



Anthropic, il country manager Mattia Gamberoni: «Trasparenti sui rischi senza precedenti dell'AI»•

Descrizione

(Adnkronos) Anthropic annuncia la nomina di Mattia Gamberoni a Country manager per l'Italia, dopo l'annuncio dell'apertura di un ufficio a Milano, che è già operativo. Il manager guiderà il team in Italia supportando imprese, startup e ricercatori. In precedenza Gamberoni ha lavorato in Salesforce e Cisco, ha guidato Stripe nell'Europa meridionale e ha contribuito all'avvio di AWS nella penisola iberica. Mentre gli allarmi sulla possibile apocalisse provocata dall'AI aumentano, il manager all'Adnkronos dice che «siamo sempre stati trasparenti sul fatto che l'AI porterà con sé enormi benefici, ma anche rischi senza precedenti». «Non commentiamo», dice il manager eventuali accessi futuri a Mythos 5, la famiglia più avanzata di modelli di Anthropic a cui è appena stato garantito l'accesso all'Agenzia europea per la cybersecurity.

Quali progetti sono attualmente in corso in Italia e chi sono i vostri clienti più importanti?

Anthropic collabora con importanti imprese italiane nei settori dei servizi finanziari, delle life sciences, dell'energia, dell'automotive e della tecnologia. Generali, Intesa Sanpaolo, Poste Italiane e Unipol sono tra le aziende che stanno estendendo l'utilizzo di Claude Code ai propri sviluppatori, affiancandolo alle competenze ingegneristiche interne per modernizzare il modo in cui sviluppano il software.

Di cosa si occuperà il team ospitato nell'ufficio di Milano? Ci sarà una componente di ricerca?

Guiderà un team focalizzato sul mercato italiano, con competenze trasversali che spaziano dal rapporto con le grandi imprese alle partnership, dal supporto tecnico e a diretto contatto con i clienti allo sviluppo dell'ecosistema. L'obiettivo è combinare una conoscenza approfondita del mercato italiano con l'accesso alle competenze e alle risorse di Anthropic a livello europeo e globale, in ambiti che comprendono prodotto, ricerca, sicurezza e implementazione. Il team italiano non opererà come un'unità autonoma, ma lavorerà in stretto coordinamento con i team di Anthropic nel mondo. Questo modello permette a clienti e partner locali di beneficiare sia di un supporto diretto sul territorio,

sia del collegamento con i team che sviluppano Claude e con il più ampio lavoro dell'azienda su prodotto e sicurezza.

Cosa ha convinto ad accettare un ruolo in Anthropic?

Il mio obiettivo è aiutare le imprese italiane a trasformare il potenziale dell'AI in valore duraturo, attraverso sistemi di AI di cui possano fidarsi. Molte aziende, istituzioni e sviluppatori in Italia stanno già sperimentando e utilizzando tecnologie di frontiera, e io ho trascorso la mia carriera a costruire mercati e a far crescere piattaforme tecnologiche in Europa. È proprio l'attenzione di Anthropic a sicurezza, affidabilità e creazione di valore nel lungo periodo ad avermi convinto ad accettare questo ruolo.

Cosa rappresenta per voi il mercato italiano in termini di numeri e obiettivi? Per Anthropic questa presenza è puramente commerciale o apre anche possibilità strategiche, o persino di M&A?

Vediamo l'Italia come un Paese con un forte potenziale nell'AI, dove imprese, istituzioni e sviluppatori stanno già utilizzando tecnologie di frontiera per sperimentare nuovi modi di lavorare. L'apertura dell'ufficio di Milano e la nomina di un Country Manager si inseriscono nel più ampio percorso di crescita europea di Anthropic e riflettono la nostra volontà di lavorare fianco a fianco con aziende, founder, università e istituzioni pubbliche italiane per promuovere un'adozione responsabile dell'AI. Il nostro obiettivo è costruire una presenza solida e affidabile, capace di supportare nel lungo periodo i clienti e l'intero ecosistema italiano.

A quali clienti contate di rivolgervi in prima battuta? Ci sono già pubbliche amministrazioni o ministeri tra i vostri clienti?

L'approccio di Anthropic in Italia parte dalle esigenze dei clienti, non da un singolo segmento. Lavoriamo con imprese nei settori dei servizi finanziari, delle life sciences, dell'energia, dell'automotive e della tecnologia, che stanno utilizzando Claude principalmente in tre modi. I team utilizzano Claude per fare ricerca, analizzare informazioni, redigere contenuti e scrivere codice. Le aziende stanno integrando Claude nei processi operativi, affidandogli attività ad alto volume e permettendo così ai propri specialisti di concentrarsi sulle decisioni che richiedono esperienza e giudizio umano. Le imprese stanno integrando Claude nei prodotti e nei servizi che offrono ai propri clienti. Oltre al mondo enterprise, stiamo lavorando con università e istituzioni pubbliche sul tema dell'adozione responsabile dell'AI in Italia.

Sono previste partnership con attori come Cdp o iniziative simili per l'ecosistema?

Avere un ufficio a Milano significa essere presenti sul territorio per supportare le imprese italiane, la ricerca e l'ecosistema tecnologico italiano nel percorso verso un'adozione sicura dell'AI. Più di recente abbiamo anche lanciato a Milano il programma Claude Community Ambassador, che riunisce founder e sviluppatori attraverso meetup, workshop e hackathon e crea un canale attraverso cui la community italiana degli sviluppatori può contribuire a plasmare il futuro di Claude.

Gli allarmi sulla sicurezza dell'AI aumentano mentre si avvicinano le grandi IPO. Qual è la sua visione sull'approccio di Anthropic alla sicurezza?

Siamo sempre stati trasparenti sul fatto che l'AI porterà con sé enormi benefici, ma anche rischi senza precedenti. Per affrontarli, continuiamo a sviluppare modelli dotati di alcune delle più solide salvaguardie del settore. Anthropic è stata pioniera nella mechanistic interpretability, un campo di ricerca che mira a comprendere cosa accade all'interno dei modelli di AI e come questi funzionano, e che oggi viene utilizzato per analizzare e prevenire potenziali fenomeni di disallineamento dell'AI. Siamo stati il primo laboratorio di AI a pubblicare una Responsible scaling policy, un framework pubblico dedicato alla mitigazione dei rischi catastrofici associati ai modelli di AI. Continuiamo inoltre a sottoporre i nostri modelli a test rigorosi per individuare capacità potenzialmente pericolose, in ambiti come la cybersecurity e la biologia, e a condividere ciò che impariamo per favorire il confronto e la ricerca nel settore. È anche per questo che riteniamo che il mondo trarrebbe beneficio dall'adozione, da parte dell'intero settore, di modalità legali e verificabili per collaborare alla definizione di criteri condivisi per lo sviluppo e il rilascio di modelli sempre più potenti.

Qual è l'attuale utilizzo in Italia dei vostri modelli più avanzati, parte della famiglia Mythos? Mentre a Enisa è stato appena garantito l'accesso a Mythos 5, avete avuto o avrete colloqui con le istituzioni italiane sul tema sicurezza?

Claude viene utilizzato in Italia in diversi ambiti, dallo sviluppo software al knowledge work, dalla ricerca e analisi alla trasformazione dei flussi di lavoro. Claude Code, in particolare, rappresenta un esempio concreto. Stiamo coinvolgendo l'ecosistema più ampio, comprese università e istituzioni pubbliche, nell'ambito di un confronto più ampio sull'adozione responsabile dell'AI. In questa fase non commentiamo eventuali accessi futuri a Mythos 5.

??

economia

webinfo@adnkronos.com (Web Info)

Categoria

1. Comunicati

Tag

1. Ultimora

Data di creazione

Settembre 14, 2026

Autore

redazione