



Pontificia Universidad
JAVERIANA
Cali

**PREDICCIÓN DE LA FINURA DEL CEMENTO UTILIZANDO PARÁMETROS
OPERACIONALES: UN ENFOQUE BASADO EN MACHINE LEARNING**

Laura Valencia Romero

Código 9026328

*Proyecto aplicado para optar al título de
Magister en Ciencia de Datos*

Director
Julián Gil González

FACULTAD DE INGENIERÍA Y CIENCIAS
MAESTRÍA EN CIENCIA DE DATOS
SANTIAGO DE CALI, MAYO 25 DE 2026

FICHA RESUMEN

TÍTULO DEL PROYECTO: Predicción de la finura del cemento utilizando parámetros operacionales: un enfoque basado en Machine Learning.

1. ÁREA DE TRABAJO: Ciencia de datos

2. TIPO DE PROYECTO: Aplicado

3. ESTUDIANTE(S): Laura Valencia Romero

4. CORREO ELECTRÓNICO: lauravalencia0826@javerianacali.edu.co

5. DIRECCIÓN Y TELEFONO: Carrera 1 B4 #58 A1 – 30, Cali. Cel. 3117956206

6. DIRECTOR: Julián Gil González

7. VINCULACIÓN DEL DIRECTOR: Profesor del Departamento de Electrónica y Ciencias de la Computación, Pontificia Universidad Javeriana Cali.

8. CORREO ELECTRÓNICO DEL DIRECTOR: julian.gil@javerianacali.edu.co

9. PALABRAS CLAVE: aprendizaje automático, finura, molinos verticales, optimización, control de calidad.

10. FECHA DE INICIO: Julio de 2025

11. FECHA DE FINALIZACIÓN: Mayo de 2026

12. RESUMEN: Este proyecto desarrolla un modelo de aprendizaje automático para predecir la finura del cemento, específicamente el porcentaje de retenido en malla 325, a partir de variables operacionales de un molino vertical. La motivación surge ante la necesidad crítica de superar las limitaciones de los métodos tradicionales de control de calidad que generan demoras y afectan tanto la eficiencia energética como la estabilidad del proceso. Para ello, se evaluaron cinco enfoques supervisados, encontrando que los modelos no lineales superaron de manera consistente al modelo lineal. Aunque el proceso gaussiano obtuvo el mejor desempeño predictivo, XGBoost fue seleccionado como modelo final por ofrecer el mejor equilibrio entre precisión, interpretabilidad y costo computacional, con un R^2 cercano a 0.60 y un RMSE aproximado de 0.46% en el conjunto de prueba. Además, el análisis del modelo permitió identificar la relevancia de variables asociadas a la composición de la alimentación, la carga interna del molino y la velocidad del separador sobre la finura del producto final.

TABLA DE CONTENIDO

INTRODUCCIÓN	6
1. DEFINICIÓN DEL PROBLEMA	7
1.1. PLANTEAMIENTO DEL PROBLEMA	7
1.2. FORMULACIÓN DEL PROBLEMA	8
2. OBJETIVOS DEL PROYECTO	8
2.1. OBJETIVO GENERAL	8
2.2. OBJETIVOS ESPECÍFICOS	8
3. MARCOS DE REFERENCIA	8
3.1. MARCO TEÓRICO	8
3.1.1. PROCESO DE MOLIENDA EN LA INDUSTRIA CEMENTERA	8
3.1.2. FINURA: SU IMPACTO Y MEDICIÓN	10
3.1.3. APRENDIZAJE AUTOMÁTICO EN PROCESOS INDUSTRIALES	11
3.1.4. MODELOS SUPERVISADOS DE PREDICCIÓN	12
3.1.5. MÉTRICAS DE EVALUACIÓN	16
3.2. ANTECEDENTES	16
4. PREPARACIÓN Y ANÁLISIS EXPLORATORIO	19
5. CONSTRUCCIÓN Y OPTIMIZACIÓN DE LOS MODELOS	27
6. EVALUACIÓN DE LOS MODELOS PREDICTIVOS	31
6.1. ESTABILIDAD DE LOS MODELOS EN VALIDACIÓN CRUZADA	32
6.2. CAPACIDAD DE GENERALIZACIÓN EN EL CONJUNTO DE PRUEBA	33
6.3. TIEMPOS DE EJECUCIÓN DE LOS MODELOS	34
6.4. ANÁLISIS ESPECÍFICO DE LOS MODELOS	35
6.4.1. REGRESIÓN LINEAL MÚLTIPLE: VERIFICACIÓN DE SUPUESTOS	35
6.4.2. MODELOS BASADOS EN ÁRBOLES: INTERPRETACIÓN MEDIANTE SHAP	36
6.4.3. PERCEPTRÓN MULTICAPA: CONTROL DE SOBREAJUSTE	39
6.4.4. PROCESO GAUSSIANO: PREDICCIÓN E INCERTIDUMBRE	40
6.5. SELECCIÓN DEL MODELO FINAL	42
7. CONCLUSIONES Y TRABAJOS FUTUROS	43

7.1. CONCLUSIONES	43
7.2. TRABAJOS FUTUROS	44
8. REFERENCIAS BIBLIOGRÁFICAS	44

LISTA DE FIGURAS

- Figura 1.** Esquema de un molino vertical de rodillos.
- Figura 2.** Potencia del molino y detección de paros.
- Figura 3.** Distribución de la finura medida como retenido en malla 325.
- Figura 4.** Correlaciones de las variables operativas con la variable de salida: (a) coeficientes de Pearson y Spearman; (b) correlación por distancia (dCor).
- Figura 5.** Distribución del retenido según niveles del índice composicional de la mezcla.
- Figura 6.** Densidad conjunta de dp_{molino} y $rpm_{\text{separador}}$ respecto al retenido.
- Figura 7.** Tendencia del retenido según la posición promedio de los rodillos.
- Figura 8.** Diagrama del flujo de trabajo para los modelos.
- Figura 9.** Valores reales vs. Valores predichos para cada modelo.
- Figura 10.** Verificación de supuestos de la regresión lineal múltiple.
- Figura 11.** Resumen de interpretabilidad del modelo XGBoost mediante valores SHAP.
- Figura 12.** Efecto de la posición promedio de los rodillos sobre la predicción del retenido según valores SHAP.
- Figura 13.** Efecto de la velocidad del separador sobre la predicción del retenido según valores SHAP.
- Figura 14.** Curva de aprendizaje del perceptrón multicapa en términos de RMSE.
- Figura 15.** Predicción media del GP con intervalo de confianza del 95% en el conjunto de prueba.

LISTA DE TABLAS

- Tabla 1.** Descripción de las variables.

Tabla 2. Resumen de la limpieza de variables.

Tabla 3. Resumen de búsqueda de hiperparámetros.

Tabla 4. Resumen de mejor configuración de hiperparámetros para cada modelo.

Tabla 5. Desempeño promedio y desviación estándar de los modelos en la validación cruzada K-Fold.

Tabla 6. Desempeño de los modelos en el conjunto de prueba.

Tabla 7. Tiempos de ejecución de los modelos.

Tabla 8. Resumen de la incertidumbre latente y total estimada en el conjunto de prueba.

Tabla 9. Length-scales para las variables en el proceso gaussiano con ARD.

LISTA DE ANEXOS

Anexo A: Cálculo de la incertidumbre del ensayo de retenido en malla 325.

INTRODUCCIÓN

En la industria cementera, parámetros de calidad como la finura (tamaño de las partículas del producto final) necesitan un control exhaustivo pues de ello depende el cumplimiento de los estándares de desempeño del material mediante propiedades como la resistencia y la durabilidad. No obstante, el método utilizado en la actualidad para la medición de la finura presenta limitaciones críticas pues los resultados tardíos de los ensayos de laboratorio impiden la realización de ajustes operativos oportunos, generando ineficiencia energética y variabilidad en la producción. Esta brecha evidencia la necesidad de tener un sistema predictivo para la estimación de la finura que permita optimizar el proceso de molienda.

El aprendizaje automático surge como una herramienta viable para aprovechar los datos históricos disponibles de variables operacionales y mediciones en el laboratorio para entrenar modelos robustos que permitan estimaciones precisas. Estudios recientes han demostrado el potencial de estas técnicas en la predicción de propiedades del cemento, aunque con enfoques que no siempre consideran la complejidad de las interacciones entre las múltiples variables operacionales y que, en términos de finura, se concentran en métricas de medición como el Blaine sin considerar otras comúnmente utilizadas como el retenido en malla 325.

Por ello, este proyecto desarrolló un modelo predictivo basado en técnicas de aprendizaje automático para estimar la finura del cemento utilizando datos operativos de un molino vertical. Para ello, se evaluaron distintos enfoques supervisados con el fin de comparar su capacidad predictiva y su pertinencia en el contexto del proceso industrial. Los resultados obtenidos permitieron confirmar la naturaleza no lineal del problema y seleccionar un modelo con un desempeño adecuado para la estimación operativa de la finura, no solo en términos de precisión, sino también de practicidad e interpretabilidad para una eventual implementación.

1. DEFINICIÓN DEL PROBLEMA

1.1. PLANTEAMIENTO DEL PROBLEMA

En las últimas décadas, la industria cementera ha experimentado un crecimiento significativo debido al incremento de la demanda de infraestructura y desarrollo urbano a nivel global, lo que exige a las cementeras trabajar en una optimización continua de los procesos de producción con el objetivo de garantizar eficiencia energética, sostenibilidad y calidad del producto final [1]. Dentro de la fabricación del cemento, uno de los procesos más críticos es la molienda, donde se homogeneizan y pulverizan las diferentes materias primas, concentrando más del 60 % del consumo eléctrico total y constituyendo uno de los mayores costos operativos [2]. Además, el grado de pulverización o tamaño de partícula del producto resultante de este proceso, conocido como finura, es un parámetro clave en la calidad, pues está directamente relacionada con propiedades como la velocidad de reacción de hidratación, el desarrollo de resistencia mecánica, la manejabilidad y la durabilidad del material [3].

Actualmente, el control de la finura del cemento se realiza mediante ensayos de laboratorio utilizando técnicas que incluyen la medida del Blaine (superficie específica por permeabilidad al aire), el retenido en malla 325 (tamizado de partículas mayores a $45\ \mu\text{m}$) y la granulometría láser (distribución de tamaño de partículas) [4]. Sin embargo, este enfoque basado en el análisis de muestras periódicas presenta varias limitaciones derivadas principalmente del tiempo que transcurre entre la toma de la muestra, su análisis y la obtención de resultados, generando problemas como desviaciones en la calidad del producto final por el retraso en la implementación de los ajustes operativos necesarios para mantener el cemento dentro de las especificaciones. Además, este desfase temporal deriva en ineficiencia energética, ya que el molino puede funcionar en condiciones subóptimas, incrementando el consumo específico de energía y la variabilidad de la producción, ya que la estabilidad del proceso recae en gran medida en la experiencia del operador mientras espera los resultados del laboratorio [5].

Por ello, contar con una estimación de esta medida se presenta como una necesidad crítica para optimizar el proceso de molienda. No obstante, esta tarea es compleja debido a la naturaleza no lineal de la relación entre la finura y los múltiples parámetros operacionales que se miden en línea, como el flujo de alimentación de los materiales, la presión hidráulica, la velocidad del separador, el caudal de aire, entre otros [6]. Esta complejidad dificulta el desarrollo de modelos matemáticos precisos y, en consecuencia, surge la necesidad de adoptar técnicas avanzadas de modelado predictivo capaces de procesar estos datos operativos, permitiendo hacer estimaciones confiables de la finura.

1.2. FORMULACIÓN DEL PROBLEMA

Reconociendo el potencial del modelado predictivo mediante técnicas de aprendizaje automático para superar las limitaciones de la medición de finura mediante métodos tradicionales, surge la pregunta central de esta investigación: ¿Cómo desarrollar un modelo de aprendizaje automático que permita predecir la finura del cemento en un molino vertical utilizando variables operacionales? A partir de esta, surgen otros interrogantes: ¿Cuáles son las variables operacionales con mayor influencia en la finura del cemento?, ¿Qué técnicas de aprendizaje automático son las más adecuadas y presentan el mejor desempeño para esta predicción, dadas las características del conjunto de datos y las relaciones no lineales entre las variables?, y, por último, ¿qué métricas de evaluación y técnicas de validación permiten garantizar la robustez estadística del modelo predictivo propuesto?

2. OBJETIVOS DEL PROYECTO

2.1. OBJETIVO GENERAL

Desarrollar un modelo de aprendizaje automático que estime la finura del cemento a partir de variables operacionales, con el propósito de contribuir a la optimización del proceso de molienda en términos de eficiencia energética y calidad.

2.2. OBJETIVOS ESPECÍFICOS

- Procesar el conjunto de datos históricos (variables de operación y mediciones de laboratorio de retenido en malla 325) mediante técnicas de tratamiento de *outliers*, datos faltantes y manejo de inconsistencias, para posteriormente identificar las variables operacionales clave que influyen significativamente en la finura del cemento mediante análisis estadísticos de correlación.
- Entrenar modelos de aprendizaje automático predictivos que permitan capturar las relaciones no lineales entre las variables de entrada y la finura.
- Validar los modelos predictivos mediante la evaluación comparativa de métricas de desempeño y técnicas de validación, asegurando su robustez y precisión para su potencial implementación en la estimación de la finura del cemento.

3. MARCOS DE REFERENCIA

3.1. MARCO TEÓRICO

3.1.1. PROCESO DE MOLIENDA EN LA INDUSTRIA CEMENTERA

La molienda de cemento es la etapa final en la producción de este material, donde las diferentes materias primas que lo conforman son reducidas a un polvo fino. Este proceso

tiene como finalidad transformar estos componentes en partículas del orden de micrómetros, aumentando su área superficial y permitiendo así su reacción con el agua lo que genera la resistencia característica del cemento. Para esta tarea, tradicionalmente se han utilizado molinos de bolas, que se destacan por su fácil operación, sin embargo, debido a su baja eficiencia, otras alternativas más eficientes en términos de energía y productividad han ganado relevancia, como es el caso de los molinos verticales de rodillos (VRM). [2] La Figura 1 presenta un esquema general de un molino vertical de rodillos y sus principales componentes:

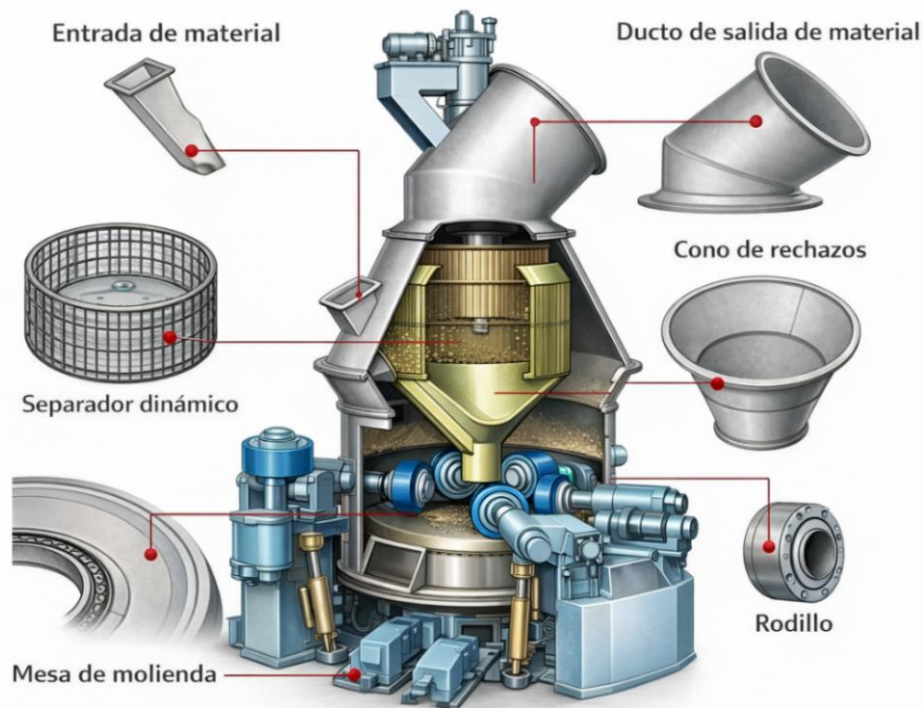


Figura 1. Esquema de un molino vertical de rodillos [7].

El proceso de molienda en un molino vertical de rodillos integra varias tareas en una sola unidad: molienda, secado, clasificación y transporte. Este consiste, en primer lugar, en el ingreso de las materias primas que caen sobre una mesa giratoria en la que, mediante fuerza centrífuga, se distribuye el material hacia la zona de molienda, donde se ubican unos rodillos que ejercen presión sobre el material mediante un sistema hidráulico. La presión y fricción a la que son sometidas las partículas en esta zona causa la reducción de su tamaño mientras que un anillo de retención, que se ubica alrededor del perímetro de la mesa, mantiene el material en la zona de molienda el tiempo necesario para alcanzar la finura deseada [2].

Simultáneamente, un flujo ascendente de aire caliente recorre el interior del molino cumpliendo dos funciones: secar el material y transportar las partículas ya molidas a la parte

superior del molino. En esta zona se encuentra un separador dinámico, compuesto por un rotor de paletas con velocidad ajustable, que tiene como función clasificar las partículas según su tamaño, rechazando las más gruesas que son retornadas a la mesa de molienda, mientras que las más finas son arrastradas hacia los filtros de colección del producto final [2].

Con lo anterior, es notable que variables operacionales como la velocidad del separador dinámico, la presión de los rodillos, el flujo de aire a través del molino, entre otras, determinan las características finales del producto en términos de finura [6].

3.1.2. FINURA: SU IMPACTO Y MEDICIÓN

La distribución del tamaño de las partículas en el cemento es una propiedad física clave en términos de la calidad, pues este parámetro incide directamente en su comportamiento hidráulico y propiedades mecánicas [3]. Cuando se presenta una mayor finura, es decir, partículas más pequeñas, estas cuentan con una mayor área superficial, lo que acelera el proceso de hidratación ya que se facilita su contacto con el agua. Este incremento en la velocidad de hidratación se traduce en un mayor desarrollo de resistencia a edades más tempranas, característica valiosa en aplicaciones donde se requieren altas resistencias iniciales [8].

No obstante, esta mayor reactividad genera un mayor calor de hidratación, que, si no es controlado adecuadamente, puede conducir a problemas de fisuración térmica, especialmente en estructuras masivas. Además, un cemento con partículas más gruesas, al tener una hidratación más lenta, tiene tiempos de fraguado más prolongados, lo que puede ser beneficioso para mantener la trabajabilidad del material durante periodos más extensos si así se requiere, aunque esto puede implicar una disminución en el desarrollo de resistencia [8], [9].

Por esto, es necesario implementar sistemas de control de la finura según los requerimientos y las propiedades deseadas del producto final. Para ello, en la industria se utilizan diferentes técnicas de medición como:

- Blaine, según la Norma Técnica Colombiana 33: este método mide el área superficial específica (m^2/kg), para ello se compacta una capa de cemento en una celda de permeabilidad y se mide la cantidad de tiempo que tarda una cantidad fija de aire en pasar a través de la muestra. Este tiempo se relaciona con la resistencia al flujo de aire, que depende del tamaño y distribución de las partículas [4], [10].
- Retenido en malla 325, según la Norma Técnica Colombiana 294: en este ensayo se dispersa una cantidad conocida de muestra sobre una malla estandarizada con aperturas de $45\ \mu m$ y se filtra con agua a presión y caudal controlados. El material

retenido se seca, se pesa y se calcula su porcentaje, el cual se relaciona con la cantidad de partículas gruesas en la muestra [4], [11].

- Granulometría láser: consiste en la medición de la distribución completa del tamaño de partículas suspendiendo una muestra en un líquido, sobre el que incide un láser. Las partículas que constituyen la muestra dispersan la luz en diferentes ángulos dependiendo de su tamaño, señal que es captada por un detector y traducida en la cuantificación porcentual de partículas para diferentes rangos de tamaño [12].

Este tipo de mediciones experimentales están sujetas a una incertidumbre que caracteriza la dispersión de los valores que podrían atribuirse razonablemente al mensurando [13]. En el caso del retenido en malla 325, la incertidumbre de la medida proviene principalmente de dos fuentes:

- Incertidumbre tipo A, asociada a la variabilidad estadística observada en ensayos repetidos bajo condiciones de repetibilidad.
- Incertidumbre tipo B, asociada a las características instrumentales, como la calibración y resolución de la balanza analítica utilizada en los pesajes.

Teniendo en cuenta ambas contribuciones se puede estimar el valor de una incertidumbre combinada, que permite establecer un límite de precisión de los valores de entrada que debe considerarse al evaluar los modelos predictivos construidos a partir de estos datos.

3.1.3. APRENDIZAJE AUTOMÁTICO EN PROCESOS INDUSTRIALES

El aprendizaje automático (*machine learning*) es un campo de la ciencia de datos que se encarga de desarrollar algoritmos y modelos que permiten a las máquinas aprender a partir de conjuntos de datos con el objetivo de reconocer patrones, hacer predicciones y extraer información valiosa para la toma de decisiones. Estas técnicas se pueden clasificar en tres categorías: aprendizaje no supervisado en el cual el algoritmo identifica patrones y relaciones entre datos no etiquetados, aprendizaje por refuerzo, que consiste en entrenar un algoritmo para tomar acciones óptimas recibiendo retroalimentación en forma de recompensas o penalizaciones y aprendizaje supervisado en el cual el algoritmo aprende a partir de datos etiquetados para hacer predicciones [14].

En el aprendizaje supervisado, el algoritmo aprende a partir de un conjunto de datos de entrenamiento donde cada observación contiene tanto las variables de entrada como las salidas conocidas (etiquetas). Durante el proceso de entrenamiento, el modelo ajusta sus parámetros internos mediante algoritmos de optimización que le permiten minimizar el error entre sus predicciones y los valores de salida reales, con el objetivo de adquirir la capacidad de generalizar lo aprendido para inferir los valores de salida correspondientes a

nuevos datos de entrada no vistos. Para evaluar su capacidad predictiva, se utiliza un conjunto de datos independiente (conjunto de prueba), que permite evaluar su desempeño y capacidad de generalización [15].

Estas técnicas se han convertido en herramientas útiles para abordar diversas problemáticas en la industria. Entre ellas se encuentra la predicción de fallas en equipos, lo que permite ahorrar y optimizar los tiempos de mantenimiento; el monitoreo continuo y reconocimiento de patrones en variables clave del proceso, lo que contribuye a garantizar la calidad del producto final; y la identificación de oportunidades para maximizar el rendimiento energético y mejorar la eficiencia operativa. Asimismo, los algoritmos de aprendizaje automático han contribuido a la optimización de los procesos logísticos, mediante la predicción basada en datos históricos de la demanda, lo que permite una mejor gestión del inventario y la planificación de recursos. Esta transformación digital impulsa la evolución hacia entornos industriales más inteligentes, sostenibles y adaptables a las demandas del mercado actual [16].

3.1.4 MODELOS SUPERVISADOS DE PREDICCIÓN

Para abordar el reto de predecir la finura, se utilizaron 5 modelos de aprendizaje automático:

3.1.4.1 REGRESIÓN LINEAL MÚLTIPLE (RLM)

La regresión lineal múltiple es una extensión del modelo de regresión lineal que busca modelar la relación entre una variable de respuesta Y y un conjunto de variables predictoras $X_1, X_2, \dots, X_n$ [17]. Para ello, se busca estimar coeficientes β que cuantifican la influencia que tiene cada predictor sobre el resultado utilizando el método de mínimos cuadrados, el cual consiste en encontrar la combinación de parámetros que mejor se ajusta a los datos observados, reduciendo al mínimo la diferencia entre los valores reales y las predicciones del modelo.

Es importante recalcar que este modelo en particular posee limitaciones pues asume que las relaciones entre la respuesta y los predictores son lineales e independientes lo que le impide capturar patrones o interacciones más complejas. Además, la validez de sus resultados depende del cumplimiento de algunos supuestos fundamentales como homocedasticidad de los errores (varianza constante), la normalidad de los residuos y la ausencia de multicolinealidad entre los predictores.

3.1.4.2 RANDOM FOREST (RF)

Esta técnica de modelado se basa en la construcción de árboles de decisión, los cuales segmentan los datos en regiones específicas mediante nodos. En esta estructura, cada rama

representa una regla de decisión basada en las características de los datos, mientras que cada hoja final asigna un valor objetivo [17]. El Random Forest potencia este concepto mediante el ensamblaje de múltiples árboles a través de la técnica de *bagging*, la cual consiste en generar múltiples subconjuntos de entrenamiento aleatorios y con reemplazo a partir del *dataset* original y posteriormente, entrenar un árbol independiente en cada uno de ellos. Al promediar los resultados de cada árbol, el modelo logra reducir la varianza al evitar que el error de un solo árbol afecte el resultado final. Para optimizar el funcionamiento del modelo y evitar el sobreajuste, es fundamental controlar hiperparámetros como:

- Número de árboles (*n_estimators*): Define la cantidad de árboles que compondrán el bosque. Un número mayor suele aumentar la estabilidad del modelo, aunque incrementa el costo computacional.
- Número de predictores por división (*max_features*): Limitar cuantas variables se consideran en cada nodo, lo que permite tener mayor diversidad entre árboles y reducir la correlación entre ellos.
- Profundidad máxima (*max_depth*): Determina el nivel máximo de ramificación de los árboles para evitar que el modelo capture ruido de los datos de entrenamiento (sobreajuste).
- Número mínimo de muestras para dividir un nodo (*min_samples_split*): controla la complejidad de los árboles al limitar un número mínimo de observaciones para realizar las divisiones internas.
- Número mínimo de muestras en una hoja (*min_samples_leaf*): establece un tamaño mínimo de datos para las hojas terminales (predicciones).

3.1.4.3 XGBOOST (XGB)

Al igual que Random Forest, este modelo se basa en la construcción de árboles de decisión, con la diferencia de que los árboles en este caso se entrenan de manera secuencial mediante un proceso llamado *boosting*. En este modelo cada árbol nuevo se construye específicamente para corregir los errores residuales de árboles anteriores mediante el *descenso de gradiente*, un algoritmo de optimización que busca reducir el error de forma iterativa, identificando la dirección en la que la función de pérdida (medida de la discrepancia entre los valores reales y las predicciones del modelo) disminuye más rápidamente para ajustar las predicciones en cada paso [18]. Además, XGBoost permite incorporar términos de regularización que actúan como un mecanismo de control que penaliza la complejidad excesiva de los árboles para obligar al modelo a mantener una estructura más simple y evitar el sobreajuste. Para controlar y optimizar este modelo se

deben ajustar algunos hiperparámetros similares a los definidos en la sección de Random Forest como la profundidad máxima o el número de árboles, y adicionalmente, algunos específicos del algoritmo como:

- Tasa de aprendizaje (`learning_rate`): Determina la magnitud del ajuste realizado por cada nuevo árbol.
- Fracción de muestras (`subsample`): Cantidad de observaciones utilizadas para entrenar cada árbol.
- Fracción de variables (`colsample_bytree`): Define la proporción de variables que se considera al entrenar cada árbol.
- Parámetros de regularización (`lambda` y `alpha`): penalizan la complejidad del modelo mediante regularización L1 y L2, respectivamente, que controlan el tamaño de los pesos de las hojas.

3.1.4.4 PERCEPTRÓN MULTICAPA (MLP)

Consiste en una arquitectura de red neuronal que se basa en una extensión del perceptrón, el cual consiste en una combinación lineal de las características de entrada, ponderadas mediante pesos y ajustadas por un término de sesgo, cuya salida se transforma mediante una función de activación no lineal [18]. Cuando estos perceptrones se organizan en múltiples capas: una capa de entrada con las características, una o varias capas ocultas y una capa de salida con la predicción, se denomina una red neuronal de tipo MLP.

El proceso de aprendizaje de esta red ocurre en dos etapas: primero, se realiza la propagación hacia adelante (*forward propagation*), donde los datos de entrada atraviesan la arquitectura de la red para generar una predicción basada en los pesos y sesgos actuales. Después, ocurre la propagación hacia atrás (*backward propagation*), donde el modelo analiza las diferencias entre la predicción y el valor real utilizando una función de pérdida, y esta información se utiliza en un proceso de optimización basado en gradientes que actualizan cada peso y sesgo, para minimizar el error en las próximas predicciones. Para regular el aprendizaje y complejidad de la red se ajustan algunos hiperparámetros como:

- Optimizador: Algoritmo basado en gradiente que se encarga de actualizar los pesos.
- Tasa de aprendizaje (`learning_rate`): Define la magnitud del ajuste aplicado a los pesos en cada iteración del proceso de optimización.
- Tamaño del lote (`batch_size`): Número de muestras que se pasan por la red en cada iteración del entrenamiento.

- Épocas (epochs): Número de iteraciones que se realiza sobre el conjunto de datos completo.
- Arquitectura de la red: Define el número de capas y neuronas de la red, para controlar su complejidad.
- Función de activación: Transformación no lineal que se aplica a las neuronas que permite aprender patrones complejos.
- Parámetros de regularización (dropout y L1/L2): al igual que en XGBoost, penalizan la complejidad de la red o, en el caso del dropout, desactivan aleatoriamente neuronas durante el entrenamiento para evitar el sobreajuste.

3.1.4.5 PROCESO GAUSSIANO (GP)

Es un modelo probabilístico que construye una distribución sobre funciones, es decir, que en lugar de asumir una única ecuación para describir la relación entre las variables de entrada y de salida, considera un conjunto de funciones posibles y asigna mayor probabilidad a los que resultan más coherentes con los datos observados en el entrenamiento [19]. De tal manera que, este modelo no solo estima una predicción media, sino que también la incertidumbre asociada a dicha observación. En regresión, este modelo se define por una función de media y una función de covarianza o *kernel*, que cuantifica la similitud entre las entradas, por lo que, puntos de entrada cercanos tienden a producir salidas similares [20]. El comportamiento del modelo depende de la elección del *kernel* y de sus hiperparámetros, los cuales se ajustan mediante un optimizador durante el entrenamiento, maximizando la *log-verosimilitud marginal*, la cual premia un buen ajuste a los datos mientras penaliza la complejidad excesiva para evitar el sobreajuste. Entre los hiperparámetros que pueden ajustarse tanto para el modelo como para el optimizador se encuentran [21]:

- *Kernel*: define la estructura de similitud entre observaciones y, por tanto, la forma de las funciones que el modelo puede aprender.
- Length-scale: controla qué tan rápido cambia la función respecto a las variables de entrada. Valores pequeños permiten variaciones más bruscas y valores grandes generan funciones más suaves.
- Varianza de la señal: regula la amplitud máxima de las oscilaciones de la función aprendida en el eje de la variable objetivo.
- Nivel de ruido: representa la variabilidad no explicada por la función subyacente y permite modelar el error de observación.
- Parámetro ν en Matérn: controla la suavidad de la función estimada cuando se utiliza este *kernel*.

- Tasa de aprendizaje (`learning_rate`): Controla el tamaño del paso que da el optimizador hacia el óptimo de la función de *log-verosimilitud marginal* en el entrenamiento.
- Número de iteraciones (`training_iterations`): Número de pasos que realiza el optimizador para ajustar los hiperparámetros.

3.1.5 MÉTRICAS DE EVALUACIÓN

La evaluación del desempeño de los modelos es una etapa crítica que permite cuantificar qué tan cerca están las predicciones de la realidad observada. Para los modelos de regresión, se utilizan métricas que miden el error o la desviación de las predicciones, permitiendo comparar la eficacia de los modelos [22]. Las métricas seleccionadas para este proyecto son:

- Error Cuadrático Medio (MSE): Consiste en el promedio de los cuadrados de las diferencias entre los valores reales y los predichos. Al elevar los residuos al cuadrado, esta métrica penaliza de manera más severa las desviaciones grandes. Cuanto más cercano a cero sea este valor, mejor es el ajuste del modelo
- Raíz del Error Cuadrático Medio (RMSE): Es la raíz cuadrada del MSE. Su principal ventaja es que devuelve el error en las mismas unidades que la variable objetivo, facilitando su interpretación.
- Error Absoluto Medio (MAE): A diferencia del MSE, esta métrica calcula el promedio de las diferencias absolutas entre la predicción y el valor real por lo que es menos sensible a los valores atípicos.
- Coeficiente de Determinación (R^2): Esta métrica indica la proporción de la varianza de la variable de respuesta que es explicada por los predictores del modelo. Un R^2 cercano a 1 sugiere que el modelo captura gran parte de la variabilidad de los datos, mientras que un valor cercano a 0 indica un ajuste pobre.

3.2. ANTECEDENTES

Ante la necesidad de enfrentar los desafíos provenientes del control de parámetros de calidad, como es el caso de la finura, diversos investigadores han realizado aproximaciones de modelos de predicción aplicando técnicas de aprendizaje automático para hacer estimaciones robustas y precisas. A continuación, se analizan diferentes estudios relevantes en este campo:

- Pani & Mohanta [6] evaluaron tres modelos de redes neuronales y Least Squares Support Vector Regression (LS-SVR) para predecir el Blaine en un molino vertical de

cemento, utilizando datos históricos de tres variables operacionales: la tasa de alimentación del material de entrada, el flujo de aire a través de molino y la velocidad del separador. Utilizaron un enfoque riguroso en el preprocesamiento de los datos aplicando tres técnicas de identificación de *outliers* (desviación estándar, diagramas de caja y método de Hampel), seguido de una división representativa de los datos en entrenamiento y prueba mediante el algoritmo CADEX. El modelo de redes neuronales con *Back Propagation Resiliente* obtuvo el mejor desempeño alcanzando un error absoluto porcentual medio (MAPE) de 6.4% en el conjunto de validación. Este estudio evidencia el potencial de las técnicas de aprendizaje automático para estimar finura a partir de variables operacionales, sin embargo, consideró únicamente tres variables de entrada, dejando por fuera otras condiciones del proceso que pueden ser influyentes en la molienda.

- Lange & Van Zyl [23] evaluaron múltiples modelos para predecir el Blaine en una planta china, incluyendo modelos lineales (regresión lineal y Lasso) y no lineales (redes neuronales, ANFIS y LSTM). La regresión lineal con reentrenamiento *online* logró el mejor R^2 (0.481) demostrando que modelos simples pueden obtener resultados aceptables en circuitos complejos. No obstante, debido a su menor rendimiento en comparación con otros estudios, se resalta que diferencias en el diseño del molino, calidad de las materias primas, ajustes operativos o incluso el desgaste de los equipos, hacen que sea difícil generalizar los modelos que funcionan bien en una planta a otros contextos. Este antecedente es relevante pues muestra que el desempeño de los modelos depende fuertemente del contexto operativo y que una mayor complejidad del algoritmo no siempre implica mejores resultados.
- Stanišić et al. [24] desarrollaron un sensor virtual para estimar la finura del cemento en un proceso de molienda con molino de bolas, utilizando variables operacionales del proceso con diferentes rezagos temporales, como alimentación, rechazos del molino, presiones, velocidad del separador y temperatura del producto. El modelo propuesto se basó en una red neuronal de funciones de base radial (RBF ANN), cuya estructura fue seleccionada de forma automática mediante técnicas de agrupamiento y criterios estadísticos (*subtractive clustering*, *k-means* y criterio de información de Akaike, *AIC*). Los resultados mostraron un buen desempeño predictivo, con un MSE de 0.0103 y un R^2 de 0.7953, superando a una red neuronal *feed-forward* usada como referencia. Este estudio aporta una perspectiva importante sobre el uso de sensores virtuales para estimar finura en línea, aunque su aplicación se desarrolla en un molino de bolas y su énfasis principal está en automatizar la construcción de la arquitectura del sensor.

- Belmajdoub & Abderafi [5] desarrollaron un sensor virtual para predecir la finura de harina cruda en un molino vertical, utilizando redes neuronales y Support Vector Regression (SVR). Su enfoque destacó la importancia de incluir variables de calidad del material para mejorar la precisión del modelo. Entre los modelos evaluados, el mejor desempeño se obtuvo con una red neuronal entrenada mediante *Bayesian Regularization* con un valor de R^2 de 0.9698. Este estudio demuestra que las características del material pueden aportar información valiosa para la estimación de la finura, sin embargo, se enfoca en harina cruda, correspondiente a una etapa previa a la molienda de cemento, por lo que sus resultados no son directamente transferibles a la predicción en el producto final.
- Topaloğlu et al. [25] desarrollaron modelos de aprendizaje automático para predecir la finura del producto final en la producción de cemento, utilizando datos industriales de control de calidad provenientes de un sistema LIMS. A diferencia de otros estudios basados únicamente en variables operativas, este trabajo incorporó variables de composición química obtenidas por XRF, módulos químicos como LSF, MS y MA, descriptores granulométricos de difracción láser y una medición intermedia de finura denominada Blaine-P. Los autores evaluaron modelos como Ridge Regression, Random Forest, XGBoost, LightGBM y CatBoost, ajustados mediante *RandomizedSearchCV* con validación cruzada de cinco particiones y evaluados en un conjunto de prueba independiente. Los mejores resultados fueron obtenidos por XGBoost, con un R^2 de 0.754 y un RMSE de 76.9 cm^2/g . Este antecedente refuerza la utilidad de los modelos basados en árboles para capturar relaciones no lineales en procesos de molienda y además evidencia el aporte que pueden tener las variables químicas y granulométricas cuando están disponibles de forma sistemática.
- Fatahi et al. [26] desarrollaron un modelo para la predicción de la presión de trabajo en un molino vertical de cemento utilizando variables de operación como tasas de alimentación, presión diferencial, la inyección de agua, entre otras variables del circuito. Para ello, compararon el desempeño de dos modelos basados en árboles: Random Forest y XGBoost, con este último demostrando el mejor desempeño con un R^2 cercano a 0.98 en el conjunto de prueba. Adicionalmente, el estudio implementó el método SHAP (*SHapley Additive exPlanations*) para la cuantificación del impacto de cada variable operativa en la predicción final, lo que aportó interpretabilidad al modelo y brindó información útil para la toma de decisiones operativas. Aunque este estudio no predice directamente la finura, constituye un

antecedente relevante, ya que muestra el potencial de modelos como XGBoost para capturar relaciones complejas y modelar la dinámica real del proceso de molienda.

En conjunto, estos antecedentes evidencian el potencial de los modelos basados en datos para estimar variables de calidad u operación en procesos de molienda. Sin embargo, varios se enfocan en Blaine o harina cruda, utilizan molinos de bolas o incorporan variables de calidad del material que no siempre están disponibles en línea. Por ello, este proyecto se diferencia al estimar el retenido en malla 325 del cemento usando variables operacionales medidas por sensores en planta, con énfasis en desempeño predictivo, interpretabilidad y aplicabilidad operativa.

4. PREPARACIÓN Y ANÁLISIS EXPLORATORIO

El conjunto de datos utilizado en este estudio corresponde a los registros históricos del proceso de molienda de cemento obtenidos entre el 1 de enero de 2022 y el 30 de abril de 2025. El *dataset* final integra las variables operacionales, capturadas automáticamente con resolución horaria por el sistema de control del molino vertical, con las mediciones de laboratorio de finura (retenido en malla 325).

Para la consolidación de la información, se implementó una estrategia de sincronización que consideró que los datos del laboratorio (finura) podían ser horarios o muestras compuestas de dos o más horas, por lo cual se realizó un procedimiento de *join* que identificó las características de cada registro de finura para alinearlos correctamente con los intervalos horarios correspondientes de las variables operativas. Adicionalmente, se filtraron los datos correspondientes a la producción de cemento de tipo UG, para enfocar el estudio en este tipo de molienda predominante y no sesgar el modelo con datos de producción de cementos especiales que se hacen con baja frecuencia. Como resultado de esta integración, se dispone de un total de 18 090 observaciones en el *dataset* inicial.

Para el análisis se incluyeron 23 variables: información del proceso (presiones, temperaturas, velocidades y corrientes eléctricas), de producción (toneladas alimentadas de cada material, recirculación, inyección de agua) y de calidad (retenido). La base de datos se limitó a variables de sensores disponibles en tiempo real, ya que el objetivo es construir un estimador de finura independiente de mediciones analíticas adicionales. Aunque variables como composición química, dureza o granulometría podrían fortalecer el modelado, estas no se registran con la misma periodicidad, por lo que su inclusión condicionaría el sistema a fuentes externas y limitaría su uso como herramienta de apoyo operativo inmediato. La Tabla 1 resume cada variable, su tipo de dato y una breve descripción de su significado dentro del proceso.

Tabla 1. Descripción de las variables.

Variable	Tipo	Descripción
fecha_hora	Temporal	Fecha y hora del registro.
dp_molino	Numérica	Diferencial de presión a través del molino como un indicador de carga interna del molino (mbar).
dp_filtro	Numérica	Diferencial de presión en el filtro principal, como un indicador de carga interna en el filtro (mbar)
t_entrada_molino_r1, t_entrada_molino_r2	Numérica	Temperaturas del gas de entrada al molino por cada ramificación (°C)
t_salida_molino	Numérica	Temperatura del gas a la salida del molino (°C)
rpm_ventilador_tiro	Numérica	Velocidad de rotación del ventilador principal de tiro (rpm)
i_ventilador_tiro	Numérica	Corriente eléctrica del ventilador principal (A)
p_motor_molino	Numérica	Potencia del motor principal del molino (kW)
rpm_separador	Numérica	Velocidad del separador dinámico (rpm)
i_elevador	Numérica	Corriente eléctrica del elevador de rechazos (A)
ton_pesador	Numérica	Flujo total de alimentación al molino (ton/h)
ton_clinker, ton_caliza, ton_yeso, ton_scm1, ton_scm2	Numérica	Toneladas por hora de cada material alimentado (ton/h)
ton_recirculacion	Numérica	Flujo de material recirculado (ton/h)
ppm_aditivo	Numérica	Dosis de aditivo líquido (ppm).
iny_agua	Numérica	Caudal de agua inyectada al molino (m ³ /h)
posicion_m1, posicion_m2	Numérica	Posiciones de los rodillos (mm)
retenido	Numérica	Porcentaje de material retenido en malla 325 (%)

Los valores faltantes, presentes en un porcentaje aproximado de 1.8% de los registros, fueron eliminados por su baja proporción, sin afectar la representatividad del conjunto de datos. Por otro lado, el tratamiento de valores atípicos se orientó inicialmente a la identificación de las observaciones con niveles bajos en variables como `p_motor_molino` que podrían corresponder a paros o condiciones de operación inestables del molino. Para ello, se utilizó como umbral el percentil 5% y se seleccionó aquellas observaciones por debajo de este límite que coincidían simultáneamente con caídas en otras variables operativas tales como la alimentación del molino (`ton_pesador`) o la carga interna (`dp_molino`). La Figura 2 muestra los datos identificados como paro:

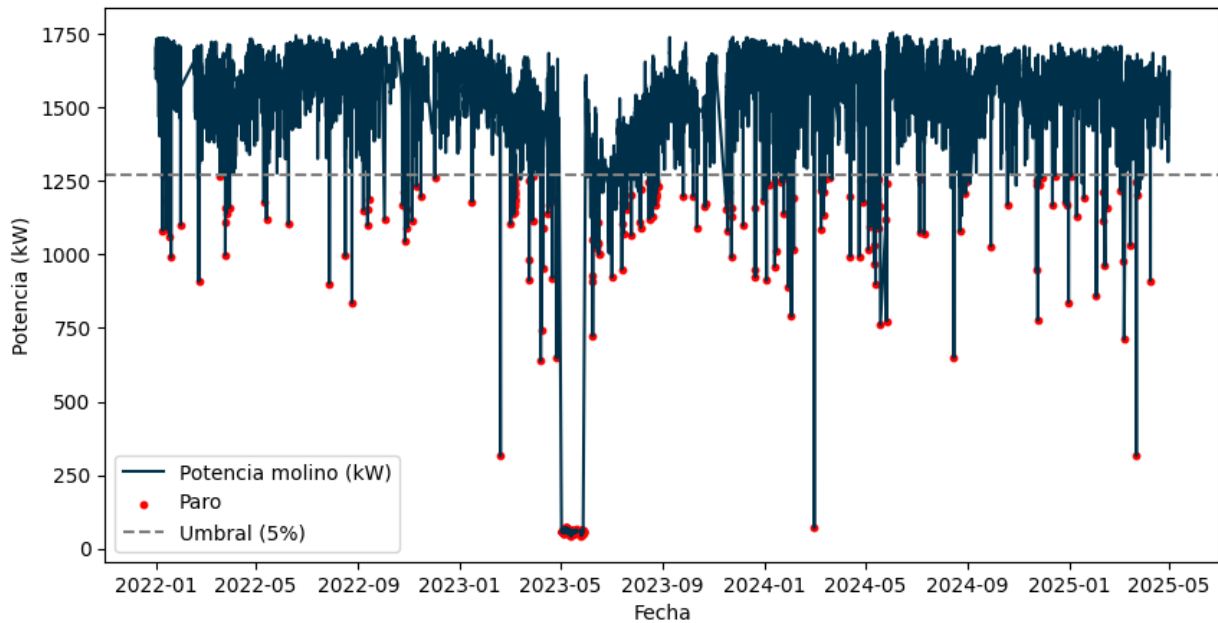


Figura 2. Potencia del molino y detección de paros.

En estos casos, los valores operacionales y la finura medida podrían no representar adecuadamente la operación normal de molienda, por lo que tales registros fueron eliminados para evitar sesgos en el modelado.

Durante la limpieza de atípicos, se detectaron registros negativos en variables como `ton_recirculacion`, `ppm_aditivo` e `iny_agua`. Desde el punto de vista operativo, estos sistemas funcionan de manera intermitente, activándose únicamente cuando las condiciones de molienda o el tipo de producto lo requieren, por lo tanto, los valores negativos se atribuyeron a desviaciones de los sensores cuando los equipos se encontraban detenidos, lo que llevó a realizar un ajuste igualando estos registros a 0. Adicionalmente, para el tratamiento de valores atípicos, se consideró que el uso exclusivo de criterios estadísticos tradicionales, como el rango intercuartílico o reglas basadas en desviaciones

estándar, podían no ser suficientes para diferenciar entre datos realmente anómalos y condiciones operativas válidas. En procesos industriales, los sensores suelen mantenerse dentro de márgenes estrechos debido a la estabilidad buscada en la planta, por ello, estas pequeñas desviaciones que pueden ser identificadas por métodos estadísticos como atípicas pueden en realidad reflejar cambios válidos como ajustes para retornar el proceso a especificación o transiciones propias de la operación, por lo que eliminarlas podría reducir la capacidad del modelo para aprender el comportamiento del molino ante diferentes condiciones reales. Por esta razón, el tratamiento de atípicos se apoyó en criterios operativos, para eliminar únicamente eventos no representativos como fallas de sensores, paros, lecturas físicamente no realistas o estados anormales del molino. La siguiente tabla resume los criterios de ajuste aplicados:

Tabla 2. Resumen de la limpieza de variables.

Variable	Umbral aplicado	Justificación operativa
rpm_ventilador_tiro	≥ 1150 rpm	Valores menores no representan una condición real de operación en molienda UG
dp_molino	≥ 27 mbar	Lecturas muy bajas indican falla del sensor o estado anormal del molino
retenido	≤ 8 %	Valores atípicos de laboratorio que pueden corresponder a errores de digitación
ton_recirculacion	Negativos $\rightarrow 0$	Valores negativos aparecen cuando no hay recirculación de material
ppm_aditivo	Negativos $\rightarrow 0$	Valores negativos aparecen cuando no se dosifica aditivo
iny_agua	Negativos $\rightarrow 0$	Valores negativos aparecen cuando no se inyecta agua por humedad suficiente del material de entrada

Una vez tratados los atípicos, se realizó un análisis de Factor de Inflación de la Varianza (VIF) para evaluar la multicolinealidad entre las variables del conjunto de datos y evitar redundancias que pudieran afectar la estabilidad del modelo. Los resultados evidenciaron una alta correlación entre las variables de posición de los rodillos (posicion_m1 y posicion_m2), con valores VIF de 17.13 y 18.77 respectivamente. Para mitigar esta

colinealidad, se decidió combinarlas en una nueva variable sintética denominada *posicion_promedio*. Adicionalmente, se eliminó la variable de alimentación total (*ton_pesador*), dado que esta representa la suma de la cantidad de cada material (como *ton_clinker*, *ton_caliza*, etc.) y dicha información ya se encontraba detallada individualmente en el *dataset*, volviendo innecesaria su inclusión. Si bien la regresión lineal múltiple es el modelo más sensible a este problema debido a la inestabilidad que puede generar en la estimación e interpretación de los coeficientes, se decidió aplicar la misma depuración de variables para todos los modelos evaluados con el fin de mantener un conjunto de datos homogéneo y permitir una comparación más justa entre los diferentes enfoques. Además, aunque modelos como los basados en árboles o redes neuronales suelen ser más robustos frente a variables correlacionadas en términos predictivos, la presencia de variables redundantes puede afectar la interpretación de la importancia de variables, distribuyendo artificialmente la influencia entre predictores que contienen información similar. Finalmente, como resultado de este proceso de depuración, el conjunto de datos final para el modelado se consolidó en 16 491 registros.

Para comenzar el análisis exploratorio, se analizó el comportamiento de la variable objetivo mediante un gráfico de distribución:

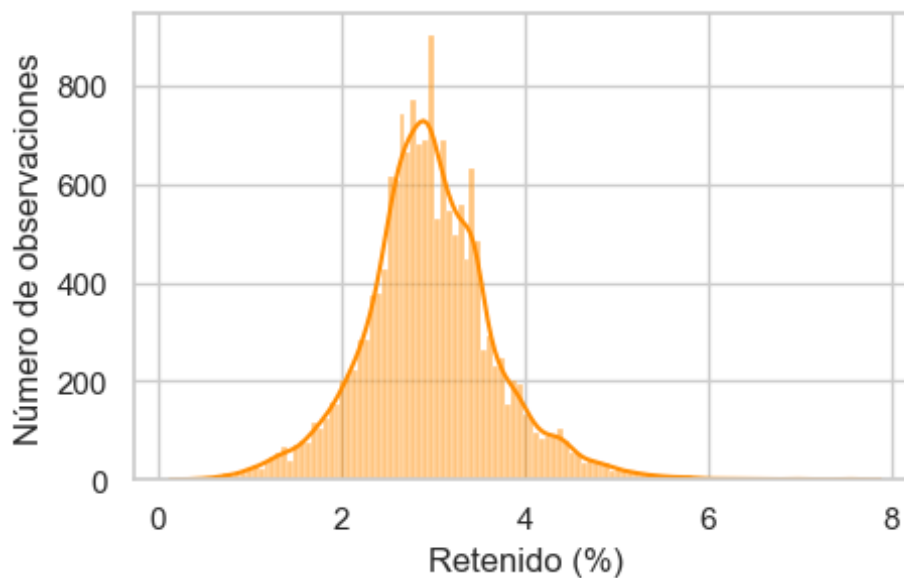


Figura 3. Distribución de la finura medida como retenido en malla 325.

Como se observa en la figura 3, los datos siguen una distribución unimodal con forma de campana, concentrándose la mayoría alrededor de un promedio de 2.96%. Esto indica que el molino opera generalmente de manera estable, aunque la cola que se extiende hacia la

derecha (valores altos) revela los momentos específicos en los que la molienda fue menos eficiente.

Posteriormente, se llevó a cabo un análisis de las variables predictoras mediante los coeficientes de correlación de Pearson, Spearman y correlación por distancia (Figura 4).

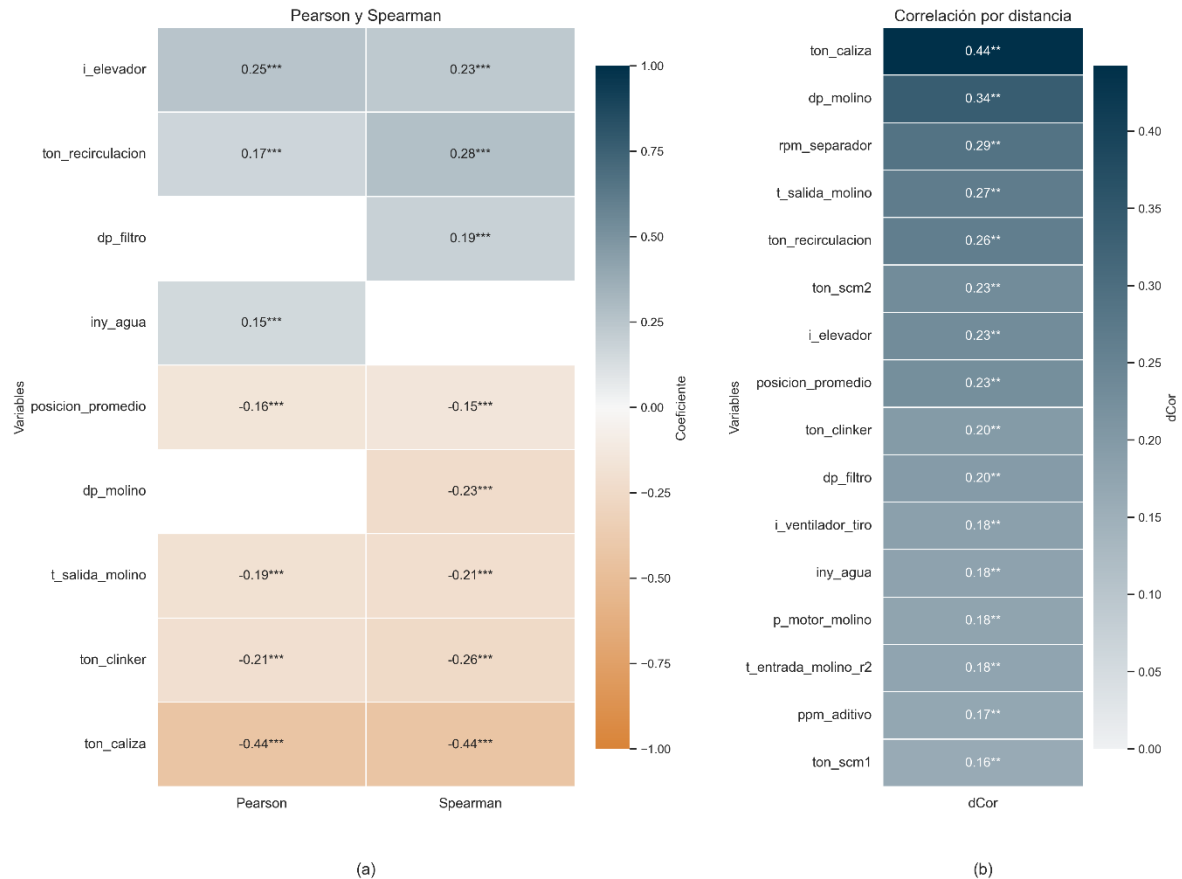


Figura 4. Correlaciones de las variables operativas con la variable de salida: (a) coeficientes de Pearson y Spearman; (b) correlación por distancia (dCor).

Inicialmente, los resultados de Pearson y Spearman mostraron asociaciones significativas del retenido con algunas variables, entre las que sobresale ton_caliza con una relación inversa, junto con otras variables con correlaciones bajas a moderadas como ton_clinker, t_salida_molino, i_elevador y ton_recirculacion. Sin embargo, la mayoría de las variables analizadas mostraron correlaciones bajas, lo que sugiere que su relación con la finura no se puede describir únicamente mediante patrones lineales o monotónicos simples. Por su parte, la correlación por distancia permitió captar dependencias más complejas, mostrando valores significativos para algunas variables relevantes en la operación como rpm_separador, dp_molino y dp_filtro, cuya importancia no era tan evidente en las métricas tradicionales. En general, estos resultados respaldan la hipótesis de que la finura

está influida principalmente por interacciones complejas y potencialmente no lineales entre variables de composición de la alimentación, como `ton_caliza`, y variables operativas del proceso, como la carga interna del molino (`dp_molino`), la velocidad del separador (`rpm_separador`), entre otras.

Debido a la alta correlación que presenta la variable `ton_caliza` con la finura, se examinó en detalle la influencia de la composición de la mezcla sobre el producto final. En este caso, con base en el conocimiento de la naturaleza de los diferentes materiales durante la molienda, se calculó un índice composicional definido como la razón entre la caliza y los materiales que comúnmente presentan mayor resistencia a la molienda, cuya distribución se muestra en la siguiente gráfica:

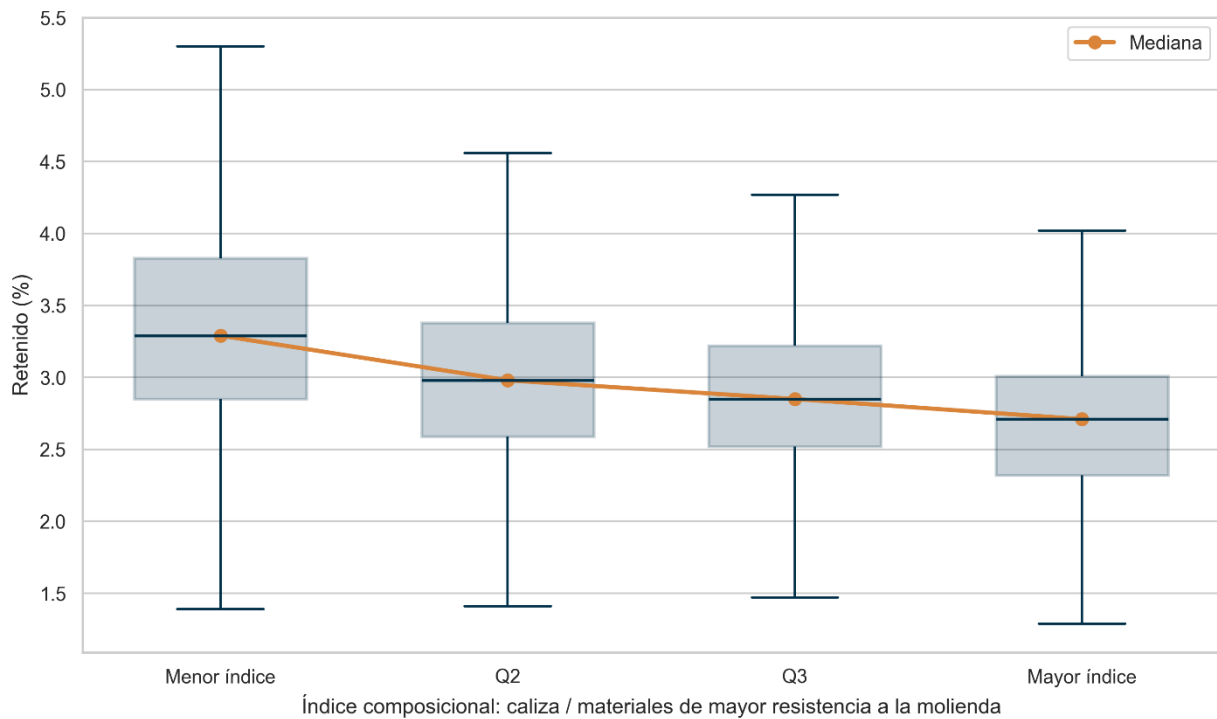


Figura 5. Distribución del retenido según niveles del índice composicional de la mezcla.

En la figura 5, se observa una disminución progresiva de la mediana del retenido a medida que aumenta la razón caliza/materiales de mayor resistencia a la molienda, lo que sugiere que mezclas relativamente más ricas en caliza y con menor proporción de componentes de baja molturabilidad tienden a asociarse con menores valores de retenido. Aunque se observa una alta variabilidad dentro de cada grupo atribuida a influencia de otras variables operativas del proceso, la tendencia general respalda la importancia de la composición de la alimentación en el comportamiento de la finura.

Complementariamente al análisis de la composición, se evaluó el impacto de las variables operativas críticas, como la presión diferencial y la velocidad del separador, mediante gráficos de densidad conjunta:

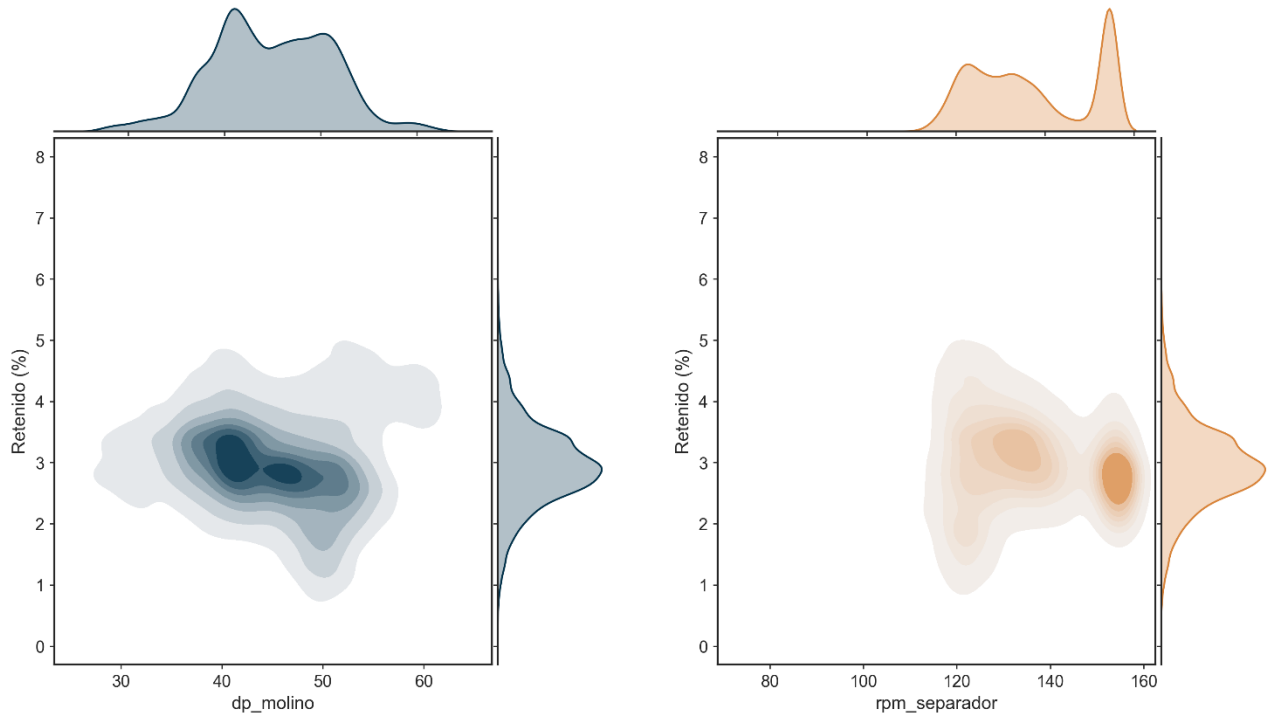


Figura 6. Densidad conjunta de dp_molino y rpm_separador respecto al retenido.

Como se observa en la gráfica 6, tanto la presión diferencial del molino (dp_molino) como la velocidad del separador (rpm_separador) presentan patrones de asociación con el retenido que no parecen responder a un comportamiento lineal simple. Específicamente, para la presión diferencial se aprecia una mayor concentración de observaciones en un rango aproximado entre 35 y 50 mbar, dentro del cual el retenido muestra una dispersión menor en comparación con valores fuera de este intervalo. Por otra parte, el comportamiento del separador sugiere que mayores velocidades de rotación favorecen la finura del producto, aunque se detectó una mayor dispersión en los datos a velocidades bajas, en el rango entre 110 y 130 rpm. En conjunto, estas observaciones validan que el molino no responde linealmente a los ajustes, sino que requiere operar dentro de rangos específicos que favorecen la estabilidad del proceso.

Finalmente, se evaluó el comportamiento del retenido respecto a otra variable operativa clave, la posición de los rodillos:

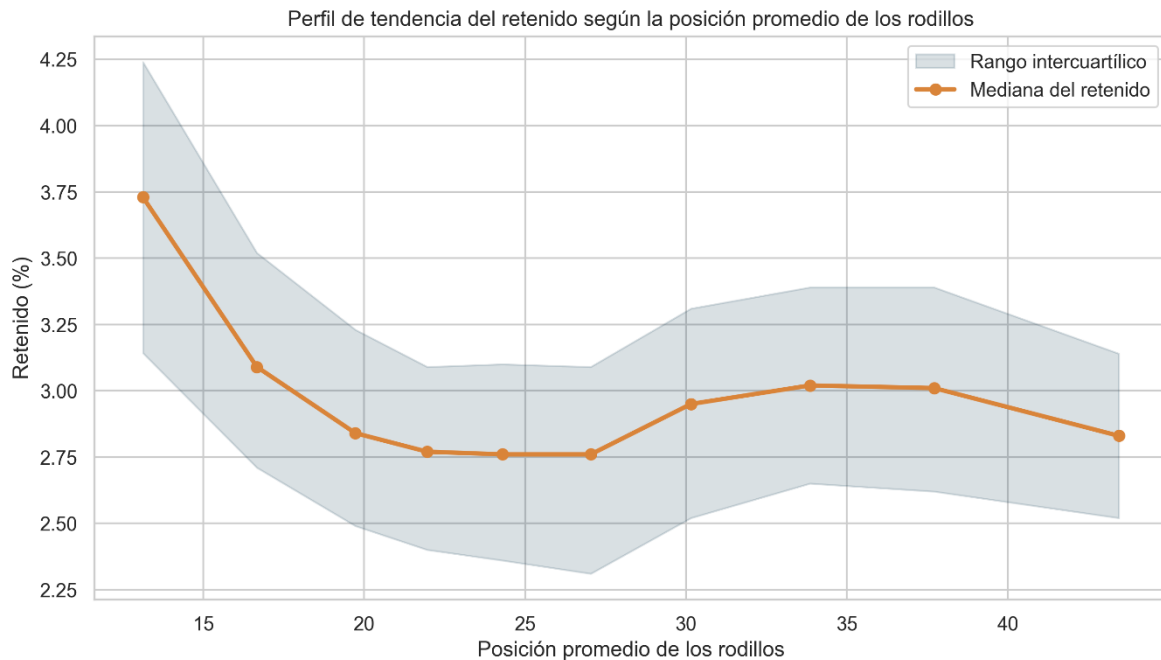


Figura 7. Tendencia del retenido según la posición promedio de los rodillos.

En la figura 7, se observa que la relación entre la variable `posicion_promedio` y el retenido no sigue un patrón lineal simple: En posiciones bajas de los rodillos se presentan mayores valores de retenido, mientras que al avanzar hacia un rango intermedio el retenido disminuye de forma apreciable para finalmente mostrar un leve incremento en posiciones medias-altas, lo que sugiere la existencia de rangos operativos diferenciados. Esta tendencia se puede explicar desde el punto de vista operativo, pues la posición promedio de los rodillos puede interpretarse como un indicador indirecto del estado de la cama de molienda: los valores más bajos de posición podrían asociarse con una baja cama de molienda y mayor inestabilidad, generando una pérdida de eficiencia y por consiguiente el retenido tiende a aumentar, mientras que en rangos intermedios se evidencian condiciones más favorables para la molienda y en los rangos más altos, un posible exceso de material en el molino podría dificultar el proceso de molienda, haciendo que el retenido aumente nuevamente.

5. CONSTRUCCIÓN Y OPTIMIZACIÓN DE LOS MODELOS

Con base en los hallazgos del análisis exploratorio, se procedió a la construcción de cinco modelos supervisados de regresión para obtener una estimación de la finura del cemento: regresión lineal múltiple (RLM), Random Forest (RF), XGBoost (XGB), perceptrón multicapa (MLP) y proceso gaussiano (GP). Se seleccionaron estos modelos para comparar enfoques de distinta complejidad, desde un modelo lineal base hasta métodos no lineales de ensamble, redes neuronales y modelos probabilísticos capaces de representar

incertidumbre. El desarrollo computacional del proyecto se realizó en Python, empleando librerías especializadas para el análisis de datos y el aprendizaje automático.

Para todos los modelos se definió como variable objetivo el retenido y como variables predictoras el conjunto de variables operacionales resultante tras la etapa de preparación y limpieza de los datos. Para todos los modelos se utilizó una partición aleatoria de 80% para entrenamiento y 20% para prueba, con semilla fija (`random state = 42`) para garantizar reproducibilidad, con el fin de reservar un conjunto independiente para la evaluación final del desempeño y un conjunto de entrenamiento para el ajuste y optimización de cada algoritmo. Adicionalmente, para todos los modelos se aplicó una técnica de validación cruzada K-Fold con 5 particiones sobre el conjunto de entrenamiento, con el fin de evaluar la estabilidad de los modelos y reducir la dependencia de los resultados frente a una sola partición de los datos. En los modelos con búsqueda de hiperparámetros, los resultados de la validación cruzada se utilizaron como criterio para la selección de la mejor combinación de parámetros, priorizando el RMSE como criterio principal de comparación debido a su interpretabilidad por estar expresado en las mismas unidades de la variable objetivo. Una vez elegida la mejor configuración, cada modelo se entrenó nuevamente con todo el conjunto de entrenamiento, para finalmente evaluar su capacidad de generalización en el conjunto de prueba mediante métricas como MAE, MSE, RMSE y R^2 . En general, el flujo de trabajo utilizado para cada modelo se resume en la figura 8:



Figura 8. Diagrama del flujo de trabajo para los modelos.

Al ser modelos de distinta naturaleza, su construcción y optimización tienen especificidades diferentes que se detallan a continuación:

- **Regresión lineal múltiple:** Se implementó como modelo base de referencia. Dado que no requiere una búsqueda iterativa de hiperparámetros, se ajustó directamente sobre los datos de entrenamiento. Como complemento para la verificación del modelo, se utilizó la librería `statsmodels` para examinar gráficamente el cumplimiento de los supuestos clásicos de la regresión lineal (diagramas de residuos versus valores ajustados, gráfico Q-Q, *scale-location* y

residuos versus *leverage*). Este análisis sirvió para diagnosticar las limitaciones de este enfoque lineal frente a la dinámica real del molino.

- **Random Forest:** La optimización de este modelo de ensamble se llevó a cabo mediante la estrategia RandomizedSearchCV. El espacio de búsqueda definido para optimizar la estructura y el crecimiento del bosque se detalla en la Tabla 3. Con el objetivo de extraer información sobre la influencia de las diferentes variables operativas en la variable objetivo, el modelo final se analizó utilizando la técnica de valores SHAP (*Shapley Additive exPlanations*), permitiendo aislar e interpretar detalladamente el efecto individual de cada variable en la predicción.
- **XGBoost:** Al igual que en Random Forest, la optimización de los hiperparámetros de crecimiento y regularización del algoritmo de *boosting* se ejecutó utilizando RandomizedSearchCV (los rangos explorados se presentan en la Tabla 3). La interpretabilidad del modelo ajustado se desarrolló también mediante el análisis de valores SHAP.
- **Perceptrón multicapa:** La arquitectura de la red neuronal se implementó mediante TensorFlow/Keras. Esta consistió en una capa de entrada, dos capas ocultas densas con activación ReLU y regularización dropout, y una neurona de salida lineal. Debido a la alta sensibilidad de este modelo frente a la escala de los datos, el modelo fue entrenado dentro de un Pipeline junto con StandardScaler. El espacio de hiperparámetros (Tabla 3) se exploró fijando un límite máximo de 500 épocas. El entrenamiento fue regulado mediante el mecanismo de *early stopping*, configurado con una paciencia de 15 épocas sobre un subconjunto de validación interno del 10 %, para prevenir el sobreajuste.
- **Proceso gaussiano:** La implementación probabilística se desarrolló en la librería GPyTorch. En este caso, se empleó un *kernel Matérn* ajustado mediante determinación automática de relevancia (ARD), precedido por una etapa de estandarización de los datos (StandardScaler). En el algoritmo se incorporó de manera explícita la incertidumbre conocida del ensayo de laboratorio (FixedNoiseGaussianLikelihood), con el objetivo de aislar matemáticamente el error del modelo respecto a la variabilidad de la medición física. Dada su arquitectura interna, la optimización no empleó RandomizedSearchCV, sino una validación cruzada estructurada sobre una malla definida para optimizar la tasa de aprendizaje, el número de iteraciones y la topología del *kernel* (Tabla 3). El ajuste final de GP permitió extraer la desviación estándar latente y total, con resultados de las predicciones y sus respectivos intervalos de confianza del 95 %.

Los hiperparámetros considerados, los métodos de búsqueda y el número de configuraciones evaluadas en la etapa de optimización se resumen en la siguiente tabla:

Tabla 3. Resumen de búsqueda de hiperparámetros.

Modelo	Método de búsqueda	Hiperparámetros ajustados	Espacio de búsqueda	Número de configuraciones evaluadas
Regresión Lineal Múltiple (RLM)	No aplica	No aplica	No aplica	No aplica
Random Forest (RF)	RandomizedSearchCV	n_estimators	{200, 400, 600}	20
		max_depth	{10, 20, 30}	
		min_samples_split	{2, 5, 10}	
		min_samples_leaf	{1, 2, 4}	
		max_features	{sqrt, 0.3, 0.5}	
XGBoost (XGB)	RandomizedSearchCV	n_estimators	{200, 400, 600}	20
		max_depth	{3, 5, 7, 10}	
		learning_rate	{0.01, 0.05, 0.1, 0.2}	
		subsample	{0.7, 0.8, 0.9, 1.0}	
		colsample_bytree	{0.5, 0.7, 0.9, 1.0}	
		min_child_weight	{1, 3, 5}	
Perceptrón multicapa (MLP)	RandomizedSearchCV	n_hidden_1	{16, 32, 64}	20
		n_hidden_2	{8, 16, 32}	
		dropout_rate	{0.1, 0.2, 0.3}	
		learning_rate	{0.0005, 0.001, 0.005}	
		batch_size	{16, 32, 64}	
Proceso Gaussiano (GP)	Búsqueda manual en malla reducida	lr	{0.1, 0.2}	12
		training_iterations	{100, 150}	
		nu	{0.5, 1.5, 2.5}	

Como resultado del procedimiento de optimización descrito anteriormente, se seleccionó para cada algoritmo la configuración que obtuvo el menor RMSE promedio en validación cruzada. La Tabla 4 resume las configuraciones finalmente adoptadas para el reentrenamiento de los modelos sobre todo el conjunto de entrenamiento y su posterior evaluación en el conjunto de prueba.

Tabla 4. Resumen de mejor configuración de hiperparámetros para cada modelo.

Modelo	Mejor configuración seleccionada
RLM	No aplica
Random Forest (RF)	n_estimators = 600; max_depth = 30; min_samples_split = 2; min_samples_leaf = 2; max_features = 0.5
XGBoost (XGB)	n_estimators = 600; max_depth = 10; learning_rate = 0.01; subsample = 0.9; colsample_bytree = 0.9; min_child_weight = 5
Perceptrón multicapa (MLP)	n_hidden_1 = 64; n_hidden_2 = 32; dropout_rate = 0.1; learning_rate = 0.0005; batch_size = 64
Proceso Gaussiano (GP)	lr = 0.2; training_iterations = 150; nu = 1.5

6. EVALUACIÓN DE LOS MODELOS PREDICTIVOS

En esta etapa se evaluó el desempeño de los modelos construidos a partir de tres criterios: estabilidad, capacidad de generalización e interpretación práctica de los resultados. La estabilidad se analizó a partir del comportamiento de las métricas obtenidas durante la validación cruzada K-Fold sobre el conjunto de entrenamiento, mientras que la capacidad de generalización se examinó mediante el desempeño final en el conjunto de prueba. Adicionalmente, se consideraron aspectos como el tiempo de ejecución y la interpretabilidad de cada enfoque, con el fin de realizar una comparación integral. Aunque se reportaron varias métricas de evaluación, se priorizó el RMSE como criterio principal de comparación por estar expresado en las mismas unidades de la variable objetivo. Cabe recalcar que en este caso las mediciones de retenido empleadas para entrenar y evaluar los modelos tienen cierta variabilidad inherente cuantificada como la incertidumbre del ensayo de laboratorio (ver Anexo A), y métricas como el RMSE no distinguen entre esta dispersión aleatoria de los datos de entrada y las desviaciones asociadas al modelo [27] por lo que, para efectos prácticos, se utilizó el valor conocido de la incertidumbre del ensayo de

laboratorio como una referencia para contextualizar el orden de magnitud del error: cuando el RMSE se aproxima a dicho valor de incertidumbre, las mejoras adicionales en el ajuste de los modelos deben analizarse con cautela, ya que podrían reflejar un mayor ajuste al ruido experimental más que una mejora en la capacidad predictiva del modelo.

6.1. ESTABILIDAD DE LOS MODELOS EN VALIDACIÓN CRUZADA

Como primer criterio de evaluación, se analizó la estabilidad de los modelos mediante los resultados obtenidos en la validación cruzada K-Fold aplicada sobre el conjunto de entrenamiento. Un resumen de los resultados obtenidos se muestra en la siguiente tabla:

Tabla 5. Desempeño promedio y desviación estándar de los modelos en la validación cruzada K-Fold.

Modelo	RMSE (media $\pm$ Desv. estándar)	R ² (media $\pm$ Desv. estándar)	MSE (media $\pm$ Desv. estándar)	MAE (media $\pm$ Desv. estándar)
Proceso gaussiano (GP)	0.4633 $\pm$ 0.0121	0.5858 $\pm$ 0.0148	0.2148 $\pm$ 0.0114	0.3389 $\pm$ 0.0065
XGBoost (XGB)	0.4727 $\pm$ 0.0121	0.5688 $\pm$ 0.0151	0.2236 $\pm$ 0.0117	0.3453 $\pm$ 0.0065
Random Forest (RF)	0.4745 $\pm$ 0.0135	0.5656 $\pm$ 0.0163	0.2254 $\pm$ 0.0130	0.3455 $\pm$ 0.0072
Perceptrón multicapa (MLP)	0.4912 $\pm$ 0.0117	0.5345 $\pm$ 0.0138	0.2414 $\pm$ 0.0117	0.3636 $\pm$ 0.0081
Regresión lineal múltiple (RLM)	0.6084 $\pm$ 0.0123	0.2857 $\pm$ 0.0203	0.3703 $\pm$ 0.0151	0.4573 $\pm$ 0.0079

En términos de estabilidad y desempeño promedio, el proceso gaussiano presentó el mejor comportamiento global, al registrar los menores valores de MAE, MSE y RMSE, y el mayor coeficiente de determinación. Muy cerca de este se ubicaron XGBoost y Random Forest, cuyos resultados fueron bastante similares entre sí y confirmaron un buen desempeño dentro del conjunto de entrenamiento, mientras que el perceptrón multicapa mostró un comportamiento intermedio, con errores mayores que los modelos basados en árboles. Por su parte, la regresión lineal múltiple presentó el menor desempeño promedio, lo que confirma que un enfoque estrictamente lineal no resulta suficiente para representar adecuadamente la dinámica del proceso. En general, las desviaciones estándar obtenidas fueron bajas para todos los modelos, lo que indica que los resultados fueron consistentes entre particiones.

6.2. CAPACIDAD DE GENERALIZACIÓN EN EL CONJUNTO DE PRUEBA

Una vez analizada la estabilidad interna de los modelos, se evaluó la capacidad de cada modelo de transferir lo aprendido en el entrenamiento a observaciones no vistas utilizando las métricas de desempeño en el conjunto de prueba, cuyos resultados se muestran a continuación:

Tabla 6. Desempeño de los modelos en el conjunto de prueba.

Modelo	RMSE	R ²	MSE	MAE
Proceso gaussiano (GP)	0.4419	0.6217	0.1953	0.3295
XGBoost (XGB)	0.4552	0.5987	0.2072	0.3368
Random Forest (RF)	0.4552	0.5986	0.2072	0.3338
Perceptrón multicapa (MLP)	0.4704	0.5714	0.2213	0.3538
Regresión lineal múltiple (RLM)	0.5979	0.3075	0.3575	0.4536

Los resultados obtenidos mantuvieron el mismo orden general observado en la validación cruzada: el proceso gaussiano conservó el mejor desempeño global, con el menor RMSE y el mayor coeficiente de determinación, seguido muy de cerca por XGBoost y Random Forest, cuyos resultados fueron prácticamente equivalentes. El perceptrón multicapa mostró un desempeño intermedio, superior al modelo lineal pero inferior al proceso gaussiano y a los modelos basados en árboles. En conjunto, estos resultados muestran que los enfoques no lineales ofrecen una mejor capacidad de generalización.

La Figura 9 presenta la relación entre los valores reales y los valores predichos por cada modelo, lo que permite visualizar de manera comparativa el ajuste alcanzado en distintos rangos de la variable objetivo. En ella se observa que, en términos generales, los modelos no lineales muestran una mayor cercanía a la diagonal de referencia, especialmente en el centro donde se concentra la mayor parte de las observaciones. Sin embargo, también se observa que en los extremos de la distribución se presentan mayores diferencias entre valores reales y predichos, donde valores bajos de retenido tienden a ser sobreestimados y valores altos tienden a ser subestimados, lo cual es consistente con la menor frecuencia de observaciones y con la mayor variabilidad del proceso en dichas regiones.

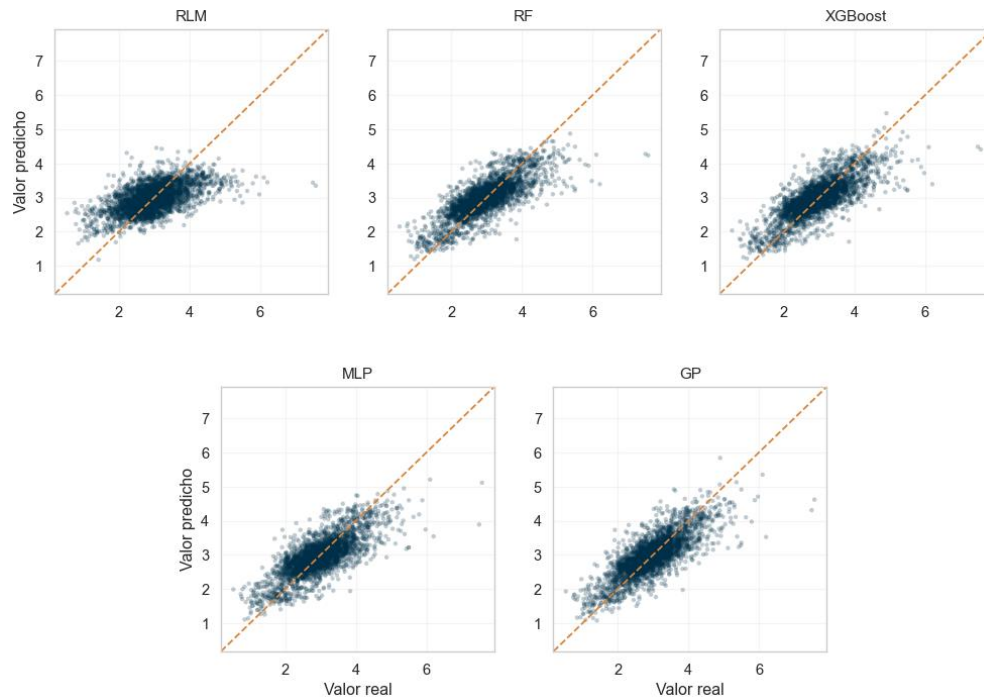


Figura 9. Valores reales vs. Valores predichos para cada modelo.

6.3. TIEMPOS DE EJECUCIÓN DE LOS MODELOS

Como criterio complementario para la selección del modelo final, se analizaron los tiempos de ejecución requeridos por cada enfoque durante el proceso de optimización y entrenamiento. En la Tabla 7, se presenta el tiempo destinado a la búsqueda de hiperparámetros, el reentrenamiento final con todo el conjunto de entrenamiento y el tiempo de predicción sobre el conjunto de prueba.

Tabla 7. Tiempos de ejecución de los modelos.

Modelo	Optimización / validación (s)	Reentrenamiento final (s)	Predicción en test (s)	Tiempo total (s)
Regresión lineal múltiple (RLM)	0.0519	0.0036	0.00075	0.056
Random Forest (RF)	167.50	45.00	1.19	213.69
XGBoost (XGB)	83.92	3.58	0.026	87.53
Perceptrón multicapa (MLP)	1950.74	26.11	0.67	1977.52
Proceso gaussiano (GP)	13113.46	1665.56	24.18	14779.02

En términos de costo computacional, la regresión lineal múltiple fue, con amplia diferencia, el modelo más eficiente, seguida por XGBoost, que presentó un tiempo total de modelado considerablemente menor al de Random Forest, mientras que el proceso gaussiano fue el modelo más costoso computacionalmente, tanto en la etapa de optimización como en el reentrenamiento y la predicción final.

6.4. ANÁLISIS ESPECÍFICO DE LOS MODELOS

Además de la comparación cuantitativa entre métricas de desempeño, se realizó un análisis específico de los modelos con el fin de examinar sus principales fortalezas, limitaciones y aportes al entendimiento del proceso.

6.4.1. REGRESIÓN LINEAL MÚLTIPLE: VERIFICACIÓN DE SUPUESTOS

La regresión lineal múltiple se utilizó como modelo base de referencia, sin embargo, su validez depende del cumplimiento de varios supuestos estadísticos, que se examinaron mediante los gráficos diagnósticos presentados en la Figura 10:

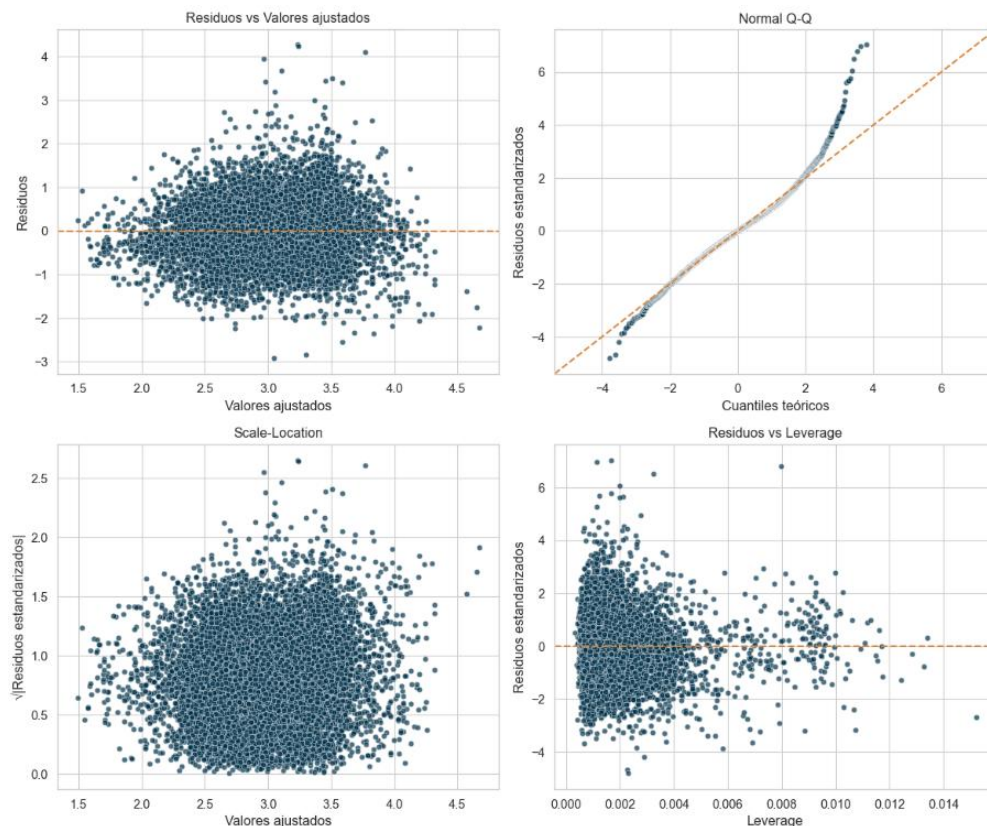


Figura 10. Verificación de supuestos de la regresión lineal múltiple.

Aunque los residuos se encuentran centrados alrededor de cero, la dispersión observada en función de los valores ajustados no es completamente homogénea, lo que sugiere que la

relación entre las variables predictoras y el retenido no puede describirse mediante una estructura lineal simple. Asimismo, el gráfico Q-Q muestra desviaciones notables respecto a la recta teórica en las colas, indicando incumplimiento del supuesto de normalidad de los residuos. Por su parte, el gráfico *Scale-Location* evidencia variaciones en la dispersión de los errores a lo largo del rango de predicción, lo que apunta a la presencia de heterocedasticidad moderada. Finalmente, aunque no se observa una concentración severa de observaciones con alto *leverage*, sí existen algunos casos aislados que podrían ejercer cierta influencia sobre el ajuste. En conjunto, estos resultados confirman que la regresión lineal múltiple presenta limitaciones para capturar adecuadamente la dinámica del proceso de molienda.

6.4.2. MODELOS BASADOS EN ÁRBOLES: INTERPRETACIÓN MEDIANTE SHAP

Dado que los modelos basados en árboles presentaron uno de los mejores desempeños predictivos, se profundizó en su interpretación mediante valores SHAP, para la identificación de variables influyentes y la dirección y forma de su efecto sobre la finura. Al ser XGBoost el modelo con mejor desempeño tanto en el conjunto de entrenamiento como en el de prueba, se realizó un énfasis en los resultados de este modelo específico, aunque posteriormente se verificó que estos resultados concordaban con los obtenidos en Random Forest. Primero, se graficó la influencia de las diferentes variables operativas, ordenándolas de mayor a menor influencia:

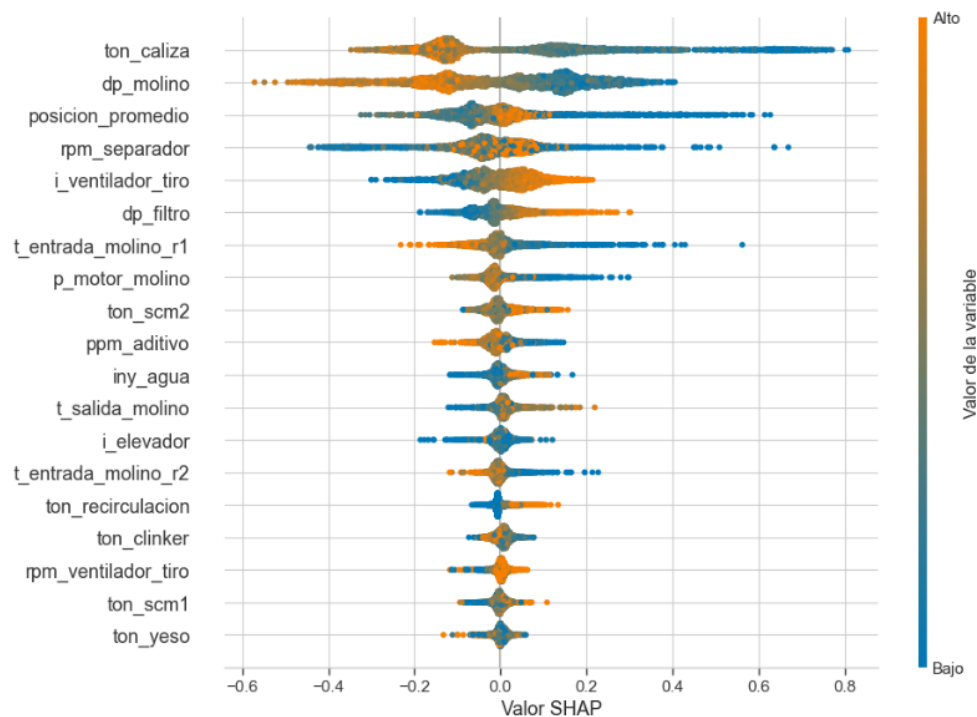


Figura 11. Resumen de interpretabilidad del modelo XGBoost mediante valores SHAP.

Este análisis permite identificar como variables más influyentes la composición de caliza y la presión diferencial del molino. En el caso de la caliza, tal como se evidenció en el análisis exploratorio, al ser un material relativamente más blando que el resto de componentes de la mezcla, tiene una relación inversa con el retenido. Por otro lado, valores altos de presión diferencial del molino tienden a relacionarse con retenidos más bajos, pues estos valores representan una mayor carga de material sobre la mesa de molienda. Tener una cama de material estable favorece la molienda pues la disminución de tamaño de partícula no ocurre únicamente por la presión ejercida por los rodillos, sino que también por la interacción entre partículas.

El gráfico 11 también permitió identificar a `posicion_promedio` y `rpm_separador` entre las variables más relevantes del modelo, sin embargo, en ambos casos no se observó una tendencia monotónica claramente definida, es decir, su efecto sobre la predicción del retenido no aumenta ni disminuye de forma uniforme a lo largo de todo su rango de valores. Por esta razón, se realizó un análisis más detallado mediante gráficos de dependencia SHAP, los cuales permiten examinar la contribución de cada variable a la predicción y, adicionalmente, observar su comportamiento conjunto con la variable de mayor interacción estimada por el modelo. En el caso de la variable `posicion_promedio`, la variable de mayor interacción elegida fue la presión diferencial del molino:

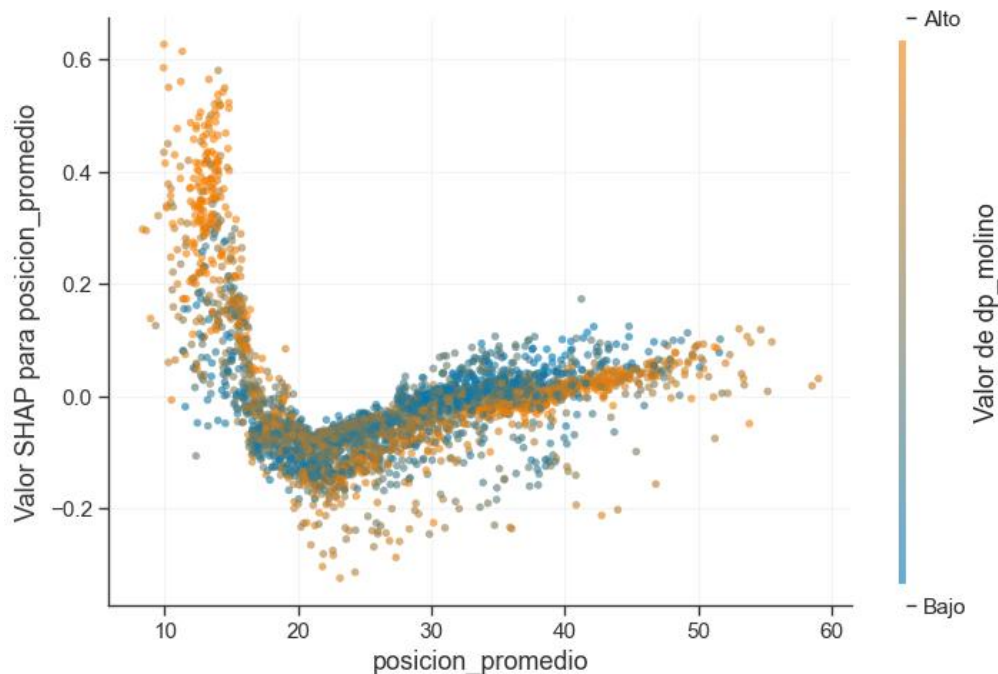


Figura 12. Efecto de la posición promedio de los rodillos sobre la predicción del retenido según valores SHAP.

Como se observa en la figura 12, valores bajos de `posicion_promedio` tienden a asociarse con contribuciones positivas a la predicción del retenido, mientras que en un rango intermedio su efecto se vuelve predominantemente negativo, para luego acercarse nuevamente a valores cercanos a cero en posiciones más altas, lo que es consistente con los resultados observados para esta variable en el análisis exploratorio. Además, esta gráfica muestra una dependencia más compleja de estos cambios con la variable `dp_molino`, donde se observa que para una misma `posicion_promedio`, condiciones de menor presión diferencial (menor carga interna de molino) pueden asociarse con mayores valores predichos de retenido pues en estos casos hay una cama de molienda más inestable y por lo tanto una molienda más ineficiente. En contraste, los valores intermedios de posición, aproximadamente entre 15 y 35 mm, parecen representar una condición operativa más favorable, en la cual la molienda tiende a generar un producto final más fino.

Este análisis de dependencia se realizó de manera análoga con la velocidad del separador dinámico, donde la variable con mayor interacción escogida fueron las toneladas de caliza:

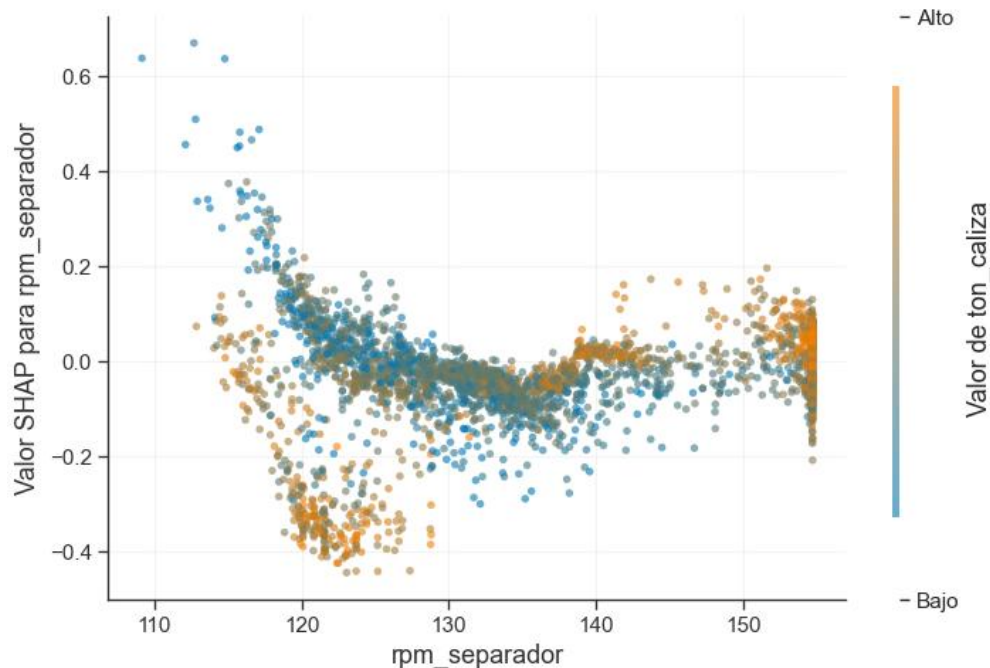


Figura 13. Efecto de la velocidad del separador sobre la predicción del retenido según valores SHAP.

En este caso, observamos que para valores bajos de `rpm_separador` se presentan valores de retenido más altos, pues hay una clasificación menos eficiente en la que la menor velocidad del separador facilita el paso de partículas más gruesas al producto final. Adicionalmente, se muestra la interacción con la variable `ton_caliza` donde, particularmente en los rangos

bajos e intermedios de `rpm_separador`, se observa que una mayor presencia de caliza parece atenuar el efecto desfavorable de trabajar a estas bajas velocidades debido a la mayor facilidad de molienda de este material. A medida que la velocidad del separador aumenta, el efecto de esta variable sobre la predicción tiende a acercarse a valores cercanos a cero y la dispersión se reduce, lo que sugiere una operación de clasificación más estable y menos sensible a la composición de la mezcla.

6.4.3. PERCEPTRÓN MULTICAPA: CONTROL DE SOBREAJUSTE

En el caso del perceptrón multicapa, además del desempeño predictivo final, resulta relevante analizar la evolución de la función de pérdida durante el entrenamiento para obtener información sobre la convergencia del modelo y evaluar la presencia de sobreajuste.

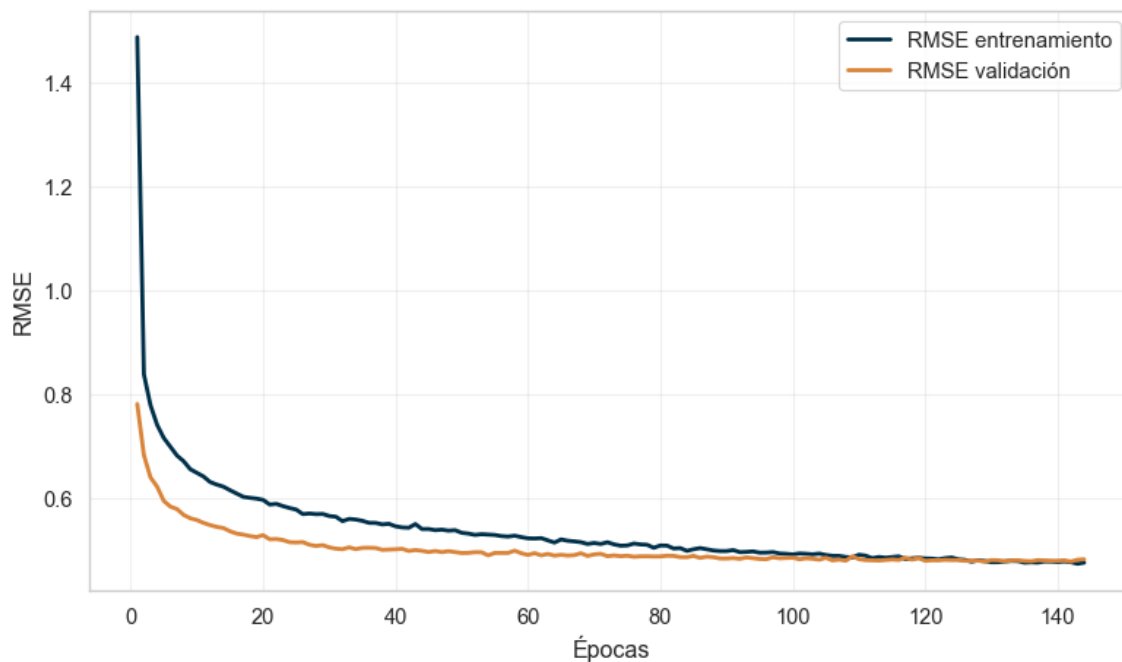


Figura 14. Curva de aprendizaje del perceptrón multicapa en términos de RMSE.

En la figura 14, la curva de aprendizaje muestra que tanto el RMSE de entrenamiento como el de validación disminuyen rápidamente durante las primeras épocas y posteriormente convergen de manera gradual hacia valores estables. La cercanía entre ambas curvas a lo largo del entrenamiento y la ausencia de una separación creciente entre ellas, sugiere que el modelo logró aprender patrones sin mostrar un sobreajuste significativo, evidenciando que el uso de regularización (dropout) y `early stopping` fue adecuado para controlar la complejidad del modelo.

6.4.4. PROCESO GAUSSIANO: PREDICCIÓN E INCERTIDUMBRE

A diferencia de los demás enfoques evaluados, el proceso gaussiano ofrece no solo una predicción puntual, sino también una estimación explícita de la incertidumbre asociada a cada observación, esto se representa en la siguiente figura donde se observa la media predictiva del modelo en el conjunto de prueba, junto con su intervalo de confianza del 95% construido a partir de la incertidumbre total estimada:



Figura 15. Predicción media del GP con intervalo de confianza del 95% en el conjunto de prueba.

En la figura 15, se observa que la predicción media dada por el modelo sigue razonablemente la tendencia general de los valores reales en el conjunto de prueba, especialmente en el rango central, donde se concentra la mayor parte de las observaciones. De igual forma, la mayoría de los valores reales permanece dentro del intervalo de confianza del 95%, lo que sugiere que el modelo no solo presenta buen desempeño en las predicciones, sino también una cuantificación coherente de la incertidumbre asociada. Sin embargo, en los extremos de la distribución, donde hay menos observaciones y el retenido presenta mayor variabilidad, se aprecian diferencias marcadas entre los valores reales y la media predictiva.

Este modelo también permitió estimar tanto la incertidumbre latente del modelo como la incertidumbre total asociada a cada observación:

Tabla 8. Resumen de la incertidumbre latente y total estimada en el conjunto de prueba.

Medida	Incertidumbre latente	Incertidumbre total
Media	0.4789	0.6353
Mediana	0.4583	0.6149
Mínimo	0.2209	0.4657
Máximo	0.8882	0.9783
Desv. estándar	0.1239	0.0960

En todos los estadísticos descriptivos, la incertidumbre total resultó mayor que la latente, lo cual era esperable, ya que además de la variabilidad propia de la función aprendida por el modelo, también incorpora la incertidumbre conocida del ensayo de laboratorio ($u = 0.412\%$), incluida explícitamente en el modelado mediante un *likelihood* de ruido fijo. En promedio, la incertidumbre latente fue de 0.4789 mientras que la incertidumbre total alcanzó 0.6353, lo que confirma que una parte significativa de la variabilidad observada en las predicciones no depende únicamente del modelo, sino también del propio proceso de medición.

De manera complementaria, los length-scales estimados en el proceso gaussiano mediante el esquema ARD permitieron cuantificar la sensibilidad relativa del modelo frente a las variables de entrada, pues valores más pequeños de length-scale se asocian con una mayor variación de la función predictiva ante cambios en esa variable, lo que puede interpretarse como un indicio de mayor relevancia relativa. Los resultados obtenidos para el mejor modelo se muestran a continuación:

Tabla 9. Length-scales para las variables en el proceso gaussiano con ARD.

Variable	Length-scale
rpm_separador	1.1431
posicion_promedio	1.7651
dp_molino	1.9492
t_entrada_molino_r2	2.2500
ton_yeso	2.5742
ton_scm2	2.6229
ton_caliza	2.6976

Variable	Length-scale
ppm_aditivo	2.9562
dp_filtro	3.2221
i_ventilador_tiro	3.4267
rpm_ventilador_tiro	3.5553
ton_scm1	3.9687
t_salida_molino	4.0391
ton_clinker	4.0487
ton_recirculacion	4.4601
iny_agua	4.4754
i_elevador	4.6194
t_entrada_molino_r1	4.6348
p_motor_molino	4.6582

Si bien esta medida no es directamente comparable con los valores SHAP obtenidos para XGBoost, se observó una coincidencia clara en la relevancia de variables como rpm_separador, posicion_promedio, dp_molino y ton_caliza. El hecho de que ambos modelos resalten las mismas variables sugiere una buena consistencia en los resultados y valida la importancia de estos parámetros en el proceso.

6.5. SELECCIÓN DEL MODELO FINAL

Respecto a las métricas de desempeño de los diferentes modelos, los valores de RMSE alcanzados por Random Forest, XGBoost y el proceso gaussiano en el conjunto de prueba se encuentran en el mismo orden de magnitud que la incertidumbre del ensayo de laboratorio, lo que sugiere que, dentro de las limitaciones impuestas por la variabilidad inherente de la medición y la complejidad del proceso industrial modelado, estos enfoques alcanzan un desempeño predictivo razonable para su aplicación como estimación operativa del retenido. Por su parte, aunque el proceso gaussiano obtuvo el menor error, la ventaja frente a los modelos basados en árboles es relativamente pequeña, por lo que la selección final del modelo no solo considera las mejores métricas de desempeño, sino también aspectos como la interpretabilidad y el costo computacional. En ese sentido, aunque el proceso gaussiano obtuvo el mejor desempeño global, XGBoost ofreció una relación más favorable entre precisión, costo computacional e interpretabilidad, al requerir un tiempo total de modelado muy inferior al del proceso gaussiano y permitir, adicionalmente, un análisis detallado del

efecto de las variables operativas mediante valores SHAP, lo cual aporta información útil para que los operadores comprendan qué condiciones del molino influyen en los resultados obtenidos. Por esta razón, se seleccionó XGBoost como modelo final del proyecto.

7. CONCLUSIONES Y TRABAJOS FUTUROS

7.1. CONCLUSIONES

El desarrollo de este proyecto permitió demostrar que la finura del cemento producido en un molino vertical de rodillos puede estimarse a partir de variables operacionales mediante técnicas de *machine learning*, con un nivel de ajuste útil para fines de seguimiento operativo.

El análisis exploratorio evidenció que la variable objetivo no puede representarse únicamente con tendencias lineales y no se ve afectada con una única variable dominante, sino que sus cambios obedecen a una serie de interacciones complejas entre las condiciones de operación y las características de la alimentación al molino. En particular, la composición de la mezcla, la acción del separador dinámico, la carga interna del molino y las características de la cama de molienda surgieron como variables clave para entender los cambios en el retenido.

En la construcción, optimización y comparación de modelos, se observó que los enfoques no lineales representaron de manera más acertada la realidad física del proceso, obteniendo en general mejor desempeño que el análisis mediante regresión lineal múltiple. El desempeño general de los modelos se evaluó teniendo en cuenta la variabilidad inherente de los datos de medición utilizados como entrada por lo que la comparación entre modelos no solo se concentró en identificar el menor error del modelo, sino en ponderar qué alternativa resultaba más conveniente dentro de un escenario de uso real.

En ese sentido, aunque el proceso gaussiano logró obtener las mejores métricas de desempeño y añadió la ventaja de la cuantificación de la incertidumbre para sus predicciones, al comparar de manera integral se seleccionó como mejor modelo el XGBoost pues tuvo un desempeño similar en términos de error en el modelado y se caracterizó por su practicidad al tener un costo computacional considerablemente menor y una mejor interpretabilidad, pues permitió evaluar el sentido y la magnitud de la influencia de cada variable en las predicciones.

En base a los resultados obtenidos, se concluye que el uso de modelos de aprendizaje automático en este contexto no solo es técnicamente viable, sino que puede convertirse en

una herramienta valiosa para mejorar la estabilidad del proceso, anticipar desviaciones de calidad y orientar ajustes operativos con mayor oportunidad.

7.2. TRABAJOS FUTUROS

Como trabajo futuro, se plantea avanzar en la implementación y despliegue del modelo seleccionado en el entorno operativo en planta, de manera que pueda integrarse con las herramientas de monitoreo actualmente utilizadas por los operadores y brindar estimaciones del retenido de manera automatizada. Esta etapa contempla el desarrollo de la solución informática para la extracción de las variables operativas, la generación de predicciones y su visualización, así como la definición de un esquema de mantenimiento periódico y reentrenamiento del modelo.

Por otro lado, debe considerarse que el presente estudio se concentró en la producción de cemento tipo UG, por ello, una línea de trabajo futura consiste en extender el proyecto a otros tipos de cemento fabricados con menor frecuencia, en los que generalmente se tienen exigencias de calidad específicas, se manejan condiciones de operación diferentes y se tiene una disponibilidad de datos significativamente menor a la del presente estudio. Este análisis permitiría evaluar hasta qué punto la metodología desarrollada puede generalizarse o si resulta más conveniente la construcción de modelos específicos para cada tipo de producto.

De manera complementaria, también sería pertinente explorar su aplicación en procesos de molienda de crudo, donde la configuración del molino, las características del material alimentado y los criterios de calidad difieren de los de la molienda de cemento. Sin embargo, considerando que este tipo de proceso cuenta con variables operativas generalmente similares a las utilizadas en la molienda de cemento, también podría ser útil contar con herramientas predictivas basadas en dichas variables medidas en línea.

Finalmente, se considera de interés ampliar el estudio a otros métodos de ensayo para la evaluación de finura, como la granulometría láser, que no fue considerada en el proyecto actual por falta de datos históricos. La incorporación de este tipo de medición permitiría una caracterización más completa de la finura, permitiendo un control más riguroso en términos de la distribución granulométrica completa del producto final.

8. REFERENCIAS BIBLIOGRÁFICAS

[1] E. Adiguzel, "Global cement industry outlook: Trends and forecasts," *World Cement Assoc. Blog*, Mar. 19, 2024. [Online]. Available: <https://www.worldcementassociation.org/blog/news/global-cement-industry-outlook-trends-and-forecasts>. [Accessed: May 22, 2025].

- [2] P. Pareek and V. S. Sankhla, "Review on vertical roller mill in cement industry & its performance parameters," *Mater. Today: Proc.*, vol. 44, no. 6, pp. 4621–4627, 2021, doi: 10.1016/j.matpr.2020.10.916.
- [3] P. C. Hewlett and M. Liska, *Lea's Chemistry of Cement and Concrete*, 5th ed. Oxford, UK: Butterworth-Heinemann, 2019, pp. 131–134.
- [4] J. F. Lamond and J. H. Pielert, Eds., *Significance of Tests and Properties of Concrete and Concrete-Making Materials*. West Conshohocken, PA: ASTM Int., 2006, pp. 438–439.
- [5] F. Belmajdoub and S. Abderafi, "Efficient machine learning model to predict fineness in a vertical raw meal of Morocco cement plant," *Results Eng.*, vol. 17, p. 100833, Mar. 2023, doi: 10.1016/j.rineng.2022.100833.
- [6] A. K. Pani and H. K. Mohanta, "Online monitoring and control of particle size in the grinding process using least square support vector regression and resilient back propagation neural network," *ISA Trans.*, vol. 56, pp. 206–221, May 2015, doi: 10.1016/j.isatra.2014.11.011.
- [7] *Vertical Roller Mill for Cement Grinding* [Online Image]. Cement-plants.com, 2024. [Online]. Available: <https://cement-plants.com/raw-material-production/cement-mill/vertical-mill/>. [Accessed: May 25, 2025].
- [8] S. O. Ehikhuenmen, F. I. Irabor, and E. E. Ikponmwosa, "The influence of cement fineness on the structural characteristics of normal concrete," *IOP Conf. Ser.: Mater. Sci. Eng.*, vol. 640, no. 1, p. 012043, 2019, doi: 10.1088/1757-899X/640/1/012043.
- [9] D. P. Bentz et al., "Effects of cement particle size distribution on performance properties of Portland cement-based materials," *Cem. Concr. Res.*, vol. 29, no. 10, pp. 1663–1671, Oct. 1999, doi: 10.1016/S0008-8846(99)00163-5.
- [10] Instituto Colombiano de Normas Técnicas y Certificación, *NTC 33: Cementos. Método de ensayo para determinar la finura del cemento hidráulico por medio del aparato Blaine de permeabilidad al aire*, 4.ª actualización. Bogotá, Colombia: ICONTEC, 2019.
- [11] Instituto Colombiano de Normas Técnicas y Certificación, *NTC 294: Cementos. Método de ensayo para determinar la finura del cemento hidráulico utilizando el tamiz de 45 μm (No. 325)*, 5.ª actualización. Bogotá, Colombia: ICONTEC, 2018.
- [12] C. F. Ferraris, M. Peltz, and B. Toman, *Measurement of Particle Size Distribution of Portland Cement and Related Materials: Accuracy and Precision* (NIST Special Publication 260-190). Gaithersburg, MD: NIST, Aug. 2018, doi: 10.6028/NIST.SP.260-190.

- [13] JCGM, *Evaluación de datos de medición—Guía para la expresión de la incertidumbre de medida*, JCGM 100:2008, GUM 1995, ed. digital 1 en español, Centro Español de Metrología, Sep. 2008, pp. 4–6.
- [14] K. B. Prakash, Ed., *Machine Learning for Industrial Applications*. Hoboken, NJ: Wiley, 2024, pp. 2–4.
- [15] J. Lu, *Machine Learning Foundations: Supervised, Unsupervised, and Advanced Learning*, 1st ed. Cham, Switzerland: Springer, 2021, pp. 11–12, doi: 10.1007/978-3-030-65900-4.
- [16] N. L. Rane, Ö. Kaya, and J. Rane, *Artificial Intelligence, Machine Learning, and Deep Learning for Sustainable Industry 5.0*. London, UK: Deep Sci. Publ., 2024, pp. 19–21.
- [17] G. James, D. Witten, T. Hastie, and R. Tibshirani, *An Introduction to Statistical Learning: with Applications in R*, 1st ed. New York, NY, USA: Springer, 2013, pp. 59–75, 303–320.
- [18] T. Gulli, A. Kapoor, and S. Pal, *Hyperparameter Tuning with Python*, 1st ed. Birmingham, UK: Packt Publishing, 2022, pp. 222–224, 229–231.
- [19] C. E. Rasmussen and C. K. I. Williams, *Gaussian Processes for Machine Learning*. Cambridge, MA, USA: MIT Press, 2006, pp. 13–17.
- [20] scikit-learn developers, "1.7. Gaussian Processes," *scikit-learn 1.8.0 documentation*. [Online]. Available: https://scikit-learn.org/stable/modules/gaussian_process.html. [Accessed: Apr. 02, 2026].
- [21] COMSOL, "More on Hyperparameters for Gaussian Processes," *Learning Center, Surrogate Modeling Theory*. [Online]. Available: <https://www.comsol.com/support/learning-center/course/surrogate-modeling-theory-271/more-on-hyperparameters-for-gaussian-processes-96781>. [Accessed: Apr. 02, 2026].
- [22] A. V. Tatachar, "Comparative Assessment of Regression Models Based on Model Evaluation Metrics," *Int. Res. J. Eng. Technol. (IRJET)*, vol. 8, no. 9, pp. 853–860, Sep. 2021.
- [23] R. Lange, T. Lange, and T. L. van Zyl, "Predicting particle fineness in a cement mill," in *Proc. 2020 IEEE 23rd Int. Conf. Inf. Fusion (FUSION)*, Rustenburg, South Africa, 2020, pp. 1–8, doi: 10.23919/FUSION45008.2020.9190236.
- [24] D. Stanišić, L. Mejić, B. Jorgovanović, V. Ilić, and N. Jorgovanović, "An algorithm for soft sensor development for a class of processes with distinct operating conditions," *Sensors*, vol. 24, no. 6, Art. no. 1948, Mar. 2024, doi: 10.3390/s24061948.

- [25] M. T. Topaloğlu, C. K. Macit, U. U. Uçar, and B. Tanyeri, "Prediction of Blaine fineness of final product in cement production using industrial quality control data based on chemical and granulometric inputs using machine learning," *Appl. Sci.*, vol. 16, no. 4, Art. no. 2046, Feb. 2026, doi: 10.3390/app16042046.
- [26] R. Fatahi, H. Abdollahi, M. Noaparast, and M. Hadizadeh, "Modeling the working pressure of a cement vertical roller mill using SHAP-XGBoost: A 'conscious lab of grinding principle' approach," *Powder Technol.*, vol. 457, Art. no. 120923, 2025, doi: 10.1016/j.powtec.2025.120923.
- [27] M. Huber, T. Müller, and R. H. Schmitt, "Differentiation between model error and measurement uncertainty in machine learning modeling for measurement processes," *Meas. Sens.*, vol. 38, Art. no. 101608, 2025, doi: 10.1016/j.measen.2024.101608.

ANEXO A: CÁLCULO DE LA INCERTIDUMBRE DEL ENSAYO DE RETENIDO EN MALLA 325

La finura del cemento mediante tamiz 325 se reporta como el porcentaje de material retenido respecto a la masa inicial de muestra. Para el ensayo realizado, la masa de muestra utilizada es aproximadamente 1 g, por lo que el modelo de medición se expresa como:

$$F = \frac{m_r}{m_0} \times 100$$

donde F es la finura reportada como porcentaje de retenido, m_r es la masa retenida en la malla y m_0 es la masa inicial de la muestra. De acuerdo con la Guía para la Expresión de la Incertidumbre de Medición (GUM), la contribución tipo A se estimó a partir de resultados experimentales obtenidos bajo condiciones de repetibilidad. Para ello, se realizaron ensayos sobre una misma muestra con 4 analistas, en los que cada uno realizó 5 determinaciones, para un total de 20 resultados. Por razones de confidencialidad no se presentan los datos individuales, sin embargo, a partir de estos datos se obtuvo una desviación estándar experimental de 0.411%. En consecuencia, para un resultado individual del ensayo, la incertidumbre estándar tipo A se toma como:

$$u_A = s = 0.411 \%$$

Además de la variabilidad experimental, se consideró la contribución de la balanza analítica empleada en el ensayo. En el certificado de calibración, se reporta una incertidumbre expandida de 0.00032 g con un factor de cobertura de 2.23. Por tanto, la incertidumbre estándar asociada a la balanza es:

$$u_b = \frac{U_b}{k} = \frac{0.00032}{2.23} = 0.000144 \text{ g}$$

Como la finura se expresa en porcentaje, esta incertidumbre debe propagarse desde gramos hasta la magnitud de salida utilizando la ley de propagación de incertidumbres aplicada al modelo de medición:

$$u_{\text{bal}}(F) = \sqrt{\left(\frac{\partial F}{\partial m_r} u(m_r)\right)^2 + \left(\frac{\partial F}{\partial m_0} u(m_0)\right)^2}$$

con:

$$\frac{\partial F}{\partial m_r} = \frac{100}{m_0}$$

$$\frac{\partial F}{\partial m_0} = -\frac{100 m_r}{m_0^2}$$

Asumiendo que ambas pesadas se realizan con la misma balanza, se tiene que $u(m_r) = u(m_0) = u_b$, y, por lo tanto:

$$u_{\text{bal}}(F) = \sqrt{\left(\frac{100}{m_0} u_b\right)^2 + \left(-\frac{100 m_r}{m_0^2} u_b\right)^2}$$

Dado que en este ensayo $m_0 \approx 1 \text{ g}$ y que el porcentaje retenido es bajo, la segunda contribución resulta despreciable frente a la primera, por lo que puede aproximarse como:

$$u_{\text{bal}}(F) \approx 100 u_b = 100 (0.000144) = 0.0144 \%$$

La incertidumbre estándar combinada del ensayo se obtiene mediante la suma cuadrática de la contribución tipo A y la contribución de la balanza:

$$u_c = \sqrt{u_A^2 + u_{\text{bal}}^2}$$

$$u_c = \sqrt{(0.411)^2 + (0.0144)^2} = 0.412 \%$$

En consecuencia, la incertidumbre del ensayo de finura está dominada por la variabilidad experimental estimada mediante el análisis tipo A, pues la contribución de la balanza resulta comparativamente pequeña.