

**TOOLBOX SESIÓN 20****CONGRESO DE
MANTENIMIENTO
& CONFIABILIDAD
M É X I C O****19^a
EDICIÓN**

Large Language Models para Mantenimiento Industrial

Enrique López Droguett

Profesor Titular, Civil and Environmental Engineering
Department, UCLA

1

Contenido



1. Fundamentos del Large Language Models (LLMs).
2. Predictivo en base a Análisis de Aceite con:
 - a. IA en la Nube: LLMs Comerciales,
 - b. IA Local: Gemma 4.

2

Fundamentos del Large Language Models

3

LLMs y Mantenimiento Industrial



4

Qué es un Large Language Model

- Modelo entrenado para **procesar y generar lenguaje** a partir de grandes volúmenes de texto.



Tokens

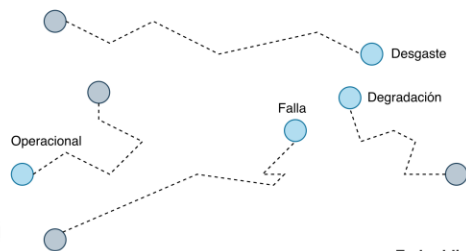
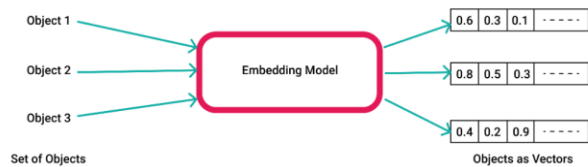
- LLMs no procesan directamente palabras completas.
- Procesan unidades llamadas **tokens**.
- Una frase técnica se transforma en **una secuencia de tokens** antes de entrar al modelo:

“Continuar monitoreo tribológico según plan de mantenimiento”

Large Language Models (LLMs), such as GPT-3 and GPT-4, utilize a process called tokenization. Tokenization involves breaking down text into smaller units, known as tokens, which the model can process and understand. These tokens can range from individual characters to entire words or even larger chunks, depending on the model. For GPT-3 and GPT-4, a Byte Pair Encoding (BPE) tokenizer is used. BPE is a subword tokenization technique that allows the model to dynamically build a vocabulary during training, efficiently representing common words and word fragments. Although the core tokenization process remains similar across different versions of these models, the specific implementation can vary based on the model's architecture and training objectives.

Embeddings

- Cada token se representa como un **vector numérico**.
- Estos vectores permiten que el modelo capture **relaciones** entre conceptos técnicos.
- Conceptos relacionados tienden a quedar cerca entre sí:
 - Desgaste,
 - Degradación,
 - Falla.



Embedding Space

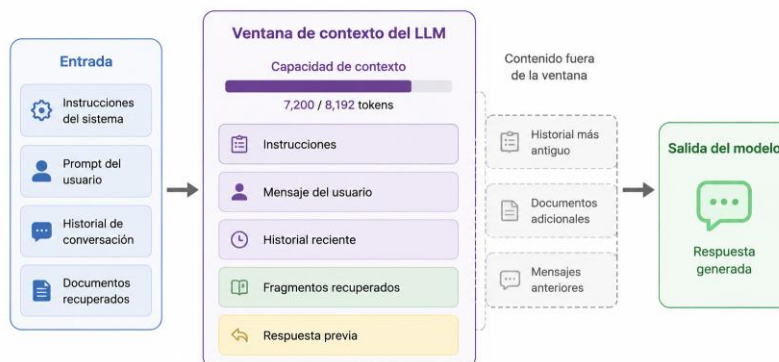
UCLA

<https://www.linkedin.com/pulse/how-do-embeddings-work-large-language-model-llm-onfinanceofficial-54awc/>

7

Ventana de Contexto (Context Window)

- Es la cantidad de información que el modelo puede considerar en una interacción.



El LLM solo puede usar directamente la información que cabe dentro de su ventana de contexto.

UCLA

8

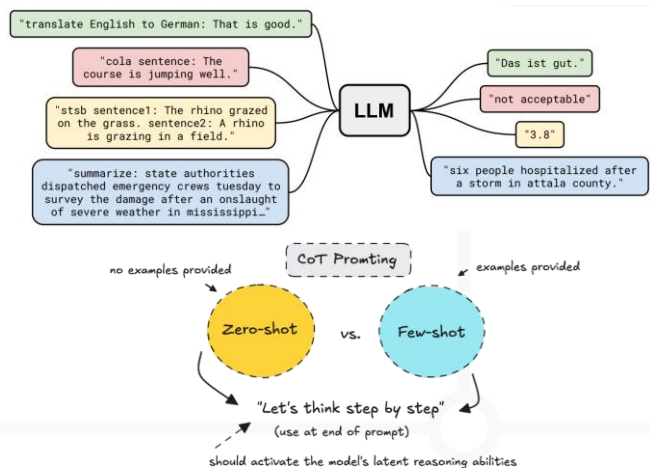
Pre-Training

- Antes de ser usados en aplicaciones específicas, los LLMs pasan por una etapa de **pre-training**:
 - Aprenden patrones generales de lenguaje, estructura y conocimiento a partir de grandes corpus.



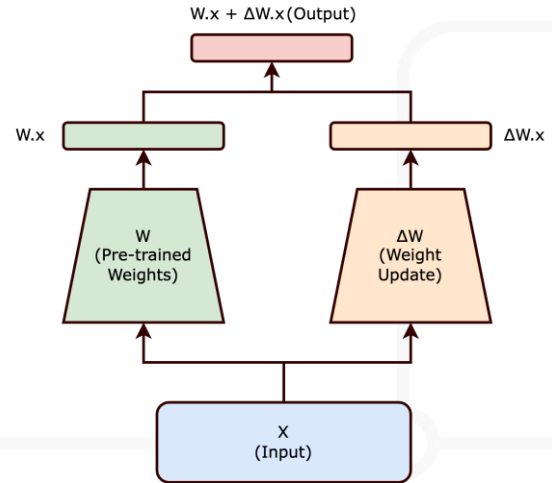
Prompting

- Consiste en guiar al modelo mediante instrucciones, contexto y ejemplos.
- No modifica los pesos del modelo.
- Estrategias comunes:
 - **Zero-shot**: instrucción directa sin ejemplos,
 - **Few-shot**: instrucción con ejemplos resueltos,
 - **Razonamiento estructurado**: respuesta organizada en etapas.



Fine-Tuning

- Adapta un modelo base a una tarea o dominio específico usando datos del problema.
- Con LLMs, una alternativa eficiente es entrenar **adaptadores** en vez de modificar todo el modelo.
- **LoRA** permite adaptar un modelo con menor costo computacional.



Hu EJ, Shen Y, Wallis P, Allen-Zhu Z, Li Y, Wang S, et al. LoRA: Low-Rank Adaptation of Large Language Models 2021. <https://doi.org/10.48550/arXiv.2106.09685>.

<https://towardsdatascience.com/understanding-lora-low-rank-adaptation-for-finetuning-large-models-936bce1a07c6/>.

11

Prompting x Fine-Tuning

Aspecto	Prompting	Fine-Tuning
Modifica Pesos	No	Sí
Esfuerzo Inicial	Bajo	Mayor
Personalización	Limitada	Alta
Control Local	Bajo o Ninguno (Depende del Proveedor)	Alto (con Open Source)



FACULTAD DE
CIENCIAS FÍSICAS
Y MATEMÁTICAS
UNIVERSIDAD DE CHILE

12

Modelos Comerciales y Open Source

Existen dos grandes formas de usar LLMs en ingeniería:

1. Modelos Comerciales:

- Ejecución en la nube,
- Alta capacidad general,
- Uso inmediato,
- Acceso mediante web o API.



Modelos Comerciales y Open Source

2. Modelos Open Source:

- Ejecución local,
- Mayor control del modelo,
- Adaptación al dominio,
- Mejor compatibilidad con datos sensibles.

Open Source LLMs

BigScience



Hugging Face



Grok



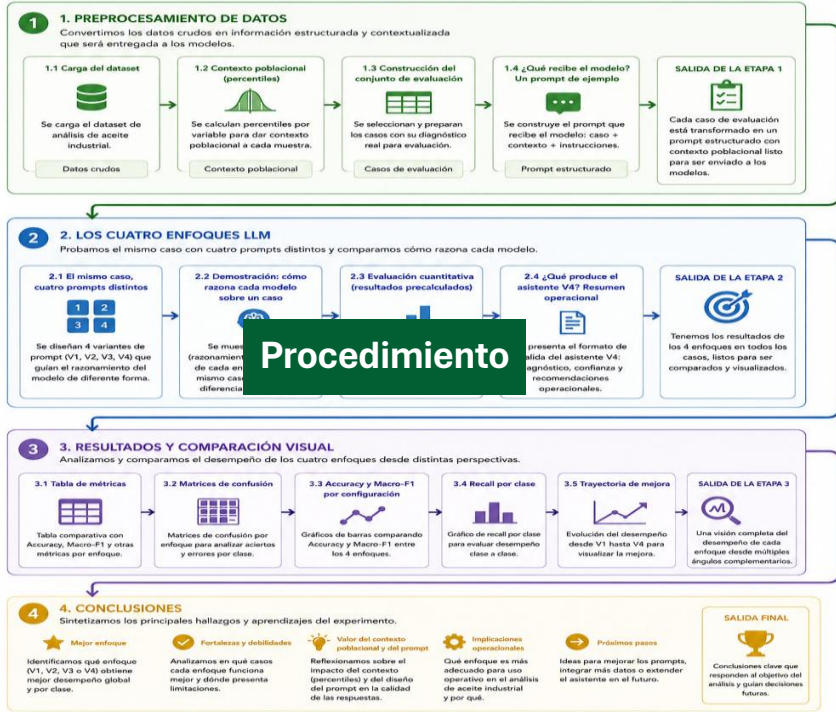
Meta

Relevancia para Mantenimiento Industrial

Los LLMs conectan datos, estructuran salidas y apoyan decisiones de mantenimiento.



IA en la Nube: LLMs Comerciales



17

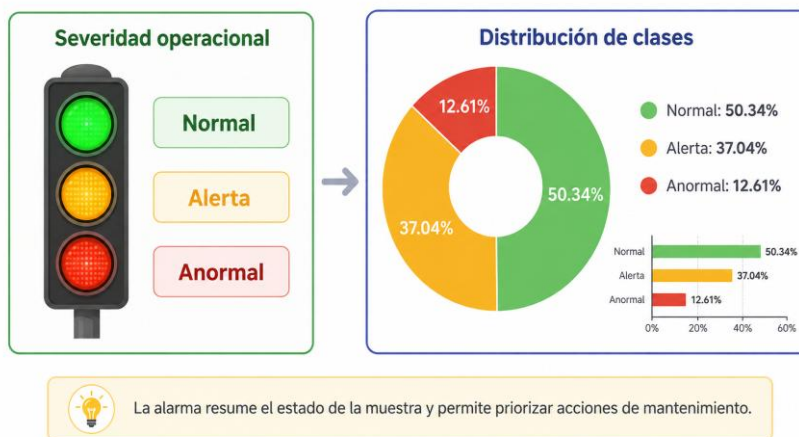
Dataset Industrial de Análisis de Aceite



20

Estimar el Estado de Alarma

Objetivo: clasificar la severidad operacional.



UCLA

21

No Todos los Errores Cuestan lo Mismo

Dirección del error en la clasificación de severidad

		Alarma predicha			Leyenda Clasificación correcta Error conservador Error riesgoso Subestimación crítica
		Normal	Alerta	Anormal	
Alarma real	Normal	Correcto	Conservador	Error	
	Alerta	Riesgoso	Correcto	Error	
	Anormal	Error crítico	Error	Correcto	

Una subestimación puede ser más grave que una falsa alarma.

UCLA

22

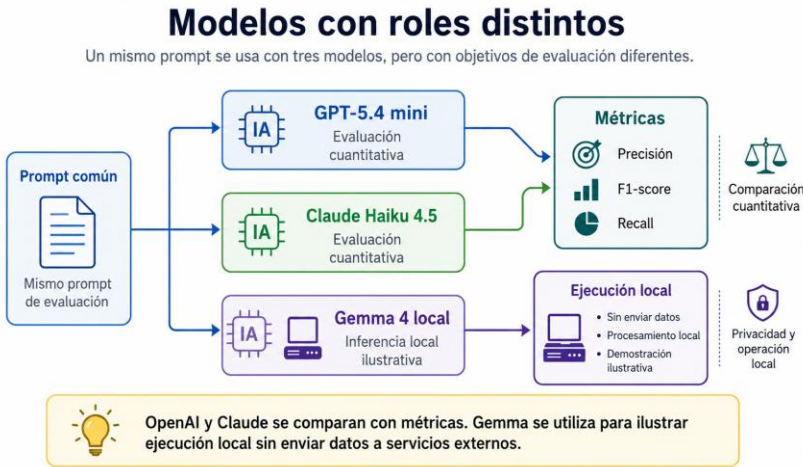
Tres Modos de Uso del LLM



UCLA

23

Modelos Usados



UCLA

24

Propiedad Intelectual y Confidencialidad

Gobernanza en el uso de modelos

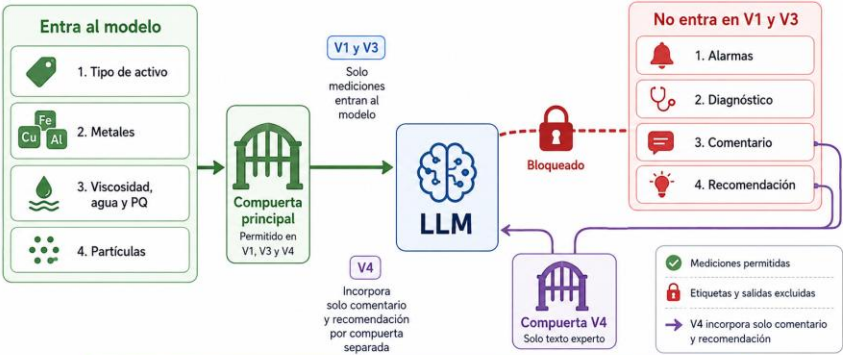
La elección entre API comercial y modelo local no es solo técnica, también implica confidencialidad, control y políticas organizacionales.



Separar Mediciones, Alarmas y Texto Experto

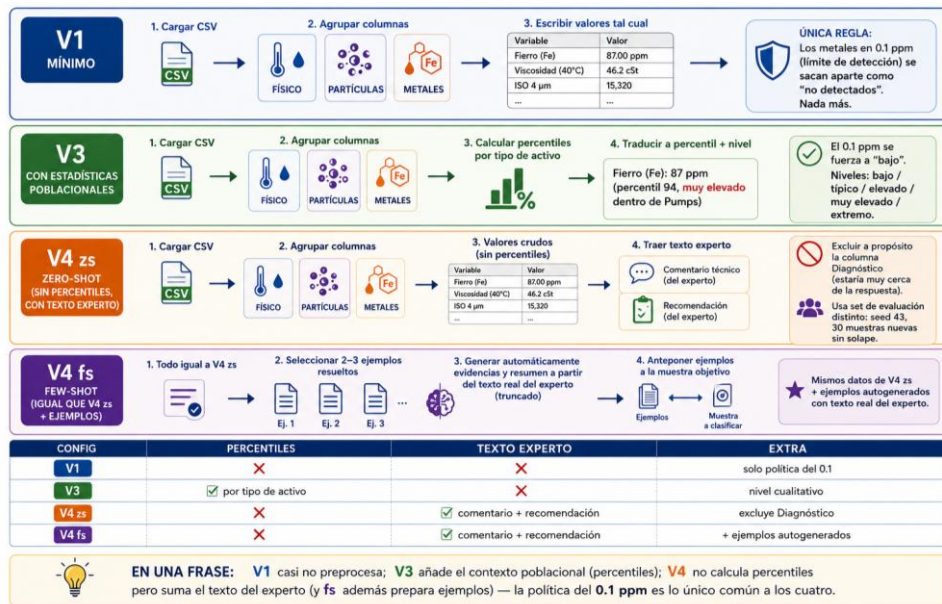
Control de información para evitar leakage

Cada modo controla qué información entra al modelo



El diseño de entrada controla el leakage. V1 y V3 usan solo mediciones, mientras V4 abre una compuerta separada para comentario y recomendación.

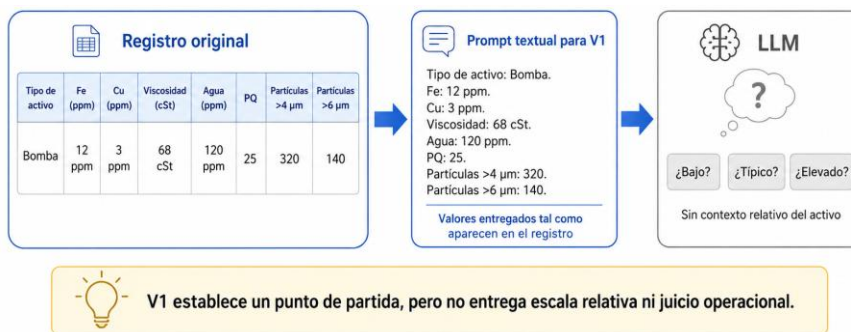
Preprocesamiento de la Data Numérica y Texto



27

V1: Valores Crudos

- V1 entrega al modelo los valores tal como aparecen en el registro:
 - Nombres de variables,
 - Valores numéricos,
 - Unidades básicas,
 - Sin escala relativa.



28

V3: Criterio Operacional

V3 agrega contexto operacional al dato

El modelo interpreta señales en conjunto, no como valores aislados.

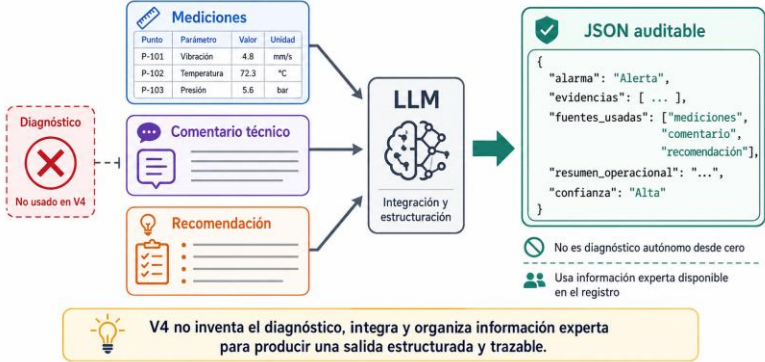


29

V4 Zero-Shot: Mediciones + Texto Experto

V4 integra información experta disponible

Integra mediciones, comentario técnico y recomendación, sin usar Diagnóstico.



30

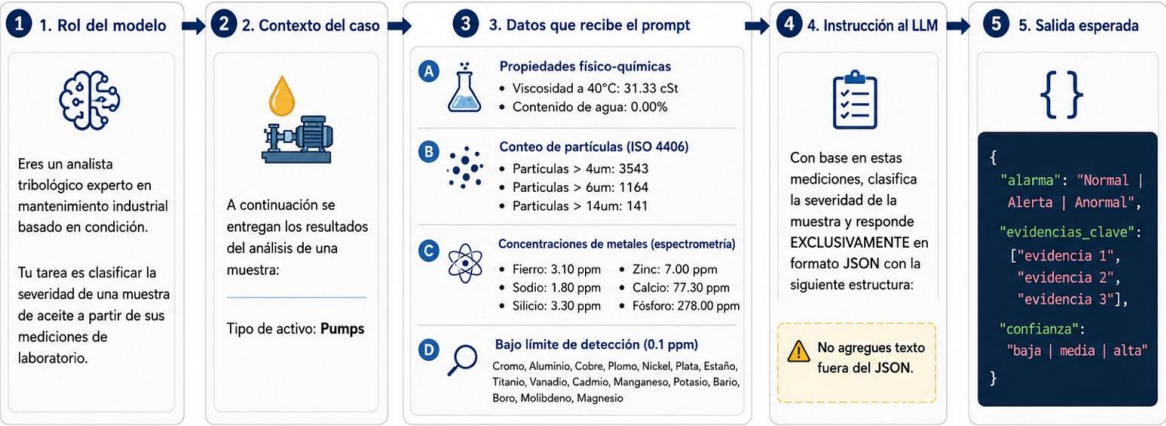
V4 Few-Shot: Ejemplos Resueltos (Observaciones)

V4 Few-shot agrega ejemplos resueltos

El modelo aprende la estructura de severidad a partir de ejemplos resueltos y luego procesa una muestra nueva con el mismo formato JSON.



Prompt de Entrada para LLMs

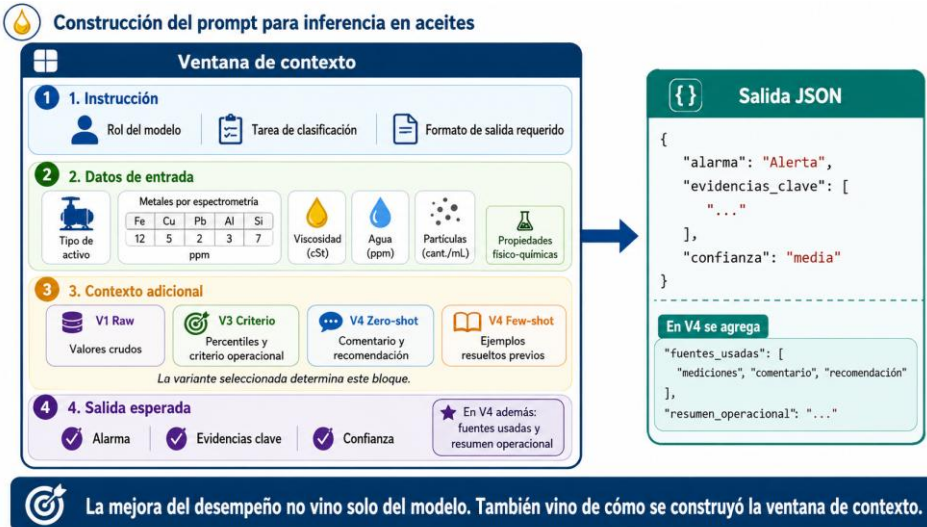


En una frase

El prompt define quién es el modelo, qué datos recibe, qué debe decidir y en qué formato exacto debe responder.



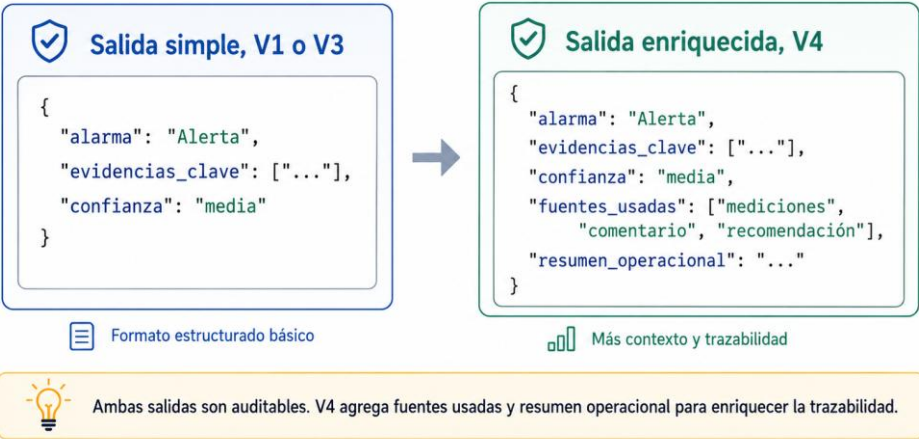
Context Window



34

Salidas Estructuradas y Auditables

Todas las configuraciones responden en formato JSON. V4 agrega mayor trazabilidad y contexto operacional.



35

Demo en Vivo

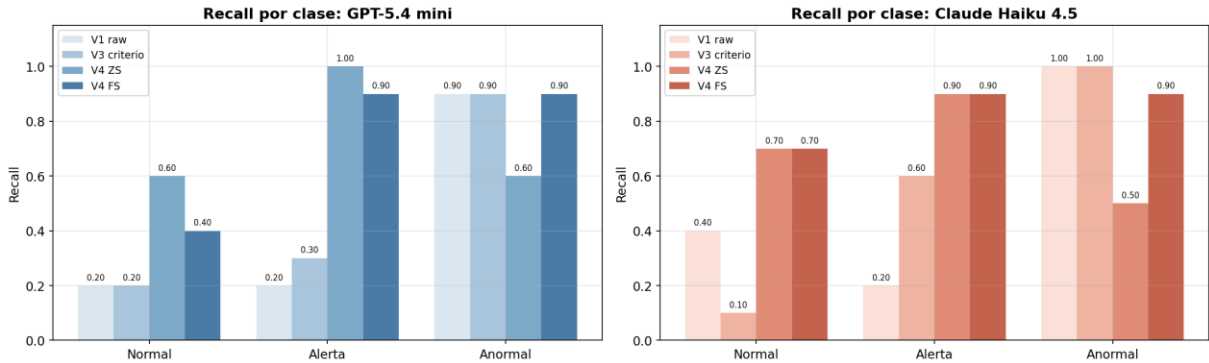
- 1. Desarrollar las soluciones predictivas
- 2. Interpretar los resultados
- 3. Esto se hace cargando el notebook:
“02_oil_analysis_llm_demo.ipynb”
- 1. El notebook y sus librerías se encuentran disponibles en:
“CMC_MX_2026_Toolbox_Code_Nube.zip”



36

Resultados: Recall por Clase (Alarma)

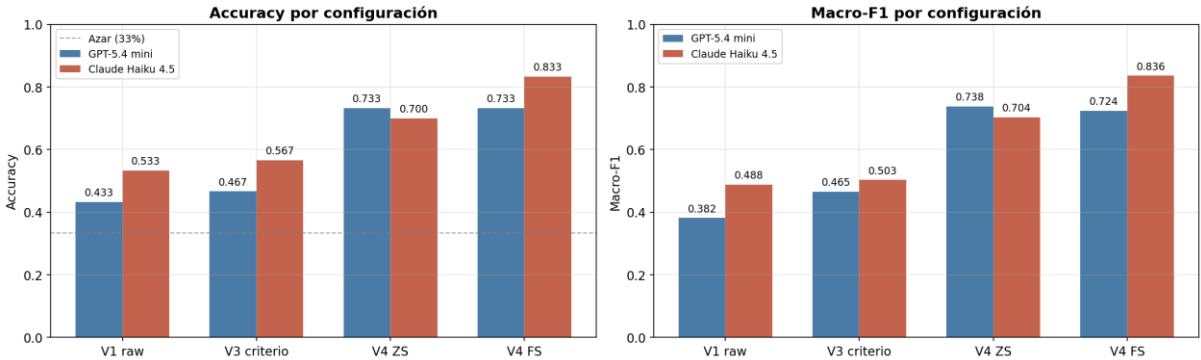
Recall por clase: evolución a través de las cuatro configuraciones



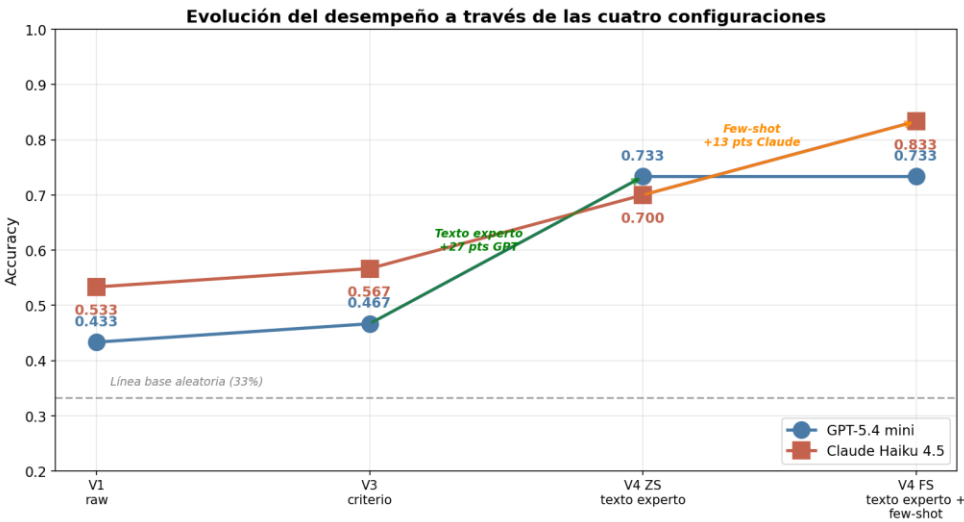
37

Resultados: Métricas de Performance

Comparación de las cuatro configuraciones sobre las mismas 30 muestras independientes



Resultados: Evolución de las Distintas LLMs Predictivas



Mejor LLM Predictiva Desarrollada



Riesgo Operacional: Claude Haiku 4.5



Qué Aprendimos de los Cuatro Modos de LLMs



Aprendizajes por configuración

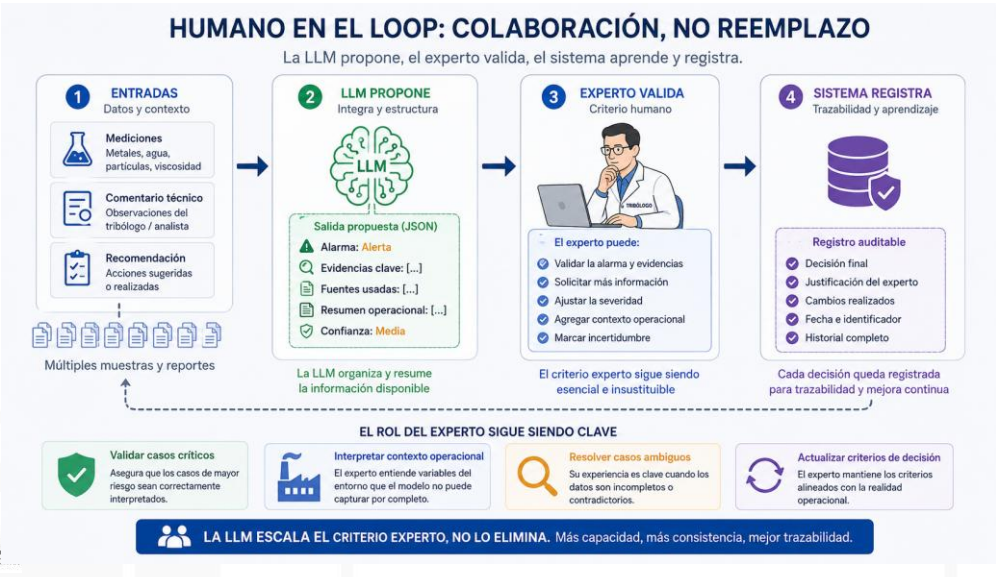
Cada configuración aporta una mejora metodológica distinta en el flujo de información.



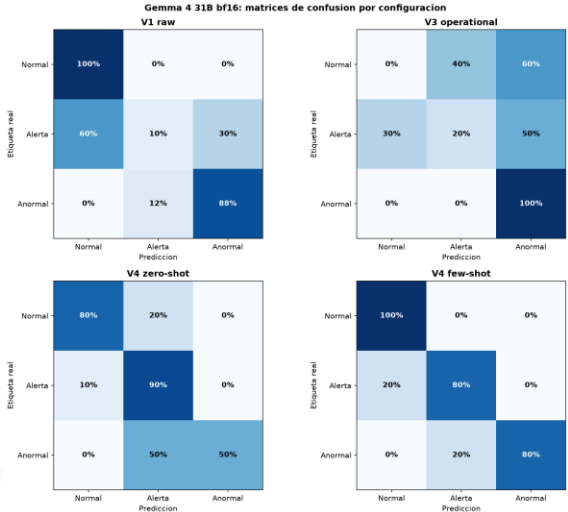
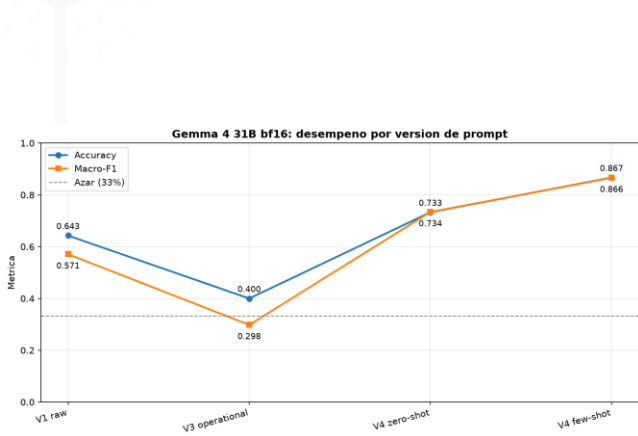
Qué Significa Esto para Mantenición



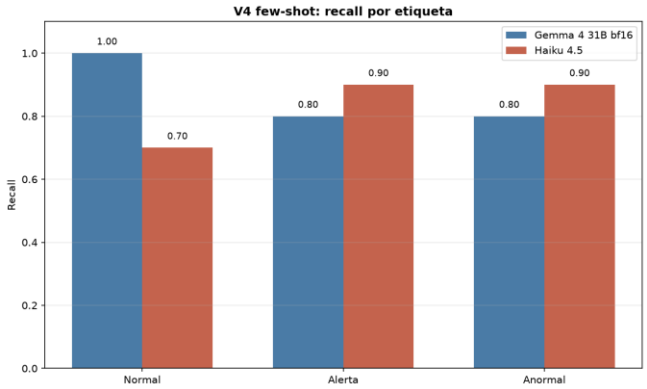
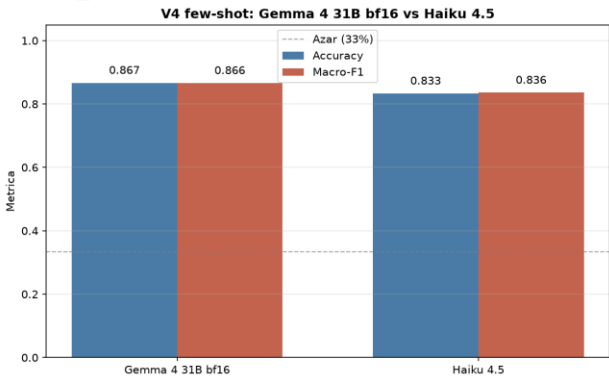
El Experto Sigue en el Sistema



Qué Pasa Si Hacemos lo Mismo con un LLM Local?

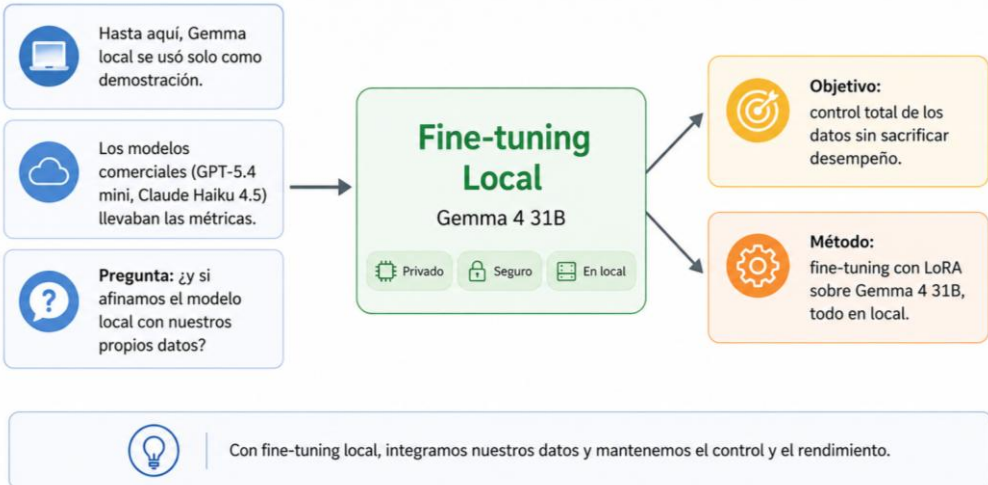


Resultados Comparables al Mejor Modelo Comercial



IA Local: Gemma 4

De Ilustrativo a Competitivo: Gemma Local Afinado



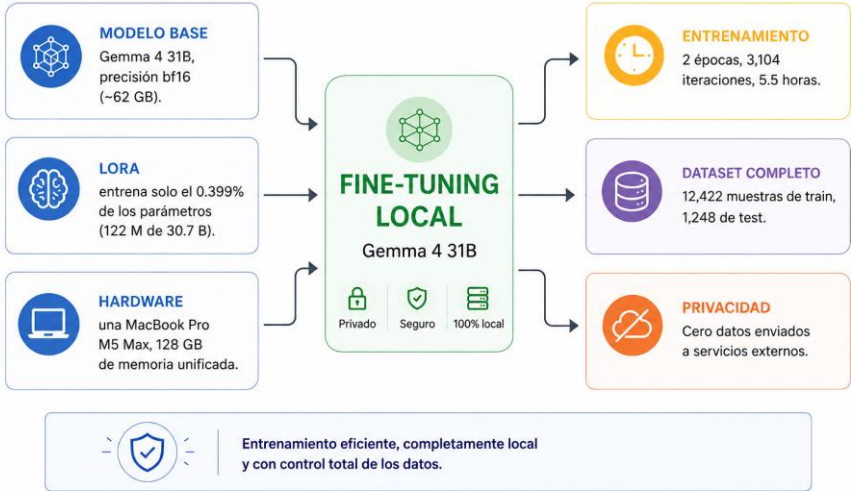
48

Procedimiento del Fine Tuning



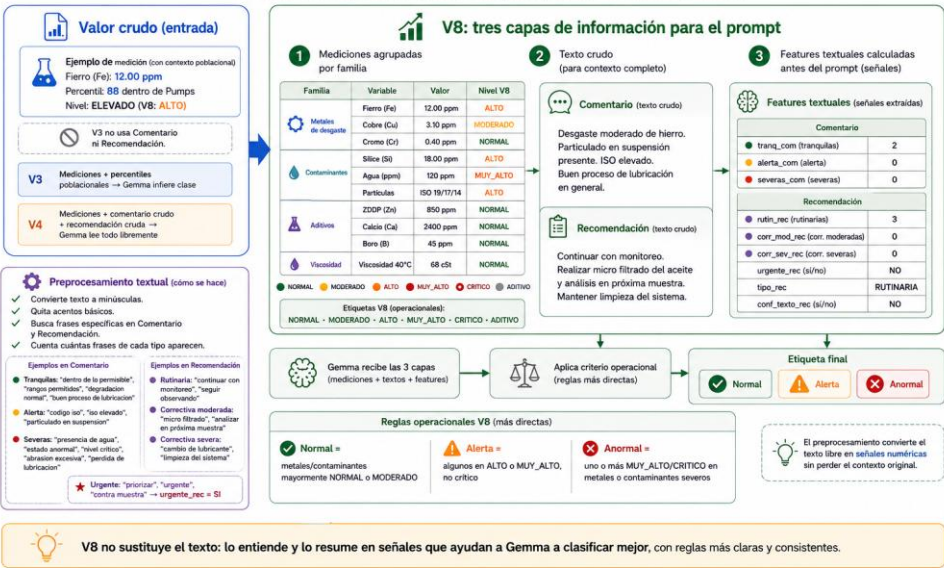
49

Cómo se Afinó: LoRA Local sobre Gemma 4-31B



Importancia del Pre-Procesamiento V8 agrega señales textuales estructuradas al prompt

El modelo interpreta señales en conjunto, no como valores aislados.



Tipo y Flujo de la Información

Aprendizajes por configuración

Cada configuración aporta una mejora metodológica distinta en el flujo de información.



Demo en Vivo: Fine-Tuning Gemma 4

1. Desarrollar la solución predictiva
2. Interpretar los resultados
3. Esto se hace cargando el notebook:
"09_v8_finetuning_demo.ipynb"
1. El notebook y sus librerías se encuentran disponibles en:
"CMC_MX_2026_Toolbox_Code_AI_Local.zip"

Solución en base a AI Local Supera a los Comerciales

Resultados de evaluación

Comparación de desempeño en el test completo

Enfoque	Accuracy	Recall Anormal
Red neuronal (baseline del advisor)	85.00%	90.6%
Mejor comercial: Claude Haiku 4.5 (V4 Few-shot)	83.3%	90%
Gemma 4 31B local, afinado (V8)	91.03%	96.2%



Gemma local alcanza 91.03% de accuracy sobre el test completo.



Supera al mejor comercial por 7.7 puntos de accuracy.



En la clase crítica (Anormal), 96.2% de recall: 6.2 puntos por encima.



Balanced Accuracy: 91.63%. F1 macro: 0.8819.



Nota de evaluación: el comercial se midió sobre 30 muestras balanceadas; Gemma sobre las 1,248 del test real desbalanceado.



Gemma 4 31B afinado localmente ofrece el mejor desempeño global, con mayor precisión y mejor detección de anomalías.



UCLA

54

Dónde Acierta y Dónde Falla

RENDIMIENTO DEL MODELO: MATRIZ DE CONFUSIÓN

Filas = clase real | Columnas = predicción

Real \ Predicho	MATRIZ DE CONFUSIÓN		
	Normal	Alerta	Anormal
Normal (627)	599 (95.5%)	28 (4.5%)	0 (0.0%)
Alerta (464)	10 (2.2%)	386 (83.2%)	68 (14.7%)
Anormal (157)	0 (0.0%)	6 (3.8%)	151 (96.2%)

HALLAZGOS CLAVE



Anormal a Normal

0 casos. Cero subestimaciones críticas.



Normal a Anormal

0 casos. Cero falsas alarmas extremas.



Detección de Anormales

De 157 Anormales reales, detecta 151; pierde solo 6.



Errores residuales

Los errores residuales son conservadores, no peligrosos.

DE LA INFORMACIÓN A LA DECISIÓN



1. Comprender
Analizar la matriz e identificar patrones clave.



2. Evaluar
Validar que los errores sean conservadores.



3. Confiar
Tomar decisiones con base en resultados confiables.



4. Mejorar
Aprender y optimizar para seguir elevando el rendimiento.



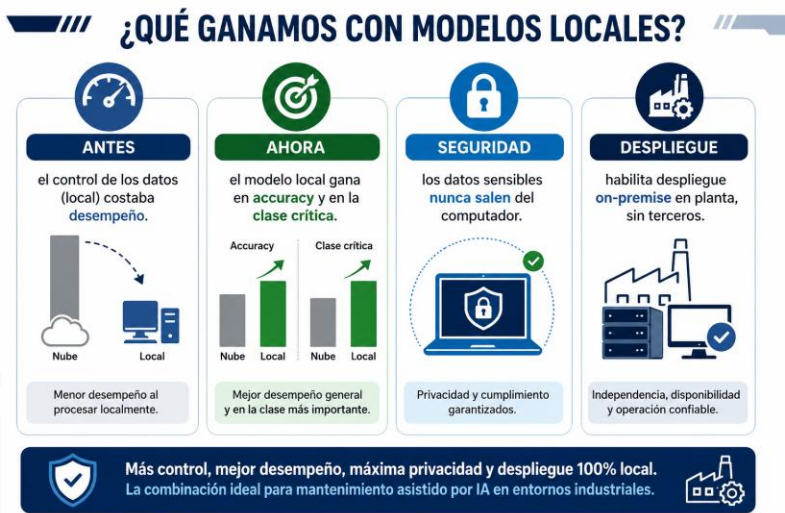
Los resultados muestran un modelo confiable y seguro: cero subestimaciones críticas y cero falsas alarmas extremas.



UCLA

55

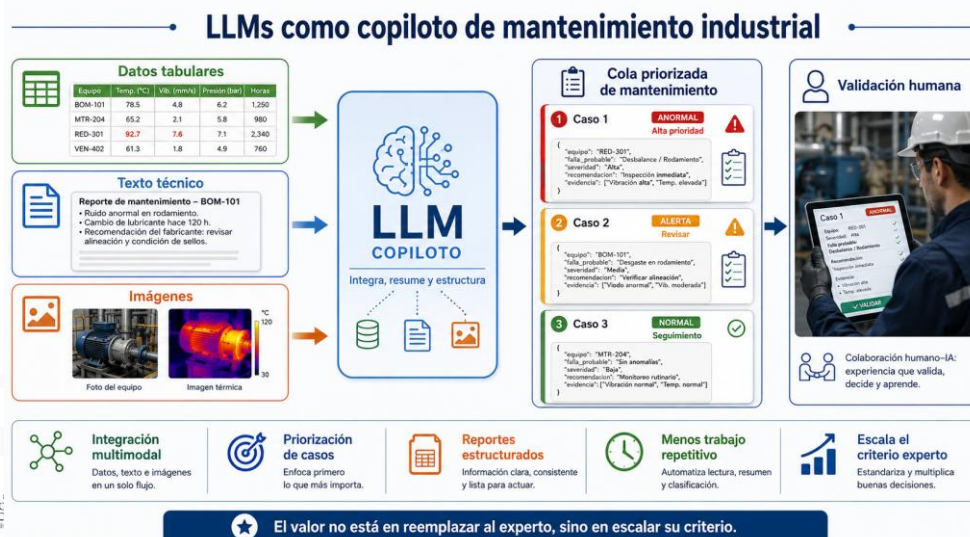
Privacidad y Desempeño: Ya No Hay que Elegir



Conclusiones

Qué Demostramos

- LLMs ya pueden apoyar el mantenimiento industrial:



UCLA

58

Por Qué Aprender Esto Ahora ?



UCLA

59

SESIÓN 20

**ESCANEA EL
CÓDIGO QR**



**RESPONDE UNA
BREVE ENCUESTA**



CONGRESO DE
MANTENIMIENTO
& CONFIABILIDAD
M É X I C O

19^a
EDICIÓN

¡GRACIAS!

Enrique López Droguett

eald@ucla.edu