



FACULTAD DE FARMACIA

Grado en Farmacia

REVISIÓN SISTEMÁTICA SOBRE LA APLICACIÓN DE MACHINE LEARNING EN FARMACOCINÉTICA SOBRE POBLACIÓN PEDIÁTRICA

Memoria de Trabajo Fin de Grado

Sant Joan d'Alacant

Junio 2025

Autor: Sergio Medina Moya
Modalidad: Revisión bibliográfica
Tutor: Amelia Ramón López

ÍNDICE

Tabla de contenido

1. RESUMEN:.....	2
2. INTRODUCCIÓN:.....	3
2.1. FARMACOCINÉTICA:.....	4
2.2. FARMACOCINÉTICA POBLACIONAL:.....	6
2.3. MACHINE LEARNING:.....	9
3. OBJETIVO:.....	11
4. MATERIALES Y MÉTODOS:.....	11
4.1. Diseño:.....	11
4.2. Fuente de la obtención de datos:.....	11
4.3. Tratamiento de la información:.....	11
4.4. Selección final de artículos:.....	12
5. RESULTADOS:.....	13
6. DISCUSIÓN:.....	16
7. CONCLUSIONES:.....	39
8. BIBLIOGRAFÍA:.....	40

1.RESUMEN:

Introducción: La farmacocinética es una disciplina clave para personalizar y optimizar los tratamientos farmacológicos de los pacientes. Este enfoque resulta especialmente útil en aquellas poblaciones con una alta variabilidad interindividual como la población pediátrica. Con el avance de la tecnología, se han desarrollado nuevas técnicas prometedoras como el caso del “machine learning” el cual, se posiciona como un claro sucesor a los modelos farmacocinéticos poblacionales clásicos. Por este motivo, es importante analizar el impacto de dicha tecnología en el ámbito farmacocinético.

Objetivos: Realizar una búsqueda sistemática sobre la aplicación de “machine learning” en farmacocinética en población pediátrica, mediante la búsqueda de artículos en las principales bases de datos científicas.

Material y métodos: Realizar una revisión sistemática de los artículos recuperados en bases de datos científicas MEDLINE (PubMed) y Scopus. Se generó una ecuación de búsqueda con los descriptores “Machine learning”, “Pharmacokinetics” y “Drug monitoring” aplicando el filtro de edad (<18 años).

Resultados: Tras obtener un total de 173 artículos tras la revisión, solamente 7 estudios fueron aceptados para su revisión y análisis. Los artículos incluidos cumplían con los objetivos del presente TFG, redactados en inglés y recuperados en texto completo.

Conclusiones: La aplicación de *Machine Learning* en farmacocinética pediátrica se postula como la evolución natural a los modelos clásicos debido a su mejora en la eficacia y precisión para ajustar dosis. Pese a esto se necesita una mayor validación clínica y desarrollo para su implementación final.

2.INTRODUCCIÓN:

A lo largo de la historia, el avance en la tecnología ha supuesto grandes cambios en multitud de áreas, incluida la farmacología. Este trabajo analiza el papel del *machine learning* (ML) y cuál es su impacto en farmacocinética y farmacometría, disciplinas clave en la optimización del tratamiento farmacológico.

Estas dos disciplinas permiten determinar con exactitud las dosis de fármacos, interacciones entre fármacos y otros parámetros clave, permitiendo conseguir un tratamiento farmacológico personalizado.

La **farmacocinética** es la ciencia que describe cómo un fármaco se comporta en el organismo (procesos de liberación, absorción, distribución metabolismo y expresión (LADME)). Por otra parte, la farmacodinamia es la ciencia que se encarga de estudiar el conjunto de procesos que describen el efecto que el fármaco provoca en el organismo.

En el contexto de este trabajo, la **farmacometría** es la ciencia encargada de implementar modelos matemáticos que permiten cuantificar los efectos de la farmacocinética y farmacodinamia en una población grande de pacientes y así poder definir y simular aquellos parámetros que permitan conocer cómo se comportan los fármacos en el organismo en la práctica clínica.

Esto permite alcanzar una terapia personalizada con una mejor eficacia y seguridad, ajustando los regímenes de dosificación de forma más efectiva para cada paciente (1)

La aplicación de la farmacocinética resulta clave en poblaciones donde la variabilidad a la respuesta del fármaco es muy alta como es el caso de la **población pediátrica**.

Esta población requiere un amplio conocimiento de las variables farmacocinéticas (LADME), puesto que estas se ven modificadas considerablemente con la edad, el grado de maduración del organismo y con el desarrollo general del paciente.

Debido a esto, muchos fármacos pueden variar su efecto en esta población aun cuando se haya calculado correctamente la dosis y los parámetros necesarios, observándose grandes diferencias con la misma terapia en adultos.

Ciertos factores, como la inmadurez de las enzimas hepáticas, la unión a proteínas plasmáticas u otros procesos, están limitados por la edad. Esto juega un papel clave farmacocinética y, por lo tanto, en la eficacia y seguridad de la terapia. (2)

Sin embargo, debido al aumento en el número de datos y su dificultad para analizar cómo en el caso de la población pediátrica, ha provocado la necesidad de implementar nuevas técnicas como el **machine learning** con el objetivo de mejorar este ámbito de la farmacología.

2.1. FARMACOCINÉTICA:

La farmacocinética estudia la evolución de las concentraciones de un fármaco en el organismo a lo largo del tiempo, permitiendo obtener un modelo matemático que permita optimizar la seguridad y eficacia de los tratamientos (3)

Los procesos que componen la farmacocinética se resumen con el acrónimo LADME: **liberación, absorción, distribución, metabolismo y excreción** del fármaco, es decir, todas las etapas desde la entrada del fármaco en el organismo hasta su eliminación.

- **Liberación:** paso del fármaco desde su forma farmacéutica hasta el lugar de absorción. (4)
- **Absorción** paso del fármaco desde el sitio de administración hasta la circulación sistémica. En el proceso de absorción influyen las propiedades fisicoquímicas del fármaco, la forma farmacéutica, características del lugar de absorción y vía de administración. (3)
- **Distribución:** dispersión o distribución del fármaco por la circulación sistémica generando un gradiente de concentración que favorece el paso al fármaco desde el espacio intravascular hacia el extravascular. Este proceso se ve influenciado por factores como el flujo sanguíneo (a mayor flujo sanguíneo el fármaco se distribuye más rápido), unión a proteínas

plasmáticas (limita el paso a los tejidos), así como el tamaño, propiedades de solubilidad y grado de ionización del fármaco. (3)

- **Metabolismo:** biotransformación del fármaco en estructuras más simples para su posterior eliminación, principalmente en el hígado puesto que las enzimas encargadas de este proceso se encuentran allí. Este proceso se ve influenciado por factores como la genética de la persona, enfermedades y posibles interacciones farmacológicas. (3)
- **Excreción:** proceso de eliminación de los fármacos o productos metabolizados, las principales vías son a través de los riñones e hígado. Pese a esto, algunos productos pueden eliminarse por vías secundarias como los pulmones, leche materna o sudor. (3)

En la población pediátrica, estos procesos se ven modificados debido a la inmadurez fisiológica propia de esta población, afectado a:

- Absorción: afectada por diferencias en el pH, motilidad intestinal de esta población.
- Distribución: mayor cantidad de agua corporal y extracelular, cambios en la unión a proteínas plasmáticas.
- Metabolismo: Menor concentración de enzimas encargadas de la biotransformación.
- Excreción: los mecanismos de excreción no se encuentran completamente desarrollados en esta población

Todos estos factores hacen imprescindible ajustar la dosis en niños adecuadamente para evitar toxicidad o ineficacia (2). La farmacocinética permite gracias al control de los procesos LADME, determinar la dosis óptima, la frecuencia de administración y la selección de la vía de administración correctamente. Esto permite optimizar el tratamiento de un paciente, alcanzando el objetivo clínico y evitando el riesgo de reacciones adversas y/o toxicidad.

Pese a esto, el enfoque individualizado de la farmacocinética clásica presenta limitaciones especialmente cuando no se puede evaluar a cada paciente por separado.

La implementación de **machine learning** y la **farmacocinética poblacional** mejora este límite ya que no sólo considera la variabilidad interindividual, sino que permite abarcar una población más amplia.

La evolución natural desde el enfoque clásico a un nuevo enfoque se representa en la figura 1 (5).

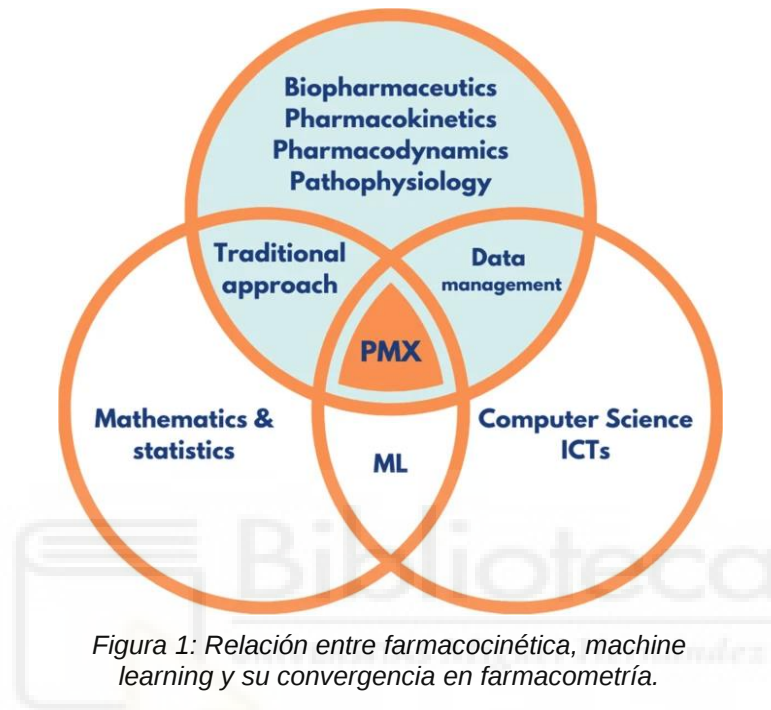


Figura 1: Relación entre farmacocinética, machine learning y su convergencia en farmacometría.

2.2. FARMACOCINÉTICA POBLACIONAL:

La farmacocinética poblacional tiene como objetivo estudiar las fuentes y correlaciones que afecta a las concentraciones plasmáticas de los fármacos en una población representativa de la que recibirá el tratamiento clínico (6). Es decir, busca entender las causas por las que las concentraciones pueden variar entre individuos y cómo se relacionan con los cambios.

Existen distintos factores que pueden afectar a la variabilidad:

- Demográficos: edad, peso, altura, género.
- Fisiopatológicos: insuficiencia renal, hipertensión.
- Genéticos: polimorfismos en enzimas encargadas de metabolizar los fármacos (citocromos)
- Otros: medicación concomitante, forma farmacéutica, vía de administración. (7)

Una vez identificados estos factores, se puede estimar la magnitud de su impacto a la dosis-respuesta del fármaco y ajustar la dosis de manera adecuada.

Un punto clave de la farmacocinética poblacional es que permite obtener resultados desde un conjunto de datos escaso o analizar grupos que normalmente, por su complejidad, no se pueden estudiar cómo la población pediátrica. A diferencia con la farmacocinética poblacional, en el enfoque clásico, los participantes de los estudios son elegidos bajo criterios muy estrictos buscando homogeneidad ya que, la variabilidad es percibida como un problema a solucionar y no a estudiar.

Esto supone un gran problema pues no se consigue una representación real de cómo funciona el fármaco puesto que la población de estudio no se asemeja a una población real sometida a mucha variabilidad.

A diferencia del enfoque tradicional, la farmacocinética poblacional permite obtener información de pacientes que sí son relevantes para el objetivo terapéutico ya que integra la variabilidad como una característica importante a estudiar.

En definitiva, la farmacocinética poblacional es importante para orientar y definir estrategias óptimas de dosificación en poblaciones o de forma individual y en la monitorización terapéutica ya que la eficacia y seguridad pueden verse afectadas por la variabilidad. Para poder comprender el impacto de esos factores en la variabilidad de las concentraciones plasmáticas, se hace uso de un modelo poblacional.

Los modelos más usados son el **método en dos etapas**, el **modelo de efectos mixtos no lineales** y **métodos bayesianos** (6) (8):

- **Método en dos etapas:** enfoque tradicional dividido que se divide en dos etapas. Primero, se estima de forma individual los parámetros farmacocinéticos de cada uno de los individuos a partir de sus concentraciones en el tiempo. En la segunda etapa, se agrupan estos parámetros obteniendo un enfoque general de toda la población estudiada. Este método es útil cuando se dispone de muchos datos por

individuo, por el contrario, no es eficiente y tiende a sobrestimar la variabilidad que presentan los individuos con información limitada.

- **Modelo de efectos mixtos no lineales:** este modelo estudia directamente los datos de una población de manera conjunta, lo que resulta útil cuando hay datos incompletos, o con características poco definibles. En este modelo, se modela directamente sobre los datos de una población por lo que la variabilidad está presente. De este modo es posible conseguir patrones sin perder la información individual. Este modelo es útil cuando se tienen pocos datos individuales o estos son difíciles de interpretar.
- **Métodos bayesianos:** se trata de una estrategia que combina la información previa, conocimientos existentes y resultados de estudios anteriores sobre una población con los datos observados en un paciente concreto, utilizando el teorema de Bayes. Este parámetro consigue una estimación individualizada de los parámetros farmacocinéticos cuando no se disponen de suficientes datos individuales ya que se compensa con los datos poblacionales (8).

La farmacocinética poblacional y los modelos son posibles gracias a la farmacocinética y la farmacometría, el desarrollo de estos modelos para su aplicación clínica depende de ambas disciplinas: la farmacometría aporta las herramientas matemáticas y estadísticas necesarias para generar el propio modelo mientras que la farmacocinética proporciona los datos necesarios de la acción del fármaco en el organismo.

Sin embargo, a la hora de optimizar tratamientos, nuevas técnicas como el mencionado “machine learning” resultan clave en el desarrollo y optimización clínico, mejorando la precisión de los modelos predictivos y automatizando el proceso.

2.3. MACHINE LEARNING:

El **Machine learning (ML)**, o aprendizaje automático, es una rama de la inteligencia artificial (IA), que capacita a las máquinas para aprender de una forma similar a los humanos, permitiendo realizar las tareas de forma autónoma y mejorar continuamente con la exposición a nueva información (9). En el ámbito de la **farmacocinética poblacional**, el ML puede integrarse en modelos matemáticos o software para procesar datos de una forma más eficiente, facilitando el aprendizaje autónomo y mejorar progresivamente con el tiempo.

El funcionamiento de ML se basa en 3 partes principales:

- **Proceso de decisión:** determina la capacidad predictiva o clasificatoria del algoritmo, según datos de entrada, que pueden estar o no definidos.
- **Función de error:** evalúa el desempeño del algoritmo (que tanto se equivoca) respecto a valores conocidos.
- **Optimización del modelo:** ajusta la diferencia que existe entre el modelo y los valores reales, repitiendo el proceso de forma iterativa hasta alcanzar un umbral de precisión aceptable. (9)

Dentro del campo de *machine learning*, destacan subcampos como el “**Deep learning**” (**aprendizaje profundo**) y las **redes neuronales**. Las redes constituyen un subcampo del machine learning, y el *Deep learning*, a su vez un subcampo de las redes neuronales.

La diferencia *entre machine learning y Deep learning* reside en la forma en que el algoritmo aprende y el grado de intervención humana requerido. El Deep learning es capaz de procesar grandes volúmenes de datos no estructurados¹, identificando las principales características de estos de forma autónoma, reduciendo la intervención humana. En contraste, el machine learning tradicional requiere de la intervención humana para poder definir las características más importantes antes de entrenar el modelo.

Las **redes neuronales** están compuestas capas o nodos, emulando la estructura de un cerebro humano, es por este motivo que recibe el nombre de “neuronal”.

¹ datos directos que no han sido procesados

Las redes neuronales incluyen:

- **Capa de entrada:** Recibe los datos iniciales.
- **Capas ocultas:** Una o varias capas que procesa los datos y aprende.
- **Capa de salida:** Propone el resultado final.

Cada nodo (en orden) determinará si la información se transmite al siguiente, permitiendo un aprendizaje jerárquico.

Dentro del machine learning, existen diversos algoritmos con aplicaciones específicas:

- **Regresión lineal:** Utilizado para predecir valores numéricos, relaciona una variable independiente con otras dependientes.
- **Regresión logística:** Empleado para predecir variables categóricas (por ejemplo, hay respuesta a un fármaco; si/no)
- **“Clustering”:** Agrupa un conjunto de datos y descubrir patrones sin necesidad de información previa. Puede ser útil para determinar poblaciones con perfiles similares.
- **Árboles de decisiones:** se emplea para predicciones numéricas o para clasificar mediante una secuencia de decisiones ramificadas, la cual funciona como una estructura jerárquica.
- **“Random forest”:** Mediante la combinación de varios árboles de decisión, permite predecir valores reduciendo el sobreajuste. Es útil cuando hay muchas variables complejas. (9)
- **“Extreme gradient boosting (XGBoost):** Algoritmo basado en árboles de decisión. Trabaja de una forma óptima y precisa por lo que, permite corregir errores de modelos anteriores (10).
- **“Catboost”:** Algoritmo que permite manejar muchas variables categóricas de forma eficiente sin procesamiento (10).

La diversidad de algoritmos y modelos permiten que esta técnica se amolde y ajuste a diversas necesidades y objetivos propuestos. Esto la consolida como una herramienta con un gran potencial para el desarrollo farmacocinético.

3.OBJETIVO:

El objetivo de este trabajo consiste en realizar una revisión sistemática sobre la aplicación del machine learning en farmacocinética, concretamente en su aplicación sobre farmacoterapia y farmacocinética poblacional en pacientes pediátricos.

4.MATERIALES Y MÉTODOS:

Este trabajo ha sido autorizado por la oficina de investigación responsable (OIR) encargada de la evaluación de proyectos de la UMH con el código: **TFG.GFA.ARL.SMM.250215.**

4.1. Diseño:

Consiste en el estudio crítico de los artículos obtenidos tras la revisión sistemática pertinente.

4.2. Fuente de la obtención de datos:

Los datos de este trabajo se obtuvieron tras una consulta directa y acceso, vía internet, de las siguientes bases de datos: MEDLINE (vía PubMed), [SCOPUS (vía Elsevier)

4.3. Tratamiento de la información:

Para definir los términos de la ecuación de búsqueda final, se planteó una estructura PICO y Con la ayuda de DeCS se obtuvieron los descriptores de los términos clave del proyecto:

- “Machine learning”
- “Pharmacokinetics”
- “Drug monitoring”

Planteando las siguientes ecuaciones de búsqueda:

MEDLINE: ((Machine Learning [MeSH Terms]) OR (Machine Learning [Title/Abstract])) AND ((pharmacokinetics [MeSH Terms]) OR (pharmacokinetic*[Title/Abstract]) OR (drug monitoring [MeSH Terms]) OR (drug monitoring [Title/Abstract])), Junto al filtro Child: birth-18 years

SCOPUS: (TITLE-ABS-KEY ("machine learning") AND TITLE-ABS-KEY ("drug monitoring") AND TITLE-ABS-KEY (pharmacokinetics)) AND (LIMIT-TO (DOCTYPE,"ar")) AND (LIMIT-TO (EXACTKEYWORD,"Child") OR EXCLUDE (EXACTKEYWORD,"Adult"))

4.4. Selección final de artículos:

Para realizar este trabajo, se seleccionaron una serie de artículos que se adecuaban a los criterios de búsqueda establecidos: uso de machine learning en la monitorización de fármacos en pacientes pediátricos, que además estuvieran escritos en inglés y/o castellano junto a acceso abierto al texto completo.

Los **CRITERIOS DE INCLUSIÓN** fueron los siguientes:

- Artículos que cumple con los objetivos del presente TFG
- Población pediátrica
- Artículos publicados en revistas originales revisadas por pares.

Los **CRITERIOS DE EXCLUSIÓN** fueron los siguientes:

- Artículos recuperados a texto completo.
- Artículos en otro idioma distinto del castellano o inglés.

5.RESULTADOS:

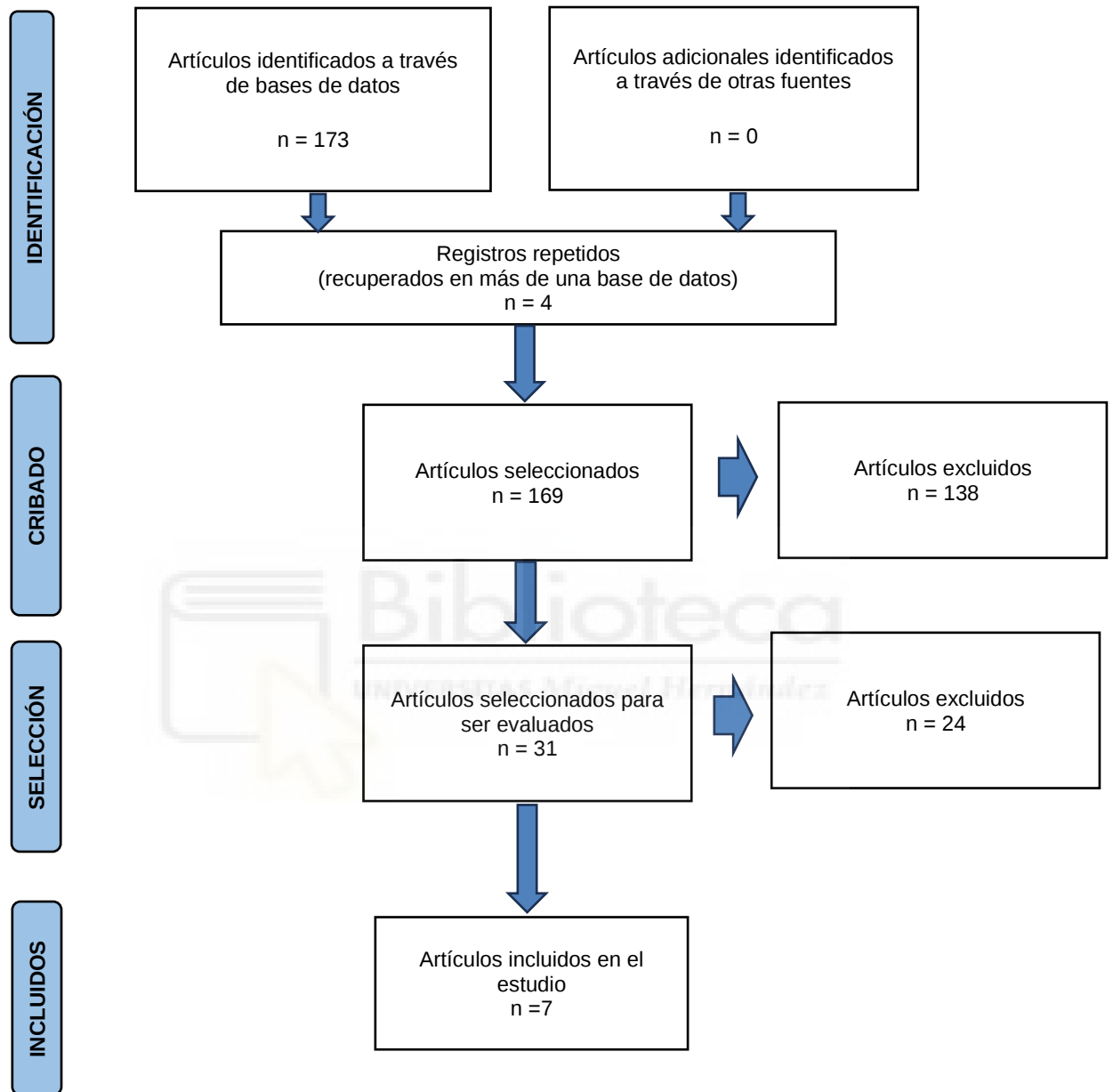


Figura 2. Diagrama de decisión de los artículos seleccionados en el trabajo

TABLA 1. ANÁLISIS DE LOS 7 ARTÍCULOS ESTUDIADOS SOBRE LA APLICACIÓN DEL MACHINE LEARNING EN FARMCOMETRIA

AUTOR, AÑO	DISEÑO DEL ESTUDIO	FÁRMACO/S ESTUDIADO/S	POBLACIÓN (EDAD Y NÚMERO)	TÉCNICA DE ML UTILIZADA	OBJETIVO PRINCIPAL DEL ESTUDIO	RESULTADOS	APLICACIÓN CLÍNICA
Tang et al 2023 (11)	Estudio retrospectivo y análisis de datos clínicos con desarrollo y validación de un nuevo modelo predictivo basado en machine learning	Vancomicina	n= 1631 Edad: 23,3-52,4 semanas	CatBoost	Evaluación del uso de ML para el cálculo de C_0 y AUC de vancomicina en neonatos	La validación del modelo CatBoost C_0 muestra una mejora general frente a los modelos clásicos en 42,5%. En cuanto al modelo AUC, presenta una precisión de 80,3%	El modelo puede ser aplicado para establecer la dosis a administrar y para llevar a cabo individualización posológica
Ponthier et al 2022 (12)	Simulación farmacocinética con base en un modelo poblacional, con una validación clínica en pacientes reales de un nuevo modelo predictivo basado en machine learning	Vancomicina	n=82 Edad: 27,43-34,54 semanas	XGBoost	Desarrollo de un modelo ML para optimizar la dosis inicial de vancomicina en neonatos	El modelo XGBoost mejora significativamente la predicción de la dosis frente a modelos clásicos validados	El modelo permite establecer dosis de vancomicina de forma más precisa reduciendo el riesgo de eventos adversos en neonatos
Ponthier et al 2024 (13)	Simulación farmacocinética basada en un modelo poblacional junto a una validación clínica en pacientes reales de un nuevo modelo predictivo basado en machine learning	Ganciclovir y valganciclovir	Ganciclovir: n= 31 Edad: 2,8-10,1 años Valganciclovir: n= 34 Edad: 0,5-11,5 años	Combinación de XGBoost, redes neuronales y bosque aleatorio	Desarrollo de un modelo ML para optimizar la dosis de ganciclovir y valganciclovir en población pediátrica y comparación de los resultados con modelos previamente publicados	El modelo ML propuesto obtuvo una mayor tasa de éxito en ambos fármacos (36,8 y 25,8%) en alcanzar el AUC deseado frente a las fórmulas o modelos existentes	El modelo podría ser aplicado con el objetivo de optimizar las dosis de ganciclovir y valganciclovir

Huang et al 2021 (14)	Estudio retrospectivo y analítico de datos clínicos con desarrollo y validación de un nuevo modelo predictivo basado en machine learning	Vancomicina	n= 407 Edad: 0,63-5 años	Combinación de XGBoost, GBRT, Bagging y ExtraTree	Desarrollo de un modelo ML para predecir las concentraciones valle de vancomicina en población pediátrica	El modelo combinado presentó una mayor precisión a la hora de predecir la concentración valle frente al modelo clásico (51,22% frente a 36,59%)	Optimización de la dosis de vancomicina en pediatría con el objetivo de reducir riesgo de toxicidad y problemas de dosificación
Yin et al 2024 (15)	Estudio observacional retrospectivo de datos clínicos y desarrollo de un nuevo modelo predictivo basado en machine learning	Vancomicina	n= 112 Edad: 1 mes - <4 años	XGBoost	Desarrollo de un modelo ML para predecir las concentraciones valle de vancomicina en población pediátrica	El modelo XGBoost fue el que mejores resultados mostró frente a los otros 6 modelos propuestos	Herramienta de optimización en la dosificación de vancomicina en población pediátrica.
Ponthier et al 2024 (16)	Simulación farmacocinética basada en un modelo poblacional, con desarrollo y validación de un nuevo modelo predictivo basado en machine learning	Ganciclovir y valganciclovir	Ganciclovir: n= 17 Edad: 1,5-10,2 años Valganciclovir: n= 9 Edad: 0,8-8,8 años	XGBoost	Desarrollo y posterior validación de un modelo ML de muestreo limitado (solo dos puntos de muestra) para estimar el AUC de ganciclovir y valganciclovir en población pediátrica	El modelo XGBoost mostró una mayor precisión a la hora de estimar la exposición de ambos fármacos, en comparación con el método clásico	El modelo propuesto puede ser aplicado para estimar la exposición a ganciclovir y valganciclovir con solo dos muestras de sangre, lo que permite un monitoreo terapéutico más sencillo y menos invasivo en la población pediátrica
Mo et al 2022 (17)	Estudio retrospectivo observacional de datos clínicos, con desarrollo y validación de un nuevo modelo predictivo basado en machine learning	Tacrolimus	n= 171 Edad: 1,9-8,9 años. Divididos en subgrupos según expresión de CYP3A5	GBDT para el grupo general, y subgrupo que no expresa. ET para el grupo que expresa	Desarrollo y validación de un 3 modelo ML para predecir la concentración valle con ajuste dosis/peso (C ₀ /D) según variables clínicas y genéticas	El modelo GBDT superó a los modelos clásicos en términos generales, con una mejor precisión y menor error. Se generó un modelo específico para cada subgrupo	Los modelos propuestos permiten optimizar la dosificación de tacrolimus en pediatría, minimizando riesgos y mejorando la eficacia terapéutica

6.DISCUSIÓN:

La aplicación de “machine learning” en la farmacometría pediátrica es un gran avance y ha demostrado tener un gran potencial, arrojando mejores resultados que análisis farmacocinéticos clásicos a la hora de caracterizar el comportamiento de los fármacos.

Destacan modelos como XGBoost los cuales claramente han superado en rendimiento a modelos bayesianos clásicos.

En esta revisión, se va a exponer por que la implementación de esta técnica es beneficiosa sobre la farmacometría en población pediátrica.

Según Tang et al (11) en 2023, los modelos de machine learning basados en la predicción de C_0^2 y AUC_{0-24h} permiten individualizar de forma precisa y segura la dosis de vancomicina en neonatos, la cual presenta una alta variabilidad interindividual en neonatos. Este enfoque permite un abordaje terapéutico más preciso, optimizando la dosis inicial de vancomicina y posteriormente el régimen terapéutico.

El planteamiento de este estudio resulta innovador, ya que integra el desarrollo independiente de ambos modelos, mejorando la aplicabilidad clínica frente a modelos farmacocinéticos poblacionales (PopPK)

- **MODELO DE PREDICCIÓN DE C_0 :**

El modelo se construyó a partir de datos obtenidos de un estudio retrospectivo, que incluyó datos de 1631 neonatos y 4894 mediciones de la concentración plasmática de vancomicina, con el objetivo de predecir la concentración mínima en rango terapéutico mediante el modelo desarrollado.

² En este estudio, C_0 hace referencia a la concentración plasmática mínima de fármaco en estado estacionario, es decir la más baja justo antes de administrar la siguiente dosis.

El primer paso consistió en un procesamiento de los datos, siguiendo unos criterios de inclusión:

- Concentraciones obtenidas de administraciones intermitentes de vancomicina
- Concentraciones tomadas en el estado estacionario³.
- Concentraciones dentro de un rango temporal específico comprendido entre el 80-120% después de haber sido administrada la última dosis.
- Concentraciones por encima del umbral mínimo detectable del ensayo clínico.

Seguidamente, se seleccionaron 15 fisiológicamente significativas⁴, mediante un modelo poblacional (PopPK) llamado Jacqz-Aigrain. Tras esto, se reemplazaron datos faltantes (como el peso al nacer, el cual tenía una falta de datos en un 24% de los casos) y se eliminaron covariables categóricas con una distribución muy desequilibrada⁵.

Finalmente, para evaluar la importancia de cada covariable y su aportación en al modelo, se utilizó una técnica llamada **SHAP**, identificando la concentración sérica de creatinina, tiempo de muestreo, dosis única por unidad de peso corporal, frecuencia de dosificación entre otras.

El conjunto de datos procesados comprendía 1017 concentraciones plasmáticas de 648 neonatos y sufrió una transformación estadística para mejorar el rendimiento del modelo que se dividieron en un grupo de entrenamiento (80%) y un grupo de prueba (20%).

Con los datos procesados, se llevó a cabo la construcción del modelo C_0 . Para ello, se probaron 8 algoritmos de machine learning⁶ y se evaluaron mediante una serie de métricas de rendimiento:

- RMSE (error cuadrático medio): cuantifica la desviación en promedio respecto al valor objetivo, valores más bajos indican mejor precisión.

³ Los pacientes alcanzaban dicho estado al haber recibido al menos tres dosis del fármaco.

⁴ Aquellas que pueden afectar a cómo se comporta el fármaco en el cuerpo.

⁵ Si una covariable aparecía más veces que otra, para evitar sesgo se eliminaban.

⁶ Árbol de decisión de refuerzo de gradiente (GBDT), Refuerzo de categóricos (CatBoost), Refuerzo de gradiente extremo (XGBoost), Máquina de refuerzo de gradiente ligera (LBGM), Regresión logística (LR), Regresión soporte vectorial (SVR), Red neuronal Tabular Learning (Tabnet) y Vecino más cercano aproximado (ANN)

- MAPE (Error porcentual absoluto medio): mide el porcentaje en el que la predicción se desvía del valor real, valores más bajos indican mejor precisión.
- R^2 (Coeficiente de determinación): mide la variación de los datos reales con el modelo, cuanto más cercano a 1 el valor, el modelo es mejor.⁷
- MPE (Error porcentual medio): Mide el porcentaje en el que la predicción se desvía del valor real, es decir, si el modelo tiende a sobreestimar los valores o subestimarlos⁸.

El modelo que arrojó mejores resultados en comparación al resto fue “CatBoost” y, por tanto, fue el modelo elegido.

Finalmente, se llevó a cabo la validación del modelo, comparando la eficacia frente a un modelo clásico PopPK. Para ello, se hizo uso de los datos de un estudio externo e independiente y se hizo una comparación de sus valores estadísticos y de sus resultados en escenarios clínicos reales.

Tabla 2. Evaluación de los modelos basada en indicadores estadísticos para C0.

Métrica	Modelo ML	Modelo PopPK
RMSE	5,02 mg/l	6.18 mg/l
MPE	-4,20%	12.4%
MAPE	29,50%	53.0%
%MPE (20%) ⁹	46,30%	28.5%
%MPE (50%)	83,60%	66.4%

Tabla 3. Evaluación de los modelos basado en escenarios clínicos para C0.

Escenario clínico	Modelo ML	Modelo PopPK
% dentro de rango terapéutico (10-20 mg/l)	62.8%	48.7%
% fuera de rango terapéutico ¹⁰	56.8%	31.8%
% precisión global	60.7%	42,6%

⁷ Mide que tan bien puede replicar el modelo los valores reales.

⁸ Si el valor es 0, el modelo no tiene sesgo en cambio, si es positivo tiende a sobreestimar y negativo a subestimar.

⁹ Porcentaje en los que el error porcentual fue menor o igual a 20% en valor absoluto.

¹⁰ Situaciones en las que el rango terapéutico es diferente a 10-20mg/l.

En comparación, el modelo ML mostró resultados más precisos que el modelo PopPK, mejorando la precisión relativa de 42,5%¹¹. Esto implica que un mayor número de pacientes recibirán correctamente la primera dosis, siendo crucial para minimizar el riesgo de toxicidad o ineficacia terapéutica.

Los escenarios clínicos reales, se distinguen según la precisión respecto al rango terapéutico:

- Precisión dentro del rango terapéutico:
 - El valor real está dentro del rango, pero el modelo predice un valor inferior → aumento innecesario de la dosis; toxicidad.
 - El valor real está dentro del rango, pero el modelo predice valor superior → posible reducción de la dosis; ineficacia terapéutica.
- Precisión fuera del rango terapéutico:
 - El valor real está por debajo del rango, pero el modelo predice que está dentro → se interpreta que el paciente está bien tratado; Sin eficacia.
 - El valor real está por encima del rango, pero el modelo predice que está dentro → no se detecta exceso de fármaco; toxicidad

Para evitar este tipo de escenarios peligrosos, se necesita una buena precisión y el modelo ML mejora significativamente a los modelos clásicos.

- **MODELO DE PREDICCIÓN DE AUC_{0-24h}:**

Este modelo, tiene como objetivo ajustar la dosis de forma más precisa. Se desarrollo utilizando estimaciones bayesianas “post hoc”¹² para obtener valores de AUC con los que construir el modelo. para ello se necesitó:

- Modelo PopPK de Jacqz-Aigrain.
- Datos de los pacientes (creatinina, peso, edad, etc).
- Concentraciones plasmáticas de vancomicina.¹³
- Régimen de dosificación.

¹¹ $(60,7\%-42,6\%) / 42,6\% \times 100 = 42,5\%$ de mejora del modelo ML frente al PopPK.

¹² Se realiza una combinación de datos poblacionales previos y datos de los pacientes.

¹³ Concentraciones valle y pico.

Este modelo se construyó siguiendo los mismos pasos mencionados en el desarrollo del modelo C₀. Para ellos, se incluyeron 1017 mediciones de AUC, siendo finalmente el modelo CatBoost el que arrojó mejores resultados.

Una vez se construyó el modelo, planteo realizar una validación con 122 estimaciones individuales de AUC, comparando las mediciones obtenidas con el modelo frente a el AUC estimado.¹⁴

Tabla 4. Evaluación del modelo ML basada en indicadores estadísticos para AUC0-24h.

Métrica	Modelo ML
RMSE	60,1 mg/l
MPE	2,2%
MAPE	9,3%
%MPE (20%)	94,7%
%MPE (50%)	100%

Tabla 5. Evaluación del modelo ML basada en escenarios clínicos para AUC0-24h.

Escenario clínico	Modelo ML
% dentro de rango terapéutico (400-600 mg*h/l)	79,7%
% fuera de rango terapéutico	81,3%
% precisión global	80,3%

Pese a que el modelo presenta una precisión global del 80,3% (alcanzado al menos el mismo nivel de precisión que los modelos clásicos), no se puede afirmar su superioridad frente a los modelos clásicos debido a la falta de datos suficientes.

Teniendo en cuenta que el modelo planteado reduce significativamente la dificultad de cálculo y que la predicción pueda obtenerse solo introduciendo las covariables necesarias, hace que el modelo sea una opción más cómoda.

Finalmente, se llevó a cabo un ensayo virtual para predecir la respuesta de los pacientes a diferentes regímenes de dosificación, utilizando el modelo validado de Jacqz-Aigrain, el cual permite simular la concentración inicial de vancomicina mediante un conjunto de datos externos también validados.

¹⁴ Esto significa que el modelo ML se comparó con valores de AUC, pero a través de estimaciones ya que no se disponía de suficientes valores reales medidos. Dichas estimaciones se consideraron AUC "referencia".

Se realizaron 100 simulaciones y se compararon los resultados con el enfoque de las guías clínicas y los modelos ML propuestos, que mostraron el mejor rendimiento a la hora de calcular la dosis dentro del rango terapéutico.

Los modelos desarrollados permiten simular diferentes esquemas de dosificación para nuevos pacientes, esto permite a nivel clínico, poder elegir el plan más adecuado antes de administrar la dosis. Tras la administración y consiguiente caracterización del fármaco, el modelo AUC permite ajustar de manera más precisa la dosis del tratamiento en función de la respuesta del paciente.

A pesar de sus ventajas, no es posible afirmar que los modelos ML son superiores a modelos clásicos, los cuales se fundamentan en mecanismos fisiológicos y dan pie a una mayor interpretación de los resultados.

Además, el estudio cuenta con una serie de limitaciones como el tamaño reducido en su validación, ausencia de algunas covariables asociadas a neonatos y, por último, la falta de datos reales de AUC en la práctica clínica mencionado anteriormente.

El artículo publicado por Ponthier et al (12) en 2022, destaca el potencial de los modelos de *Machine Learning* entrenados con datos simulados, especialmente cuando los datos clínicos disponibles son limitados. En este estudio, los autores buscaron desarrollar un modelo o algoritmo de machine learning capaz de predecir el AUC de la vancomicina, a partir del cual poder calcular la dosis inicial óptima, mediante perfiles farmacocinéticos simulados.

Estos perfiles se obtuvieron de modelos farmacocinéticos poblacionales validados. Asimismo, evaluaron los resultados con modelos o ecuaciones previamente validadas.

La vancomicina en neonatos presenta una gran variabilidad farmacocinética, viéndose afectada por factores como el peso, edad posmenstrual y la función renal, esto convierte la individualización de la dosis en un reto clínico considerable.

Entrenar un algoritmo capaz de predecir la dosis con datos simulados a partir de la exposición, sin requerir una gran base de datos clínicos reales representa un modelo innovador.

El desarrollo del modelo comenzó con una simulación de 1900 perfiles farmacocinéticos¹⁵, dividiendo los perfiles en 19 edades gestacionales comprendidas entre 24 y 42 semanas y realizando 100 simulaciones por cada una de estas edades, partiendo de un modelo previamente publicado.

Por otra parte, dicho modelo, sufrió una serie de modificaciones en la variabilidad “entre ocasiones” (diferentes respuestas de los pacientes en distinto momento) y se generaron dos versiones del modelo con diferente nivel de error (mayor o menor “ruido”). Estos fueron comparados seleccionando dos pacientes simulados de forma aleatoria, mostrando una menor variabilidad el modelo de ruido reducido, siendo más fiable para la aplicación clínica.

Las covariables elegidas para la simulación fueron la edad gestacional, peso al nacer, edad postnatal, edad postmenstrual, aumento de peso diario, peso actual y creatinina plasmática, todas ellas simuladas de forma independiente para cada semana de edad gestacional, lo que combinaciones poco realistas.

Así mismo, se definió la dosis de referencia¹⁶ (corresponde a la utilizada en el hospital donde del estudio) y la dosis de carga¹⁷. Con dichos valores, se pudo ejecutar la simulación y obtener los perfiles simulados.

Para estimar la exposición al fármaco entre 48 y 72 horas, se calculó el AUC_{48-72h} multiplicando la concentración simulada por 24 horas. Se eliminaron aquellos atípicos¹⁸ y finalmente, se obtuvieron 1717 perfiles simulados.

Una vez obtenidos los datos simulados, se llevó a cabo el análisis con machine learning, para ello se analizaron 3 algoritmos de ML: XGBoost, MARS y GLMNET.

¹⁵ Se pueden entender como pacientes virtuales.

¹⁶ Basada en la edad postmenstrual y creatinina.

¹⁷ Basada en la edad posmenstrual y el peso actual

¹⁸ Valores extremos fuera del intervalo 5% al 95% de los valores simulados.

Los datos fueron divididos:

- Grupo de entrenamiento (75%), con el objetivo de “enseñar” al algoritmo.
- Grupo de prueba (25%), con el objetivo de comprobar si el modelo funciona.

A su vez el grupo de entrenamiento se subdividió en un conjunto de análisis donde se prueban los diferentes modelos (80%) y un conjunto de evaluación para determinar el mejor modelo (20%).

Por otra parte, los datos se normalizaron, con el objetivo de que todos los valores influyan por igual y se convirtieron variables categóricas se codificaron.

Los 3 algoritmos se compararon en el conjunto de evaluación, siguiendo distintas métricas de rendimiento (RMSE, R^2 , MPE). Siendo XGBoost el que mejores resultados arrojó en términos de $rMPE^{19}$ (8,6%) y $rRMSE$ (35,7%) y por ello elegido para calcular la dosis óptima de vancomicina.

XGBoost, se refinó combinando los conjuntos de análisis y evaluación para finalmente evaluar su precisión en un conjunto de pruebas de datos nuevos no utilizados para el entrenamiento.

Las variables con mayor impacto se predijeron utilizando una técnica llamada permutación aleatoria²⁰, siendo las de mayor relevancia (en orden): Creatinina, dosis de carga, edad posmenstrual, perfusión continua, peso actual). Tras ellos, el modelo entrenado y refinado, se usó para predecir el AUC_{48-72h} con el objetivo de alcanzar un valor de $AUC_{objetivo}$ de 500mg*h/l, siguiendo la siguiente fórmula:

$$Dosis_{ML} = AUC_{48-72h\ objetivo} * \frac{Dosis_{administrada}}{AUC_{48-72h\ predicho}}$$

Para evaluar la robustez del modelo ML y su ajuste con el rango terapéutico óptimo, se generó un nuevo conjunto de datos simulados con una semilla²¹ distinta. Para ello, se comparó la dosis de referencia del hospital, la dosis generada por el algoritmo y la dosis estimada por una ecuación derivada de un modelo PopPK. El rango terapéutico se definió entre 400-600mg*h/l.

¹⁹ Error medio de predicción relativo.

²⁰ Si la eliminación de la variable causa una gran disminución de la precisión, esta es importante, la técnica consiste en modificar aleatoriamente los valores de estas variables.

²¹ Hace referencia a la capacidad de generación, garantizado que sean datos aleatorios y que garantice la variabilidad.

Para cada enfoque, se calculó el AUC extrapolado, el cual permite comprar el impacto de cada dosis obtenida. Los resultados se agruparon en 3 grupos; por debajo del objetivo ($<400\text{mg}\cdot\text{h/l}$), en objetivo ($400\text{--}600\text{mg}\cdot\text{h/l}$) y por encima ($>600\text{mg}\cdot\text{h/l}$). Para determinar la capacidad de los 3 enfoques, se realizó una prueba de chi-cuadrado.

XGBoost propuso dosis más baja cuando el AUC estaba por encima de rango y dosis similares cuando estaba dentro del rango y, por otra parte, arrojó la mejor tasa de alcance del objetivo en un 46,9% en comparación a la ecuación del modelo PopPK (41,4%) y dosis de referencia (34,5%).

Finalmente, utilizando un conjunto de 82 pacientes de dos hospitales francés, sometidos a monitoreo terapéutico, se compararon los 3 enfoques aportando así una validación externa real.

Para cada método, se evaluó su AUC extrapolado, clasificándose de igual forma que en la prueba de robustez (por debajo, en rango, por encima de rango) y analizando su capacidad con una prueba chi-cuadrado.

En esta prueba, la tasa de alcance del modelo ML fue mejor, que la ecuación previa (35,5% frente a 28%) y similar a la dosis de referencia (35,5%), aunque sin diferencias estadísticamente significativas ($p=1$).

En definitiva, el modelo ML es capaz de predecir la primera dosis con mayor precisión que los métodos previos, aunque la diferencia en la práctica clínica no fue extremadamente significativa. Esto puede ser debido a la falta de una muestra representativa de mayor tamaño.

La fórmula derivada de un modelo PopPK, fue elegida ya que mostró buen rendimiento, covariables relevantes y tenía una validación previa, por lo tanto, la mejora de rendimiento del modelo ML con respecto a dicha fórmula es muy significativa, tanto en la práctica clínica como en simulación. Es importante destacar que el rendimiento de la fórmula fue mucho peor que en el estudio original (28% frente a 70%) debido al cambio en el rango terapéutico objetivo.

Puesto que la exposición a vancomicina en pacientes pediátricos puede dar lugar a reacciones adversas, alcanzar de forma óptima la dosis correcta y el rango terapéutico es clave.

El porcentaje de sobreexposición a la vancomicina fue menor con la dosis obtenida mediante el modelo ML que con la formula, lo que resulta en una menor incidencia de reacciones adversas.

Pese a que hay una mejora evidente con el modelo ML propuesto, su tasa de alcance sigue siendo baja (35%) por lo que hay un gran margen de mejora debido a los muchos factores de variabilidad que afectan y no han sido considerados.

El estudio presentó ciertas limitaciones como son el pequeño tamaño muestral de pacientes reales, y la falta de consideración de ciertas variables, a pesar de esto el algoritmo mejora la tasa de alcance y arroja datos prometedores, con una menor probabilidad de sobreexposición.

Según Ponthier et al en 2024 (13), al igual que en el artículo de 2022 (12), se destaca la opción de entrenar algoritmos de machine learning con datos simulados para estimar la mejor dosis inicial de ganciclovir (GCV) y valganciclovir (VGCV) en población pediátrica.

Para lograr dicho objetivo, se entrenaron los algoritmos con perfiles farmacocinéticos generados mediante simulaciones Monte Carlo, lo que permitió una mayor diversidad de perfiles simulados. El objetivo terapéutico propuesto para el algoritmo fue alcanzar un AUC_{0-24h} de 40-60mg*h/l, un rango extrapolado de estudios previos en adultos. los resultados del algoritmo se compararon con una formula derivada de un modelo PopPK previamente validado en una población real de pacientes.

El primer paso para construir el modelo consistió en crear 300 perfiles farmacocinéticos por sexo y por cada año entre 1 y 18 años, obteniendo 10.800 perfiles empleando diferentes covariables, elegidas en base a los modelos poblacionales previos para GCV y VGCV. Tras filtra los valores extremos, se obtuvieron 9873 perfiles para GCV y 9809 para VGCV, válidos para la simulación.

Las covariables empleadas fueron la edad (1-18 años con distribución uniforme²²), peso y talla (según estándar de crecimiento de la OMS con distribución normal truncada²³), creatininemia (se hizo de manera independiente a las demás covariables), área de superficie corporal (siguiendo la fórmula de Mosteller), aclaramiento de creatinina (siguiendo formula de Schwartz), sexo y tipo de trasplante (distribución uniforme).

A continuación, de acuerdo con los modelos PopPK ya publicados, aplicando las dosis recomendadas para profilaxis de cada fármaco²⁴, se simuló las concentraciones de los medicamentos en cada perfil cuando se alcanzó el estado estacionario. Tras esto, se ajustó el error al mínimo posible para evitar el “ruido” y se calculó mediante la regla del trapecio el AUC0-24h. Para todos los pacientes se eliminaron valores extremos²⁵ para evitar sesgos.

En el desarrollo de los algoritmos ML, se analizaron los datos de GCV y VGCV por separado en un grupo de entrenamiento (75%) y prueba (25%) seleccionados de forma aleatoria y se procesan las variables numéricas y categóricas. Se creó una variable categórica (sí/no) que indicaba si el AUC0-24h alcanzaba el objetivo de exposición entre 40-60mg*h/l. y los algoritmos entrenados fueron XGBoost, redes neuronales, Random Forest y un modelo combinado llamado “*stacking*” que junta los tres modelos mencionados.

Sobre el conjunto de entrenamiento, se ajustaron los parámetros de los algoritmos y se procedió a evaluar utilizando diferentes métricas como la sensibilidad y especificidad y Valores predictivo negativo (VPN) y positivo (VPP) El modelo con mayor precisión a la hora de alcanzar el objetivo terapéutico. fue el modelo “*stacking*”.

A continuación, se estimó la probabilidad que tendría cada paciente simulado en alcanzar el rango objetivo con el modelo elegido, examinando diferentes dosis que iban desde 0.1 hasta 5 veces la dosis guía.

²² Misma probabilidad para cada variable.

²³ Evita valores extremos.

²⁴ GCV: 5mg/kg una vez al día intravenoso, VGCV: 7xBSAxCrCL con un máximo de 900mg por dosis vía oral.

²⁵ Percentiles <5 y >95

La dosis asociada a la mayor probabilidad de alcanzar el objetivo se seleccionó como dosis inicial óptima de ML. Posteriormente, se comparó la propuesta de dosis de ML con la dosis propuesta por la fórmula clásica y la dosis guía.

En general, el modelo ML arrojó dosis más altas para aquellos pacientes cuyo AUC estaba por debajo del objetivo, mejorando la proporción de pacientes que alcanzarían dicho objetivo, para aquellos pacientes en rango el modelo ML propuso dosis similares y aquellos casos por encima del rango objetivo, el modelo propuso dosis menores.

Para finalizar, se llevó a cabo una validación externa con pacientes reales (31 para GCV y 34 para VGCV), comparando la dosis inicial propuesta tanto por el modelo ML como la dosis propuesta por la fórmula clásica y la dosis guía. Los resultados fueron:

Tabla 6. Evaluación del porcentaje de éxito de distintos métodos.

Método	GCV (% éxito)	VGCV (% éxito)
ML	25,8%	35,5%
Fórmula poblacional clásica	22,6%	35,5%
Dosis Guía	12,9%	20,6%

Esto implica, que para GCV la tasa incremento un 3,2 % respecto a la fórmula clásica, lo que podría evitar el fracaso terapéutico en un mayor número de pacientes. Para VGCV, la tasa fue igual, esto puede deberse a que la mayoría de los pacientes utilizados para la evaluación externa pertenecían al mismo centro que los utilizados para desarrollar el modelo PopPK de referencia.

Ambos fármacos presentan una alta variabilidad interindividual y por ello, alcanzar el objetivo terapéutico suele ser difícil, por este motivo, se alcanza una tasa de objetivo tan baja.

Entre las fortalezas que presenta este estudio, cabe destacar el uso de datos simulados para suplir la escasez de datos reales, la transparencia en la metodología (el código y los parámetros fueron publicados) y el enfoque en algoritmos ML los cuales superan las fórmulas tradicionales.

Sin embargo, el estudio presenta una serie de limitaciones, la más importante es que en el estudio se asumió, una relación completamente lineal entre la dosis administrada y el AUC, lo cual puede no ser representativo de lo que ocurre realmente. Otra imitación fue el rango del objetivo terapéutico, el cual se basó en estudios realizados en adultos, no pudiendo afirmar si el rango es igualmente válido para la población pediátrica por falta de información específica en niños. Por último, otra limitación que afectó al estudio fue el uso de pocas variables a la hora de entrenar el algoritmo ML.

Pese a todo esto, el estudio ha conseguido desarrollar un modelo ML capaz de estimar la dosis de GCV y VGCV, mejorando las tasas de éxito a la hora de alcanzar el objetivo terapéutico, aunque requiero de futuras investigaciones para confirmar su aplicabilidad en la práctica clínica.

En otro artículo de Ponthier et al, del año 2024 (16), los autores buscan desarrollar un modelo de ML capaz de determinar el AUC0-24h partiendo de un muestreo limitado²⁶, mediante una metodología similar al artículo (13), anteriormente explicado.

Para ello, se simularon 10.800 perfiles siguiendo la misma metodología que el anterior, con 300 perfiles por sexo y edad entre 1 y 18 años, manteniendo la misma distribución de variables, dosis de referencia y variables clave mediante simulaciones Monte Carlo.

A diferencia del anterior, el filtro para eliminar los valores atípicos se amplió de un intervalo de 5-95% a 1-99%, permitiendo una mayor inclusión de perfiles. Por otra parte, el conjunto de entrenamiento (75%) se subdividió en un conjunto de análisis (80%) y un conjunto de evaluación (20%).

Los datos fueron procesados de igual forma, aplicando estrategias de normalización y codificación para las variables y comparando el rendimiento de los modelos mediante los mismos parámetros estadísticos (rMPE y rRMSE y su valor absoluto).

²⁶ Partiendo de muy pocos datos de concentración.

Un aspecto diferencial entre ambos artículos es la elección de algoritmos a estudiar, este estudio comparó XGBoost, MARS, GLMNET, Random Forest y SVM, a diferencia del artículo anterior que hizo uso de XGBoost, redes neuronales, Random Forest y un modelo combinado llamado “*stacking*”. En este nuevo artículo, la validación externa se utilizaron 17 pacientes para GCV y 9 para VGCV.

El enfoque de muestreo limitado se analizó mediante distintas combinaciones de muestras tomadas en momentos distintos; 0/2, 0/3, 0/4, 1/3, 1/4, 2/4, 0/6, 1/6 y 2/6 horas. Para cada algoritmo y muestreo limitado, se ajustaron hiperparámetros con una validación cruzada.

Las mejores combinaciones fueron 2/6 para VGCV y 0/2 para GCV, con XGBoost como algoritmo ya que presentaron un menor RMSE y MPE.

EL modelo ML se comparó con un método de referencia llamado “MAP-BE”²⁷ basado en los modelos farmacocinéticos clásicos mencionados en ambos estudios. Los resultados estadísticos de XGBoost en el grupo de prueba fueron:

Tabla 7. Resultados estadísticos en el grupo de entrenamiento para VGCV.

Método	rMPE	rRMSE	% valores fuera del 20% del valor referencia
XGBoost	0,4	5,7	0,8%
Modelo clásico “Fracnk”	-1,9	46,6	43,9%
Modelo clásico “Facchin”	1,1	12,9	7%
Modelo clásico “Nguyen”	28,9	97,7	36,4%

²⁷ Maximun-a-posteriori Bayesian Estimation.

Tabla 8. Resultados estadísticos en el grupo de entrenamiento para GCV.

Método	rMPE	rRMSE	% valores fuera del 20% del valor referencia
XGBoost	0,9	12,4	10,4%
Modelo clásico "Fracnk"	-5,6	44,4	61,3%
Modelo clásico "Nguyen"	4,4	61,8	75,7%

En el caso de los pacientes externos, los tiempos de la toma de muestra oscilaba en el caso de la muestra de 2 horas (1,8-2,3 horas) y para la de 6 (5,9-6,3 horas) y, al igual que en el grupo de pruebas, el algoritmo XGBoost con esta combinación de muestras arrojó los mejores resultados en comparación con el modelo clásico. Los resultados estadísticos de XGBoost en el grupo de validación externo fueron:

Tabla 9. Resultados estadísticos en pacientes externos reales para VGCV.

Método	rMPE	rRMSE	% valores fuera del 20% del valor referencia
XGBoost	0,2	16,5	36,3%
Modelo clásico "Fracnk"	37,3	56,4	71,7%
Modelo clásico "Facchin"	-7,7	35,2	72,7%
Modelo clásico "Nguyen"	108	123	90,9%

Tabla 10. Resultados estadísticos en pacientes externos reales para GCV.

Método	rMPE	rRMSE	% valores fuera del 20% del valor referencia
XGBoost	-9,7	17,2	28,6%
Modelo clásico "Fracnk"	4,1	56,4	95,5%
Modelo clásico "Nguyen"	105,7	125,8	100%

El modelo ML demostró un rendimiento superior a los modelos clásicos en pacientes reales para VGCV mientras que para GCV pese a una buena precisión, se aprecia cierto sesgo probablemente debido a la alta variabilidad farmacocinética de la población.

Por otra parte, los autores también analizaron una estrategia de 3 muestras en lugar de 2, observando que no mejoraba significativamente la predicción, lo que da importancia a las estrategias simplificadas.

En cuanto a las limitaciones del estudio, la información clínica de los pacientes procede de un estudio retrospectivo lo que puede generar sesgos. Por otra parte, el tiempo real de muestreo no siempre coincide con los tiempos teóricos lo que puede afectar a la precisión de la estimación y factores como la omisión del consumo de alimento para VGCV que puede afectar a su absorción, así como la alta variabilidad de la población estudiada.

En definitiva, los autores desarrollaron un modelo capaz de calcular el AUC a partir de dos muestras de concentración para poder individualizar el tratamiento en población pediátrica, esto reduce el número de extracciones necesarias lo que implica ser un método menos invasivo. El modelo mostro ser una herramienta precisa y generalizable la cual, necesita de futuros estudios que confirme su impacto clínico.

En el artículo de Huang et al (14) publicado en el año 2021, se concluye que los modelos de *Machine Learning (ML)* son útiles para predecir la concentración mínima (valle) de vancomicina en población pediátrica.

La población del estudio incluyó datos de 407 pacientes²⁸ tratados con vancomicina, cuyas concentraciones en sangre de vancomicina se midieron en el hospital mediante una técnica cromatográfica tras la cuarta dosis de vancomicina, justo antes de la siguiente dosis, permitiendo obtener la concentración valle del medicamento en sangre.

Tras esto, los autores analizaron las variables objetivo del estudio, así como las covariables potencialmente influyentes²⁹. Para ello se completaron los datos faltantes de estas mediante de manera estadísticamente significativa mediante un método llamado “MICE”³⁰, imputando los valores de forma robusta y minimizando el riesgo de sesgo durante el análisis. Esta técnica permite aumentar la precisión del modelo ajuste del modelo con los datos³¹.

El resultado fue un conjunto de 24 variables para cada individuo, siendo la principal la concentración mínima de vancomicina en sangre.

El conjunto de datos se dividió en:

- Grupo de entrenamiento (80%): para crear y ajustar el modelo
- Grupo de prueba (20%): para analizar el rendimiento del modelo

Para generar el modelo ML final, se aplicaron 8 algoritmos (Árbol de decisiones, regresión de vectores de soporte, bosque aleatorio, “Adaboost”, “Bagging”, “Extratree”, “GBRT” y “XGBoost”). Se evaluaron con parámetros estadísticos (R^2 , MAE, MSE, RMSE) y el porcentaje predictivo dentro de un margen del $\pm 30\%$ del valor real.

Los cinco algoritmos con mejor rendimiento en general fueron XGBoost, GBRT, ExtraTree, Bagging y árbol de decisión. Para mejorar la precisión y la predicción, los autores combinaron estos 5 algoritmos en uno único. Cada uno de estos algoritmos tiene un peso determinado en el algoritmo final. El aporte de cada uno

²⁸ Siguiendo una serie de criterios de inclusión y exclusión.

²⁹ Demográficas, analíticas, clínicas.

³⁰ Imputación múltiple por ecuaciones encadenadas.

³¹ alto valor de R^2

fue: XGBoost (4 partes) GBRT (4 partes), ExtraTree (0.5 partes), Bagging (1 parte), Arbol de decisiones (0.5 partes)).

Esto implica que los algoritmos con mayor peso son XGBoost y GBRT, siendo los más precisos. Tras esto, los autores analizaron que variables tenían mayor peso sobre el algoritmo XGBoost al ser de los 5 algoritmos el que presenta un mejor ajuste con los datos. Las variables más relevantes fueron: aclaramiento de creatinina, procalcitonina (PTC), ácido úrico y peso³².

Teniendo el modelo ensamblado y las variables con mayor peso asignadas, se analizó el modelo en un conjunto de prueba y se obtuvo un $R^2 = 0,614$ y una precisión de predicción de la concentración del 51,22%.

En definitiva, la elección de un modelo conjunto por parte de los autores es debido a que esta, es la opción óptima para un uso clínico ya que, presenta mayor predicción frente a XGBoost (51,22% frente a 42,68%) y un R^2 parecido (0,614 frente a 0,657). Pese a que, por separado, XGBoost tuvo un MSE menor que el conjunto³³, el modelo conjunto sigue siendo la mejor opción.

Posteriormente, se tomaron 20 pacientes del hospital como grupo de validación independiente y se analizó el modelo conjunto obteniendo un $R^2 = 0,622$ y una precisión del 72,69%, valores más altos que en el conjunto de prueba. Esto indica que el modelo trabaja correctamente con nuevos pacientes lo que indica que tiene una buena capacidad de predicción.

Finalmente, los datos obtenidos del modelo conjunto se compararon con los obtenidos mediante un modelo poblacional clásico, evaluando como de cerca están las predicciones de la concentración real mínima³⁴ de ambos modelos. Los valores obtenidos fueron de una precisión del 51,22% para el modelo ML y 36,59% para el modelo clásico a la hora de predecir la concentración.

Los resultados del modelo ML fueron superiores y, además, el modelo clásico presenta un R^2 menor que el modelo conjunto (0,3 frente a 0,614).

³² PTC y ácido úrico se relacionan positivamente con la concentración (si aumenta la variable, aumenta la concentración y el aclaramiento de creatinina se relaciona negativamente. Esto es clave para ajustar el modelo.

³³ MSE menor implica menos error cuadrático.

³⁴ Con un intervalo de $\pm 30\%$ de la concentración real.

Los modelos clásicos no son capaces de incluir para el cálculo de la concentración de vancomicina, la cantidad de parámetros que un modelo ML sí que puede, esto pone de manifiesto que el modelo ML es claramente más preciso y permite un mejor ajuste.

Un aspecto clave de los modelos ML es su escalabilidad, ya que pueden mejorar automáticamente con nuevos datos, por ello, elegir el modelo más preciso es clave. El modelo conjunto combina las ventajas de todos los algoritmos por separado, reduciendo el error que estos puedan generar.

Aunque el estudio se centra en estimar la concentración mínima de vancomicina, las guías clínicas recomiendan que sea el AUC el parámetro utilizado para ajustar la dosis ya que refleja mejor la exposición al fármaco. Los autores no abordaron dicho parámetro debido a la falta de concentraciones pico y valle necesarias para calcular con precisión el AUC.

A modo de conclusión, los autores presentan las limitaciones que tiene el estudio, siendo la principal el pequeño tamaño de muestra, lo que limita el desarrollo y aprendizaje del modelo ML. A pesar de esto, el modelo arroja resultados que demuestran ser una potencial mejora frente a los modelos clásicos a la hora de predecir la concentración de vancomicina en población pediátrica.

El estudio de Minghui et al (15) de 2024, tiene como objetivo desarrollar un modelo ML capaz de predecir la concentración mínima de vancomicina en población pediátrica. Pese a que el índice AUC/MIC³⁵ sea el principal parámetro recomendado en la monitorización de vancomicina, resulta complicado determinarlo. Por este motivo se elige la concentración mínima como objetivo.

El estudio incluyó pacientes menores de 4 años del hospital infantil de Shanghai que se encontraba bajo monitorización farmacéutica de vancomicina. Los parámetros de inclusión fueron administración intermitente, estado estacionario, media hora antes de la siguiente dosis y de exclusión menores de un mes, antecedentes de lesión renal aguda, entre otros, incluyendo finalmente 120

³⁵ MIC= concentración inhibitoria máxima.

mediciones de 112 pacientes. Los datos se dividieron en un grupo de entrenamiento (70%) y un grupo de prueba (30%).

El resto de los datos necesarios se obtuvo de los registros de cada paciente incluyendo regímenes de dosificación, información demográfica y resultados clínicos de los pacientes. En general, los datos eran parecidos tanto en el grupo de prueba como entrenamiento, pero se encontraron algunas diferencias en los niveles de creatinina y en el número de administraciones del fármaco.

Los datos recopilados se procesaron, ajustando los valores para que estuvieran en la misma escala numérica y que así, el modelo no aportase mayor peso a una variable con valores extremos.

Se analizaron 7 algoritmos ML; Regresión lineal (LR), árboles de decisión potenciados por gradiente (GBDT), máquinas de vectores de soporte (SVM), árboles de decisión (DT), bosques aleatorios (RF), “Bagging” y XGBoost. En primer lugar, se ajustaron los hiperparámetros para cada uno de los algoritmos con el objetivo de obtener el mejor rendimiento posible. Además, se analizó la covariable más importante a la hora de predecir la concentración mínima mediante la técnica SHAP.

Una vez planteado el modelo, con el conjunto de prueba se obtuvo la concentración mínima predicha por dichos modelos y se comparó con los valores medidos reales mediante parámetros estadístico (R^2 , MSE, RMSE y MAE). De los 7 modelos, XGBoost arrojó los mejores resultados (MAE = 2,55, RMSE = 4,13, MSE = 17,12 y R^2 = 0,59). Para este algoritmo, las covariables con mayor peso fueron el nitrógeno ureico en sangre (BUN), creatinina sérica (SCR) y el aclaramiento de creatinina (CCR), valores clave en la función renal, demostrando el papel fundamental de la función renal en la eliminación de vancomicina.

Este estudio, demuestra que es posible desarrollar un algoritmo ML capaz de predecir la concentración mínima de vancomicina. XGBoost mostró resultados superiores a otros algoritmos gracias a su gran capacidad de manejar datos con una relación no lineal y múltiples variables, lo que resulta relevante en la población pediátrica ya que presenta una alta variabilidad interindividual.

Finalmente, cabe mencionar que el estudio presenta una serie de limitaciones, la principal de ellas es que se realizó en un único centro hospitalario. Por otra

parte, carece de validación con un grupo externo y la exclusión de ciertos individuos como los menores de un mes y pacientes con antecedentes de lesión renal aguda, lo que limita la extrapolación del modelo. Estos factores pueden afectar al comportamiento del modelo en términos de precisión y seguridad.

Sin embargo, el estudio muestra un modelo con mucho potencial clínico y monetario ya que podría reducir los costes de los hospitales en la monitorización.

Finalmente, el artículo de Mo et al (17) (2022) busca generar un modelo ML para predecir y ajustar la concentración de tacrolimus³⁶ (TAC) por dosis/peso (Co/D) en pacientes pediátricos con síndrome nefrótico refractario, combinando datos clínicos y genéticos, obteniendo una herramienta que ayude a optimizar la dosificación individualizada. El TAC es un fármaco esencial que presenta un estrecho margen terapéutico y alta variabilidad interindividual, atribuido principalmente a polimorfismos genéticos de enzimas metabolizadoras, como el polimorfismo de CYP3A5*3.

El estudio amplía el enfoque incluyendo factores fisiológicos como la función renal, aunque no se explica por completo toda la variabilidad. En estudios previos, se han desarrollado modelos farmacocinéticos poblacionales, con una escasa integración de estos factores genéticos. Debido a esto, el enfoque de este estudio resulta novedoso.

Para el desarrollo del modelo, se seleccionaron pacientes de un estudio previo, incluyendo pacientes menores de 16 años, tratados al menos 3 meses con TAC, incluyendo finalmente un total de 171 pacientes. Los criterios de exclusión fueron pacientes con síndrome nefrótico sensible a corticoesteroides, mala adherencia al tratamiento o pacientes con medicación concomitante que altere los niveles de TAC. Finalmente, se recopilaron 82 variables clínicas (edad, sexo, concentración mínima de TAC, entre otras) y se incluyeron 244 polimorfismos genéticos, seleccionadas por su relevancia farmacocinética y frecuencia alélica ($\geq 5\%$) en la población del estudio.

³⁶ Inmunosupresor con biodisponibilidad baja y variable, metabolizado por el citocromo P450 y eliminado casi al 100% a través de la bilis.

Las concentraciones de TAC se midieron antes de la siguiente administración de la dosis. Por otra parte, los polimorfismos se analizaron mediante PCR y la técnica MassARRAY, para evitar sesgo del genotipo CYP3A5, se dividieron los pacientes según la expresión de este, desarrollando un modelo específico para el grupo que expresaba el gen y los que no.

Durante el procesamiento de datos, se eliminaron variables con más de un 10% de datos faltantes, se imputaron datos cuando fue posible, se transformaron las variables categóricas y finalmente se normalizaron las variables mediante un proceso llamado Z-score. La selección de características y variables se realizó usando 4 algoritmos de ML (XGBoost, Extremely Randomized Trees, Random Forest y GBDT), generando 16 modelos basados en combinaciones y normalización de estos algoritmos. Esto permitió seleccionar las variables más relevantes y predictiva usando principalmente XGBoost.

Finalmente, con este conjunto de característica seleccionado³⁷, se entrenaron cinco algoritmos ML (Lasso Regresión, XGBoost, Extremely Randomized Trees, Random Forest y GBDT). El conjunto de datos se dividió en un conjunto de entrenamiento y otro de prueba y se evaluaron los modelos usando parámetros estadísticos como R^2 , RMSE y MSE. Con esto, se determinó que GBDT fue el más eficaz en el grupo general con un $R^2= 0,444$; MSE = 591,03; MAE $\approx 18,29$. El subgrupo que no expresaba CYP3A5, GBDT mostró un $R^2=0.264$ y Por su parte, Extremely Randomized Trees (ET) fue el algoritmo más preciso en el grupo que expresaba el genotipo CYP3A5 con un $R^2= 0,380$.

Para validar externamente los modelos, se reclutó a 30 pacientes de los cuales se incluyeron 25, aplicando los mismos criterios de inclusión. Los modelos mostraron una correlación estadísticamente significativa entre los valores predichos y los observados en términos de C_0/D , confirmando su aplicabilidad clínica.

En definitiva, el estudio destaca la utilidad de integrar variables clínicas y genéticas en fármacos como TAC, el cual está altamente influenciado por esto. Los modelos planteados demuestran ser capaces de identificar relaciones entre variables clínicas y polimorfismos genéticos, permitiendo una predicción precisa.

³⁷ Variables clínicas como la albúmina basal, edad y género y por otra parte polimorfismos.

Sin embargo, el estudio presenta limitaciones importantes, como el tamaño muestral, y la inclusión de datos de un único centro y la ausencia de validación clínica real. En conclusión, el trabajo demuestra que el uso de ML en farmacocinética en población pediátrica puede superar las limitaciones de métodos tradicionales, mediante 3 modelos, adaptados al genotipo específico, útiles a la hora de predecir la dosificación de TAC en pediatría mejorando la seguridad y eficacia del tratamiento.

A modo general, la evidencia sigue siendo limitada ya que la mayoría de los estudios carecen de una validación externa con pacientes reales y aquellos que, si presentan, lo hacen sobre un grupo reducido y homogéneo, lo que dificulta que los resultados sean escalables a la población y puedan ser aplicados de inmediato. La principal limitación corresponde a la falta de grandes conjuntos de datos pediátricos, que dificulta el desarrollo de un modelo preciso ya que, en sí, el ML es una herramienta que, con un correcto desarrollo, presenta una superioridad aplastante en términos de precisión sobre los modelos clásicos.

Otro punto importante es la falta de conocimiento general sobre esta tecnología, el desarrollo de nuevas técnicas debe ir acompañado de una correcta adaptación y comprensión por parte de los sanitarios, siendo imprescindible aprender a usar dicha herramienta. Por este motivo, destaco aquellos artículos donde se presenta el modelo de forma abierta, con una interfaz clara y en código abierto, fomentando la mejora constante del modelo. Una vez superadas estas limitaciones, se abre un horizonte de posibilidades muy amplio sobre el uso de ML en farmacocinética, sobre todo en poblaciones vulnerables como la pediátrica.

En resumen, el ML tiene un potencial inmenso para transformar la farmacocinética poblacional tal y como se conoce, pero su implementación en la práctica clínica debe ser gradual y metodológica. Sin duda alguna, el futuro de la farmacocinética poblacional y clínica pasa por la implementación de esta tecnología, siempre y cuando se desarrolle bajo una evidencia clara.

7.CONCLUSIONES:

En base al objetivo del trabajo (aplicación del *Machine Learning* en farmacocinética para población pediátrica), se puede concluir que el ML es una herramienta innovadora y eficaz, que representa una evolución a los métodos farmacocinéticos tradicionales, especialmente cuando nos enfrentamos a una población con una alta variabilidad interindividual y expuesta a cambios constantes.

En base a los estudios revisados, el *Machine Learning* se puede desarrollar como la herramienta principal para la predicción de dosis iniciales y para poder llevar a posteriori, una monitorización terapéutica sobre los pacientes pediátricos. Esto resulta altamente relevante en fármacos como la Vancomicina, Ganciclovir y Valganciclovir los cuales presentan un estrecho margen terapéutico y una precisión extrema en el ajuste de la dosificación para evitar riesgo de ineficacia y/o toxicidad.

Sin embargo, los estudios no demuestran una aplicación concluyente debido al gran número de limitaciones que presenta el desarrollo de una herramienta de tal magnitud.

8. BIBLIOGRAFÍA:

1. Aragón M. Periódico UNAL - Farmacometría, clave en la optimización de tratamientos y en el desarrollo de medicamentos [Internet]. Edu.co. [citado 2025 mayo]. Disponible en: <https://periodico.unal.edu.co/articulos/farmacometria-clave-en-la-optimizacion-de-tratamientos-y-en-el-desarrollo-de-medicamentos>
2. SAAVEDRA S. I, QUIÑONES S. L, SAAVEDRA B. M, SASSO A. J, LEÓN T. J, ROCO A. A. Farmacocinética de medicamentos de uso pediátrico, visión actual. Revista chilena de pediatría. 2008 Jun [citado 2025 Marzo] ; 79(3): p. 249-258. Disponible en: http://www.scielo.cl/scielo.php?script=sci_arttext&pid=S0370-41062008000300002&lng=es . <http://dx.doi.org/10.4067/S0370-41062008000300002>
3. Robles Piedras AL. Importancia de la farmacocinética en la práctica clínica. Educación y Salud Boletín Científico Instituto de Ciencias de la Salud. 2019 [citado 2025 Marzo];8(15): p. 123-126. Disponible en: <https://repository.uaeh.edu.mx/revistas/index.php/ICSA/issue/archive>
4. Sociedad Española de Farmacia Hospitalaria (SEFH). Escuela de Pacientes. Liberación y absorción de los medicamentos [Internet]; 2024 [citado 2025 Marzo 25]. Disponible en: <https://www.sefh.es/escuela-de-pacientes-conoce-tus-medicamentos-detalle.php?mdl=4&tm=49> .
5. Ibarra M, Lorier M, Trocóniz IF. Pharmacometrics: Definition and History. In A T, editor. The ADME Encyclopedia [Internet]; Springer; 2022. p. 933–939. Disponible en: http://dx.doi.org/10.1007/978-3-030-84860-6_171
6. Sun H, Fadiran E, Jones C, et al. Population Pharmacokinetics. A regulatory perspective. Clin Pharmacokinet [Internet]. 1999;37. p. 41-58. Disponible en: <http://dx.doi.org/10.2165/00003088-199937010-00003>
7. Barranco Garduño LM, Neri Salvador C, León Molina H, et al. La farmacocinética poblacional y su importancia en la terapéutica. Medicina Interna de México [Internet]. 2011 [citado 2025 Abril];27(4):370–377. Disponible en: <https://www.mediagraphic.com/pdfs/medintmex/mim-2011/mim114h.pdf>
8. Universidad de Salamanca. METODOLOGÍA DE LA MONITORIZACIÓN. Curso. Salamanca : Universidad de Salamanca, Farmacia y Tecnología Farmacéutica; 2022. Disponible en: <https://cursotdmsalamanca.es/01/MA%20-%20LUNES%20-%202001.pdf>
9. What is machine learning (ML)? [Internet]. Ibm.com. 2025 [citado 2025 Abril]. Disponible en: <https://www.ibm.com/think/topics/machine-learning>
10. Kavlakoglu E, Russi E. ¿Qué es XGBoost? [Internet]. Ibm.com. 2025 [citado 2025 Abril]. Disponible en: <https://www.ibm.com/es-es/think/topics/xgboost>
11. Tang BH, Zhang JY, Allegaert K, Hao GX, Yao BF, Leroux S, et al. Use of Machine Learning for Dosage Individualization of Vancomycin in Neonates. Clin Pharmacokinet. [Internet]. 2023;62(8):1105–1116. Disponible en: <http://dx.doi.org/10.1007/s40262-023-01265-z>
12. Ponthier L, Ensuque P, Destere A, et.al. Optimization of Vancomycin Initial Dose in Term and Preterm Neonates by Machine Learning. Pharm Res [Internet]. 2022;39(10):2497–2506. Disponible en: <http://dx.doi.org/10.1007/s11095-022-03351-6>
13. Ponthier L, Autmizguine J, Franck B, et.al. Optimization of Ganciclovir and Valganciclovir Starting Dose in Children by Machine Learning. Clin Pharmacokinet [Internet]. 2024;63(4):539–550. Disponible en: <http://dx.doi.org/10.1007/s40262-024-01362-7>
14. Huang X, Yu Z, Bu S, et.al. An Ensemble Model for Prediction of Vancomycin Trough Concentrations in Pediatric Patients. Drug Des Devel Ther [Internet]. 2021;15:1549–1559. Disponible en: <http://dx.doi.org/10.2147/DDDT.S299037>
15. Yin M, Jiang T, Yuan Y, et.al. Optimizing vancomycin dosing in pediatrics: a machine learning approach to predict trough concentrations in children under four years of age. Int J Clin Pharm [Internet]. 2024;46(5):1134–1142. Disponible en: <http://dx.doi.org/10.1007/s11096-024-01745-7>
16. Ponthier L, Bénédicte F, Autmizguine J, et.al. Application of machine-learning models to predict the ganciclovir and valganciclovir exposure in children using a limited sampling strategy. Antimicrob Agents Chemother [Internet]. 2024;68(10):e0086024. Disponible en: <http://dx.doi.org/10.1128/aac.00860-24>
17. Mo X, Chen X, Wang X, et.al. Prediction of Tacrolimus Dose/Weight-Adjusted Trough Concentration in Pediatric Refractory Nephrotic Syndrome: A Machine Learning Approach Pharmgenomics Pers Med [Internet]. 2022;15:143–155. Disponible en: <http://dx.doi.org/10.2147/PGPM.S339318>