

# Efficient Knowledge Base Population with LLMs for Industrial Applications: A Space Mission Case Study

Cristian Berrío Aroca – cberrio@expert.ai

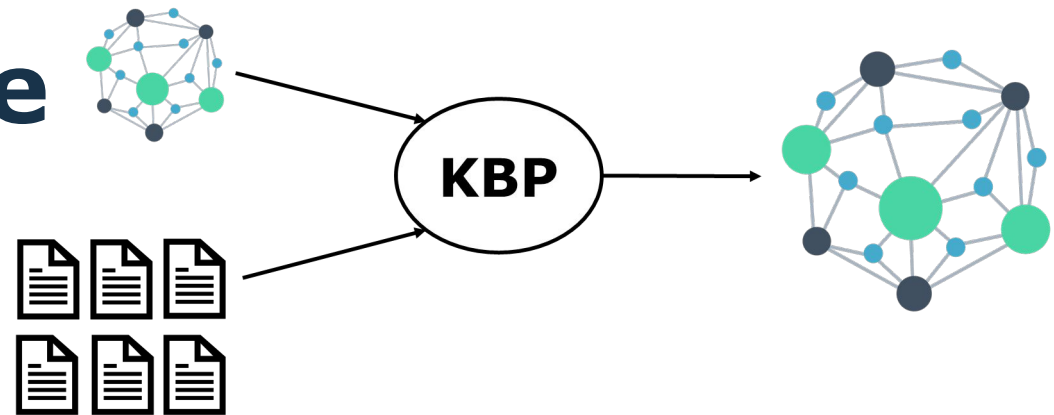
José Manuel Gómez-Pérez - jmgomez@expert.ai



# The Problem

**Structured data from text is still hard**

- **Pipelines are complex**
- **Errors propagate**
- **Scaling is expensive**



# Why LLMs Changed the Game

**Before: multi-step pipelines**



**Now: one model**



**BUT:**  
**Expensive**  
**Slow**  
**Hard to deploy**

# Our Idea

**Use large LLMs once**

**Train small specialized models**

**SSLM**

- **Faster**
- **Cheaper**
- **On-prem**

# How It Works (Intuition)



# Key Result

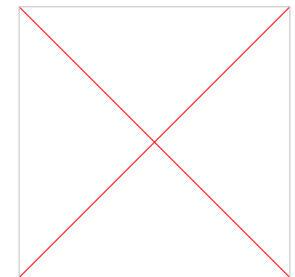
**$\sim 1B \approx 12B$  performance**

**Runs on 1 GPU**

# Rethink LLM usage

**Large → data**

**Small → production**



# Thanks



[Linkedin.com/company/expert-ai](https://www.linkedin.com/company/expert-ai)



[Twitter.com/expertdotai](https://twitter.com/expertdotai)



[info@expert.ai](mailto:info@expert.ai)

*Efficient Knowledge Base Population with Large  
Language Models for Industrial Applications: A  
Space Mission Case Study*

Cristian Berrío Aroca – [cberrio@expert.ai](mailto:cberrio@expert.ai)  
José Manuel Gómez-Pérez - [jmgomez@expert.ai](mailto:jmgomez@expert.ai)