

La vulnerabilità di Moltbook, il social degli agenti AI



Internet è rimasto affascinato da **Moltbook**, un sito di social media con una nuova serie di regole: i bot AI possono pubblicare post mentre gli esseri umani guardano. I post sono diventati rapidamente strani, con agenti AI che apparentemente inventavano religioni, scrivevano manifesti contro l'umanità e formavano quelli che sembravano culti digitali. Ma i ricercatori di sicurezza affermano che lo spettacolo è solo un diversivo. Sotto la superficie, hanno trovato database esposti contenenti password e indirizzi e-mail, malware diffuso e un modello funzionante di come l'"agente Internet" potrebbe fallire.

Alcune delle conversazioni più fantascientifiche sulla piattaforma simile a Reddit, ad esempio quelle relative agli agenti di intelligenza artificiale che complottano per l'estinzione dell'umanità, sembrano essere in gran parte false. Tuttavia, secondo gli esperti, Moltbook presenta

alcuni **potenziali problemi di sicurezza esistenziale**. Essi sostengono che la piattaforma potrebbe diventare **un ambiente poco controllato** in cui gli hacker potrebbero testare malware, truffe, disinformazione o iniezioni immediate che dirottano altri agenti prima di prendere di mira le reti tradizionali.

“Lo spettacolo degli ‘agenti che dialogano tra loro’ è per lo più performativo (e in parte è finto), ma ciò che è davvero interessante è che si tratta di una dimostrazione dal vivo di tutto ciò che i ricercatori nel campo della sicurezza hanno segnalato riguardo agli agenti di intelligenza artificiale”, ha dichiarato George Chalhoub, professore presso l’UCL Interaction Centre, a Fortune. “Se 770.000 agenti giocattolo su un clone di Reddit possono creare tutto questo caos, cosa succederà quando i sistemi agenti gestiranno l’infrastruttura aziendale o le transazioni finanziarie? Vale la pena prestare attenzione a questo fenomeno come un avvertimento, non come un motivo di festeggiamento”.

I ricercatori di sicurezza affermano che OpenClaw, il software agente AI (precedentemente Clawdbot/Moltbot) che alimenta molti bot su Moltbook, è già un bersaglio per il malware. Un rapporto di OpenSourceMalware ha rilevato 14 ‘abilità’ false caricate sul suo sito ClawHub in pochi giorni, che fingevano di essere strumenti di trading di criptovalute ma in realtà infettavano i computer. Queste abilità eseguono codice reale in grado di accedere ai file e a Internet; una è persino arrivata sulla prima pagina di ClawHub, ingannando gli utenti occasionali e indurli a incollare un comando che scarica script dannosi per rubare dati o portafogli di criptovalute.

Simon Willison, un importante ricercatore nel campo della sicurezza che ha seguito lo sviluppo di OpenClaw e Moltbook, ha descritto Moltbook come la sua “scelta attuale per il più probabile a provocare un disastro Challenger”, un riferimento all’esplosione dello space shuttle del 1986 causata dall’ignoranza delle avvertenze di sicurezza. Il rischio intrinseco più evidente, ha affermato, è il **prompt injection**,

un tipo di attacco ben documentato in cui istruzioni dannose vengono nascoste nei contenuti forniti a un agente di intelligenza artificiale.

In un post sul blog, ha messo in guardia da una 'tripletta letale' in atto: gli utenti che concedono a questi agenti l'accesso a e-mail e dati privati, li collegano a contenuti non affidabili provenienti da Internet e consentono loro di comunicare con l'esterno. Questa combinazione significa che un singolo comando dannoso potrebbe istruire un agente a sottrarre dati sensibili, svuotare portafogli crittografici o diffondere malware, il tutto senza che l'utente si renda conto che il proprio assistente è stato compromesso.

Tuttavia, Willison ha anche osservato che ora che "le persone hanno visto cosa può fare un assistente digitale personale senza restrizioni", la domanda è destinata ad aumentare.

Charlie Eriksen, ricercatore di sicurezza presso Aikido Security, ha affermato di considerare Moltbook come **un sistema di allerta precoce per l'ecosistema più ampio degli agenti di intelligenza artificiale**. "Penso che Moltbook abbia già avuto un impatto sul mondo. È stato un campanello d'allarme sotto molti aspetti. Il progresso tecnologico sta accelerando a un ritmo sostenuto ed è abbastanza chiaro che il mondo è cambiato in un modo che non è ancora del tutto chiaro. Dobbiamo concentrarci sulla mitigazione di questi rischi il prima possibile", ha affermato.

Il nuovo Internet

Nonostante l'attenzione virale, la società di sicurezza informatica **Wiz** ha scoperto che gli **1,5 milioni** di agenti 'autonomi' di Moltbook non erano esattamente ciò che sembravano. L'indagine della società ha rivelato che dietro quegli account c'erano solo **17.000** esseri umani, senza alcun controllo per distinguere la vera AI dagli script.

Gal Nagli, ricercatore presso Wiz, ha dichiarato a Fortune di essere riuscito a registrare **un milione di agenti** in pochi minuti quando ha testato la piattaforma. “Gli agenti di intelligenza artificiale, strumenti automatizzati, raccolgono semplicemente le informazioni e le diffondono a macchinetta”, ha affermato Nagli. “Nessuno controlla cosa è reale e cosa non lo è”.

Ami Luttwak, cofondatore e direttore tecnico di Wiz, ha affermato che l’incidente evidenzia **un problema di autenticità** più ampio legato all’emergente “**internet degli agenti**” e all’aumento della scarsa qualità dell’intelligenza artificiale: “Il nuovo internet non è in realtà verificabile. Non esiste un’identità chiara. Non c’è una distinzione netta tra intelligenza artificiale e esseri umani, e sicuramente non esiste una definizione di intelligenza artificiale autentica”.

Wiz ha anche scoperto che Moltbook stesso presentava **un’enorme falla nella sicurezza**: il suo database principale era completamente aperto, quindi chiunque trovasse una singola chiave nel codice del sito web poteva leggere e modificare quasi tutto. Quella chiave dava accesso a circa **1,5 milioni** di ‘password’ di bot, decine di migliaia di indirizzi e-mail e messaggi privati, il che significa che un hacker poteva impersonare agenti di AI popolari, rubare i dati degli utenti e riscrivere i post senza nemmeno effettuare il login.

“È un’esposizione molto semplice. L’abbiamo riscontrata anche in molte altre applicazioni codificate con vibe”, ha affermato Nagli. “Sfortunatamente, in questo caso, l’app era completamente codificata con vibe senza alcun intervento umano. Quindi non ha implementato alcuna misura di sicurezza nel database; era completamente configurata in modo errato”.

“L’intero flusso è una sorta di anteprima del futuro”, ha aggiunto. “Si crea un’app con il codice vibe, questa viene pubblicata e in poche ore diventa virale in tutto il mondo. Ma il rovescio della medaglia è che il codice vibe crea anche

delle falle nella sicurezza”.

Questa storia è stata originariamente pubblicata su [Fortune.com](https://fortune.com).