

PERCORSO FORMATIVO — AI GENERATIVA APPLICATA

Dispensa 8

Ambienti Simulati e Agenti AI

Comprendere, senza timore, i sistemi capaci di eseguire compiti in autonomia, e come sperimentarli in sicurezza.

A cura di Marco Consiglio

Formatore e Course Director — Contenuti aggiornati al 2026

Indice

1. Introduzione e obiettivi didattici	3
2. Che cos'è un agente AI, oltre la sola conversazione.....	4
3. La differenza tra un assistente conversazionale e un agente	6
4. Perché un ambiente simulato è il modo più sicuro per sperimentare.....	9
5. Un esempio guidato di ambiente simulato	11
6. I rischi specifici degli agenti autonomi.....	13
7. Un principio di supervisione proporzionata al livello di autonomia.....	15
8. Le applicazioni professionali più promettenti nel breve periodo	17
9. Caso pratico guidato.....	18
10. Glossario dei termini chiave.....	20
11. Checklist finale e autoverifica.....	21
12. Riferimenti e risorse aggiornate al 2026	22

1. Introduzione e obiettivi didattici

Questa dispensa prosegue il percorso AI Generativa Applicata, affrontando un tema più avanzato rispetto alle applicazioni già presentate nelle dispense precedenti: gli agenti AI, sistemi capaci non solo di rispondere a una richiesta con un testo, ma di eseguire in autonomia una sequenza di azioni per raggiungere un obiettivo, e gli ambienti simulati che permettono di sperimentare con questi sistemi in condizioni controllate e sicure.

Applicando quanto già presentato nella Dispensa 1 di questo percorso a proposito del secondo luogo comune sfatato, questa dispensa affronta con lo stesso approccio equilibrato già più volte richiamato in questo percorso un tema che genera spesso un misto di curiosità e preoccupazione, offrendo strumenti concreti per comprenderlo e sperimentarlo con la dovuta cautela.

A chi si rivolge questa dispensa

Si rivolge a chi ha seguito le prime sette dispense di questo percorso e desidera comprendere, partendo dalle basi già costruite, che cosa sono gli agenti AI e come sperimentarli in sicurezza attraverso ambienti simulati.

Obiettivi di apprendimento

1. Comprendere che cos'è un agente AI, oltre la sola conversazione testuale;
2. Distinguere chiaramente un assistente conversazionale da un agente autonomo;
3. Comprendere perché un ambiente simulato rappresenta il modo più sicuro per sperimentare;
4. Conoscere un esempio guidato di ambiente simulato;
5. Riconoscere i rischi specifici degli agenti autonomi;
6. Comprendere un principio di supervisione proporzionata, e conoscere le applicazioni professionali più promettenti nel breve periodo.

2. Che cos'è un agente AI, oltre la sola conversazione

Applicando quanto già presentato nella Dispensa 1 di questo percorso a proposito della definizione operativa di AI generativa, iniziamo definendo con precisione che cos'è un agente AI, un concetto distinto ma correlato agli strumenti conversazionali già presentati nelle dispense precedenti di questo percorso.

Una definizione operativa

Un agente AI è un sistema costruito a partire da uno strumento di AI generativa, capace non solo di rispondere a una domanda con un testo, ma di pianificare ed eseguire in autonomia una sequenza di azioni concrete (cercare un'informazione, compilare un modulo, inviare una comunicazione, interagire con altri strumenti software) per raggiungere un obiettivo assegnato, con un livello di autonomia superiore rispetto a un semplice scambio conversazionale.

Un'analogia utile per comprendere questa distinzione

Un'analogia utile: uno strumento conversazionale già presentato nelle dispense precedenti di questo percorso può essere paragonato a un consulente che risponde a domande e fornisce indicazioni, mentre un agente AI può essere paragonato a un assistente a cui si delega l'esecuzione concreta di un compito, che procede autonomamente attraverso più passaggi operativi fino al completamento dell'obiettivo assegnato, riportando poi il risultato finale, secondo un livello di autonomia operativa distinto dal solo scambio di informazioni già presentato negli strumenti conversazionali di base.

Un elemento di continuità con quanto già presentato in questo percorso

Applicando quanto già presentato nella Dispensa 7 di questo percorso a proposito dell'automazione parziale, gli agenti AI possono essere considerati un'estensione più avanzata degli stessi principi di automazione già presentati in quella dispensa, con la differenza che un agente non si limita a generare una proposta di testo da rivedere, ma esegue direttamente azioni concrete, un elemento che richiede, come sarà approfondito nel capitolo 6 di questa dispensa, un livello di attenzione e supervisione ancora maggiore rispetto a quanto già presentato per l'automazione parziale di compiti testuali.

Perché questo tema merita una dispensa dedicata

Applicando lo stesso principio di proporzionalità già più volte richiamato in questa collana, gli agenti AI meritano una trattazione dedicata, distinta dalle applicazioni pratiche già presentate nelle Dispense 5, 6 e 7 di questo percorso, poiché introducono una categoria di rischi e opportunità qualitativamente diversa: non più la sola generazione di contenuti da rivedere prima dell'uso, ma l'esecuzione autonoma di azioni con conseguenze dirette, un salto concettuale che giustifica pienamente l'approfondimento specifico presentato in questa dispensa, prima di

passare, nelle dispense successive di questo percorso, alle ulteriori applicazioni pratiche della creatività multimediale e dell'integrazione nei flussi di lavoro aziendali.

3. La differenza tra un assistente conversazionale e un agente

Approfondiamo ora sistematicamente la distinzione già introdotta nel capitolo precedente, un elemento centrale per comprendere correttamente le potenzialità e i rischi specifici degli agenti AI rispetto agli strumenti conversazionali già presentati nelle dispense precedenti di questo percorso.

Un confronto sistematico

Caratteristica	Assistente conversazionale	Agente AI
Output	Un testo, un'immagine, un contenuto	Un'azione concreta eseguita
Autonomia	Nessuna, risponde a una singola richiesta	Pianifica ed esegue più passaggi in sequenza
Supervisione tipica	Verifica del contenuto finale	Verifica di ciascun passaggio o del processo

Un esempio concreto di questa differenza

Un esempio pratico chiarisce questa distinzione: chiedere a uno strumento conversazionale già presentato nelle dispense precedenti di questo percorso di scrivere un'email di sollecito a un cliente produce un testo che un professionista rivede, personalizza secondo quanto già presentato nella Dispensa 5 di questo percorso, e invia personalmente. Un agente AI, in un contesto opportunamente configurato, potrebbe non solo scrivere l'email, ma anche identificare autonomamente quali clienti richiedono un sollecito sulla base di criteri predefiniti, e inviare effettivamente le comunicazioni senza un intervento manuale per ciascun singolo invio, un livello di autonomia che comporta implicazioni pratiche molto diverse.

Un principio di gradualità nell'autonomia concessa

Applicando lo stesso principio di gradualità già più volte richiamato in questo percorso, gli agenti AI non operano secondo un'unica modalità di autonomia fissa, ma possono essere configurati con diversi livelli di supervisione richiesta, da un agente che richiede un'approvazione esplicita prima di ciascuna azione concreta, a un agente che opera con maggiore autonomia entro limiti predefiniti, un tema che sarà approfondito con maggiore dettaglio nel capitolo 7 di questa dispensa.

Un elemento pratico: la trasparenza sulle azioni intraprese

Un ultimo elemento pratico, distinto dai livelli di autonomia già presentati in questo capitolo, riguarda la trasparenza del processo seguito da un agente per raggiungere il proprio obiettivo: un

agente ben configurato dovrebbe fornire, insieme al risultato finale, un resoconto chiaro delle azioni intraprese e delle decisioni prese lungo il percorso, secondo un principio analogo al ragionamento esplicito già presentato nella Dispensa 4 di questo percorso, ma applicato a una sequenza di azioni concrete piuttosto che a un solo ragionamento testuale. Questa trasparenza permette a un professionista di verificare non solo il risultato finale, ma anche la logica e il percorso seguiti per raggiungerlo, un elemento essenziale per costruire fiducia informata nel comportamento dell'agente nel tempo.

4. Perché un ambiente simulato è il modo più sicuro per sperimentare

Applicando quanto già presentato nel capitolo 3 di questa dispensa a proposito dei diversi livelli di autonomia, presentiamo ora un concetto centrale per chi desidera avvicinarsi a questi strumenti più avanzati con la dovuta prudenza: l'ambiente simulato, uno spazio protetto in cui sperimentare senza il rischio di conseguenze reali indesiderate.

Una definizione operativa

Un ambiente simulato è uno spazio di sperimentazione in cui un agente AI può operare ed eseguire azioni senza che queste abbiano conseguenze reali sui sistemi effettivi di un'organizzazione: un'email che sembra essere inviata ma resta in realtà confinata all'ambiente simulato, una modifica a un documento che non altera il file reale, un'azione che permette di osservare il comportamento dell'agente senza alcun rischio concreto.

Perché questo strumento è particolarmente prezioso per chi si avvicina per la prima volta

Applicando quanto già presentato nella Dispensa 1 di questo percorso a proposito della curiosità sperimentale come approccio efficace, un ambiente simulato permette di applicare pienamente questo stesso principio anche a strumenti più avanzati e potenzialmente più rischiosi come gli agenti AI, senza il timore che un errore di configurazione o un comportamento inatteso dell'agente possa causare un danno reale, un elemento che rende questo strumento particolarmente prezioso per costruire familiarità e fiducia prima di un eventuale utilizzo in un contesto realmente operativo.

Un principio pratico: iniziare sempre da un ambiente simulato per compiti nuovi

Applicando lo stesso principio di gradualità già più volte richiamato in questo percorso, un professionista che desidera configurare un agente AI per un nuovo compito, anche quando ha già maturato esperienza con altri agenti in passato, dovrebbe idealmente testarlo prima in un ambiente simulato, verificando che il comportamento dell'agente sia effettivamente allineato alle proprie aspettative, prima di autorizzarlo a operare su sistemi e dati reali.

Un elemento pratico: che cosa fare quando un ambiente simulato non è disponibile

Applicando lo stesso principio di proporzionalità già più volte richiamato in questo percorso, non tutti gli strumenti o i contesti offrono un ambiente simulato pienamente strutturato secondo quanto già presentato in questo capitolo: quando questa opzione non è disponibile, un professionista può comunque avvicinarsi allo stesso principio attraverso approssimazioni pratiche, come testare un agente su un sottoinsieme di dati non sensibili e facilmente verificabili, o limitare inizialmente l'autorizzazione dell'agente a compiti con conseguenze reversibili e di

scarsa rilevanza economica, secondo la stessa logica di rischio contenuto già presentata in questo capitolo, pur senza la piena protezione offerta da un ambiente simulato formale.

5. Un esempio guidato di ambiente simulato

Presentiamo ora un esempio concreto e semplificato di come un ambiente simulato potrebbe essere utilizzato per testare il comportamento di un agente AI, applicando gli strumenti concettuali già presentati nei capitoli precedenti di questa dispensa.

Il contesto dell'esempio

Un professionista desidera configurare un agente AI per gestire autonomamente una prima risposta alle richieste di preventivo ricevute via email, secondo un compito concettualmente simile all'automazione della classificazione già presentata nella Dispensa 7 di questo percorso, ma con un livello di autonomia superiore, poiché l'agente dovrebbe anche predisporre e inviare autonomamente una prima risposta.

Le fasi tipiche di un test in ambiente simulato

1. Configurare l'agente con istruzioni chiare sul compito da svolgere, applicando gli stessi principi di prompt engineering già presentati nella Dispensa 3 di questo percorso;
2. Fornire all'agente, all'interno dell'ambiente simulato, un insieme di email di esempio rappresentative delle situazioni reali che dovrà affrontare;
3. Osservare il comportamento dell'agente su ciascuna email di esempio, verificando se le azioni intraprese siano effettivamente coerenti con le aspettative;
4. Identificare eventuali comportamenti inattesi o problematici, affinando le istruzioni fornite prima di considerare un eventuale utilizzo in un contesto reale.

Un principio pratico: testare anche i casi limite

Applicando quanto già presentato nella Dispensa 7 di questo percorso a proposito della gestione dei casi ambigui, un test efficace in ambiente simulato non dovrebbe limitarsi a verificare il comportamento dell'agente sui casi più tipici e prevedibili, ma dovrebbe includere esplicitamente anche casi limite o ambigui (una richiesta di preventivo incompleta, una comunicazione che non è realmente una richiesta di preventivo mascherata da tale), per verificare come l'agente gestisce situazioni che si discostano dallo schema più comune, un elemento cruciale per costruire fiducia nel comportamento dell'agente prima di un utilizzo su dati reali.

6. I rischi specifici degli agenti autonomi

Applicando quanto già presentato nella Dispensa 1 di questo percorso a proposito dei reali limiti di questi strumenti, approfondiamo ora sistematicamente i rischi specifici associati agli agenti AI, distinti e per certi versi amplificati rispetto ai rischi già presentati per gli strumenti puramente conversazionali.

Il rischio di azioni irreversibili

Applicando quanto già presentato nel capitolo 3 di questa dispensa a proposito della differenza tra un testo generato e un'azione eseguita, un rischio specifico e particolarmente rilevante degli agenti AI riguarda la possibilità che un'azione eseguita autonomamente, se erronea, non sia facilmente correggibile: un'email inviata per errore, un ordine effettuato in modo inappropriato, o una modifica applicata a un documento reale comportano conseguenze concrete difficilmente reversibili, a differenza di un testo generato ma non ancora inviato, che un professionista può sempre correggere o scartare prima di un utilizzo effettivo.

Il rischio di propagazione degli errori

Applicando quanto già presentato nel capitolo 2 di questa dispensa a proposito della sequenza di azioni eseguite da un agente, un secondo rischio specifico riguarda la propagazione di un errore iniziale attraverso l'intera sequenza di azioni successive: se un agente commette un errore nella prima fase di un compito articolato, questo errore può propagarsi e amplificarsi nelle fasi successive, un rischio che non si presenta nello stesso modo per una singola risposta testuale isolata già presentata nelle dispense precedenti di questo percorso.

Il rischio di un'interpretazione eccessivamente letterale dell'obiettivo

Un terzo rischio, meno intuitivo ma ugualmente rilevante: un agente potrebbe interpretare l'obiettivo assegnato in modo eccessivamente letterale, intraprendendo azioni tecnicamente coerenti con l'istruzione ricevuta ma non allineate con l'intento reale del professionista che l'ha formulata, un rischio che richiede istruzioni particolarmente chiare e complete, secondo gli stessi principi già presentati nella Dispensa 3 di questo percorso, ma applicati con un rigore ancora maggiore data la natura operativa, non solo comunicativa, delle azioni intraprese da un agente.

Un quarto rischio: l'interazione con sistemi esterni non completamente prevedibili

Un ultimo rischio, distinto dai tre già presentati in questo capitolo, riguarda le situazioni in cui un agente interagisce con sistemi esterni (un altro software, una piattaforma di terzi) il cui comportamento non è completamente sotto il controllo diretto del professionista che ha configurato l'agente: un cambiamento nell'interfaccia di un sistema esterno, un'interruzione temporanea di un servizio, o una risposta inattesa da parte di un sistema terzo possono generare comportamenti imprevedibili nell'agente stesso, un rischio che si aggiunge a quelli già presentati e

che richiede una particolare attenzione quando un agente è configurato per interagire con più sistemi esterni contemporaneamente.

7. Un principio di supervisione proporzionata al livello di autonomia

Applicando quanto già presentato nel capitolo 6 di questa dispensa a proposito dei rischi specifici degli agenti autonomi, presentiamo ora un principio pratico centrale: la supervisione dovrebbe essere calibrata in proporzione al livello di autonomia concesso e alla rilevanza delle conseguenze di un eventuale errore.

Una scala di supervisione proporzionata

Livello di rischio	Livello di supervisione richiesto
Basso (azioni facilmente reversibili)	Supervisione a campione, verifica periodica
Medio	Approvazione esplicita prima di ciascuna azione rilevante
Alto (azioni irreversibili o di rilevanza economica)	Supervisione umana diretta e continuativa, nessuna autonomia concessa

Un principio pratico di introduzione graduale dell'autonomia

Applicando lo stesso principio di gradualità già più volte richiamato in questo percorso, un professionista dovrebbe introdurre progressivamente un livello crescente di autonomia per un agente AI, partendo da un livello di supervisione elevato secondo quanto già presentato in questo capitolo, e riducendo gradualmente questa supervisione solo dopo aver costruito, attraverso un utilizzo prolungato e positivamente validato secondo gli strumenti già presentati nel capitolo 5 di questa dispensa, una fiducia solida nel comportamento affidabile dell'agente per quello specifico compito.

8. Le applicazioni professionali più promettenti nel breve periodo

Concludiamo la parte teorica di questa dispensa con una panoramica delle applicazioni professionali degli agenti AI che appaiono più promettenti e mature nel contesto attuale, applicando lo stesso principio di equilibrio tra entusiasmo e cautela già più volte richiamato in questo percorso.

Compiti di ricerca e raccolta di informazioni

Applicando quanto già presentato nella Dispensa 6 di questo percorso a proposito dell'analisi di documenti, gli agenti AI si dimostrano particolarmente efficaci in compiti di ricerca strutturata (racogliere informazioni da più fonti su un argomento specifico, preparare una prima bozza di ricerca comparativa), un ambito in cui il rischio di azioni irreversibili già presentato nel capitolo 6 di questa dispensa è tipicamente contenuto.

Compiti di prima elaborazione, con revisione umana sistematica

Applicando quanto già presentato nella Dispensa 7 di questo percorso a proposito dell'automazione parziale, gli agenti AI si prestano bene a compiti di prima elaborazione che restano sempre sottoposti a una revisione umana sistematica prima di qualunque conseguenza reale, secondo lo stesso principio di automazione parziale già presentato in quella dispensa, applicato a compiti più articolati e multi-fase rispetto alla sola generazione di un testo singolo.

Compiti di organizzazione e preparazione, in contesti a basso rischio

Un terzo ambito applicativo promettente, distinto dai due già presentati in questo capitolo, riguarda compiti di organizzazione e preparazione che non comportano conseguenze dirette verso l'esterno di un'organizzazione: la preparazione di una prima bozza di agenda per una riunione a partire da materiali già disponibili, l'organizzazione preliminare di documenti in cartelle secondo criteri predefiniti, o la predisposizione di un primo elenco di attività da svolgere sulla base di comunicazioni ricevute. Questi compiti, per la loro natura puramente organizzativa e interna, presentano tipicamente un livello di rischio contenuto anche in caso di errore, poiché le loro conseguenze restano generalmente circoscritte all'ambito interno dell'organizzazione e facilmente correggibili, secondo gli stessi criteri di valutazione del rischio già presentati nel capitolo 7 di questa dispensa.

Un principio di sintesi per l'intera dispensa

Gli agenti AI e gli ambienti simulati, in tutte le dimensioni presentate in questa dispensa — la distinzione dagli strumenti conversazionali, il valore della sperimentazione protetta, i rischi specifici dell'autonomia operativa, la supervisione proporzionata, le applicazioni più promettenti — rappresentano un'evoluzione più avanzata degli strumenti già presentati in questo percorso, da avvicinare con la stessa curiosità sperimentale già presentata nella Dispensa 1, ma con un livello di cautela proporzionalmente maggiore rispetto alla maggiore autonomia operativa che questi sistemi comportano.

Un ultimo elemento: un ambito in rapida evoluzione

Un'ultima riflessione conclude questa dispensa: gli agenti AI rappresentano, tra tutti gli strumenti presentati in questo percorso, l'ambito probabilmente più giovane e in più rapida evoluzione, con capacità e limiti che potrebbero cambiare significativamente nel tempo rispetto a quanto presentato in questa dispensa. Applicando lo stesso principio di aggiornamento continuativo che sarà presentato con maggiore dettaglio nell'ultima dispensa di questo percorso, un professionista dovrebbe mantenere un aggiornamento periodico su questo specifico ambito, mantenendo sempre, indipendentemente da come questi strumenti evolveranno, i principi fondamentali di prudenza e supervisione proporzionata già presentati in questa dispensa come guida costante nel loro utilizzo.

9. Caso pratico guidato

Concludiamo la parte teorica con un caso pratico semplificato, che applica gli strumenti presentati in questa dispensa.

Il contesto

Riprendiamo il caso già presentato nelle dispense precedenti di questo percorso: il responsabile amministrativo, dopo aver automatizzato con successo la classificazione delle email già presentata nella Dispensa 7 di questo percorso, valuta l'introduzione di un agente AI più autonomo, capace non solo di classificare ma anche di predisporre autonomamente una prima risposta per le richieste più semplici e ricorrenti.

Passo 1 — Configurazione dell'agente in ambiente simulato

Applicando gli strumenti del capitolo 5, l'agente viene configurato e testato inizialmente in un ambiente simulato, con un insieme di email di esempio rappresentative, incluse alcune situazioni limite già presentate in quel capitolo.

Passo 2 — Osservazione dei risultati e identificazione di un problema

Applicando gli strumenti del capitolo 6, durante il test emerge un rischio specifico: l'agente, di fronte a una richiesta di preventivo con specifiche tecniche non chiare, tenta comunque di formulare una risposta basata su un'interpretazione autonoma delle specifiche, un comportamento potenzialmente problematico secondo il rischio di interpretazione eccessivamente letterale già presentato in quel capitolo.

Passo 3 — Correzione delle istruzioni

Applicando gli strumenti del capitolo 3, le istruzioni fornite all'agente vengono corrette, specificando esplicitamente che, in caso di specifiche tecniche non chiare, l'agente dovrebbe segnalare la richiesta per un intervento umano diretto, invece di procedere autonomamente con un'interpretazione non verificata.

Passo 4 — Introduzione graduale con supervisione elevata

Applicando gli strumenti del capitolo 7, l'agente viene introdotto in un contesto operativo reale con un livello di supervisione elevato, richiedendo un'approvazione esplicita del responsabile amministrativo prima dell'invio effettivo di ciascuna risposta, secondo la scala di supervisione proporzionata già presentata in quel capitolo.

Passo 5 — Riduzione progressiva della supervisione

Applicando gli strumenti del capitolo 7, dopo un periodo di utilizzo con esiti costantemente positivi, il livello di supervisione viene gradualmente ridotto per le sole richieste più semplici e

standardizzate, mantenendo invece un controllo diretto per le situazioni più complesse o ambigue, secondo il principio di gradualità già presentato in quel capitolo.

Che cosa insegna questo caso

Il caso mostra come un approccio prudente e graduale, che parte dalla sperimentazione in ambiente simulato e introduce l'autonomia dell'agente solo progressivamente, permetta di beneficiare delle potenzialità di questi strumenti più avanzati, riducendo significativamente i rischi specifici già presentati in questa dispensa.

10. Glossario dei termini chiave

Come per le dispense precedenti di questo percorso, riportiamo un glossario di sintesi dei principali termini introdotti.

Agente AI	Sistema capace di pianificare ed eseguire in autonomia una sequenza di azioni per raggiungere un obiettivo assegnato.
Ambiente simulato	Spazio di sperimentazione protetto in cui un agente può operare senza conseguenze reali sui sistemi effettivi.
Azione irreversibile	Azione eseguita da un agente le cui conseguenze non sono facilmente correggibili una volta completata.
Supervisione proporzionata	Principio secondo cui il livello di controllo umano dovrebbe essere calibrato in base al rischio dell'azione intrapresa.

11. Checklist finale e autoverifica

Come per le dispense precedenti, dedichiamo un momento all'autoverifica personale prima di proseguire.

Domande di autoverifica

5. Sapresti definire che cos'è un agente AI, distinguendolo da un assistente conversazionale?
6. Ricordi perché un ambiente simulato è particolarmente prezioso per chi si avvicina per la prima volta a questi strumenti?
7. Sapresti descrivere le fasi tipiche di un test in ambiente simulato?
8. Ricordi almeno due rischi specifici degli agenti autonomi rispetto agli strumenti conversazionali?
9. Sapresti spiegare il principio di supervisione proporzionata al livello di autonomia?
10. Ricordi almeno un'applicazione professionale promettente nel breve periodo?

Come procedere

Se le risposte sono arrivate con sicurezza, si può proseguire con la Dispensa 9, dedicata alla creatività multimediale con l'AI. In caso di incertezza, si consiglia di rileggere i capitoli corrispondenti, in particolare quelli relativi ai rischi specifici degli agenti autonomi e alla supervisione proporzionata.

12. Riferimenti e risorse aggiornate al 2026

Per approfondimenti autonomi, si segnalano alcune categorie di risorse utili:

- Le funzionalità di ambiente simulato offerte dai principali fornitori di strumenti di AI generativa, spesso documentate nelle relative guide ufficiali;
- Pubblicazioni divulgative sugli agenti AI e le loro applicazioni professionali;
- Comunità professionali dedicate alla condivisione di esperienze pratiche con questi strumenti più avanzati;
- La sperimentazione diretta, sempre a partire da un ambiente simulato secondo quanto presentato in questa dispensa, la risorsa più efficace per costruire una comprensione pratica e sicura di questi strumenti.

Chiusura della Dispensa 8

Con questa dispensa prosegue il percorso AI Generativa Applicata, fornendo gli strumenti per comprendere e sperimentare in sicurezza gli agenti AI. La dispensa successiva affronterà la creatività multimediale, un'ulteriore applicazione pratica che completa la panoramica degli strumenti presentati in questo percorso.

Marco Consiglio

Formatore e Course Director — 2026