

【AI Summer Camp】

9. Orquestación de Agentes Autónomos

Juan Luis Trapero Ventura



Juan Luis Trapero Ventura



- Graduado en Física Fundamental y Máster en Física Teórica, UCM
- AI Engineer & Consultant, *Nfq Advisory*
- Certificado AI Engineer / Architect por Anthropic, Azure y Databricks
- *Finisher* Maratón Madrid 2025/2026

ÍNDICE DE CONTENIDOS

1. De la IA Generativa a la IA Ejecutiva
2. Interoperabilidad y pegamento digital
3. HUMANO COMO PIVOTE (H.I.T.L.)
4. 5k-tón: aplicación a un caso de uso



MÓDULO 1

DE LA IA GENERATIVA A LA IA EJECUTIVA

【AI Summer Camp】

La (digi)evolución de los LLM

Cada fila es la misma IA subiendo un peldaño: del becario que solo habla al equipo de agentes que razonan y actúan solos.

El salto no es que piense mejor, es que ejecuta mejores acciones.

Fase	Novedad	Autonomía	Uso
LLM	Genera texto	Baja	Chat
Flujos con LLM	Automatiza generación de texto	Baja	Resumen, clasificación, ...
Structured Output	Genera datos estructurados	Baja	Extraer campos, multi-idioma, ...
Tool calling	Acciones con impacto	Media	CRUD sobre base de datos
Agentes (ReAct)	Interacción con feedback	Alta	Proceso auto gestionado
Multi-Agentes	Flujos dinámicos	Muy alta	Entorno multi-disciplinar

Componentes de un agente

LLM

El cerebro. Razona, entiende y decide cuál es el siguiente paso.
Gran impacto en desempeño, pero también en la factura.

Memoria

Recuerda lo que ya pasó. Corto / largo plazo, episódica, ...
Nos permite construir una historia, pero también saturar el contexto.

Tools (MCP)

Las manos. Le dejan interactuar con otros sistemas, no solo hablar.
5 tools curadas >> 100 tools genéricas.

H1 2026: Harness engineering

Comandos

Repetitivo, para escenarios claros y acotados.

ACI (Terminal)

Interacción directa con la máquina que lo ejecuta.

Skills

Exposición gradual a situaciones flexibles y *open-ended*.

Permisos

Límites claros y programáticos, por seguridad y cumplimiento.

Guías/Sensores

Detección de progreso o desalineamiento en el proceso.

Agent Teams

Sistemas de agentes generados de manera dinámica.

DEMO

MÓDULO 2

INTEROPERABILIDAD Y PEGAMENTO DIGITAL

【AI Summer Camp】

Taxonomía de agentes

En 2022 un LLM opinaba.
En 2026 trabaja.

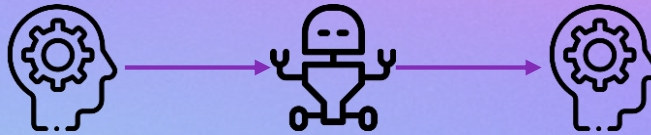
La diferencia no está en el modelo, sino en cómo se apalanca en el entorno: qué hace, qué toca y hasta dónde llega.

Tipo	Rol	Uso
Conversacional	Punto de entrada interactiva / UI	Chat
Informacional	Accede a bases de datos, con información	Investigación, contextualización, ...
Operativo	Genera eventos en el entorno del agente	Generar tareas, enviar e-mails, ...
Planificador	Descompone y acota en fases o tareas	Ejecuciones complejas
Orquestador	Gestiona la ejecución de otros agentes	Toma de decisiones dinámicas

DEMO

Workflows de agentes

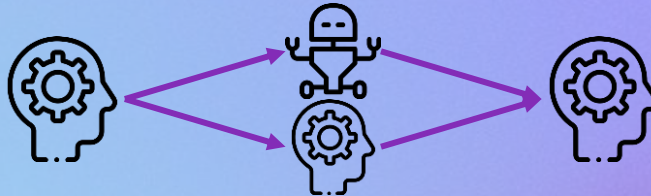
Secuencial



Cada agente recibe como input el output del agente anterior.

Pipeline de extracción y guardado de datos

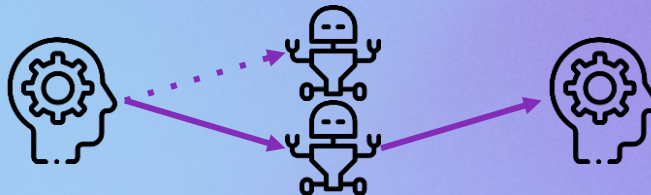
Paralelo



Varios agentes que realizan su tarea sin depender de los otros.

Generación de informes multi-disciplinares

Enrutado

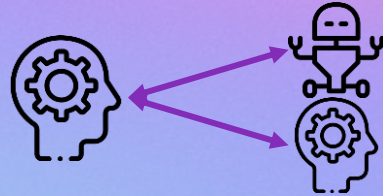


Decide que subagente es el adecuado para la tarea

RAG en conocimientos por departamentos

Equipos de agentes

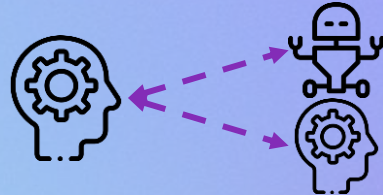
Orquestador



Organiza jerárquicamente tareas.
Los subagentes rinden cuentas
de vuelta siempre.

Atención al cliente en
casos controlados

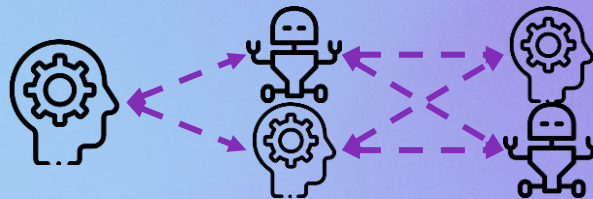
Handoff



Organiza jerárquicamente tareas.
Los subagentes toman el control
del caso.

Atención al cliente por
departamentos
especializados

Mesh



Red de agentes totalmente
distribuidos y conectados /
descubiertos dinámicamente.

Gestión de incidencias
multi-departamento
escaladas

DEMO

MÓDULO 3

HUMANO COMO PIVOTE (H.I.T.L.)

【AI Summer Camp】

Casos con dependencia humana

1. Rendición de cuentas
 2. Efectos irreversibles
 3. Entornos regulados
 4. Desviaciones detectadas
- "Una máquina nunca puede rendir cuentas; por tanto, no debe tomar la decisión." — IBM, 1979.
 - Rendición de cuentas: alguien tiene que **firmar la decisión**, y eso no se delega a un sistema.
 - Efectos irreversibles: **cancelar, borrar, pagar**: lo que no se deshace, se valida antes.
 - Entornos regulados: **salud, banca, seguros**. El negocio exige supervisión humana.
 - Desviaciones: si un **guardrail** detecta algo anómalo, el control vuelve a la persona.

Formas de interacción humana

La intervención humana tiene grados: del simple "avísame de lo que hiciste" al "no muevas nada sin mi OK".

Elegir el nivel correcto para cada acción es el diseño. Demasiado control ahoga al agente; demasiado poco, lo deja suelto.

Tipo	Nivel de intervención	Caso
Validación de resultados	Bajo	Redacción de correos
Validación de decisiones	Medio	Escritura en base de datos (pasivo, reactivo)
Solicitud de criterio experto	Alto	Ayuda con fuentes de datos para búsqueda (proactivo)
Comportamientos desviados detectados	Crítico	Detección de prompt-injection

DEMO

MÓDULO 4

5k-tón: aplicación a un caso de uso

Micro reto práctico: Project Management (abierto a propuestas)

Requisito funcional

Acceso a información sobre el proyecto

Creación, lectura, actualización y eliminación de objetivos, tareas, entregas, ...

Escalado a humano en situaciones necesarias

Opcional: distribuable a otros usuarios

Requisito técnico

Sistema de agentes conectados a fuentes de datos reales

Integración con algún sistema de notificaciones o mensajes

Desarrollo en N8N

Opcional: Integración con MCP custom

GRACIAS

Juan Luis Trapero

LinkedIn

juanlu.trp@gmail.com

【AI Summer Camp】

MADRID
INNOVATION

【MIL】

AINetwork
Artificial Intelligence Association