



Huawei presenta OceanStor M900 Context Memory Storage per accelerare l'inferenza IA nei data center hyperscale

Descrizione

COMUNICATO STAMPA - CONTENUTO PROMOZIONALE

SHANGHAI, 18 settembre 2026 /PRNewswire/ - In occasione dello HUAWEI CONNECT 2026, David Wang, Vicepresidente del Consiglio di Amministrazione e Presidente a rotazione di Huawei, ha presentato ufficialmente OceanStor M900 Context Memory Storage durante il suo keynote. Progettato per l'inferenza IA nei data center hyperscale, il prodotto fornisce ai SuperPoD uno spazio di memoria completamente condiviso che offre capacità su scala PB e prestazioni di livello TB/s. Ciò segna un cambiamento nell'infrastruttura IA da un modello basato sul calcolo a un'elevata collaborazione tra calcolo, rete e archiviazione. Questo contribuirà a liberare la potenza di calcolo dei SuperPoD.

Il 2026 ha visto l'accelerazione della transizione dell'IA dalle innovazioni tecnologiche all'implementazione su larga scala. Le applicazioni di IA si sono evolute dai chatbot agli agenti in grado di svolgere autonomamente compiti complessi. Questi agenti sono ampiamente adottati nei settori strategici, segnando l'inizio dell'era dell'IA generativa.

Man mano che i modelli di grandi dimensioni raggiungono i 10 mila miliardi di parametri, i SuperPoD stanno diventando la scelta ottimale per le infrastrutture di IA. I modelli di grandi dimensioni tradizionali supportano finestre di contesto superiori a un milione di token, l'inferenza multi-turn e le attività complesse sono diventate la norma e i dati della cache KV generati durante l'inferenza continuano a crescere. Queste tendenze hanno spinto la memoria on-chip e la DRAM oltre i loro limiti in termini di capacità e redditività. È ormai obiettivo comune nel settore costruire un sistema di archiviazione a più livelli che coordini la memoria on-chip, la DRAM e gli SSD per creare uno spazio di memoria completamente condiviso con capacità estesa.

Huawei ha presentato OceanStor M900 Context Memory Storage per superare i rallentamenti nella capacità di memoria nel contesto ultra-lungo e nell'inferenza multi-turn. OceanStor M900 utilizza la rete UnifiedBus per creare una cache KV multilivello globale su scala PB con connessioni one-hop. Ciò libera completamente il potenziale di potenza di calcolo dei SuperPoD e accelera l'inferenza IA

nei data center hyperscale. OceanStor M900 Context Memory Storage ha tre funzionalità chiave:

Superare i limiti di capacità per potenziare l'IA su larga scala con memoria estesa

Alimentata dalla rete di interconnessione UnifiedBus ad alta velocità, la cache KV realizza il pooling globale e la condivisione con archiviazione a più livelli. La cache KV dei SuperPoD viene estesa dalla memoria on-chip e dalla DRAM agli SSD, consentendo a un singolo cluster di fornire 64 PB di capacità. La capacità di cache KV disponibile per NPU viene aggiornata da gigabyte a terabyte, consentendo di memorizzare, condividere e riutilizzare più contesto. Ciò aumenta significativamente il rapporto di hit della cache KV.

Potenziare le prestazioni di inferenza per liberare completamente la potenza di calcolo

OceanStor M900 è la prima architettura del settore ad integrare la CPU, l'unità di controllo di rete e l'unità di controllo NAND. Fornisce semantica KV nativa per consentire una connessione one-hop dall'NPU del SuperPoD agli SSD. Ciò elimina la necessità di conversione del protocollo e di inoltro della CPU, riducendo drasticamente la latenza di accesso da millisecondi a 60 microsecondi, una riduzione del 90%. Un singolo cluster fornisce 40 TB/s di larghezza di banda di accesso aggregata, 1,5 volte superiore rispetto alle soluzioni peer. In tipici scenari di programmazione IA, questa architettura raddoppia il throughput dei token del cluster di inferenza e dimezza il tempo necessario al primo token, convertendo la potenza di calcolo in produttività.

Riduzione dei costi dei token per consentire un'adozione su larga scala ed economicamente efficiente dell'IA

OceanStor M900 utilizza la prima tecnologia di archiviazione adattativa KV-aware del settore, che prevede i cicli di vita della cache KV in base al valore dei dati e distribuisce in modo intelligente i dati tra i supporti di archiviazione. Questa tecnologia consente fino a 24 scritture su disco al giorno (DWPD), estendendo la resistenza degli SSD di 16 volte e garantendo la stabilità per tre anni. Riducendo i costi di sostituzione dei media e di manutenzione e gestione, riduce i costi a lungo termine delle infrastrutture di inferenza IA su larga scala e consente un'adozione più rapida dell'IA.

Man mano che l'IA si espande nei principali sistemi di produzione in tutti i tipi di industrie, l'infrastruttura dell'IA si sta evolvendo da un modello basato sul calcolo verso una collaborazione più stretta tra calcolo, rete e archiviazione. L'archiviazione della memoria di contesto sarà essenziale per migliorare continuamente la capacità e l'efficienza di accesso delle cache KV di inferenza hyperscale.

View original content: <https://www.prnewswire.com/news-releases/huawei-presenta-oceanstor-m900-context-memory-storage-per-accelerare-linferenza-ia-nei-data-center-hyperscale-302883305.html>

Copyright 2026 PR Newswire. All Rights Reserved.

COMUNICATO STAMPA - CONTENUTO PROMOZIONALE: Immediapress - un servizio di diffusione di comunicati stampa in testo originale redatto direttamente dall'ente che lo emette. Adnkronos e Immediapress non sono responsabili per i contenuti dei comunicati trasmessi

-

immediapress/pr-newswire

Categoria

1. Comunicati

Tag

1. ImmediaPress

Data di creazione

Settembre 18, 2026

Autore

redazione

default watermark