

UNA SERIE SULL'INTELLIGENZA ARTIFICIALE

Io, Gemini e altri quattro

Piano Editoriale — 10 Episodi

Scaletta completa dei contenuti con abstract per ogni episodio

ROCCO MASSIMILIANO MONDO

mondorocomassimiliano.it

INTRODUZIONE ALLA SERIE

«Io, Gemini e altri quattro» è una serie mensile dedicata all'intelligenza artificiale generativa, pubblicata sul blog <https://mondoroccomassimiliano.it>. Il progetto nasce da un'esigenza concreta: raccontare la rivoluzione dell'IA con gli occhi di un progettista elettrico che ne sperimenta gli strumenti nella propria attività professionale quotidiana.

La serie si distingue per tre caratteristiche fondamentali: una progressione tecnica deliberata (ogni episodio si appoggia sui precedenti), un'impostazione pratica (ogni concetto teorico è accompagnato da esempi reali), e uno stile personale e diretto che non rinuncia al rigore senza cadere nel gergo specialistico fine a sé stesso.

Il percorso porta il lettore dalla panoramica macroeconomica della rivoluzione IA (Parte I) fino alla costruzione di un assistente artificiale personale e specializzato (Parte X), passando per la storia dell'IA, l'architettura tecnica dei modelli, il prompt engineering, il confronto tra i principali LLM, gli agenti, il RAG, le applicazioni nel settore elettrico e l'uso locale con fine-tuning. La creazione di una IA dal DNA STEM per la ricerca borderline in campo ingegneristico e quantistico. Per concludere, un esempio di interazione con il modello A.D.A.M.

PANORAMICA DEI 10 EPISODI

N.	Titolo	Data	Tema	Stato
Parte I	La Rivoluzione dell'Intelligenza Artificiale	Gennaio 2026	Panoramica della rivoluzione IA — dal CES 2026 alla democratizzazione tecnologica	Pubblicato
Parte II	Dalla Macchina di Turing alla Robotica Neurale: Il Viaggio dell'IA	Febbraio 2026	Storia dell'intelligenza artificiale — dalle origini simboliche alla robotica neurale	Pubblicato
Parte III	IA Generativa Spiegata in Dettaglio: Un Viaggio nel Cuore della Tecnologia IA	Marzo 2026	Architettura tecnica degli LLM — Transformer, Self-Attention, allucinazioni, sicurezza dati	Pubblicato
Parte IV	Il Prompt Engineering: L'Arte di Comunicare con l'IA	Aprile 2026	Tecniche pratiche di prompt engineering — struttura, metodi, applicazione tecnica	Pubblicato
Parte V	Il Grande Confronto: GPT, Gemini, Claude, LLaMA e Grok	Maggio 2026	Analisi comparativa dei 5 principali LLM — parametri tecnici, punti di forza, scelta per caso d'uso	Pubblicato
Parte VI	Gli Agenti IA: Quando l'Intelligenza Artificiale Agisce	Giugno 2026	AI Agents, tool use, automazione di workflow — dai modelli reattivi ai sistemi proattivi	Pubblicato
Parte VII	RAG e la Memoria Contestuale: Dare un Cervello Privato all'IA	Luglio 2026	Retrieval Augmented Generation — come estendere la conoscenza di un LLM con dati propri	Pubblicato
Parte VIII	IA Locale e Fine-Tuning: Addestrare il Proprio Modello	Agosto 2026	LLM locali con Ollama, fine-tuning su dataset	Pianificato

			specializzati, hardware e privacy totale	
Parte IX	Il Progetto A.D.A.M. IA personale per la ricerca ingegneristica borderline	Settembre 2026	Sintesi della serie: il progetto A.D.A.M. , DNA di una IA votata alla ricerca scientifica	Pianificato
Parte X	Il Progetto A.D.A.M. IA personale per la ricerca ingegneristica borderline	Ottobre 2026	Sintesi della serie: un corpo per A.D.A.M. , architettura progettuale per una infrastruttura sostenibile nel tempo.	Pianificato

	Extra contenuti	Novembre 2026	Dialogando con A.D.A.M.	Pianificato
--	-----------------	---------------	--------------------------------	--------------------

Legenda: ■ Pubblicato ■ Pianificato

NOTE EDITORIALI

Cadenza: un episodio al mese. Ogni articolo è pensato per essere autonomo ma beneficia della lettura in sequenza.

Lunghezza target: 2.500–4.000 parole per episodio, con infografiche e diagrammi tecnici integrati.

Livello tecnico: progressivo. Parti I-II accessibili a tutti; Parti VII-X per lettori con conoscenze informatiche di base.

Repository: gli esempi di codice e i prompt mostrati negli episodi saranno disponibili sul blog con licenza Creative Commons.

ABSTRACT PER EPISODIO

Parte I — La Rivoluzione dell'Intelligenza Artificiale

Data: Gennaio 2026 Stato: **Pubblicato**

Tema: Panoramica della rivoluzione IA — dal CES 2026 alla democratizzazione tecnologica

Il primo episodio inquadra la rivoluzione dell'IA nel contesto del CES 2026 di Las Vegas, fotografando il 2025 come l'anno della svolta. Vengono introdotti i due paradigmi fondamentali — inferenza e addestramento — e viene esplorato il concetto di democratizzazione tecnologica: la possibilità di far girare modelli potenti localmente, su hardware personale. Non manca la lettura critica dell'altra faccia della medaglia: il vertiginoso aumento del costo delle memorie DDR5 e la tensione tra accesso aperto e controllo economico dei grandi player. La conclusione cita Dante: «Fatti non foste a viver come bruti».

Parte II — Dalla Macchina di Turing alla Robotica Neurale: Il Viaggio dell'IA

Data: Febbraio 2026 Stato: **Pubblicato**

Tema: Storia dell'intelligenza artificiale — dalle origini simboliche alla robotica neurale

Un viaggio cronologico in quattro tappe: l'IA simbolica di Alan Turing (1950-1960), il cambio di paradigma di Deep Blue (1980-2010), l'era della generazione con LLM e Agenti (2017-2025), e la frontiera della robotica neurale distribuita (2025-2030). La narrazione alterna rigore storico e curiosità divulgativa, con il consiglio di vedere il film «The Imitation Game» e con una riflessione finale sull'inevitabilità di questa rivoluzione e sulla scelta, tutta umana, tra paura e opportunità.

Parte III — IA Generativa Spiegata in Dettaglio: Un Viaggio nel Cuore della Tecnologia IA

Data: Marzo 2026 Stato: **Pubblicato**

Tema: Architettura tecnica degli LLM — Transformer, Self-Attention, allucinazioni, sicurezza dati

L'episodio più tecnico dei primi tre entra nell'anatomia di un LLM con l'architettura Transformer. Partendo dal flusso di inferenza (uomo → IA → output → valutazione umana), si descrive come un LLM non memorizzi i dati ma li «distilli» in miliardi di pesi numerici. Il meccanismo di Self-Attention viene spiegato attraverso l'esempio della parola «pesca». Si affrontano i limiti strutturali: le allucinazioni e la sicurezza dei dati privati. Una tabella comparativa chiude il confronto tra IA tradizionale e IA generativa. Conclusione: «il limite non è più la potenza di calcolo, ma la qualità dei dati e la capacità umana di porre le domande giuste».

Parte IV — Il Prompt Engineering: L'Arte di Comunicare con l'IA

Data: Aprile 2026 Stato: **Pianificato**

Tema: Tecniche pratiche di prompt engineering — struttura, metodi, applicazione tecnica

Primo episodio interamente pratico della serie. Si introduce il prompt engineering come disciplina comunicativa che determina la qualità dell'output di un LLM. Vengono analizzati i cinque componenti di un prompt efficace (Ruolo, Contesto, Compito, Formato, Vincoli) e le principali tecniche professionali: Zero-Shot, Few-Shot, Chain-of-Thought e Role+CoT. Una sezione dedicata agli errori più comuni completa la parte teorica. Il caso pratico conclusivo mostra come strutturare un prompt per il calcolo illuminotecnico di un'officina meccanica secondo la norma UNI EN 12464-1, con verifica manuale del risultato.

Parte V — Il Grande Confronto: GPT, Gemini, Claude, LLaMA e Grok

Data: Maggio 2026 Stato: **Pianificato**

Tema: Analisi comparativa dei 5 principali LLM — parametri tecnici, punti di forza, scelta per caso d'uso

Una guida sistematica ai cinque grandi protagonisti del panorama LLM. Per ogni modello si descrive l'azienda sviluppatrice, la filosofia progettuale e i punti di forza distintivi. Una tabella comparativa valuta context window, multimodalità, disponibilità open-source, accesso API e costo. La sezione pratica risponde alla domanda concreta: quale scegliere per la scrittura tecnica, per la ricerca web, per il coding, per l'uso locale. Chiude una riflessione sul paradigma open-source vs proprietario nel contesto della democratizzazione tecnologica introdotta in Parte I.

Parte VI — Gli Agenti IA: Quando l'Intelligenza Artificiale Agisce

Data: Giugno 2026 Stato: **Pianificato**

Tema: AI Agents, tool use, automazione di workflow — dai modelli reattivi ai sistemi proattivi

I modelli fin qui descritti sono sistemi reattivi: rispondono a domande. Gli Agenti IA rappresentano il salto successivo: sistemi capaci di pianificare, eseguire azioni nel mondo reale e adattarsi dinamicamente ai risultati. Si introduce il concetto di tool use (uso di strumenti: ricerca web, API esterne, scrittura di file, esecuzione di codice) e il protocollo MCP (Model Context Protocol). Casi d'uso pratici: un agente che legge bollette Edison, le analizza e produce un report in Excel; un agente che verifica la conformità normativa di un progetto elettrico. Si discutono anche i rischi: l'autonomia degli agenti richiede fiducia calibrata e supervisione umana.

Parte VII — RAG e la Memoria Contestuale: Dare un Cervello Privato all'IA

Data: Luglio 2026 Stato: **Pianificato**

Tema: Retrieval Augmented Generation — come estendere la conoscenza di un LLM con dati propri

Un LLM ha una conoscenza ferma alla data del suo training. Il RAG (Retrieval Augmented Generation) è la tecnica che permette di «iniettare» conoscenza aggiornata e privata nel processo di inferenza senza riaddestrare il modello. Si spiega il concetto di vettore di embedding e di database vettoriale (ChromaDB, FAISS, Pinecone). Caso d'uso pratico: costruire un sistema RAG che risponde a domande sulle normative CEI partendo dai documenti PDF ufficiali. Si analizzano implicazioni di privacy (dati restano locali), costo e complessità di implementazione. L'articolo apre la strada al progetto personale descritto nella Parte X.

Parte VIII — L'IA nell'Ingegneria Elettrica: Dal Calcolo al Progetto Assistito

Data: Agosto 2026 Stato: **Pianificato**

Tema: Applicazioni specifiche dell'IA nel settore dell'ingegneria elettrica e illuminotecnica

Episodio dedicato alla convergenza tra intelligenza artificiale e ingegneria elettrica, il campo professionale dell'autore. Si esplora come gli LLM possano assistere nella verifica normativa (norme CEI, IEC, UNI EN), nel calcolo illuminotecnico (raccordandosi con ElettroLux e la Guida al Calcolo Illuminotecnico v1.6), nella generazione di relazioni tecniche strutturate e nell'analisi di schemi elettrici tramite visione artificiale. Si presentano casi reali di utilizzo e si valutano i limiti: l'IA come assistente che potenzia il professionista, non come sostituto. Una riflessione sull'etica professionale nella verifica degli output chiude l'articolo.

Parte IX — IA Locale e Fine-Tuning: Addestrare il Proprio Modello

Data: Settembre 2026 Stato: **Pianificato**

Tema: LLM locali con Ollama, fine-tuning su dataset specializzati, hardware e privacy totale

La «democratizzazione tecnologica» introdotta nella Parte I qui diventa progetto concreto. Con strumenti come Ollama, LM Studio e llama.cpp è oggi possibile far girare modelli come LLaMA 3.3, Mistral o Phi-4 su hardware consumer con GPU discreta. Si analizza il processo di fine-tuning (LoRA, QLoRA) per specializzare un modello base su un dominio tecnico ristretto — in questo caso le normative di ingegneria elettrica. Sezione tecnica dedicata a: requisiti hardware, formati di dati per il training, metriche di valutazione. L'hardware necessario viene confrontato con i costi cloud equivalenti, sfatando il mito che il locale sia per forza più costoso.

Parte X — Il Progetto Definitivo: Un'IA Personale per l'Ingegneria

Data: Ottobre 2026 Stato: **Pianificato**

Tema: Sintesi della serie — costruzione di un assistente IA specializzato per ingegneria elettrica e fisica

Episodio conclusivo della serie. Partendo da quanto appreso nei nove episodi precedenti, si presenta il progetto concreto: un assistente IA personale, eseguito localmente, specializzato in ingegneria elettrica, normative CEI e calcolo illuminotecnico, con integrazione RAG su documenti tecnici propri. Si descrive l'architettura del sistema (modello base fine-tunato + RAG + interfaccia Python), le lezioni apprese nel percorso e le prospettive future. L'articolo si chiude con una riflessione personale sull'importanza della conoscenza tecnica nell'era dell'IA: la macchina amplifica il professionista che la sa usare, ma non sostituisce la competenza di chi pone le domande giuste.

