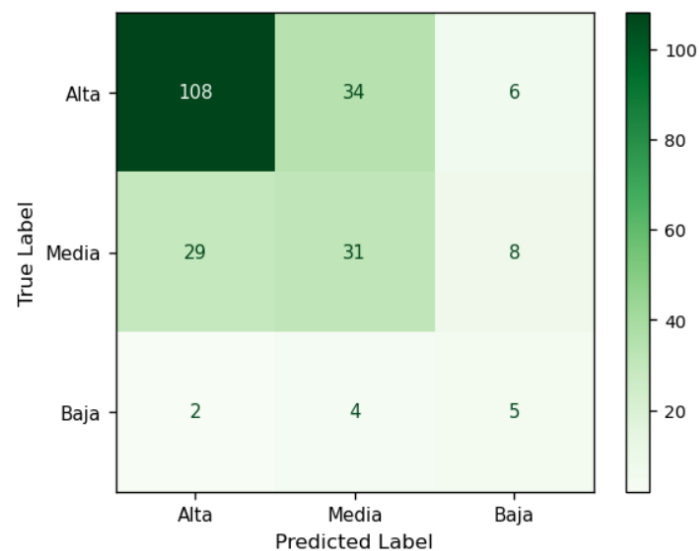


Imagen 13. Matriz de confusión del modelo de árbol de decisión. Fuente: elaboración propia mediante Python

La **matriz de confusión** presenta una mayor dificultad de clasificación precisa, no tanto en los pacientes con alta satisfacción, pero sobre todo en los pacientes con baja satisfacción, debido a la muestra muy reducida de estos pacientes.

3.2.2.3 Visualización e interpretación del árbol

La representación gráfica de los primeros niveles del árbol permite identificar visualmente las principales **reglas** de decisión generadas por el modelo.

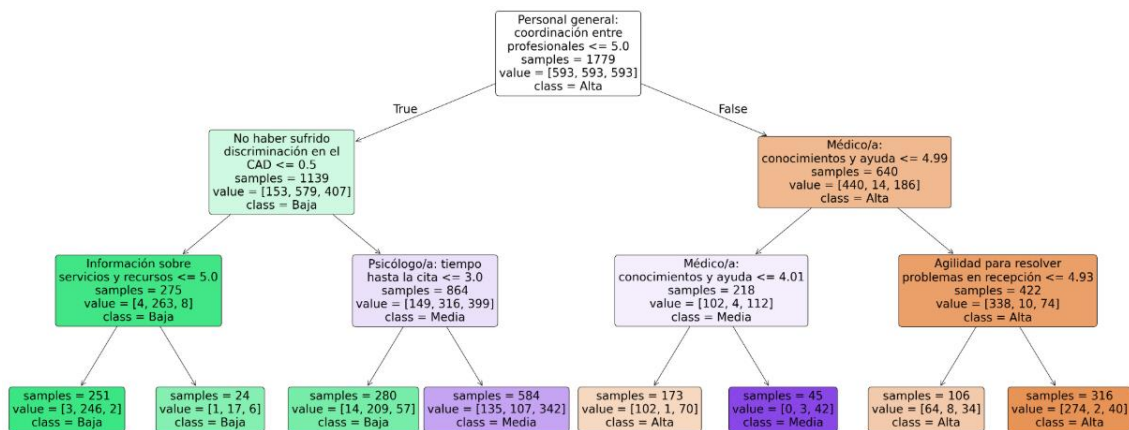


Imagen 14. Niveles del árbol de decisión e interpretación de reglas principales. Fuente: elaboración propia mediante Python

Las **primeras divisiones** se asocian a variables relacionadas con la coordinación entre los profesionales, situaciones discriminatorias, la calidad de atención recibida y la rapidez de resolución de problemas en situaciones urgentes.

3.2.3 Comparación preliminar de los modelos base

Finalmente, la **comparación** entre ambos **modelos** base desarrollados muestra un rendimiento superior de la regresión logística respecto al árbol de decisión, tanto en el Accuracy como en el F1-score.

Modelo	Accuracy	F1-score
Regresión logística	0.6564	0.6557
Árbol de decisión	0.6344	0.6424

Imagen 15. Comparación del rendimiento de los modelos predictivos base. Fuente: elaboración propia mediante Python

Ambos modelos presentan una capacidad predictiva razonable como primera aproximación, aunque se siguen visualizando esas dificultades en la **clasificación** de las categorías **minoritarias**, lo que justifica la necesidad de utilizar modelos más avanzados, además de técnicas de optimización y validación.

3.3 Modelos avanzados de Machine Learning

3.3.1 Random Forest

3.3.1.1 Entrenamiento del modelo

Para mejorar la capacidad predictiva de los modelos base, se ha realizado el estudio de un modelo **Random Forest** (Pedregosa et al., 2011), que se basa en un conjunto de árboles de decisión entrenados de manera conjunta utilizando técnicas de *bagging*, junto con una selección aleatoria de variables.

Este modelo ha sido configurado con 300 árboles de decisión, con **restricciones** de profundidad máxima y número mínimo de observaciones por nodo para reducir el *overfitting*, combinando el ajuste de **balance** de clases “*class_weight = balanced*”, con el subconjunto equilibrado mediante SMOTE.

3.3.1.2 Evaluación del modelo

Este modelo ha obtenido un Accuracy del 69,6% y un **F1-score** de 0,689, mejorando los resultados obtenidos de los modelos base.

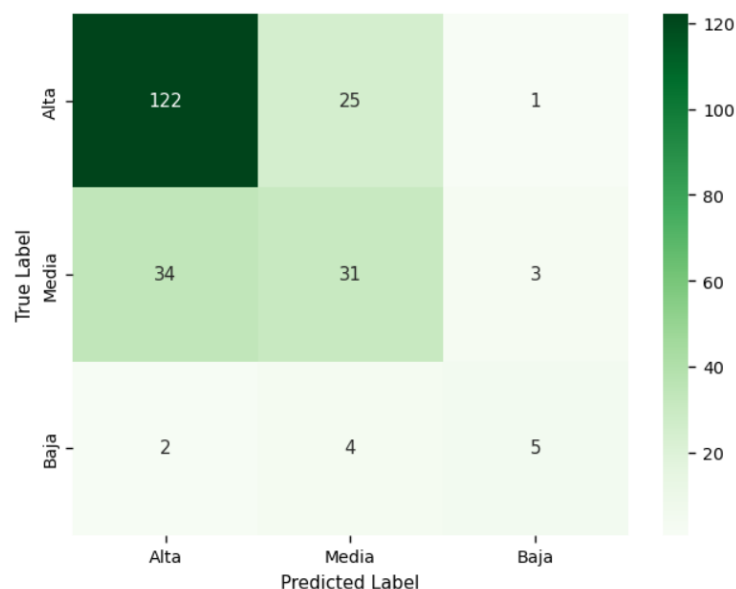


Imagen 16. Matriz de confusión del modelo Random Forest. Fuente: elaboración propia mediante Python

La **matriz de confusión** muestra una elevada capacidad para identificar la alta satisfacción, y un comportamiento ligeramente mejorado en la clasificación de los grupos de satisfacción media y baja, a pesar del desbalanceo existente.

3.3.1.3 Importancia de variables

El análisis de **importancia** de **variables** del Random Forest muestra que las variables relacionadas con la coordinación de los profesionales del centro, la rapidez de atención y la percepción de los recursos del centro son aquellas que presentan una mayor importancia y **contribución** al modelo predictivo.

Variable	Descripción	Importancia
P31_2	Personal general: coordinación entre profesionales	0.0415
P59	Opciones terapéuticas suficientes	0.0353
P61_No, nunca	P61_No, nunca	0.0311
P16	Agilidad en la primera cita	0.0308
P31_4	Personal general: rapidez ante urgencias	0.0300
P25_1	Trabajador/a social: tiempo hasta la cita	0.0287
P31_1	Personal general: trato recibido	0.0257
P26_1	Psicólogo/a: tiempo hasta la cita	0.0255
P21_2	Agilidad para resolver problemas en recepción	0.0251
P26_3	Psicólogo/a: conocimientos y ayuda	0.0248
P44	Valoración de piso de apoyo al tratamiento	0.0247
P28_1	Enfermero/a: tiempo hasta la cita	0.0236
P28_2	Enfermero/a: duración de consulta	0.0213
P36	Valoración de terapia de grupo	0.0183
P21_1	Trato recibido en recepción	0.0165

Imagen 17. Variables más importantes en el modelo Random Forest. Fuente: elaboración propia mediante Python

Resulta especialmente relevante la aparición de la variable de experiencia de **discriminación** en el CAD, como una de las más importantes.

En concreto, corresponde a la categoría en la que los pacientes indican no haber sufrido discriminación, lo que sugiere una **asociación positiva** entre la ausencia de alguna vivencia discriminatoria y mayores niveles de satisfacción.

3.3.2 XGBoost

3.3.2.1 Entrenamiento del modelo

El modelo **XGBoost**, entrena múltiples árboles de decisión de manera iterativa corrigiendo los errores de los anteriores, configurado con 300 estimadores, profundidad máxima moderada y mecanismos para regularizar y garantizar la estabilidad predictiva.

3.3.2.2 Evaluación del modelo

Este modelo ha sido el que ha obtenido el mejor rendimiento global, alcanzando un Accuracy del 70,04% y un F1-score de 0,687.

La **matriz de confusión** muestra una mejora clara en la identificación de la alta satisfacción, manteniendo un comportamiento similar en las categorías minoritarias.

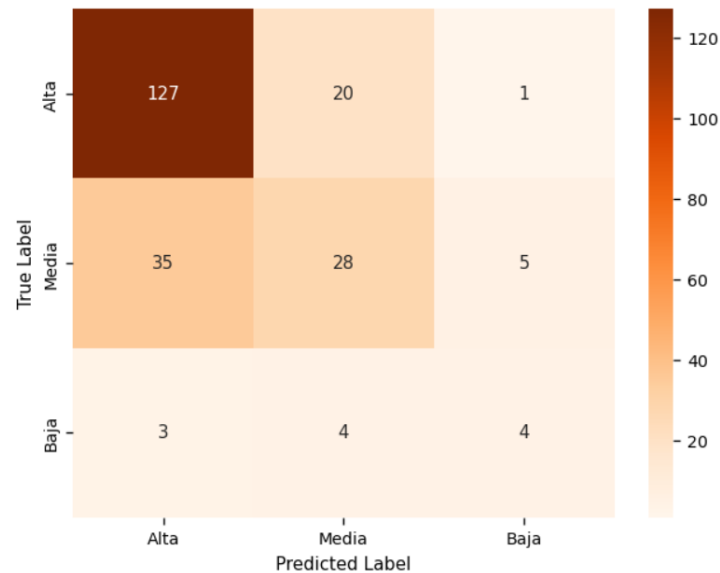


Imagen 18. Matriz de confusión del modelo XGBoost. Fuente: elaboración propia mediante Python

3.3.2.3 Importancia de variables

El análisis de **importancia** de variables del XGBoost muestra resultados bastante consistentes con el del Random Forest. Entre todas ellas destacan:

Variable	Descripción	Importancia
P59	Opciones terapéuticas suficientes	0.0521
P31_2	Personal general: coordinación entre profesionales	0.0519
P61_No, nunca	P61_No, nunca	0.0440
P31_1	Personal general: trato recibido	0.0250
P26_3	Psicólogo/a: conocimientos y ayuda	0.0221
P30_3	Terapeuta ocupacional: conocimientos y ayuda	0.0198
P25_1	Trabajador/a social: tiempo hasta la cita	0.0191
P23_3	Médico/a: conocimientos y ayuda	0.0175
P62_3	Discriminación por idioma	0.0148
P26_1	Psicólogo/a: tiempo hasta la cita	0.0144

Imagen 19. Variables más importantes en el modelo XGBoost. Fuente: elaboración propia mediante Python

La variable de experiencias discriminatorias gana en importancia en este modelo respecto al Random Forest, aumentando casi un 1,5%, un **aumento** reseñable en un contexto donde el **peso explicativo** se encuentra bastante repartido.

3.3.3 Redes Neuronales (MLP)

3.3.3.1 Arquitectura y entrenamiento

Finalmente, se ha desarrollado una **red neuronal** multicapa (MLP), con dos capas ocultas, entrenadas mediante algoritmo Adam y función de activación ReLU.

El objetivo es evaluar la capacidad predictiva de una **arquitectura neuronal** para capturar relaciones complejas y no lineales en el conjunto de datos.

3.3.3.2 Evaluación del modelo

Sobre el test, el **MLP** ha obtenido un Accuracy del 58,6% y un F1-score de 0,605, reflejando un rendimiento bastante inferior al de los modelos basados en árboles.

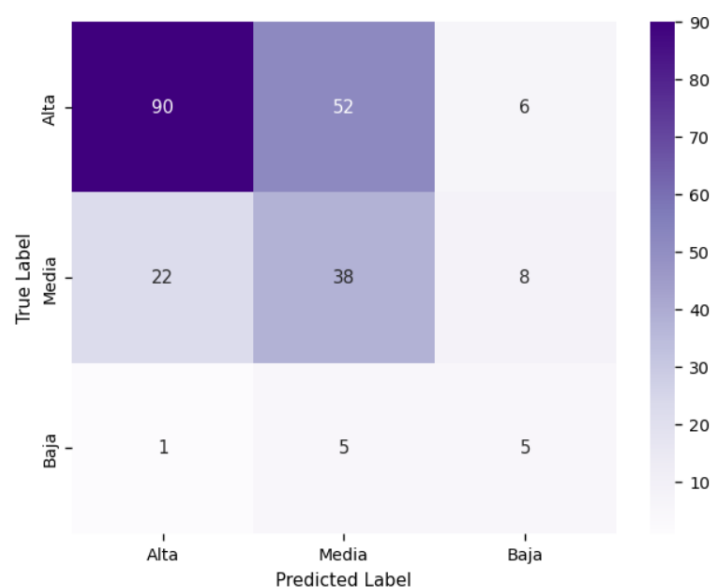


Imagen 20. Matriz de confusión del modelo de Redes Neuronales (MLP). Fuente: elaboración propia mediante Python

La **matriz de confusión** muestra una menor capacidad de generalización, especialmente en la alta satisfacción, que se reduce incluso a menos de 100, los pacientes de esa categoría identificados correctamente.

Se trata de un resultado coherente, ya que las **redes neuronales** requieren de grandes conjuntos de datos y estructuras más complejas para superar a los modelos basados en árboles.

3.3.4 Comparación global de los modelos

Los resultados muestran que el XGBoost ha obtenido el mejor rendimiento, seguido de cerca por el Random Forest. Ambos modelos superan claramente a la red neuronal y a los modelos base.

Modelo	Accuracy	F1-score
Random Forest	0.6960	0.6892
XGBoost	0.7004	0.6873
MLP	0.5859	0.6055

Imagen 21. Comparación de las métricas entre los modelos avanzados. Fuente: elaboración propia mediante Python

En general, los modelos basados en árboles muestran un mayor **rendimiento** y una mayor **capacidad** para poder capturar relaciones complejas y una mayor **estabilidad** frente al desbalance de clases.

3.4 Optimización y validación de modelos

3.4.1 Validación cruzada mediante Cross-Validation

3.4.1.1 Configuración de validación cruzada

Para evaluar la estabilidad y la capacidad de predicción de los modelos, se ha aplicado una estrategia de **validación cruzada** estratificada mediante “StratifiedkFold”.

La validación se ha configurado con 5 **particiones** ($n_splits = 5$) y aplicando un **componente aleatorio** mediante “ $random_state = 42$ ”, siendo evaluadas mediante F1-score ponderado (“ $f1_weighted$ ”), para compensar el desbalanceo entre clases.

3.4.1.2 Evaluación de estabilidad de los modelos

Esta estrategia se ha aplicado sobre los principales modelos del estudio:

- Regresión Logística
- Árbol de decisión
- Random Forest
- XGBoost

Los resultados muestran una **elevada estabilidad** general, especialmente en aquellos basados en árboles ensamblados. De todos ellos, el XGBoost fue el que obtuvo el mejor rendimiento, con un F1-score medio de 0,88, seguido del Random Forest (0,862).

Por otro lado, la regresión logística también mostró un **comportamiento** estable y sólido, mientras que el árbol de decisión fue el más limitado, con mayor **variabilidad** entre particiones (F1-score medio = 0,713, Desviación Estándar = 0,022).

3.4.1.3 Comparación de resultados de validación

La comparación final de la **validación cruzada** confirma que los modelos ensemble presentan un rendimiento superior respecto a los modelos base tradicionales.

Modelo	Media F1	Desv. Est. F1
XGBoost	0.880000	0.008000
Random Forest	0.862100	0.011900
Regresión logística	0.829600	0.013500
Árbol decisión	0.713800	0.022100

Imagen 22. Comparación de modelos mediante validación cruzada. Fuente: elaboración propia mediante Python

3.4.2 Optimización de hiperparámetros mediante Grid Search

3.4.2.1 Optimización de Random Forest

A continuación, se ha optimizado el Random Forest mediante **Grid Search** y validación cruzada estratificada de cinco particiones, para mejorar su rendimiento predictivo.

Para ello, se han evaluado distintas combinaciones de **hiperparámetros**, y la mejor fue la siguiente:

- Número de árboles (“*n_estimators*”) = 500
- Profundidad máxima (“*max_depth*”) = 20
- Mínimo de registros para dividir un nodo (“*min_samples_split*”) = 2
- Mínimo de registros permitidos en una hoja (“*min_samples_leaf*”) = 1

El **Random Forest optimizado** ha obtenido un F1-score medio en la validación cruzada de 0,879 sobre el entrenamiento. Posteriormente, sobre el conjunto de test obtuvo un Accuracy de 0,714 y un F1-score de 0,700, mejorando ligeramente el modelo sin optimizar.

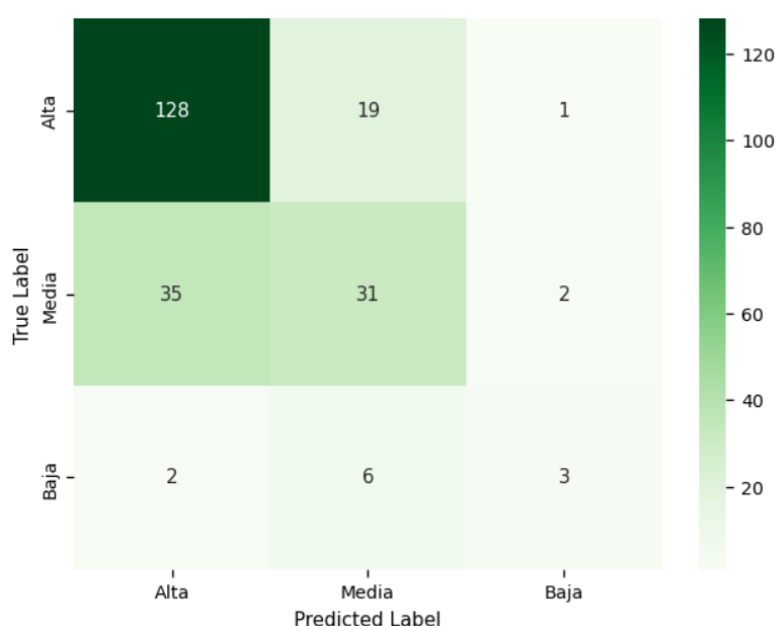


Imagen 23. Matriz de confusión del modelo Random Forest optimizado. Fuente: elaboración propia mediante Python

La **matriz de confusión** indica una mejora en la clasificación de la satisfacción alta, manteniendo una capacidad predictiva moderada para identificar las categorías media y baja.

3.4.2.2 Optimización de XGBoost

Se ha optimizado también el **XGBoost** y entre todas las combinaciones de hiperparámetros, la que obtuvo un mejor resultado fue:

- Número de árboles = 300

- Profundidad máxima = 7
- Tasa de aprendizaje (“*learning rate*”) = 0,05
- Porcentaje de registros por iteración (“*subsample*”) = 0,8

El XGBoost optimizado ha obtenido un F1-score medio en **validación cruzada** de 0,882, lo que representa el valor más elevado entre todos los modelos.

Posteriormente, sobre el test, el modelo optimizado obtuvo un **Accuracy** de 0,700 y un **F1-score** ponderado de 0,689, ligeramente inferior al de Random Forest optimizado.

3.4.2.3 Comparación del antes y después de la optimización

Finalmente, se va a comparar el rendimiento de los modelos antes y después de la optimización de los **hiperparámetros**.

Modelo	F1-score
RF original	0.6892
RF optimizado	0.7001
XGB original	0.6873
XGB optimizado	0.6886

Imagen 24. Comparación del F1-score antes y después de la optimización. Fuente: elaboración propia mediante Python

Los resultados indican una **mejora moderada** tras la optimización, especialmente en el Random Forest, cuyo F1-score aumentó de 0,689 a 0,700. En el XGBoost, la mejora fue algo menor, sugiriendo que la configuración inicial ya era cercana a la óptima.

3.4.3 Comparativa de métricas de evaluación

3.4.3.1 Comparación global de métricas

Para poder realizar una comprobación del rendimiento de los modelos, se ha construido un gráfico del **F1-score** sobre el conjunto de test.

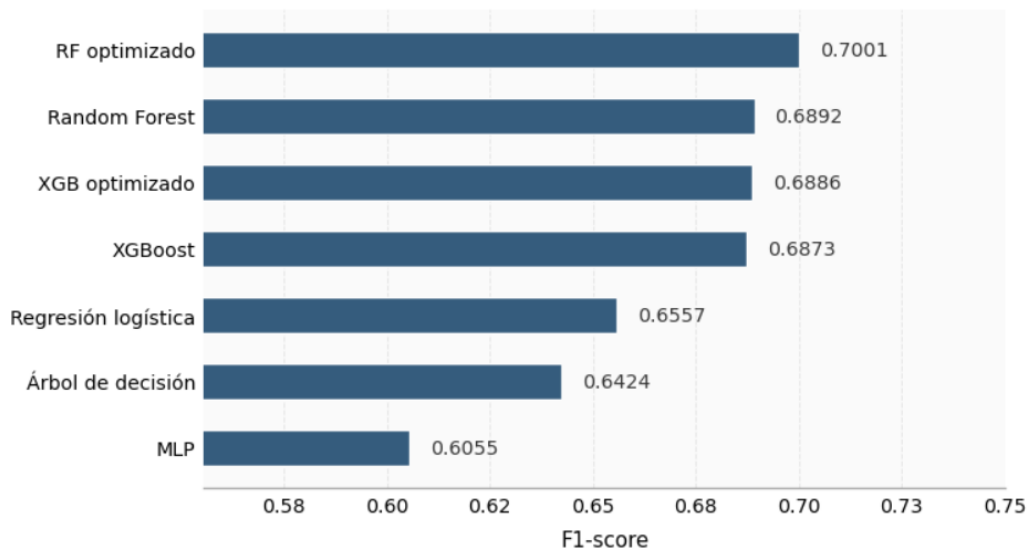


Imagen 25. Comparación global del rendimiento predictivo de los modelos evaluados mediante F1-score. Fuente: elaboración propia mediante Python

Los resultados muestran que los **modelos ensamblados** obtuvieron mejores rendimientos globales. En concreto, el Random Forest optimizado consiguió el mejor F1-score con un valor de 0,7001 y un Accuracy de 71,36%, mientras que el XGBoost obtuvo un resultado similar (F1-score = 0,6886 y Accuracy = 70,04%).

3.4.3.2 Análisis comparativo del rendimiento

A partir de los resultados, podemos observar que los algoritmos más complejos y basados en ensamblado de árboles, presentan una mayor **capacidad predictiva** para clasificar los distintos niveles de satisfacción.

La **regresión logística** ha presentado mayores limitaciones para capturar relaciones no lineales entre las variables predictivas (Hosmer et al., 2013), aunque con un rendimiento aceptable.

Por otro lado, el modelo de **árbol de decisión** muestra un rendimiento moderadamente inferior, por su mayor sensibilidad al sobreajuste (Quinlan, 1986; Loh, 2011) y una menor capacidad de generalización.

Finalmente, el modelo de **Redes Neuronales** (MLP) ha presentado el peor rendimiento de todos (Rumelhart et al., 1986; Goodfellow et al., 2016; Shickel et al., 2018), lo que confirma que esta arquitectura no es la más adecuada para un conjunto de datos pequeño y con predominio de variables categóricas (LeCun et al., 2015).

3.4.3.3 Selección preliminar del mejor modelo

Considerando las métricas de evaluación, la validación cruzada, la capacidad predictiva y la estabilidad, se han seleccionado el Random Forest y el XGBoost, ambos optimizados, para la fase final de análisis e interpretación avanzada.

Sin embargo, el modelo de Random Forest optimizado ha sido, de los dos elegidos, el modelo principal al presentar el **mejor rendimiento** final sobre el conjunto de test.

3.4.3.4 Comparación mediante curvas ROC-AUC

Para complementar el análisis, se ha calculado también el **ROC-AUC**, que se trata de una métrica capaz de evaluar la **capacidad de discriminación** de cada modelo (Fawcett, 2006) para diferenciar entre los distintos niveles de satisfacción.

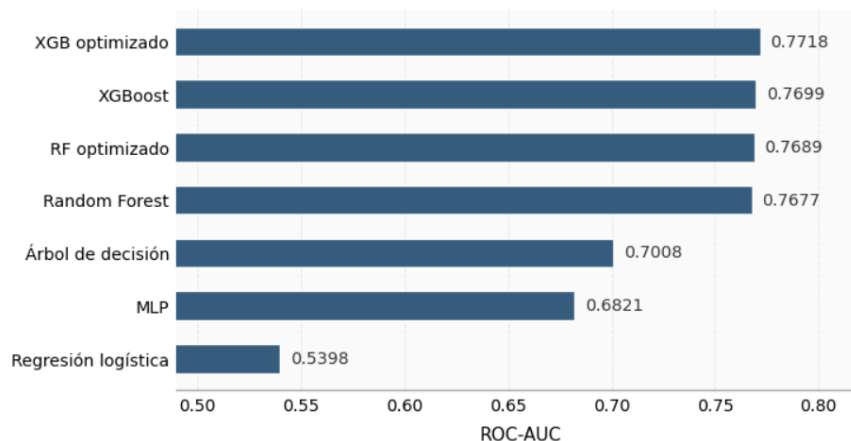


Imagen 26. Comparación global de los modelos predictivos mediante ROC-AUC. Fuente: elaboración propia mediante Python

Los resultados de los valores **ROC-AUC**, confirman que los modelos basados en ensamblado continúan siendo los que presentan mejores rendimientos predictivos.

El **XGBoost optimizado** ha alcanzado el valor más elevado, con un ROC-AUC de 0,772, seguido muy de cerca por su modelo original (0,770), y del Random Forest optimizado (0,769) y del Random Forest original (0,768).

Por el contrario, los **modelos base**, presentan una **menor capacidad** discriminativa. La regresión logística ha obtenido el valor más bajo (ROC-AUC = 0,540), mientras que el árbol de decisión y la red neuronal han alcanzado valores intermedios de 0,700 y 0,682.

Además, se ha representado las **curvas ROC** de los dos mejores modelos optimizados, tanto el Random Forest como el XGBoost. Mencionar que, al tener una variable objetivo con 3 categorías, se ha utilizado un *micro-average*, para resumir el comportamiento global de todo el modelo con respecto a todas las clases.

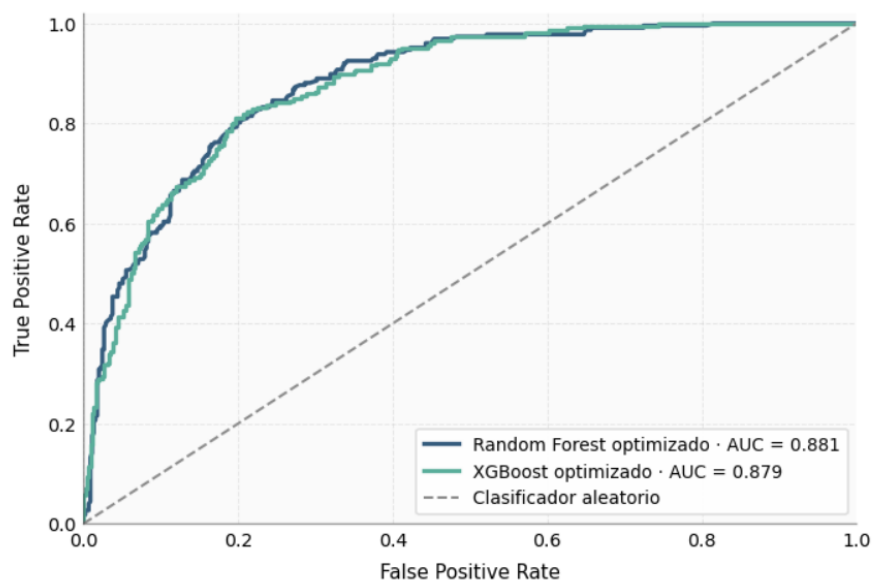


Imagen 27. Curvas ROC de los modelos Random Forest y XGBoost optimizados. Fuente elaboración propia mediante Python

La gráfica muestra unas curvas ROC bastante similares, situadas por encima de la diagonal, lo que confirma la elevada **capacidad discriminativa** de ambos algoritmos para diferenciar entre los distintos niveles de satisfacción.

Por tanto, las diferencias entre los cuatro mejores modelos son insignificantes, sugiriendo que todos ellos presentan un **rendimiento discriminatorio** muy similar.

3.4.4 Selección del modelo final

Finalmente, tras comparar los distintos modelos desarrollados, se ha procedido a seleccionar el **modelo final** para las fases posteriores de análisis avanzado e interpretación.

Los resultados han mostrado que los **modelos** basados en **ensamblados** de árboles han obtenido un rendimiento superior. En concreto, el modelo de Random Forest optimizado ha alcanzado el mejor resultado global sobre el test, con un Accuracy de 71,36% y un F1-score de 0,7001.

Aunque, el XGBoost optimizado ha obtenido el mayor ROC-AUC (0,772) (Chen & Guestrin, 2016; Friedman, 2001), la diferencia respecto al de Random Forest optimizado (0,769) es prácticamente insignificante. Por ello, se va a seleccionar finalmente el modelo optimizado de Random Forest como el **principal** del estudio, por su mejor **rendimiento** sobre el test.

Asimismo, se ha querido mantener el XGBoost optimizado como modelo de referencia para las fases de interpretación más avanzada, por su **robustez**.

3.5 Interpretabilidad y explicación de resultados

3.5.1 Importancia de variables

3.5.1.1 Importancia de variables en Random Forest Optimizado

En este apartado, se identifica los factores con mayor relevancia e influencia sobre la satisfacción general, analizando la **importancia** de las **variables** en el modelo de Random Forest optimizado.

Variable	Descripción	Importancia
P31_2	Personal general: coordinación entre profesionales	0.0339
P59	Opciones terapéuticas suficientes	0.0283
P61_No, nunca	No haber sufrido discriminación en el CAD	0.0265
P26_1	Psicólogo/a: tiempo hasta la cita	0.0239
P31_4	Personal general: rapidez ante urgencias	0.0239
P16	Agilidad en la primera cita	0.0237
P25_1	Trabajador/a social: tiempo hasta la cita	0.0236
P31_1	Personal general: trato recibido	0.0228
P21_2	Agilidad para resolver problemas en recepción	0.0217
P44	Valoración de piso de apoyo al tratamiento	0.0213
P28_1	Enfermero/a: tiempo hasta la cita	0.0212
P26_3	Psicólogo/a: conocimientos y ayuda	0.0203

Imagen 28. Variables más relevantes del modelo RF Optimizado. Fuente: elaboración propia mediante Python

Entre las más importantes destacan la **coordinación** entre **profesionales**, la percepción de disponer de alternativas terapéuticas diferentes durante el tratamiento y principalmente, la ausencia de experiencias de discriminación, junto con diversos aspectos relacionados con los profesionales del centro.

3.5.1.2 Importancia de variables en XGBoost Optimizado

Posteriormente, se ha procedido a analizar la importancia de las variables del XGBoost optimizado. Los **resultados** obtenidos son muy **consistentes** con los del RF optimizado:

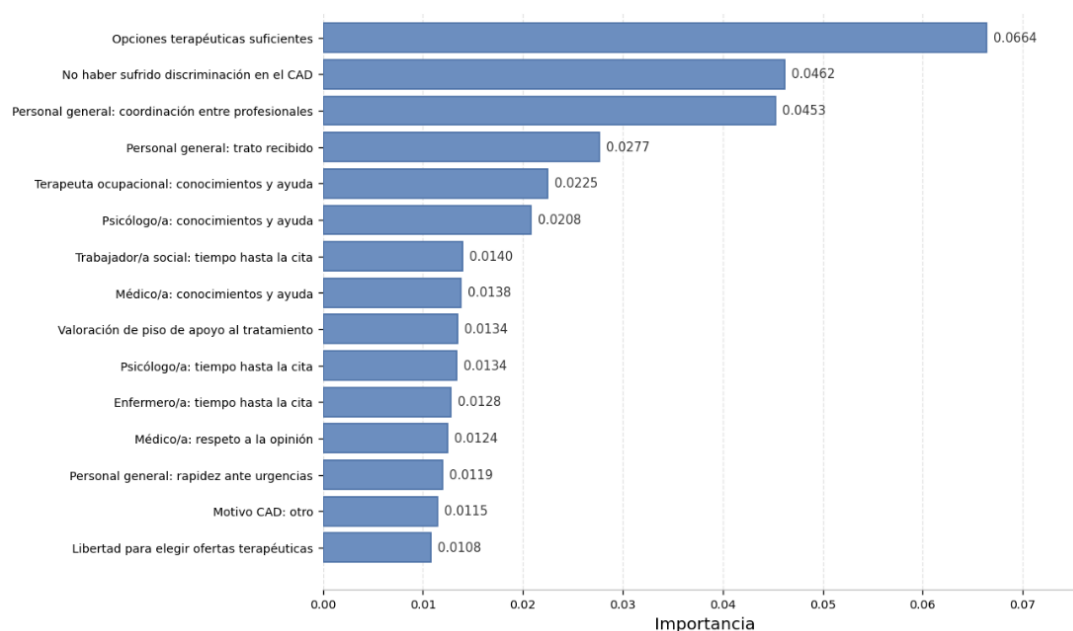


Imagen 29. Variables más relevantes del modelo XGBoost optimizado. Fuente: elaboración propia mediante Python

La variable con **mayor relevancia** ha sido la percepción de disponer de alternativas terapéuticas suficientes, seguida de la ausencia de experiencias discriminatorias en el CAD y la coordinación entre los profesionales.

3.5.1.3 Ranking combinado Random Forest y XGBoost

Finalmente, se ha construido un **ranking** conjunto de la **importancia** media de ambos modelos, para identificar aquellos factores más consistentes independientemente del modelo utilizado.

Variable	Descripción	Importancia RF	Importancia XGBoost	Importancia media
P59	Opciones terapéuticas suficientes	0.0283	0.0664	0.0474
P31_2	Personal general: coordinación entre profesionales	0.0339	0.0453	0.0396
P61_No, nunca	No haber sufrido discriminación en el CAD	0.0265	0.0462	0.0363
P31_1	Personal general: trato recibido	0.0228	0.0277	0.0252
P26_3	Psicólogo/a: conocimientos y ayuda	0.0203	0.0208	0.0205
P25_1	Trabajador/a social: tiempo hasta la cita	0.0236	0.0140	0.0188
P26_1	Psicólogo/a: tiempo hasta la cita	0.0239	0.0134	0.0186
P31_4	Personal general: rapidez ante urgencias	0.0239	0.0119	0.0179
P44	Valoración de piso de apoyo al tratamiento	0.0213	0.0134	0.0174
P16	Agilidad en la primera cita	0.0237	0.0105	0.0171
P28_1	Enfermero/a: tiempo hasta la cita	0.0212	0.0128	0.0170

Imagen 30. Ranking conjunto de variables más relevantes de ambos modelos. Fuente: elaboración propia mediante Python

Los resultados sitúan entre las 3 variables más relevantes, las ya mencionadas en ambos modelos, junto a otros factores relacionados con los profesionales del centro y la **calidad** de la **atención** recibida.

Por tanto, estos hallazgos sugieren que la satisfacción no depende únicamente de las características del tratamiento (McLellan & Hunkeler, 1998; Stallvik et al., 2020; Saloner et al., 2022), sino de otros factores, y permiten identificar áreas de mejora, y que va a servir de base para el análisis avanzado mediante técnicas **SHAP** en el siguiente apartado.

3.5.2 SHAP values

Con el objetivo de mejorar la **interpretabilidad** del modelo seleccionado, en este caso, el Random Forest optimizado, se ha aplicado la metodología **SHAP** sobre él, la cual se trata de una técnica basada en teoría de juegos (Lundberg & Lee, 2017; Molnar, 2022), permitiendo cuantificar la contribución individual de cada variable en la clasificación de los usuarios utilizando el subconjunto de prueba (test set).

3.5.2.1 Importancia global de variables mediante SHAP

En primer lugar, se ha calculado la importancia global de las variables, mediante el valor absoluto medio de los **SHAP** values, lo que va a permitir identificar aquellas que generen un mayor impacto sobre las predicciones.

Las más **relevantes** fueron la coordinación entre profesionales (Lundberg et al., 2020), el grado de suficiencia de las alternativas terapéuticas, la agilidad de recepción para resolver problemas, la rapidez del personal ante urgencias, el trato recibido por el personal general y la agilidad para la primera cita.

3.5.2.2 Interpretación direccional de los efectos

Posteriormente, se ha llevado a cabo el análisis del **gráfico resumen** de SHAP, el cual permite observar no solo la importancia de cada variable, sino también la dirección de la influencia sobre la probabilidad de pertenecer o no al grupo de alta satisfacción.

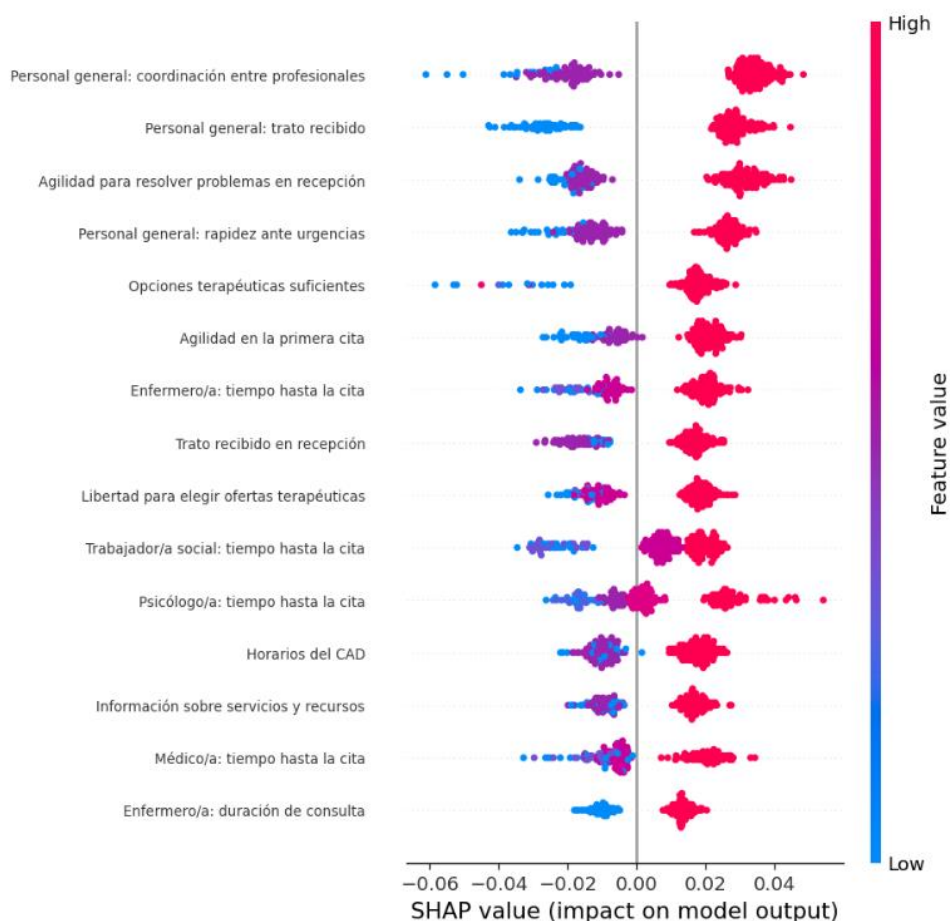


Imagen 31. SHAP Gráfico Resumen para la clase de alta satisfacción. Fuente: elaboración propia mediante Python

Los resultados muestran que valores elevados en las variables mencionadas, contribuyen **positivamente** a la **clasificación** de los usuarios dentro del grupo de alta satisfacción, mientras que las valoraciones bajas reducen dicha probabilidad.

3.5.2.3 Análisis individual de variables relevantes

Para profundizar en la interpretación de los factores más influyentes, se han analizado diferentes **gráficos de dependencia** SHAP de algunas variables más relevantes extraídas.

Los gráficos de dependencia (Imágenes 32, 33 y 34) confirman que valoraciones positivas en suficiencia de opciones terapéuticas, coordinación profesional y ausencia de discriminación generan **contribuciones SHAP positivas**, mientras que las negativas reducen la probabilidad de alta satisfacción.

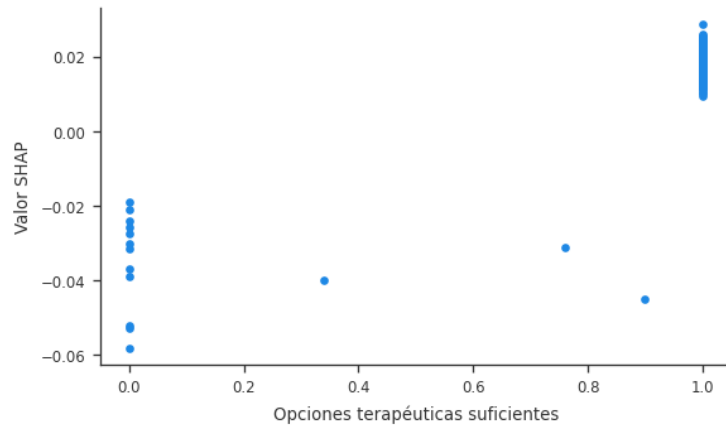


Imagen 32. Influencia de las opciones terapéuticas suficientes. Fuente: elaboración propia mediante Python

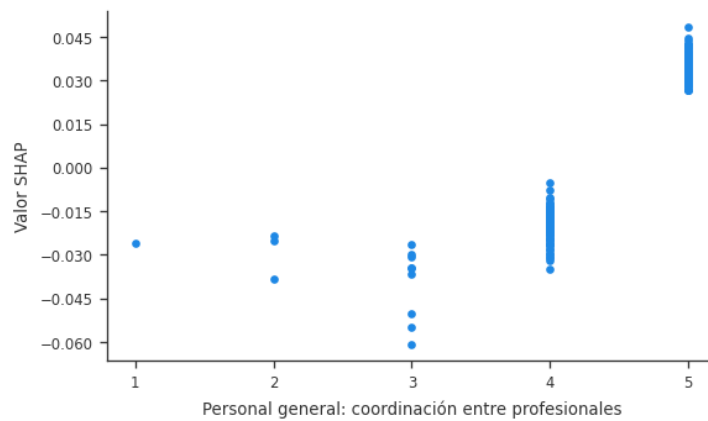


Imagen 33. Influencia de la coordinación entre profesionales. Fuente: elaboración propia mediante Python

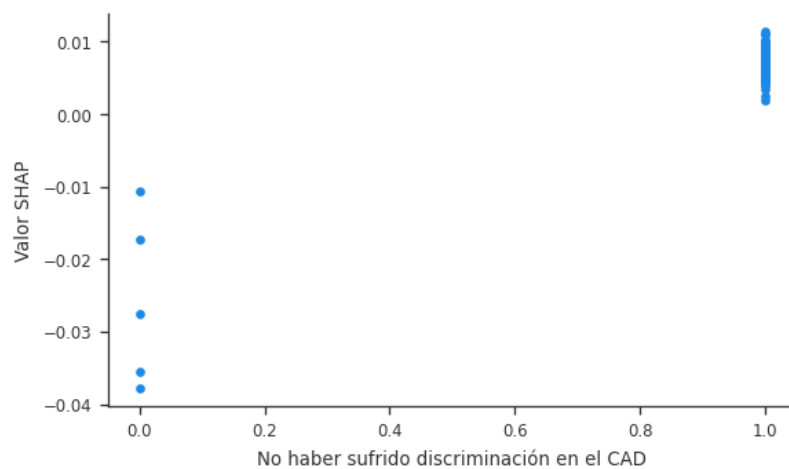


Imagen 34. Influencia de experiencias discriminatorias en el CAD. Fuente: elaboración propia mediante Python

3.5.2.4 Explicación de predicciones

Para finalizar, se ha realizado un gráfico “*Waterfall*”, que permite descomponer concretamente la predicción en las contribuciones individuales de cada variable.

Este gráfico muestra cómo determinados **factores** o variables empujan la **predicción** hacia niveles superiores o inferiores de satisfacción, ayudando a comprender mejor el proceso de decisión seguido por el modelo.

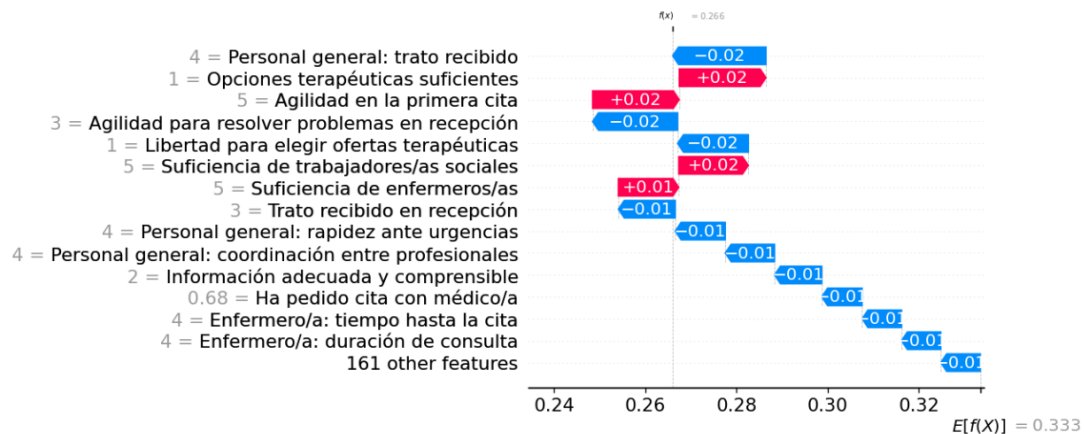


Imagen 35. Explicación individual de una predicción mediante SHAP. Fuente: elaboración propia mediante Python

Los resultados confirman de nuevo la **relevancia** de **variables** relacionadas con la calidad percibida de la atención, la coordinación profesional, la rapidez de acceso al servicio, y la valoración de los recursos terapéuticos del centro.

En conjunto, el análisis **SHAP** permite validar que los principales factores asociados a la satisfacción de los usuarios, están relacionados con los identificados en apartados anteriores, proporcionando una interpretación más transparente y consistente con las predicciones del modelo final.

3.5.3 Reglas de asociación (A-Rules) y patrones explicativos

Como análisis complementario, se ha realizado el análisis de **reglas de asociación** para identificar un conjunto de combinaciones de factores (Agrawal & Srikant, 1994; Fournier-Viger et al., 2017) que presenten mayores niveles de satisfacción.

Para ello, se han elegido algunas de las variables más relevantes, transformándolas en **binarias** (coordinación profesional, trato recibido, rapidez ante urgencias, suficiencia de alternativas terapéuticas y ausencia de experiencias de discriminación).

Los resultados muestran que las valoraciones positivas son bastantes frecuentes; el 98,1% calificó con una nota elevada el **trato recibido**, el 97,5% indica que no ha sufrido discriminaciones en el CAD, el 95,4% valora muy positivamente la coordinación de los profesionales del centro y el 90,5% considera suficientes las opciones terapéuticas.

La **regla** identificada más **relevante** asocia la presencia conjunta de todas estas variables con la pertenencia al grupo más elevado de satisfacción.

Antecedentes	Consecuente	Soporte	Confianza	Lift
Coordinacion_profesional_alta, Primera_cita_agil, Rapidez_urgencias_alta, Opciones_terapeuticas_suficientes, No_discriminacion, Trato_recibido_alto	Alta_satisfaccion	0.5472	0.7381	1.1286

Imagen 36. Regla de asociación identificada más relevante del análisis. Fuente: elaboración propia mediante Python

Esta regla presenta un soporte de 0,5472, indicando que aparece en el 54,7% de los casos y una confianza de 0,738, por lo que, cuando se cumplen conjuntamente estas condiciones, la probabilidad de alta satisfacción es del 73,8%. Finalmente, el lift es de 1,129, indicando una **asociación positiva** entre estas condiciones y la elevada satisfacción.

3.6 Comparación interanual de modelos

3.6.1 Entrenamiento independiente por año

Para analizar la **estabilidad temporal** de los modelos, se ha realizado un entrenamiento independiente para cada año (2022, 2023 y 2024), aplicando la misma metodología que en los modelos globales.

La finalidad es comprobar si la **capacidad predictiva** se mantiene estable en el tiempo o si existen diferencias significativas entre años, lo que podría indicar algún cambio en los factores asociados a la satisfacción.

Año	Modelo	Accuracy	F1-score	ROC-AUC
2022	Random Forest	0.747	0.736	0.833
2022	XGBoost	0.720	0.709	0.839
2023	Random Forest	0.737	0.731	0.844
2023	XGBoost	0.671	0.664	0.767
2024	Random Forest	0.662	0.632	0.710
2024	XGBoost	0.701	0.681	0.701

Imagen 37. Rendimiento predictivo por año y por modelo. Fuente: elaboración propia mediante Python

Los resultados muestran que ambos algoritmos, especialmente el XGBoost, mantienen una **capacidad predictiva adecuada y estable** en los tres años. No obstante, se pueden observar diferencias importantes entre años, sobre todo, en el de Random Forest.

El Random Forest obtiene sus **mejores resultados** en 2022 y 2023, con F1-score superiores al 0,73 y ROC-AUC de 0,833 y 0,844, respectivamente. Sin embargo, en 2024 se produce una caída del F1-score hasta el valor 0,632 y un ROC-AUC de 0,710, sugiriendo una mayor dificultad para identificar los patrones de satisfacción ese año.

Por otra parte, el XGBoost presenta un **comportamiento** bastante más **estable** con respecto al F1-score, aunque con una pequeña disminución progresiva del ROC-AUC, pasando de 0,839 en 2022 a 0,701 en 2024. En conjunto, ambos modelos mantienen una **capacidad predictiva adecuada**, aunque con cierta pérdida de capacidad discriminativa en el último año.

3.6.2 Evolución temporal del rendimiento

La siguiente imagen representa la evolución temporal del F1-score de ambos algoritmos.

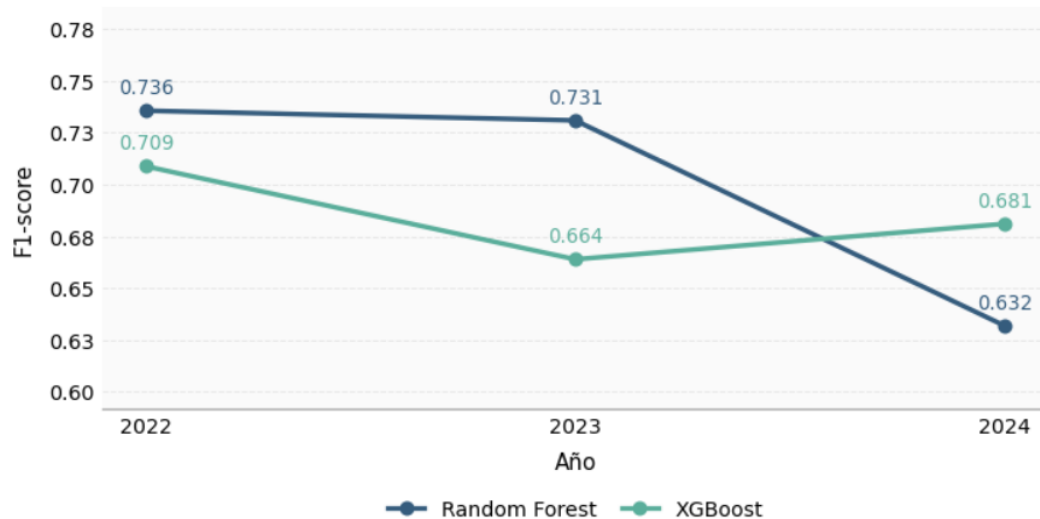


Imagen 38. Comparación interanual del F1-score de los modelos Random Forest y XGBoost.
Fuente: elaboración propia mediante Python

En líneas generales, puede observarse que el modelo de Random Forest presenta mejores resultados que el de XGBoost durante los años 2022 y 2023. Sin embargo, en 2024 se produce la **caída** de la **capacidad predictiva** del Random Forest mencionada anteriormente, lo que provoca que el XGBoost sea el modelo con mejor rendimiento en ese año.

Adicionalmente, se ha representado la **evolución temporal** del ROC-AUC, para visualizar más fácilmente la capacidad discriminativa de los modelos.

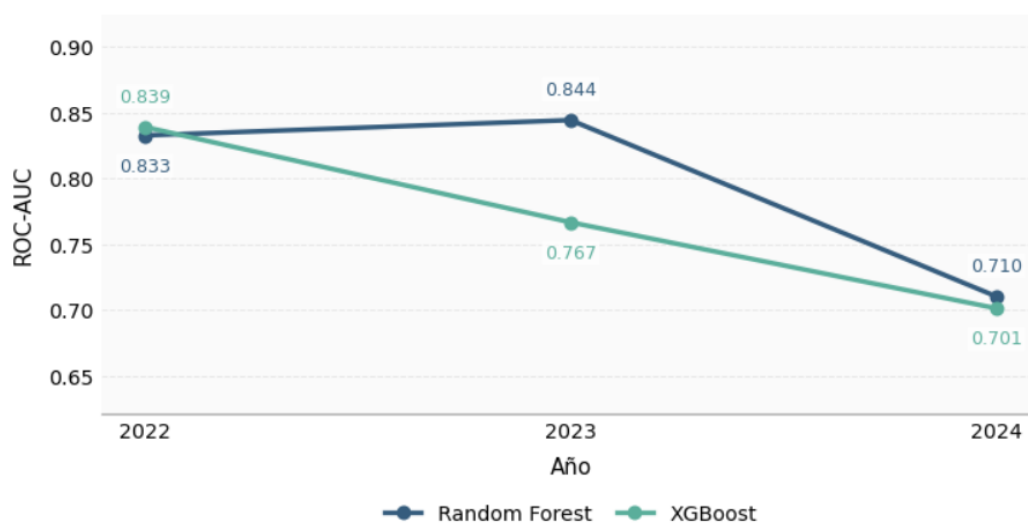


Imagen 39. Comparación interanual del ROC-AUC de los modelos Random Forest y XGBoost.
Fuente: elaboración propia mediante Python

La evolución del ROC-AUC muestra mejor **capacidad discriminativa** en 2022 y 2023, especialmente del modelo de Random Forest en 2023, con un valor de 0,844. No obstante, se puede observar una caída en 2024 en ambos algoritmos, situándose el Random Forest y el XGBoost en valores de 0,710 y 0,701, respectivamente.

Esta caída o **cambio de comportamiento** se puede explicar por algún posible cambio en el perfil de los usuarios o por la relevancia de la que disponen determinados factores predictivos a lo largo del tiempo y que cambia notablemente en el último año analizado.

3.6.3 Estabilidad de variables explicativas

El análisis de **estabilidad** de las variables identifica factores que aparecen de forma recurrente entre los más relevantes de cada modelo independientemente del año.

En el caso del Random Forest, las variables comunes que se han identificado en los tres años han sido:

Nº	Variable común	Descripción
1	P31_1	Personal general: trato recibido
2	P31_2	Personal general: coordinación entre profesionales
3	P31_4	Personal general: rapidez ante urgencias

Imagen 40. Variables comunes entre 2022, 2023 y 2024 en Random Forest. Fuente: elaboración propia mediante Python

Por otra parte, el XGBoost ha identificado solamente una variable como común durante los tres años:

Nº	Variable común	Descripción
1	P31_1	Personal general: trato recibido

Imagen 41. Variables comunes entre 2023 y 2024 en XGBoost. Fuente: elaboración propia mediante Python

Estos resultados indican que la percepción subjetiva del **trato recibido** por el personal general constituye un determinante fundamental de la satisfacción independientemente del año considerado y del modelo utilizado.

Asimismo, la **coordinación** entre profesionales y la **capacidad de respuesta** ante situaciones urgentes aparecen también como elementos de influencia constante sobre la satisfacción.

Destacar, especialmente entre estas variables, la presencia de variables relacionadas con el trato recibido, la coordinación entre profesionales y la capacidad de respuesta del servicio en situaciones urgentes, lo que las constituye en **elementos estructurales** de la satisfacción de los pacientes del CAD.

3.7 Fast Interpretable Greedy-Tree Sums (FIGS)

Para complementar los análisis de interpretabilidad, se ha incorporado un modelo basado en la metodología Fast Interpretable Greedy-Tree Sums (**FIGS**) (Tan et al., 2025), que permite el crecimiento conjunto de varios árboles de decisión, en la que la suma constituye la predicción final.

Este algoritmo combina la capacidad predictiva de los modelos basados en árboles con altos niveles de **interpretabilidad** (Rudin, 2019), permitiendo representar el proceso de clasificación de usuarios mediante un conjunto reducido de reglas de fácil comprensión.

3.7.1 Entrenamiento y evaluación del modelo FIGS

Tras el entrenamiento, el modelo obtuvo un Accuracy del 64,32% y un F1-score ponderado de 0,6537. Aunque son inferiores con respecto a los de Random Forest y XGBoost, son resultados aceptables si se considera que el objetivo principal de este modelo es **maximizar la interpretabilidad** del modelo.

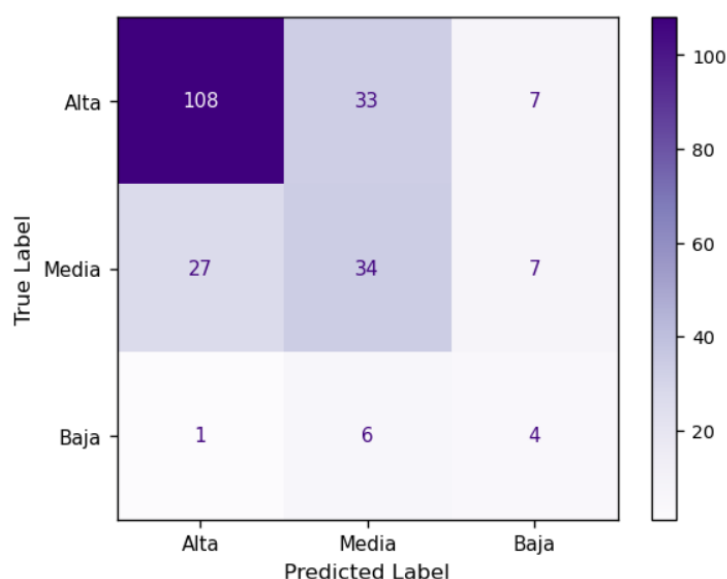


Imagen 42. Matriz de confusión del modelo FIGS. Fuente: elaboración propia mediante Python

La **matriz de confusión** muestra un comportamiento similar al de los modelos anteriores. La categoría de alta satisfacción sigue siendo la clase mejor identificada, mientras que la categoría de baja satisfacción sigue siendo la más difícil de clasificar por su baja frecuencia.

3.7.2 Reglas principales del modelo

Una de las principales **ventajas** de este **algoritmo**, es la posibilidad de representar el proceso de clasificación mediante una serie de reglas explícitas y fácilmente interpretables.

El análisis de la estructura generada por el modelo muestra que la **coordinación** entre **profesionales** aparece como una de las **reglas más relevantes**, siendo uno de los factores con mayor capacidad para diferenciar entre distintos niveles de satisfacción.

Asimismo, otras variables que intervienen en las reglas obtenidas son la ausencia de experiencias de **discriminación**, la valoración de las **opciones terapéuticas** disponibles y el **trato recibido** en la recepción.

3.7.3 Variables más relevantes identificadas por FIGS

La importancia de variables calculada por el modelo **FIGS** ha permitido identificar los factores que más contribuyen a la clasificación del nivel de satisfacción.

Variable	Descripción	Importancia
P31_2	Personal general: coordinación entre profesionales	0.2681
P26_1	Psicólogo/a: tiempo hasta la cita	0.1589
P21_1	Trato recibido en recepción	0.1187
P61_No, nunca	No haber sufrido discriminación en el CAD	0.1108
P23_3	Médico/a: conocimientos y ayuda	0.0806
P59	Opciones terapéuticas suficientes	0.0689
P27	Ha pedido cita con enfermero/a	0.0400
P32_7	Suficiencia de otro personal sanitario	0.0335
P57	Valoración del servicio de orientación laboral	0.0324
P47	Valoración de piso de apoyo a la reinserción	0.0271

Imagen 43. Variables más relevantes identificadas por FIGS. Fuente: elaboración propia mediante Python

Los resultados colocan, de nuevo, a la **coordinación** entre profesionales como la más **relevante** de todo el modelo. Tras ella, aparecen factores relacionados con el tiempo de espera para acceder a la atención psicológica, el trato recibido en recepción, la ausencia de discriminación, la valoración de algunos aspectos médicos y la percepción de disponer de opciones terapéuticas suficientes.

3.7.4 Comparación con Random Forest y XGBoost

Finalmente, se ha realizado una **comparación** del **rendimiento** del FIGS con los dos modelos de referencia en el estudio: el Random Forest y el XGBoost, ambos optimizados.

Modelo	Accuracy	F1-score
Random Forest optimizado	0.7137	0.7001
XGBoost optimizado	0.7004	0.6886
FIGS	0.6432	0.6537

Imagen 44. Comparación de rendimiento entre FIGS, Random Forest y XGBoost. Fuente: elaboración propia mediante Python

El FIGS ha obtenido un **rendimiento** predictivo **inferior** al de los modelos de ensamblado previos. Sin embargo, esa pérdida moderada en la capacidad predictiva se compensa con una **mejora** notable en la **interpretabilidad**, permitiendo explicar las decisiones mediante reglas transparentes.

CAPÍTULO III: DISCUSIÓN DE RESULTADOS

4.1 Principales hallazgos del estudio

Los resultados muestran que la **satisfacción** de los usuarios de los CAD puede explicarse a partir de los factores recogidos en las encuestas de 2022, 2023 y 2024.

Los resultados obtenidos a lo largo del estudio muestran que la satisfacción de los usuarios puede venir explicada a partir del conjunto de factores que se recogen en la encuesta. Entre los **modelos predictivos** desarrollados, los cuales todos han alcanzado niveles de precisión correctos, fueron los modelos basados en árboles de decisión y técnicas de ensamblado los que obtuvieron los mejores resultados (James et al., 2021).

Entre todos los algoritmos analizados, el modelo de **Random Forest** optimizado fue el modelo que alcanzó un mejor rendimiento global, obteniendo un Accuracy del 71,4% y un F1-score de 0,700 (Breiman, 2001). Unos resultados que ponen de manifiesto la capacidad de las técnicas avanzadas de **Machine Learning**, para la identificación de patrones más complejos en las respuestas de los usuarios y de esta manera, poder predecir su nivel de satisfacción.

Con respecto a los resultados, uno de los hallazgos más relevantes que se puede sacar del estudio es, que los factores relacionados con la **calidad** de la **atención recibida**, presentan una influencia y una relevancia mayores que otras características personales, sociodemográficas o de experiencia de tratamiento (Batbaatar et al., 2017; Parasuraman et al., 1988; Al-Abri & Al-Balushi, 2014). En concreto, la coordinación entre profesionales, la rapidez de respuesta ante situaciones urgentes, el trato recibido por parte del personal y la disponibilidad de alternativas terapéuticas suficientes (Tiffany et al., 2012; Sanghani & Moler, 2015), son algunas de las variables que aparecen de forma **recurrente** entre los principales factores que ayudan a **predecir** la satisfacción.

Otro hallazgo relevante es la importancia de la ausencia de experiencias discriminatorias dentro de los CAD. Tanto los análisis de importancia de variables de los modelos como

la interpretación que obtenemos mediante valores **SHAP**, nos muestran que aquellos usuarios que no han percibido situaciones de discriminación suelen presentar mayores niveles de satisfacción con el servicio recibido, y por ello, **mayor probabilidad** de pertenecer al grupo de alta satisfacción.

Dicho hallazgo del estudio pone de manifiesto la importancia de los aspectos relacionales, así como del clima existente dentro del centro, en la experiencia de los pacientes del CAD.

Los análisis realizados de interpretabilidad mediante SHAP values han permitido además identificar y **comprender** la dirección de los **efectos** de los factores más relevantes. En general, usuarios que valoran positivamente la coordinación profesional, el trato recibido, la rapidez de atención y la disponibilidad de opciones terapéuticas tienen una mayor probabilidad de pertenecer al grupo de alta satisfacción. Por el contrario, valoraciones negativas en estos factores **reducen** de forma notoria dicha **probabilidad**.

Por otro lado, el análisis realizado de **reglas de asociación** ha identificado una serie de patrones consistentes en los usuarios más satisfechos. Las reglas obtenidas muestran que las distintas combinaciones entre una adecuada coordinación profesional, una atención rápida, un trato positivo por parte del personal, la ausencia de experiencias discriminatorias y la suficiencia de alternativas terapéuticas se encuentran frecuentemente **asociadas** con **niveles** más **altos** de satisfacción. Unos resultados que refuerzan las conclusiones obtenidas en los modelos predictivos previos.

Finalmente, en la **comparación interanual** realizada entre 2022 y 2024, tanto el rendimiento de los modelos como en los factores explicativos identificados han obtenido una **elevada estabilidad temporal**. En particular, variables relacionadas con el trato recibido, la coordinación entre profesionales y la rapidez de respuesta se encuentran entre los factores más recurrentes y estables en los distintos años y modelos, sugiriendo tratarse de elementos estructurales de la satisfacción de los usuarios de los CAD.

En conjunto, los resultados obtenidos permiten extraer que la satisfacción depende principalmente de la **calidad percibida** y de la **experiencia asistencial** ofrecida por los CAD, más que de otros factores como los individuales o sociodemográficos.

4.2 Factores determinantes de la satisfacción de los usuarios

Una vez identificados los principales hallazgos del estudio, es necesario profundizar en aquellos **factores** que resultan **determinantes** y que presentan una **mayor influencia** sobre la satisfacción de los usuarios de los Centros de Atención a las Drogodependencias (CAD).

Los resultados obtenidos durante el estudio muestran que la satisfacción de los usuarios está estrechamente relacionada con la **calidad percibida** del servicio, asociándose con aspectos organizativos, relacionales y asistenciales que desempeñan un papel fundamental en la valoración global realizada por los usuarios encuestados.

Uno de los factores identificados como más determinantes es la **coordinación** entre profesionales. Esta variable se encuentra entre aquellas con mayor **capacidad predictiva**, reflejando la percepción de los usuarios en la que los distintos especialistas trabajan de forma coordinada durante el tratamiento.

Este resultado manifiesta la importancia que dan los usuarios a que el centro ofrezca un servicio y una **atención integrada y coherente**, sobre todo, allí donde intervienen perfiles profesionales diversos como médicos, psicólogos, enfermeros, terapeutas ocupacionales o trabajadores sociales.

Junto a la coordinación profesional, se encuentran factores también como el **trato recibido** y la **rapidez de respuesta** ante situaciones urgentes. Dichos factores actúan también como elementos clave para explicar la satisfacción de los usuarios.

La cercanía, el respeto y la atención del personal son valorados muy **positivamente** por los usuarios, reforzando como se ha mencionado anteriormente, la importancia de los **aspectos humanos y relacionales**, sobre todo, en este tipo de contextos de salud mental y adicciones, donde la confianza entre el profesional y el paciente resulta fundamental.

Asimismo, la capacidad de responder ágilmente ante situaciones de necesidad o de urgencia, puede interpretarse como una especie de **indicador de accesibilidad** y de **adaptación** del servicio o del tratamiento a las necesidades reales de los pacientes.

Otro factor relevante es la disponibilidad de **opciones terapéuticas** suficientes. Aquellos usuarios, que consideran que el centro ofrece alternativas correctas y adaptadas a sus necesidades, muestran mayores niveles de satisfacción. Por tanto, la existencia de diversos recursos puede favorecer la atención más personalizada, y a su vez, el aumento de la valoración global de los usuarios.

Por otra parte, se encuentra la ausencia de **experiencias de discriminación**, uno de los factores más relevantes identificados durante el estudio. Esta variable indica que quienes no son sometidos a situaciones de discriminación dentro del CAD presentan una mayor probabilidad de situarse en los niveles más elevados de satisfacción (Nordfjærn et al., 2010). Por tanto, pone de manifiesto la gran **importancia** de garantizar **entornos respetuosos** y libres, especialmente en aquellos colectivos que han estado normalmente expuestos a situaciones de exclusión social.

Finalmente, existen algunas actividades complementarias que son ofrecidas por los centros, como la terapia grupal, programas de orientación social o determinados **recursos de apoyo**, que también presentan capacidad explicativa en algunos años, aunque es cierto que su importancia relativa es menor que los factores anteriormente descritos.

4.3 Implicaciones para la gestión de los CAD

Los resultados obtenidos tienen **implicaciones** muy importantes para la gestión de los Centros de Atención a las Drogodependencias (CAD), ya que permiten identificar qué aspectos concretos del servicio generan mayor impacto o es necesario mejorar para incrementar la satisfacción de los usuarios.

En primer lugar, los resultados sugieren, sobre todo, la necesidad de reforzar la coordinación entre los profesionales, lo que resulta ampliamente **recomendable** fomentar nuevos mecanismos de **comunicación** entre trabajadores o mecanismos de **seguimiento** compartido para que así los distintos perfiles profesionales que intervengan estén totalmente informados.

En segundo lugar, la calidad del **trato recibido** por el personal general, que se trata de un elemento fundamental de la experiencia de los pacientes, su mejora es recomendable mediante la implantación de **programas de formación** orientados a mejorar las habilidades de comunicación y empatía o, unas pruebas de selección más duras en el aspecto relacional, lo que puede influir en el aumento de la satisfacción percibida.

Adicionalmente, la rapidez de respuesta y la agilidad para la accesibilidad al servicio también son aspectos **prioritarios de gestión**. Por tanto, reducir los tiempos de espera para la primera cita y, garantizar una atención ágil y contundente ante situaciones de urgencia puede mejorar significativamente la valoración global de los pacientes.

Por otra parte, los resultados también ponen de manifiesto la importancia de ofrecer una amplia variedad de recursos y opciones terapéuticas, con el que los usuarios estén satisfechos, ya que una mejora en estos aspectos puede **favorecer un tratamiento** más adaptado a las necesidades de cada usuario y por ello, una mejora en la percepción del servicio.

Asimismo, la **ausencia de discriminación** en los CAD es un elemento importante, por lo que resulta esencial por parte de los centros, promover entornos inclusivos, respetuosos, y libres de estigmatización, donde todas las personas sean tratadas y puedan recibir un servicio o un trato basado en la igualdad y el respeto.

En conjunto, las **implicaciones** de los resultados señalan que las mejoras en la satisfacción de los usuarios dependen principalmente de acciones a nivel organizativo y asistencial, y que tienen que ser gestionadas directamente por los responsables de los centros.

4.4 Limitaciones del estudio

A pesar de los resultados obtenidos, el estudio presenta algunas **limitaciones** que deben ser consideradas a la hora de interpretar las conclusiones que se realizarán en el siguiente apartado.

En primer lugar, el análisis procede de unas bases de datos procedentes de encuestas de satisfacción, en las que previamente se ha realizado una **homogeneización** de todas ellas y en la que sus respuestas están sujetas a la percepción subjetiva de los propios usuarios. Por tanto, las respuestas pueden estar afectadas por algún **sesgo** de respuesta o por cualquier circunstancia personal existente en el usuario en el momento de cumplimentar el cuestionario.

En segundo lugar, la distribución de la variable objetivo presenta un claro **desbalanceo** de categorías, siendo el grupo de baja satisfacción claramente el minoritario con respecto a los de satisfacción media y alta. Aunque esta situación se ha intentado corregir de forma parcial utilizando técnicas como **SMOTE** durante la modelización (Chawla et al., 2002), sigue siendo una limitación que puede influir de cierta manera en la capacidad predictiva de algunos modelos.

Asimismo, solo se dispone hasta la fecha, únicamente, de las bases de datos hasta el año 2024, por lo que, los resultados obtenidos en el estudio solo reflejan lo analizado en dicho **periodo temporal** y no podemos visualizar hoy en día los posibles cambios futuros a nivel organizacional, los **recursos disponibles** o los perfiles de usuarios atendidos, al no disponer de los datos del 2025 y del primer trimestre del 2026.

Por último, es relevante señalar que hay algunas variables potencialmente relevantes para poder explicar la **satisfacción**, pero que no se encuentran disponibles en la base de datos utilizada, y, por tanto, dichas variables no han podido ser incorporadas al análisis como determinadas características socioeconómicas, clínicas o psicológicas.

No obstante, estas **limitaciones** no invalidan, en ningún caso, los resultados obtenidos, sino que deben considerarse aspectos condicionados por la información que se tiene disponible y, interpretarse como oportunidades para futuras investigaciones y así, poder ampliar y profundizar en un futuro los factores asociados a la satisfacción de los usuarios de los CAD.

CONCLUSIONES Y RECOMENDACIONES

5.1 Conclusiones

Los resultados que se han obtenido durante el estudio permiten concluir que la satisfacción de los usuarios puede explicarse a partir de algunos **factores**, identificados como **relevantes**, que se han recogido en dichas encuestas. Además, los **modelos predictivos** que se han desarrollado han mostrado una capacidad adecuada y estable para clasificar a los usuarios en los distintos niveles de satisfacción, especialmente, aquellos algoritmos basados en técnicas de ensamblado como el Random Forest o el XGBoost.

Aunque, más allá del **rendimiento** de los modelos, el principal valor que se puede extraer del estudio trata de la identificación de los factores que influyen en la experiencia de los usuarios de forma algo más **significativa**. Entre todos ellos, han mostrado que la satisfacción no depende únicamente de la calidad del tratamiento, sino también de otros aspectos tanto organizativos como relacionales vinculados al mejor funcionamiento del centro.

Concretamente, las variables relacionadas con la **coordinación** entre profesionales, el trato recibido, la rapidez de respuesta ante urgencias, la agilidad de acceso al servicio y la disponibilidad de alternativas terapéuticas suficientes, son aquellas que aparecen de manera **recurrente** en los resultados como los **elementos claves** y más relevantes para explicar la pertenencia a los distintos niveles de satisfacción.

Asimismo, la ausencia de experiencias **discriminatorias** emerge también como otro de los factores clave identificados durante el estudio. Por ello, la importancia de promover **entornos respetuosos**, inclusivos y centrados en la salud de los pacientes.

Por otro lado, la comparación de modelos realizada entre los años disponibles ha mostrado una elevada **estabilidad** tanto en el rendimiento de los modelos como en los principales factores explicativos de la satisfacción, sugiriendo que hay determinadas dimensiones que mantienen una influencia constante sobre los usuarios a lo largo del tiempo.

Finalmente, los resultados obtenidos confirman la utilidad de las técnicas de **Business Analytics** y **Machine Learning** como herramientas de apoyo para la evaluación de servicios públicos a través de encuestas, permitiendo de esta manera transformar cantidades enormes de información en outputs útiles que sirvan para la **toma de decisiones** y la **mejora continua** de estos servicios.

5.2 Recomendaciones

A partir de los resultados y conclusiones obtenidas, en este apartado, se van a proponer una serie de **recomendaciones** orientadas a mejorar la experiencia de los pacientes y la calidad de tratamiento y servicios prestados por los CAD.

En primer lugar, se recomienda principalmente reforzar los mecanismos de coordinación entre los distintos profesionales del centro y que intervienen en el tratamiento de los usuarios, lo que puede favorecer una **atención** más **completa, integral y coherente** para los pacientes.

En segundo lugar, resulta importante y conveniente que el centro comience a impulsar **estrategias** o talleres centrados en la **mejora** de la calidad del trato y una mejor **atención**, fomentando el progreso hacia la comunicación afectiva, la empatía y el respeto hacia los pacientes.

También, los centros deberían prestar especial atención a la **reducción** de los **tiempos** de espera, y a mejorar el acceso al servicio y a los recursos del centro, ya que de esta manera se mejoraría adicionalmente, la rapidez de respuesta antes situaciones urgentes.

Por último, un aspecto bastante relevante es seguir manteniendo o incluso ampliar la **diversidad de recursos terapéuticos** que ofrece el centro, favoreciendo de esta manera un mejor itinerario o tratamiento adaptado a las necesidades específicas de cada usuario.

Por tanto, la incorporación de las herramientas de análisis de datos utilizadas en este estudio puede ayudar a monitorizar de forma anual o continua la satisfacción de los usuarios y, contribuir además con la toma de decisiones basadas en los resultados obtenidos para la mejora de los servicios de los CAD.

FUTURAS LÍNEAS DE INVESTIGACIÓN

Este estudio realizado abre diversas líneas de **futuras investigaciones** que permitirían ampliar y analizar en mayor profundidad el conocimiento de los factores asociados a la satisfacción de los usuarios de los Centros de Atención a las Drogodependencias (CAD).

En primer lugar, sería realmente interesante la incorporación de **nuevas fuentes** de información que complementen a las encuestas, entre ellas, variables clínicas, socioeconómicas o incluso algunos factores relacionados con la evolución terapéutica de los pacientes. Esto, permitiría realizar **modelos predictivos** más completos y así poder mejorar la capacidad predictiva de los resultados obtenidos, ya que hasta ahora sigue habiendo varios factores no observados.

En segundo lugar, asumiendo que estuvieran disponibles las encuestas del 2025 y primer trimestre del 2026, futuras investigaciones podrían ampliar el **horizonte temporal** del análisis no solo anteriores al 2022, sino también incorporando nuevas ediciones de encuestas del último año, e incluso del trimestre anterior (enero a marzo 2026). A partir de ahí, se facilitaría aún más el análisis de **tendencias a largo plazo** y poder evaluar los cambios en la satisfacción de los usuarios a lo largo del tiempo.

Por otra parte, en este trabajo se ha llevado a cabo el análisis de los factores asociados con una elevada **satisfacción** de los usuarios, por lo que podría ser realmente interesante profundizar en el **análisis** de estos factores de los usuarios con **niveles bajos** de satisfacción.

Es cierto que se dispone de una **clasificación** bastante **desbalanceada**, pero con la aplicación de técnicas específicas se podría identificar también con precisión los factores más asociados con este grupo y, de esta manera, realizar una **toma de decisiones** e intervenciones para mejorar su experiencia en el servicio del CAD.

Finalmente, la metodología usada en este trabajo puede aplicarse a otros servicios sanitarios o sociosanitarios, lo que permitiría comparar resultados y sobre todo, analizar

y evaluar la utilidad de las técnicas y algoritmos de Business Analytics como apoyo para comprender mejor las áreas susceptibles de mejora y elevar la calidad de los servicios públicos.

Declaración de Uso de Herramientas de Inteligencia Artificial Generativa en Trabajos Fin de Grado

ADVERTENCIA: Desde la Universidad consideramos que ChatGPT u otras herramientas similares son herramientas muy útiles en la vida académica, aunque su uso queda siempre bajo la responsabilidad del alumno, puesto que las respuestas que proporciona pueden no ser veraces. En este sentido, NO está permitido su uso en la elaboración del Trabajo fin de Grado para generar código porque estas herramientas no son fiables en esa tarea. Aunque el código funcione, no hay garantías de que metodológicamente sea correcto, y es altamente probable que no lo sea.

Por la presente, yo, Adrián Fernández Figueroa, estudiante de ADE + Business Analytics de la Universidad Pontificia Comillas al presentar mi Trabajo Fin de Grado titulado “Modelos predictivos de satisfacción en los programas de tratamiento de adicciones (2022-2024)”, declaro que he utilizado la herramienta de Inteligencia Artificial Generativa ChatGPT u otras similares de IAG de código sólo en el contexto de las actividades descritas a continuación [el alumno debe mantener solo aquellas en las que se ha usado ChatGPT o similares y borrar el resto. Si no se ha usado ninguna, borrar todas y escribir “no he usado ninguna”]:

1. Constructor de plantillas: Para diseñar formatos específicos para secciones del trabajo.
2. Sintetizador y divulgador de libros complicados: Para resumir y comprender literatura compleja.
3. Revisor: Para recibir sugerencias sobre cómo mejorar y perfeccionar el trabajo con diferentes niveles de exigencia.
4. Corrector de estilo literario y de lenguaje: Para mejorar la calidad lingüística y estilística del texto.
5. Referencias: Usado conjuntamente con otras herramientas, como Science, para identificar referencias preliminares que luego he contrastado y validado.
6. Traductor: Para traducir textos de un lenguaje a otro.

Afirmo que toda la información y contenido presentados en este trabajo son producto de mi investigación y esfuerzo individual, excepto donde se ha indicado lo contrario y se han dado los créditos correspondientes (he incluido las referencias adecuadas en el TFG y he explicitado para que se ha usado ChatGPT u otras herramientas similares). Soy consciente de las implicaciones académicas y éticas de presentar un trabajo no original y acepto las consecuencias de cualquier violación a esta declaración.

Fecha: 4 Junio de 2026

Firma: Adrián Fernández Figueroa

BIBLIOGRAFÍA

- Acion, L., Kelmansky, D., van der Laan, M., Sahker, E., Jones, D., & Arndt, S. (2017). Use of a machine learning framework to predict substance use disorder treatment success. *PLOS ONE*, 12(4), e0175383.
- Agrawal, R., & Srikant, R. (1994). Fast algorithms for mining association rules in large databases. *Proceedings of the 20th International Conference on Very Large Data Bases (VLDB)*, 487–499.
- Al-Abri, R., & Al-Balushi, M. (2014). Patient satisfaction survey as a tool towards quality improvement. *Oman Medical Journal*, 29(1), 3–7.
- Batbaatar, E., Dorjdagva, J., Luvsannyam, A., Savino, M. M., & Amenta, P. (2017). Determinants of patient satisfaction: A systematic review. *Perspectives in Public Health*, 137(2), 89–101.
- Bergstra, J., & Bengio, Y. (2012). Random search for hyper-parameter optimization. *Journal of Machine Learning Research*, 13, 281–305.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic Minority Over-sampling Technique. *Journal of Artificial Intelligence Research*, 16, 321–357.
- 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.
- Chung, J., & Teo, J. (2023). Single classifier vs. ensemble machine learning approaches for mental health prediction. *Brain Informatics*, 10(1), 1.
- Davenport, T. H., & Harris, J. G. (2007). *Competing on analytics: The new science of winning*. Harvard Business Press.
- Delegación del Gobierno para el Plan Nacional sobre Drogas. (2018). *Estrategia Nacional sobre Adicciones 2017–2024*. Ministerio de Sanidad, Consumo y Bienestar Social.

- Donabedian, A. (1988). The quality of care: How can it be assessed? *JAMA*, 260(12), 1743–1748.
- Fawcett, T. (2006). An introduction to ROC analysis. *Pattern Recognition Letters*, 27(8), 861–874.
- Fournier-Viger, P., Lin, J. C.-W., Vo, B., Chi, T. T., Zhang, J., & Le, H. B. (2017). A survey of itemset mining. *WIREs Data Mining and Knowledge Discovery*, 7(4), e1207.
- Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, 29(5), 1189–1232.
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2nd ed.). Springer.
- Hosmer, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied logistic regression* (3rd ed.). Wiley.
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An introduction to statistical learning: With applications in R* (2nd ed.). Springer.
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444.
- Liu, N., Kumara, S., & Reich, E. (2021). Gaining insights into patient satisfaction through interpretable machine learning. *IEEE Journal of Biomedical and Health Informatics*, 25(6), 2215–2226.
- Loh, W.-Y. (2011). Classification and regression trees. *WIREs Data Mining and Knowledge Discovery*, 1(1), 14–23.
- Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems (NeurIPS)*, 30, 4765–4774.
- Lundberg, S. M., Erion, G., Chen, H., DeGrave, A., Prutkin, J. M., Nair, B., Katz, R., Himmelfarb, J., Bansal, N., & Lee, S.-I. (2020). From local explanations to global understanding with explainable AI for trees. *Nature Machine Intelligence*, 2(1), 56–67.

- McKinney, W. (2010). Data structures for statistical computing in Python. Proceedings of the 9th Python in Science Conference, 51–56.
- McLellan, A. T., & Hunkeler, E. (1998). Patient satisfaction and outcomes in alcohol and drug abuse treatment. *Psychiatric Services*, 49(5), 573–575.
- Ministerio de Sanidad. (2022). Encuesta sobre alcohol y otras drogas en España, EDADES 2020. Delegación del Gobierno para el Plan Nacional sobre Drogas.
- Mira, J. J., & Aranaz, J. (2000). La satisfacción del paciente como una medida del resultado de la atención sanitaria. *Medicina Clínica*, 114(Supl 3), 26–33.
- Mira Solves, J. J., Aranaz Andrés, J. M., Rodríguez-Marín, J., Buil, J. A., Castell, M., & Vitaller, J. (1998). SERVQHOS: Un cuestionario para evaluar la calidad percibida de la asistencia hospitalaria. *Medicina Preventiva*, 4, 12–18.
- Molnar, C. (2022). *Interpretable machine learning: A guide for making black box models explainable* (2nd ed.). Christoph Molnar.
- Nordfjærn, T., Rundmo, T., & Hole, R. (2010). Treatment and recovery as perceived by patients with substance addiction. *Journal of Psychiatric and Mental Health Nursing*, 17(1), 46–64.
- Obermeyer, Z., & Emanuel, E. J. (2016). Predicting the future—Big data, machine learning, and clinical medicine. *New England Journal of Medicine*, 375(13), 1216–1219.
- Observatorio Español de las Drogas y las Adicciones (OEDA). (2023). Informe sobre alcohol, tabaco y drogas ilegales en España. Ministerio de Sanidad.
- Parasuraman, A., Zeithaml, V. A., & Berry, L. L. (1988). SERVQUAL: A multiple-item scale for measuring consumer perceptions of service quality. *Journal of Retailing*, 64(1), 12–40.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., & Duchesnay, E. (2011). Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12, 2825–2830.

- Powers, D. M. W. (2011). Evaluation: From precision, recall and F-measure to ROC, informedness, markedness and correlation. *Journal of Machine Learning Technologies*, 2(1), 37–63.
- Provost, F., & Fawcett, T. (2013). *Data science for business: What you need to know about data mining and data-analytic thinking*. O'Reilly Media.
- Quinlan, J. R. (1986). Induction of decision trees. *Machine Learning*, 1(1), 81–106.
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?": Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144.
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5), 206–215.
- Rumelhart, D. E., Hinton, G. E., & Williams, R. J. (1986). Learning representations by back-propagating errors. *Nature*, 323(6088), 533–536.
- Saloner, B., Stoller, K. B., & Alexander, G. C. (2022). Moving addiction care to the mainstream: Improving the quality of buprenorphine treatment. *New England Journal of Medicine*, 379(1), 4–6.
- Sanghani, R. M., & Moler, A. K. (2015). Improving consumer satisfaction with addiction treatment: An analysis of alumni preferences. *Journal of Addiction*, 2015, 509864.
- Saturno, P. J., Imperatori, E., & Corbella, A. (1990). Evaluación de la calidad asistencial en atención primaria. Ministerio de Sanidad y Consumo.
- Shickel, B., Tighe, P. J., Bihorac, A., & Rashidi, P. (2018). Deep EHR: A survey of recent advances in deep learning techniques for electronic health record (EHR) analysis. *IEEE Journal of Biomedical and Health Informatics*, 22(5), 1589–1604.
- Simsekler, M. C. E., Qazi, A., Ellahham, S., & Ozonze, O. (2021). Exploring drivers of patient satisfaction using a random forest algorithm. *BMC Medical Informatics and Decision Making*, 21(1), 1–11.
- Sitzia, J., & Wood, N. (1997). Patient satisfaction: A review of issues and concepts. *Social Science & Medicine*, 45(12), 1829–1843.

- Stallvik, M., Flemmen, G., Salthammer, J. A., & Nordfjærn, T. (2020). Factors predicting satisfaction in outpatient substance abuse treatment: A prospective follow-up study. *Substance Abuse Treatment, Prevention, and Policy*, 15(1), 1–11.
- Stone, M. (1974). Cross-validatory choice and assessment of statistical predictions. *Journal of the Royal Statistical Society: Series B*, 36(2), 111–133.
- Tan, Y. S., Singh, C., Nasser, K., Agarwal, A., & Yu, B. (2025). Fast Interpretable Greedy-Tree Sums. *Proceedings of the National Academy of Sciences*, 122(7), e2310151122.
- Tiffany, S. T., Friedman, L., Greenfield, S. F., Hasin, D. S., & Jackson, R. (2012). Beyond drug use: A systematic consideration of other outcomes in evaluations of treatments for substance use disorders. *Addiction*, 107(4), 709–718.
- Topol, E. J. (2019). High-performance medicine: The convergence of human and artificial intelligence. *Nature Medicine*, 25(1), 44–56.
- van Buuren, S., & Groothuis-Oudshoorn, K. (2011). mice: Multivariate imputation by chained equations in R. *Journal of Statistical Software*, 45(3), 1–67.
- Vapnik, V. N. (1995). *The nature of statistical learning theory*. Springer.
- Zhang, S., Chen, J. Y. Q., Pang, H. N., Lo, N. N., Yeo, S. J., & Liow, M. H. L. (2021). Development and internal validation of machine learning algorithms to predict patient satisfaction after total hip arthroplasty. *Journal of Orthopaedics, Surgery and Research*, 16, 511.

ANEXO

A.1. Armonización de la base de 2022 a la nomenclatura común

```
rename_2022 = {
    "Q01": "CAD", "Q02": "SEXO", "Q03": "EDAD", "Q04": "ESTUDIOS", "Q05":
    "SITUACION", "Q05_other": "P5_OTROS", "Q10_01": "P10_1", "Q10_02": "P10_2",
    "Q10_03": "P10_3", "Q10_04": "P10_4", "Q10_05": "P10_5", "Q10_06": "P10_6",
    "Q10_07": "P10_7", "Q10_08": "P10_8", "Q10_09": "P10_9", "Q10_99": "P10_99",
    "Q36": "P36", "Q37": "P36A", "Q38": "P37", "Q39": "P39A", "Q40":
    "P39_LUGAR_METADONA_2022", "Q40_other": "P39_LUGAR_METADONA_OTROS_2022",
    "Q41_Q4101": "P41_DISTANCIA_METADONA_2022", "Q41_Q4102": "P41", "Q63": "P63",
    "Q64": "P64", "Q65": "P65", # ... (resto de variables del cuestionario)}

df_2022 = df_2022.rename(columns=rename_2022)
```

A.2. Outliers: umbrales y Rango Inter cuartílico (duración del tratamiento)

```
df_model.loc[df_model["P9_num"] > 480, "P9_num"] = np.nan
df_model.loc[df_model["P9_B_num"] > 40, "P9_B_num"] = np.nan

for col in ["P9_num", "P9_B_num"]:
    serie = df_model[col].dropna()
    Q1 = serie.quantile(0.25)
    Q3 = serie.quantile(0.75)
    IQR = Q3 - Q1
    lower = Q1 - 1.5 * IQR
    upper = Q3 + 1.5 * IQR
    outliers = serie[(serie < lower) | (serie > upper)]
```

A.3. Creación de variables derivadas

```
df_model["tiempo_total_meses"] = df_model["P9_num"]
df_model["tiempo_total_meses"] = df_model["tiempo_total_meses"]
    .fillna(df_model["P9_B_num"] * 12))

df_model["tratamiento_largo"] = (df_model["tiempo_total_meses"] > 24
    ).astype(int)
```

A.4. Recodificación de escalas ordinales (Likert) y binarias

```
mapa_likert = {"Muy mal": 1, "Malo": 2, "Mal": 2, "Regular": 3, "Normal": 3,
    "Bien": 4, "Bueno": 4, "Fácil": 4, "Muy bien": 5, "Muy bueno": 5, "Muy fácil":
    5, "Difícil": 2, "Muy difícil": 1, "Nada de acuerdo": 1, "Poco de acuerdo":
```

```
2, "De acuerdo": 4, "Muy de acuerdo": 5, "No sabe/No contesta": np.nan, "No sabe / no contesta": np.nan}
```

```
for col in ordinales:
```

```
    df_model[col] = df_model[col].replace(mapa_likert)
```

```
    df_model[col] = pd.to_numeric(df_model[col], errors="coerce")
```

```
mapa_binario = {"Si": 1, "Si": 1, "No": 0, "No sabe/No contesta": np.nan, "No sabe / no contesta": np.nan}
```

```
for col in variables_binarias:
```

```
    if col in df_model.columns:
```

```
        df_model[col] = df_model[col].replace(mapa_binario)
```

A.5. Imputación de valores faltantes (Multiple Imputation)

```
from sklearn.experimental import enable_iterative_imputer
```

```
from sklearn.impute import IterativeImputer, SimpleImputer
```

```
from sklearn.ensemble import RandomForestRegressor
```

```
# Categóricas:
```

```
imputer_cat = SimpleImputer(strategy="most_frequent")
```

```
df_model[cat_cols] = imputer_cat.fit_transform(df_model[cat_cols])
```

```
# Numéricas:
```

```
imputer_num = IterativeImputer(
```

```
    estimator=RandomForestRegressor(
```

```
        n_estimators=50,
```

```
        random_state=42,
```

```
        n_jobs=-1
```

```
    ),
```

```
    max_iter=10,
```

```
    initial_strategy="median",
```

```
    random_state=42,
```

```
    skip_complete=True
```

```
)
```

```
df_model[num_cols] = imputer_num.fit_transform(df_model[num_cols])
```

A.6. Construcción de la variable objetivo (tres niveles)

```
def recodificar_satisfaccion(x):
```

```

        if x in [1, 2, 3]: return "Baja"
        elif x == 4:      return "Media"
        elif x == 5:      return "Alta"
        else:             return np.nan

df_model["target_satisfaccion"] = (
    df_model["P13"]
    .apply(recodificar_satisfaccion)
)

```

A.7. Asociación categórica: Chi-cuadrado y V de Cramér

```

from scipy.stats import chi2_contingency

for var in cat_vars:
    tabla = pd.crosstab(X[var], y)
    chi2, p, dof, expected = chi2_contingency(tabla)
    n = tabla.sum().sum()
    cramers_v = np.sqrt(chi2 / (n * (min(tabla.shape)-1)))

```

A.8. Tratamiento del desbalanceo (SMOTE, solo en entrenamiento)

```

from imblearn.over_sampling import SMOTE

smote = SMOTE(random_state=42)

X_train_smote, y_train_smote = smote.fit_resample(X_train, y_train)

```

A.9. Configuración de los modelos supervisados base

```

# Regresión Logística
log_model = LogisticRegression(solver="lbfgs",
    max_iter=5000,
    random_state=42
)

# Árbol de decision
tree_model = DecisionTreeClassifier(
    max_depth=5,
    random_state=42
)

```

```

)

# Random Forest Base
rf_model = RandomForestClassifier(
    n_estimators=300,
    max_depth=12,
    min_samples_split=10,
    min_samples_leaf=5,
    random_state=42,
    class_weight="balanced"
)

# XGBoost Base
xgb_model = XGBClassifier(
    n_estimators=300,
    max_depth=6,
    learning_rate=0.05,
    subsample=0.8,
    colsample_bytree=0.8,
    random_state=42,
    eval_metric="mlogloss"
)

# Red Neuronal (MLP)
mlp_model = MLPClassifier(
    hidden_layer_sizes=(128, 64),
    activation="relu",
    solver="adam",
    alpha=0.0001,
    learning_rate_init=0.001,
    max_iter=500,
    random_state=42
)

```

A.10. Validación cruzada y optimización de hiperparámetros

```

cv = StratifiedKFold(
    n_splits=5,
    shuffle=True,

```

```

        random_state=42
    )

# Random Forest
param_rf = {
    "n_estimators": [200, 300, 500],
    "max_depth": [5, 10, 15, 20],
    "min_samples_split": [2, 5, 10],
    "min_samples_leaf": [1, 2, 4]
}

rf_search = RandomizedSearchCV(
    estimator=RandomForestClassifier(
        random_state=42,
        class_weight="balanced"
    ),
    param_distributions=param_rf,
    n_iter=15,
    scoring="f1_weighted",
    cv=cv,
    verbose=1,
    random_state=42,
    n_jobs=-1
)

# XGBoost
param_xgb = {
    "n_estimators": [100, 200, 300],
    "max_depth": [3, 5, 7],
    "learning_rate": [0.01, 0.05, 0.1],
    "subsample": [0.8, 1.0]
}

xgb_search = RandomizedSearchCV(
    estimator=XGBClassifier(
        random_state=42,
        eval_metric="mlogloss"
    ),

```

```

    param_distributions=param_xgb,
    n_iter=15,
    scoring="f1_weighted",
    cv=cv,
    verbose=1,
    random_state=42,
    n_jobs=-1
)

```

A.11. Interpretabilidad con SHAP (clase Alta Satisfacción)

```

import shap

explainer = shap.TreeExplainer(best_rf)
shap_values_raw = explainer.shap_values(X_test)

clase_interpretar = "Alta"
idx_clase = list(modelo_shap.classes_).index(clase_interpretar)

if isinstance(shap_values_raw, list):
    shap_values_clase = shap_values_raw[idx_clase]
else:
    shap_values_clase = shap_values_raw[:, :, idx_clase]

```

A.12. Reglas de Asociación (variables binarias + algoritmo Apriori)

```

from mlxtend.frequent_patterns import apriori, association_rules

rules_data = pd.DataFrame()

# Alta satisfacción
rules_data["Alta_satisfaccion"] = (df_rules["target_satisfaccion"] == "Alta"
).astype(int)

# Opciones terapéuticas suficientes
rules_data["Opciones_terapeuticas_suficientes"] = (df_rules["P59"] == 1
).astype(int)

# No haber sufrido discriminación

```

```

rules_data["No_discriminacion"] = (df_rules["P61"] == "No, nunca")
                                ).astype(int)

# Coordinación profesional alta
rules_data["Coordinacion_profesional_alta"] = (df_rules["P31_2"] >= 4)
                                                ).astype(int)

# Trato recibido alto
rules_data["Trato_recibido_alto"] = (df_rules["P31_1"] >= 4)
                                     ).astype(int)

# Rapidez ante urgencias alta
rules_data["Rapidez_urgencias_alta"] = (df_rules["P31_4"] >= 4)
                                         ).astype(int)

# Primera cita ágil
rules_data["Primera_cita_agil"] = (df_rules["P16"] >= 4)
                                   ).astype(int)

frequent = apriori(
    rules_data.astype(bool),
    min_support=0.20,
    use_colnames=True
)

rules = association_rules(
    frequent,
    metric="confidence",
    min_threshold=0.70
)

```

A.13. Modelo FIGS

```

from imodels import FIGSClassifier

figs_model = FIGSClassifier(
    max_rules=20,
    max_trees=5,
    max_depth=3,

```

```
        random_state=42
    )

    figs_model.fit(
        X_train_smote,
        y_train_figs,
        feature_names=list(X_train.columns)
    )
```