

עובד — ומהן הטכנולוגיות המתקדמות ביותר GPU כיצד בשנת AI 2026 לפריצת חומת הזיכרון במאיצי

FLOPS תמונת המצב: הבעיה כבר איננה רק

לא מספיק GPU איננה עוד "ה" AI נכון ל-31 באוגוסט 2026, הדרך המדויקת ביותר לתאר את מגבלת הביצועים של מערכות רוחב פס לזיכרון: **בשרשרת של קירות** מודרניות נתקלות למעשה AI יחיד. מערכות "Memory Wall מהיר", ואפילו לא "יש והספק הנצרך, **(Communication Wall)** תקשורת בין מאיצים, **(Capacity Wall)** HBM קיבולת, **(Bandwidth Wall)** מקומי Training, Prefill, איזה קיר דומיננטי משתנה לפי שלב העבודה. **(Power / Data-Movement Wall)** בהעברת נתונים מגיעים לאותה LLM inference מחקרי מערכות עדכניים על Mixture-of-Experts או Decode, long-context Attention compute, memory bandwidth, memory capacity i-synchronization/ bottleneck יכול לעבור בין compute, memory bandwidth, memory capacity i-synchronization/ bottleneck יכול לעבור בין compute, memory bandwidth, memory capacity i-synchronization/ communication batch size, context length, parallelism להחמרה 1

INT4 ו-FP4, FP8 המעבר ל-. **החישוב נעשה זול וצפוף יותר מהר מהיכולת להזין אותו בנתונים** הבסיסית היא ש- עדיין צריכה KV-cache ו-activation, בצורה דרמטית, אך כל מטריצת משקולות Tensor Core מגדיל את קצב פעולות ה- נעשים מהירים יותר, הבעיה הארכיטקטונית זוהי מ"כמה Tensor Cores להגיע פיזית למקום שבו מתבצע החישוב. לכן ככל ש- אפשר להביא, לאן, באיזו אנרגיה, ובכמה שימוש חוזר?". גם סקירות עדכניות של bytes אפשר לבצע? אל "כמה MACs data movement מראות שחלק גדול מהאנרגיה במערכות חישוב מודרניות מוצא על Processing-in-Memory arithmetic 2

כדי להפריד עובדות מסוגים שונים, אשתמש לאורך המחקר בסיווג הבא:

Measured / Demonstrated — מערכת או link סיליקון, עצמאי ולא גוף עצמאי של היצרן ולא גוף עצמאי של היצרן, אצין זאת.

Research Result — תוצאה ממאמר אקדמי/מערכת מחקרית.

Vendor Claim — performance projection או benchmark, נתון מוצר.

Industry Projection — roadmap, sample תעשייתית.

Speculation — מסקנה ארכיטקטונית שלי על בסיס המידע הזמין.

PFLOPS עד 50 Rubin GPU מפרסמת ל-NVIDIA, לדוגמה. Roofline אחת הדרכים הטובות ביותר לראות את המגמה היא ב- **2,273 FLOP/byte** זה יוצר נקודת איזון תיאורטית של כ- HBM4 bandwidth של 22 TB/s inference NVFP4 של שאינו kernel כלומר **1,875 FLOP/byte** נמצא סביב 8 TB/s dense NVFP4 עם 15 Blackwell Ultra B300, 15 PFLOPS peak compute יכול להתקרב ל-HBM שמגיע מ- byte מצליח לבצע אלפי פעולות על כל 3

סטטוס	Roofline משוער knee	רלוונטי Peak שפורסם	Bandwidth מקומי	Fast memory	מאיץ, נכון ל-2026
Vendor spec / shipping-generation 4	~1,875 FLOP/B	15 PFLOPS dense NVFP4	8 TB/s	288 GB HBM3E	NVIDIA Blackwell Ultra B300

מאיץ, נכון ל-2026	Fast memory	Bandwidth מקומי	רלוונטי שפורסם Peak	Roofline משוער knee	סטטוס
NVIDIA Rubin	288 GB HBM4	22 TB/s	50 PFLOPS NVFP4 inference	~2,273 FLOP/B	Vendor Claim / preliminary 2026 specs ⁵
NVIDIA Rubin, FP8/FP6 training	288 GB HBM4	22 TB/s	17.5 PFLOPS	~795 FLOP/B	Vendor Claim / preliminary ⁶
AMD Instinct MI350/MI355X	288 GB HBM3E	8 TB/s עד	מפרסמת עד AMD 10.1 PFLOPS במצבים הכוללים sparsity	לא השוואה ישירה ל-dense	Vendor spec ⁷
Google TPU8t	216 GB HBM + 128 MB SRAM	6.528 TB/s	12.6 PFLOPS FP4	~1,930 FLOP/B	Vendor spec ⁸
Google TPU8i	288 GB HBM + 384 MB SRAM	8.601 TB/s	10.1 PFLOPS FP4	~1,174 FLOP/B	Vendor spec ⁸

והפעולות הנספרות שונים. היא כן sparsity, utilization בין המאיצים: הפורמטים, הגדרת benchmark ההשוואה איננה לכן **FLOP/byte של מאיץ 2026 מגיע למאות ואף אלפי machine balance**: ממחישה את העיקרון הארכיטקטוני Tensor Cores קטן, רחוקים מאוד מהאזור שבו עוד batch Decode נמוך, במיוחד Arithmetic Intensity בעלי workloads מועילים.

מודרני GPU מה באמת קורה בתוך

פעילים, להסתיר threads שמנסה להחזיק אלפי throughput-oriented הוא מערכת. "CPU עם הרבה cores" אינו GPU קרוב ליחידות החישוב operands ולשמור כמה שיותר multithreading באמצעות latency.

execution נמצאים SM בתוך **Streaming Multiprocessor — SM** היא NVIDIA של GPU יחידת הבנייה המרכזית של מהיר המשמש on-chip וזיכרון Tensor Cores, CUDA cores, scheduler, גדול register file, מסוגים שונים pipelines בכל Tensor Cores ארבעה, SMs מתארת 160 NVIDIA, לדוגמה Blackwell Ultra. L1/shared-memory complex. Tensor Cores. ⁹

Tensor Cores הם matrix רגילות integer operations ל-floating-point scalar/vector במטפלים ב- **CUDA Cores** של מטריצות ומבצעים פעולות tiles הם מקבלים, instruction אחד בכל multiply-add ייעודיים: במקום לבצע engines memory עלו במהירות כה גדולה ביחס ל-AI FLOPS גדולה יותר. זו אחת הסיבות ש-matrix-multiply-accumulate bandwidth. ⁹

threads של 32 warps מאוגדים ל-threads. **SIMT — Single Instruction, Multiple Threads** מודל הביצוע הוא threads branch במקביל. כאשר warp של ה-threads עבור ה-instruction מנסה להריץ warp scheduler הם חלק ממודל SIMT בגודל 32 והמודל ה-warp והניצולת יורדת. הגדרת divergence לבחור מסלולים שונים, נוצרת warp-branch. ¹⁰

execution units מקרובים מאוד ל-, on-chip הם: kernel שממוחזרים בתוך operands המקום המועדף ל- registers יכולים להיות warps פחות, registers דורש יותר thread מוגבל; ככל שכל register file עצום. אבל bandwidth ובעלי SM הפעילים ב- warps מחולקים בין ה- registers עלול לרדת. משאבי occupancy בוזמנית, ולכן resident ¹¹

נשלט במפורש על ידי התוכנה Shared memory L1 Cache. Shared memory נמצאים proximity מתחתיהם מבחינת NVIDIA בארכיטקטורות. cache מנוהל יותר כמו L1 block; באותו threads בין operands ולשיתוף tiling משמש ל- kernel CUDA עצמו בנוי מבנקים, ובמודל Shared memory. מודרניות הם חולקים משאב פיזי מאוחד שחלקו ניתן להקצאה שונה עלולים לסדר גישות שהיו יכולות להתבצע bank conflicts; successive banks; successive 32-bit words הנוכחי במקביל ¹²

reuse ומשרת את המאיץ כולו. הוא חשוב במיוחד משום שהוא המקום האחרון שבו ניתן ללכוד SMs נמצא מחוץ ל- L2 cache ה- on-chip interconnect HBM ל- memory partitions/controllers, interfaces סביבו נמצאים. HBM לפני פנייה ל- register → shared/L1 יותר הבלוקים. הפרטים הפיזיים משתנים מדור לדור, אך היררכיית I/O, caches, SMs שמחבר את ה- kernels היא הבסיס שעליו מתבצעת אופטימיזציית HBM off-chip L2 → ¹³

מסלול קריאה טיפוסי, באופן סכמטי, הוא:

HBM DRAM arrays → HBM PHY / memory controller → L2 → on-chip fabric/NoC → SM L1/ shared memory → registers → CUDA/Tensor Core

ובכתיבה:

Tensor/Core result → registers/shared → L2 → memory controller → HBM

operand אם. ואנרגיה wires, buffers, PHYs, clocks בכל מעבר כזה נצרכים. latency הנקודה החשובה איננה רק HBM כותב אותו ל- kernel משולם פעם אחת. אם HBM ונעשה בו שימוש 100 פעמים, מחיר ההבאה מ- register נשמר ב- יכולים לתת kernel fusion ו- tiling משולם שוב. זו בדיוק הסיבה ש- data movement קורא אותו שוב, מחיר ה- kernel arithmetic. אחד של transistor משמעותי בלי להוסיף אפילו speedup ¹⁴

Arithmetic Intensity - Roofline Model

מוגדר (AI) Arithmetic Intensity:

$$AI = \frac{\text{FLOPs executed}}{\text{bytes transferred from the limiting memory level}}$$

בקיבול Roofline:

$$Performance \leq \min(P_{peak}, AI \times BW)$$

משתמשת באותה NVIDIA. של רמת הזיכרון הנבדקת bandwidth הוא BW ו- compute throughput הוא P_{peak} כאשר ¹⁵ שלה Roofline הגדרה בניתוח

ridge: מנקודת ה- AI אם

$$AI_{ridge} = \frac{P_{peak}}{BW}$$

compute-bound. kernel memory-bound. יכול להפוך ל-.

יכולים להתנהג הפוך לחלוטין Tensor Core המשתמשים בדיוק באותו workloads הבדל זה מסביר מדוע שני

MACs ומשמש לעשרות או מאות SRAM/registers נטען פעם אחת ל- tile של A ב- $C = A \times B$ **גדול GEMM** -
compute-bound גדול הוא לעיתים קרובות GEMM יכול להיות גבוה מאוד ולכן Arithmetic Intensity

אחד token עשוי להיקרא פעם אחת כדי לייצר weight לעומת זאת, כל Decode של **matrix-vector multiply** -
בלבד. אפילו FLOP/byte כ-1 — multiply-add של FLOPs ומייצר בערך שני bytes דורש 2 BF16 מסוג weight, בהפשטה
ברור, FLOP/B של מעל 1,000 machine balance לעומת. metadata/scales לפני FLOP/byte נותן סדר גודל של 4-INT4
Tensor Cores אינו מסוגל "להאכיל" את batch decode מדוע.

weight אם כל GB weights כולל כ-140 BF16 parameters אפשר לראות זאת במספר פשוט. מודל היפותטי של 70
לסרוק אותם בכ-6.36, **בגבול פיזיקלי אידיאלי לחלוטין**, יכול TB/s מאיץ עם 22, batch=1, token נקרא פעם אחת עבור
ms, 35 GB raw weights FP4-ב. tokens/s או 57- ms הגבול הוא 17.5 TB/s לכל היותר. ב-8 tokens/s כלומר ~157 ms
bandwidth של lower bounds הם: benchmarks אלה אינם TB/s ב-22 ms מורידים את גבול הסריקה לכ-1.59
TB/s נתוני ה-22 ו-8 line efficiency, attention, network overhead, dequantization, synchronization, KV-cache לפני
Blackwell Ultra ו-Rubin מגיעים מ-3.

compute. לא פחות משהיא טכנולוגיית memory-wall היא טכנולוגיית quantization גם הסיבה ש-

פוגעים בכמה קירות בזמן Transformers מדוע

נכונה רק בחלק מהמקרים. בפועל, כל שלב שלהם יוצר יחס אחר בין "Transformers הם memory-bound" האמירה
compute, storage ו-communication.

ולכן במאיץ בודד חלק ניכר — QKV projections, MLPs, output projections — גדולים מאוד GEMMs כולל **Training**
activations ו-weights, gradients, optimizer state, אולם ככל שהמודל גדל. compute-bound מהעבודה יכול להיות
בשלב זה הבעיה עוברת בהדרגה. tensor/data/pipeline/expert parallelism ואז נדרשים, HBM עוברים את קיבולת ה-
של issue עדכניים מראים שהמגבלה היא מערכתית ולא רק LLM מחקרי מערכות. collectives ו-interconnect ל-HBM
arithmetic throughput. 1

היטב Tensor Cores גדולים יחסית ולכן מנצלים GEMMs MLP ו-Q/K/V שלם במקביל. ה- prompt מעבד **Prefill**
intermediate של I/O וה- sequence length גדלה ריבועית עם attention matrix, ארוך, לעומת זאת Attention ב-
partial reductions ו-Q/K/V כ- tiled attention נולד בדיוק מכאן: לבצע FlashAttention. יכול להפוך לבעיה לבעיה
על FlashAttention-3 HBM. S × S intermediates במקום לכתוב SRAM/registers on-chip שאינם ככל האפשר ב-
Hopper על TFLOPS/s עד כ-740 FlashAttention-2 של כ-1.5-2 × לעומת speedup דיווח על Hopper
Research Result. במבחני המחקר שלו FP8-ב. 16

גדולים, אך weights ה- sequence אחד חדש לכל token מגיע iteration הוא כמעט התמונה ההפוכה. בכל **Decode**
עלול לקרוא GPU נמוך. המשמעות היא שה- weights של reuse קטן; לכן batch/token במצד ה- matrix dimension
memory bandwidth אחד הגורמים לכך ש- arithmetic כדי לעשות יחסית מעט parameters של GB עשרות או מאות
Llama 3.1 8B: 16 GB BF16 weights ו-1.5 TB/s HBM bandwidth-only ceiling של 94 tokens/s בסדר גודל של
מדובר tokens/s בסדר גודל של 94 bandwidth-only ceiling גוברים HBM ו-1.5 TB/s weights BF16 GB כ-16: 3.1 8B
vender, **בניתוח**, אך העיקרון האריתמטי כללי, 17

KV Cache של b precision, head dimension d , H_{KV} KV heads, L layers עם Transformer יוצר קיר נוסף. עבור **KV Cache** sequence length S :

$$KV\ bytes \approx 2 \times L \times H_{KV} \times d \times b \times S$$

נותנים כ־320 BF16 ו־128 head dimension, KV heads, שמונה layers, לדוגמה היפותטית, 80 Key-Value ה־2 הוא עבור עוד לפני, TiB של 32 היה דורש מעל 1.2 Batch. **יחיד sequence** GiB זה כבר כ־39.1 tokens אב־128. token לכל KiB **micro-optimizations**; הם לא offload "paging, prefix sharing, KV quantization, GQA/MQA, weights. KV cache של bandwidth cost ואת footprint מתארים במפורש את serving מחקרי. capacity management מנגנוני כבעיה מרכזית. ¹⁸

PagedAttention, שאינם חייבים להיות **fixed-size blocks** KV cache מחלק את, 2023 SOSP vLLM שהוצג עם, **contiguous** **overprovisioning** **fragmentation** המטרה היא להפחית. **virtual-memory paging** בדומה ל־**contiguous** **capacity** פיזי כלל; הוא פשוט מונע בזבז **bandwidth** יעיל יותר. זו דוגמה מצוינת לפתרון שאינו מגדיל **sharing** ולאפשר **data movement**. **Research Result**. ¹⁹

אך הוא עלול, token פעילים עבור כל **experts** שרק חלק מה־**arithmetic** מפחית **Mixture-of-Experts (MoE)** **capacity** עדיין צריכים להתגורר איפשרו, כך ש־**expert weights** **expert** להחרף שני קירות אחרים. ראשית, כלל ה־**All-to-All** המתאים ולהחזיר אותם — דפוס **expert** שמחזיק את GPU ל־tokens צריך לשלוח **Expert Parallelism** **node/rack** חוצה **expert placement** כאשר **MoE** מרכזי ב־**bottleneck** הוא **All-to-All** עבודות מ־2025-2026 מראות ש־צריך לקחת בחשבון את ההבדל הגדול **EP/TP placement** **Megatron** עצמה מזהירה בפרסומי **NVIDIA** **boundaries**; **NVLink** **network bandwidth**. ²⁰

לכן, נכון יותר ב־2026 לדבר על:

דרך הטיפול העיקרית	דוגמה אופיינית	מה מוגבל	"קיר"
locality, hierarchy, software/hardware co-design	Tensor עם GPU רעבים Cores לנתונים	הפער הכללי בין data ל־arithmetic supply	Memory Wall
HBM4, quantization, SRAM, CIM	batch-1 LLM decode	ברמת bytes/s מסוימת hierarchy	Bandwidth Wall
HBM, GQA/KV quantization, DDR/CXL/pooling	ארוך KV-cache, trillion-param MoE	אפשר state כמה להחזיק בזיכרון המהיר	Capacity Wall
NVLink/UALink/Ethernet/InfiniBand, topology, optical I/O	All-Reduce, TP, MoE All-to-All	bytes/s + latency בין chips/nodes	Communication Wall
locality, 3D stacking, CIM, photonics, lower precision	דרך TB/s העברת DRAM/SerDes	והספק כולל joules/bit	Power / Data Movement Wall

bit כמה באמת עולה להזיז

משנים **packaging** ו־PHY, voltage, access width, wire length, SRAM size, process node, **anchors** **סדרי גודל והשוואת** את המספרים משמעותית. לכן טבלת האנרגיה הבאה היא **benchmark** ולא.

תנועה	סדר אנרגטי	הערה
Register / very-local SRAM	עד pJ/bit הנמוך ביותר; לרוב תת- array/path לפי low-pJ/bit	הערכות תעשייתיות; locality היתרון המרכזי הוא בתנאים מסוימים, אך pJ/bit סביב 0.1~ SRAM מציבות זיהוי ערך אוניברסלי ²¹
3D very-near memory path	במודל מחקרי pJ/bit ~0.66-0.88 2026	Research Result , לא מוצר, HBM ²²
NVLink-C2C on-package	בדור pJ/bit דיווחה ~1.3 NVIDIA Hopper	Vendor measured/spec anchor. ²³
HBM3/HBM3E	הוא סדר pJ/bit low-single-digit גודל מקובל; סקירה מ-2026 HBM3E ל- pJ/bit מצטטת ~3.4	implementation. המספר תלוי מאוד ב- ²⁴
Optical I/O של Lightmatter M1000	2.3 pJ/bit laser כולל	Vendor measured on EVK/reference platform , 114.6 Tb/s bidirectional. ²⁵
Board/rack electrical link	SerDes/גבוה יותר ככל שהמרחק גדלים retiming	אין co-packaged optics; אחת הסיבות למעבר ל- link. מספר יחיד תקף לכל ²⁶

הופכים להספק אדיר pJ/bit אפילו 2-3. **bits המרחק והמספר הכולל של** ההשוואה החשובה יותר מהמספר המדויק היא במודל bit-transfer עבור אנרגיית ה- kW הוא כ- $1.92 \text{ pJ/bit} \times 8 \text{ bits/byte} \times \text{TB/s} \times 100$ TB/s, לדוגמה, 10^2 tens of TB/s. עצמו נהפך לבעיה תרמית והספקית bandwidth: אינו פתרון חנים ל- 100 TB/s HBM פשוט. זו הסיבה ש"פשוט נגדיל" ²⁷

הטכנולוגיות שמנסות לפרוץ את הקירות

ניתן לחלק את הפתרונות לשלוש משפחות. Memory Wall אין טכנולוגיה יחידה שפותרת את

bytes להעביר יותר — HBM4, 2.5D/3D, wider interfaces, optical I/O.

bytes להעביר פחות — quantization, sparsity, compression, FlashAttention, fusion, tiling, KV optimization.

bytes מלכתחילה להפסיק להעביר את ה- Near-Memory Computing, PIM Compute-in-Memory.

טכנולוגיה	Bandwidth	Capacity	מפחיתה traffic	משנה מקום חישוב	בגרות ב-2026
HBM3E/HBM4	כן, מאוד	כן	לא	לא	Production
HBM4E	כן	כן	לא	לא	Samples / roadmap
CoWoS / 2.5D	width מאפשר גדול	מאפשר יותר stacks/chiplets	מעט, בגלל קצרים links	לא	Production

לכיוון Gb/s stable operation, scalability החברה מתארת 14 HBM4E של samples כבר סיפקה במאי 2026 Samsung
 לכן TB/s מציגות יעד של 4~ Samsung אחרות של roadmap הנוכחי; תצוגות sample-stack וכ- 16 TB/s
 3.6 TB/s **roadmap/vendor target**. 4~ sample, נכון יותר לסווג כ- 3.6 TB/s 30

SK hynix. ציבורי סופי ואמין שמצדיק טבלת נתונים כאילו מדובר במוצר specification, עדיין אין, נכון למועד המחקר **HBM5** ל- HBM4E/HBM5 custom logic/base dies גבוהים יותר ויותר hybrid bonding, stacks מאותתות על packaging וחברות fact, לא **Industry Projection**, בשלב זה הן HBM5 לכן תחזיות מספריות קשיחות ל- era. ³¹

אכלומר ~2.75 TB/s ל-Blackwell Ultra 22 TB/s קופץ מ-8 Rubin. "איננו "פתרון סופי HBM4 הבעיה העמוקה היא שגם אפוא אינו יורד — הוא Roofline knee כ-3.33x. ה-, PFLOPS עולה מ-15 ל-50 NVFP4 inference peak אבל bandwidth; עד"ן נמשך. הנתונים הם memory-ל"arithmetic זו אינדיקציה לכך שהמרצץ בין FLOP/B. אף עולה מ-1,875 לכ-2,273 vendor specs והיחס הוא חישוב מהנתונים הללו ³².

Advanced packaging: nCoWoS 3-D

3D integration pitches עם hybrid bonding משתמש ב-TSMC SoIC. במצאים ממש זה על זה dies: הולך צעד נוסף. Intel Foveros Direct ב-2025 volume production ל-10 nm stacking דיווחה ש-TSMC 3-μm מתחת ל-10-μm pitches ציבורי כולל מעבר מ-9-μm roadmap כאשר direct copper-to-copper hybrid bonding משתמש ב-5 μm. ³⁴

ליחידת vertical wires יותר orders-of-magnitude קטן יותר פירושו bump pitch מחשוב מפני ש-Hybrid bonding למשהו הקרוב יותר דולהפוך את זיכרון ה-3compute מממש מעל או מתחת ל-SRAM/DRAM logic בכך ניתן לשים. **שטח** bandwidth density מבחינת on-chip resource.

למשל — memory subsystem הוא מונח רחב יותר שבו פעולות מסוימות מתבצעות בתוך **Processing-in-Memory (PIM)** אינם יכולים לחשוף external pins שהם bandwidth כדי לנצל — banks או ליד stacked DRAM של logic layer

נשמרים SRAM CIM, weight bits. memory array עצמו אל MAC מכנים את ה- **Compute-in-Memory (CIM/IMC)** קרוב מאוד bit-serial דיגיטלי או arithmetic מבצע **Digital CIM**. משמשים גם לחישוב wordlines/bitlines SRAM בתאי באופן vector-matrix multiplication מבצע array שה- currents/charges/voltages מנצל **Analog CIM**; לתאי ולא רק רעיון mainstream research direction עדות לכך שהתחום הוא CIM, שלם ל- session הקדיש 2025 ISSCC. פיזיקלי תיאורטי. ³⁶

ADC/DAC הוא analog CIM החיסרון ב- operation בכל MAC → SRAM לא צריך לעבור weight: היתרון הפוטנציאלי עצום general- של difficulty overhead, quantization/noise, process variation, calibration, limited precision, purpose programming. SRAM CIM במחיר DRAM, הרבה פחות צפוף מ- area: SRAM גם משלם במחיר 10T/8T CIM cells trade-off עדכניות מדגישות בדיוק IEEE הצפיפות יכולה להיות גרועה עוד יותר. סקירות ³⁷

capacitor-based analog in-memory היא משתמשת ב- EnCharge AI computation עם weights SRAM. EN100 בעיקר 2025 מיועד client/edge 200 TOPS. ומפורסם כ-200 TOPS. Vendor Measurement, שנמדדו בסיליקון לעומת מוצרים מתחרים, אך זו MACs על TOPS/W אהחברה טוענת ל-20 datacenter-scale independent benchmark. ³⁸

בתוך "Performance Memory" משלב Corsair **Digital In-Memory Compute** — הולכת בכיוון אחר d-Matrix compute fabric, Performance Memory 300 עם TB/s local bandwidth, 512 עד לצד GB capacity memory; TB/s aggregate Performance Memory cards 1,200 לשרת שמונה GB bandwidth. d-Matrix חשוב במיוחד. LLM performance עצמה מסמנת את נתוני ה- d-Matrix: **preliminary and subject to change**. ³⁹

Quantization, sparsity ו-compression

tensor כלל, אלא להקטין את ה- memory היא לעיתים לא לשנות את memory wall הדרך הזולה ביותר לפרוץ.

Raw payload:

Format	Bytes/value	BF16 מול traffic הפחתת
BF16 / FP16	2	1×
FP8 / INT8	1	~2×
FP4 / INT4	0.5	~4×

אינם בדיוק פי 4. אבל ההשפעה כפולה savings ולכן, packing overhead, scales, metadata, alignment יש peak performance העובדה שמאיצי 2026 מתכננים את ה- cache/HBM נכנסים ב- weights יותר וגם קטנים יותר weights optimization הפך לחלק מהארכיטקטורה ולא low precision מוצביעה על כך ש- FP4/NVFP4/MXFP4 סביב Rubin, TPU8 הביצועים שלהם bit מתחת ל-8 formats כולם כוללים IMI350 ו- TPU8. ⁴⁰

low-bit LLM weights אחרים הראו כבר ב-2023-24 שניתן להוריד weight-only quantization ומחקרי AWQ **Research Result**. מותאמים לחומרה kernels ועם accuracy תוך שמירה טובה יחסית על representations ⁴¹

working-set יכול להיות גדול יותר מה- cache ה- long-context serving חשבו לא פחות: ב- KV-cache quantization accuracy, אף effective context capacity, יכול כמעט להכפיל/לרבע 4 FP8/INT8 ל- BF16 האחר, ולכן מעבר מ-

לא ייצור compression כפדי שה־ fused kernels מקובלים. מחקרי 2024–26 מפתחים ⁴² overhead גדול יותר מהחיסכון overhead

שלא מאוחסן Zero weight. מנצלים אותה בפועל sformat כאשר החומרה וה- compute מורידה **Sparsity**.
 לדלג גם על hardware שמאפשרת ל- memory bandwidth; structured sparsity עדיין צורך compressed format
 speedup. $\times$ אינו שקול אוטומטית ל- 2^2 sparse FLOPS אלא 2^2 Bandwidth Wall יכולה לפגוע גם ב- storage/loads

הוא חלק ממערכת הזיכרון Software

משנים את מסלול החיים של עושים דבר ארכיטקטוני: הם persistent kernels ו־FlashAttention, tiling, kernel fusion

במקום:

HBM → kernel A → HBM → kernel B → HBM → kernel C

מנסים לבצע

HBM → SRAM/registers → A → B → C → HBM

FlashAttention-3 לֹא-IO-awareness גם pipeline TMA, Tensor Cores ו־warp specialization במבחנים שלו FlashAttention-2 ביחס לֹא-improvement אִי-הואמאמר מדווח על 1.5-2 Hopper. 43

NVIDIA Rubin ממשיכה את אותו כיוון בחומרה: Tensor Memory Accelerator – asynchronous data movement **load, (overlap) לחפוף** אלא MMA הם חלק מהארכיטקטורה, משום שהאתגר כבר אינו רק לבצע movement **compute.** 44

CXL - memory pooling

switching, fabric-attached memory, sharing dynamic capacity; CXL 3.2 Dynamic Capacity Devices עם multi-head memory pooling hosts. למספר **Measured / Demonstrated ecosystem interoperability.** ⁴⁵

class. GB/s הוא 256 Rubin במערכת 6x16 PCIe בועד, TB/s HBM4 מפרסמת 22 Rubin: המספרים מדגישים את ההבדל. גבוה יותר latency ועם bandwidth היטב, מדובר בסדר גודל של עשרות פעמים פחות physical link מנצל את CXL גם אם עבור HBM הוא אינו תחליף ל-; capacity disaggregation; cold KV cache, model spill, shared pools; CXL אידיאלי ל- inner-loop GEMM/decode. Rubin I/O הטכני של NVIDIA מגיעים מהתיעוד 44

Network Wall - הופך ל-Memory Wall

יחיד GPU של 2025-26 הוא שהיחידה שאותה ספקי החומרה אופטימיזו כבר איננה AI המעבר המשמעותי ביותר בארכיטקטורת **multi-rack compute domain** או **rack** היא

במחירה את ההיררכיה בצורה כמעט מושלמת Rubin

- GPU: **22 TB/s HBM4 bandwidth.**
- GPU-to-GPU scale-up: **3.6 TB/s NVLink 6 per GPU.**
- scale-out networking: **0.4 TB/s per GPU** לפי מפרט NVL72.
- rack ברמת 72 GPUs: 1,580 TB/s aggregate HBM bandwidth, 260 TB/s NVLink Switch bandwidth ו-28.8 TB/s scale-out.

בבקירוב, GPU כלומר, פר

$$HBM : NVLink : ScaleOut \approx 22 : 3.6 : 0.4 \approx 55 : 9 : 1$$

byteאחר, ועוד יותר מ־ GPU שנמצא אצל byteהמקומי הוא משאב שונה לחלוטין מ־ HBMשנמצא ב־ Byte. זה נתון קריטי ל־ rack.שנמצא מעבר ל־

Blackwell Ultra 8 היחס דומה: TB/s HBM 1.8 מול TB/s NVLink per GPU. 4

memory hierarchy לפיכך, נהפכת להמשך של datacenter של NoC hierarchy-

Registers → SRAM/L1 → L2 → HBM → local GPU fabric → rack fabric → inter-rack network → pooled/remote memory

ואנרגיה bandwidth, latency יכול לגור", עם כולם "מקומות שבהם, compiler/system architect מנקודת המבט של שונים.

All-Reduce הוא traffic pattern ב־training. Data-parallel workers צריכים לאחד gradients; tensor-parallel layers במספר גדול של S בגודל tensor עבור, bringing all-reduce, forward/backward במהלך collectives דורשים ranks, ללא GPU FLOPS לכן הגדלת reduce-scatter + all-gather בסך הכול לאורך $2S$ bytes מעביר בערך rank כל, fabric bandwidth הגדלת communication ל־GEMM יכולה פשוט להזיז את הזמן מ־fabric bandwidth.

traffic יכול לשלוח accelerator מפורזים, ולכן כל experts – tokens routed: קשה יותר טופולוגית של MoE **All-to-All** expert parallelism, xnnv bottleneck – כל All-to-All inter-node מחקרים מ-2026 מתארים. peers למספר רב של 46 כדי להימנע ממנו node בתוך expert communication ועבודות אחרות מציעות להחזיק

NVLink ומקדמת את NVL72 generation 200GB-ל- aggregate all-to-all bandwidth מפרסמת NVIDIA 130 TB/s. scale-up domain. לשם כך, הפוך fabrics per-GPU NVLink 3.6 TB/s היא מכפילה את Rubin ב-
Vendor specs. אבל הכיוון הארכיטקטוני ברור. ⁴⁷

היא "Network Wall? הופך ל- Memory Wall לכן התשובה לשאלה "האם

זו החוצה בהיררכיה wall לא באופן אוניברסלי – אלא: אבל מבוססת היטב, Speculation,

- ־single-GPU batch-1 decode הוא HBM wall.
- ־8-GPU tensor parallelism הוא עשוי להיות NVLink wall.
- ־rack-scale dense training הוא collective wall.
- ־large MoE הוא עשוי להיות All-to-All / topology wall.
- ־multi-rack model with remote KV הוא network + capacity wall.

הגרסה הבין־שבבית שלו הוא Memory Wall; אינו מחליף את Network Wall כלומר, ה־

Lightmatter M1000 4,000 mm² photonic interposer קיבלה תאוצה עצומה silicon photonics בדיוק הסיבה ש־
laser כולל 114.6 Tb/s bidirectional optical bandwidth, 256 fibers 2.3 pJ/bit ו־
Measured/Demonstrated by vendor של החברה. זה rack-scale validation environment deployment hyperscale בלתי תלוי 48

UCIe עם TeraPHY optical I/O chiplet Ayar Labs wafer sort, EVT, DVT עברה 4 Tb/s גרסת 8 Tb/s
multi-day UCIe tests, link testing, engineering validation עברה 8 Tb/s גרסת 8
GPU mass deployment מיליוני GPUs בתוך מיליוני PowerPoint, 49

Photonic Fabric מפתחת, של \$3.5B consideration ב־2 בפברואר 2026 Marvell שנרכשה על ידי Celestial AI,
all-optical scale-up מציינת שהיישום הראשון מיועד ל־ Marvell package/system/rack; ברמת scale-up interconnect
"network transceiver" פחות כ optics רואה industry זוהי אינדיקציה חזקה לכך שה־
accelerator fabric. 50

Photonic matrix multiplication תאורטית, נמצא בשלב פחות בוגר משמעותית, Optical compute
weights/storage, ADC/DAC, non-linear functions, precision, linear algebra לבצע
silicon-photonics והחיבור לעולם הדיגיטלי נשארים מכשולים. סקירת laser efficiency
photonic computing בעוד mainstream commercialization הם כבר communications ממראה ש־
optical I/O בהרבה מ־ optical Tensor Core. 51

GPU + HBM החברות שמנסות לשבור את מודל

להלן 13 שחקנים, כולל הסטארטאפים שביקשת ושחקנים חדשים שהפכו רלוונטיים במיוחד ב־2025-26

חברה	ארכיטקטורה וזיכרון	כיצד הנתונים זזים / מה נשבר	מה הוכח בפועל	Vendor מה עדיין Claim	האם יכולה להתחרות GPU+HBM?
Cerebras	Wafer-Scale Engine; מאות cores אלפי 44 GB SRAM on-wafer WSE-3 ב־ family, fabric wafer; כל ה־ external memory systems למודלים גדולים.	מבטלת כמעט chip-to-chip traffic בתוך wafer ומספקת PB/s-class SRAM bandwidth.	הוא WSE-3 silicon deployed; ב־2026 החברה הציגה עם CS-4 WSE-3 שלושה ומערכת Turbo Nexus. 52	עד 30 inference speed 10 ⁴ throughput/W הם GPU מול Vendor Claims. 53	Speculation: כן, בנישות גדולות inference/ training עם dataflow; 44 GB SRAM אינה מספיקה למודלי frontier ולכן external model storage/ streaming חיוני.

חברה	ארכיטקטורה וזיכרון	כיצד הנתונים זזים / מה נשבר	מה הוכח בפועל	Vendor מה עדיין Claim	האם יכולה להתחרות GPU+HBM?
Groq	Deterministic LPU/TSP-style architecture, software-scheduled execution large on-chip SRAM.	weights/ activations נשמרים קרוב לחישוב; קובע compiler movement מראש במקום dynamic caches/ schedulers.	LPU deployed GroqCloud; החברה מפרסמת >80 TB/s on-chip SRAM BW. בדצמבר 2025 חתמה Groq על NVIDIA non-exclusive technology license ונשארה חברה עצמאית. ⁵⁴	performance superiority claims "10x" vendor הם claims.	Speculation: חזק low-latency inference; SRAM מחייבת scale-out למודלים גדולים ולכן interconnect נהפך לבעיה.
SambaNova	Reconfigurable Dataflow Unit. SN50: 432 MB SRAM + 64 GB HBM2E עד 512 GB DDR5 לפי chip- datasheet/site.	ממפה compiler graph dataflow; hot-SRAM operators, HBM weights/KV, large-DDR cache/model capacity. ⁵⁵	SN40 generation deployed; הוכרו SN50 בפברואר 2026 והחברה צינה shipping במחצית השנייה של 2026. ⁵⁶	performance/TCO מול B200 עד support-10T הם parameters Vendor Claims.	Speculation: מהגישות המאוזנות ביותר, כי אינה מנסה DRAM להחליף; בלבד SRAM- programmability/compiler quality הם קריטיים.
Tenstorrent	Tensix tiled architecture; local SRAM tile, NoC, RISC-V control cores ו-GDDR6 off-chip. Blackhole cards מגיעים עד 32 GB GDDR6.	explicit DMA בין NoC דרך local SRAM, DRAM tiles Ethernet; חשוף software יחסית למיקום הנתונים. ⁵⁷	Blackhole cards מוצעים בפועל ומפורסמים in-stock; Wormhole/Blackhole software פתוח stack. ⁵⁸	"infinite scalability" performance/cost הם superiority vendor positioning.	Speculation: יתרון משמעותי openness/cost; HBM-class ללא קשה כרגע BW high-end לנצח GPUs large-dense decode.

חברה	ארכיטקטורה וזיכרון	כיצד הנתונים זזים / מה נשבר	מה הוכח בפועל	Vendor מה עדיין Claim	האם יכולה להתחרות GPU+HBM?
d-Matrix	Digital In-Memory Compute (DIMC); compute משולב SRAM-like "Performance Memory", עם capacity memory חיצוני.	weight MAC נעשה ליד/בתוך storage; dual Corsair card: 4 GB Performance Memory @300 TB/s + 512 GB capacity memory. ³⁹	Corsair silicon/ platform קיים; החברה דיווחה D-גם על 3 memory silicon operational במעבדה.	LLM כמעט כל speedup/TCO באתר numbers מוגדרים preliminary; אלמשל 10 interactive speed הוא Vendor Claim. ³⁹	Speculation: אחת האלטרנטיבות הישירות והמעניינות HBM-ל- centric inference, Dבמיוחד אם 3 DRAM+DIMC Training/ general-purpose עדיון פחות מוכח.
Etched	ב-2026 החברה rack-scale inference architecture עם HBM/SRAM hybrid , Cluster Scale Memory Low Voltage Inference; פחות דגש פומבי כיום על המיתוג Sohu המקורי.	proprietary low-latency fabric יוצר shared memory pool chips; SRAM בין latency, HBM capacity.	החברה אומרת N4P A0 ש- silicon חזר TSMC מ- customer ו- rack validation ⁵⁹ מתבצע.	80%+ peak FLOP utilization, SOTA latency/ throughput \$1B demand Vendor Claims ללא public independent benchmark מפורט.	Speculation: גבוה משום upside design rack+memory מאוד אגרסיבי; גם גבוה uncertainty במיוחד עד הופעת benchmarks ומפרט מלא.
Lightmatter	כיום בעיקר photonic interconnect , לא replacement Tensor processor: Passage photonic interposer + optical I/O.	מכל I/O מוציאה שטח interposer במקום רק משפת die; 114.6 Tb/s M1000.	M1000 EVK קיים ופועל validation data center; כולל 2.3 pJ/bit laser. ⁶⁰	million-XPU scale and hyperscale הם future/vendor claims.	לא תחליף אלא complement: יכולה להיות אחת הטכנולוגיות שמאפשרות GPU/ASIC racks להמשיך לגדול.

חברה	ארכיטקטורה וזיכרון	כיצד הנתונים זזים / מה נשבר	מה הוכח בפועל	Vendor מה עדיין Claim	האם יכולה להתחרות GPU+HBM?
Celestial AI / Marvell	Photonic Fabric -optical scale-up, package-to-package/system/rack.	מקרב את optics energy/reach של remote accelerator communication -scale-up fabric.	Marvell השלימה את הרכישה בפברואר 2026; technology/platform קיים ויש hyperscaler engagements לפי Marvell. 61	commercial large-scale deployment projection. עדיין	Complement: אם תצליח, תשנה ולא Network Wall את Tensor Core עצמו.
Ayar Labs	TeraPHY UCIE optical I/O chiplet + external SuperNova laser.	electrical UCIE chiplet, WDM optics traffic מוציא במטרים עד קילומטרים.	4 Tb/s generation EVT/system integration; 8 Tb/s version עברה extensive engineering/link validation. 49	economics high-volume GPU integration טרם הוכחו בפומבי.	Complement, פוטנציאל גבוה multi-rack scale-up.
EnCharge AI	Capacitor-based analog SRAM Compute-in-Memory.	נשאר weight SRAM array משתמש MAC-charge-domain מפחית, memory↔MAC traffic.	silicon מחקרי underpinning קיימים; EN100 הוכרו עם 200 TOPS client/edge. 38	20× efficiency הם TCO 10× vendor comparisons.	Speculation: מבטיח מאוד edge/inference; datacenter frontier models capacity, precision software scaling לא הודגמו.

חברה	ארכיטקטורה וזיכרון	כיצד הנתונים זזים / מה נשבר	מה הוכח בפועל	Vendor מה עדיין Claim	האם יכולה להתחרות GPU+HBM?
FuriosaAI	Tensor Contraction Processor (TCP); 256 MB on-chip SRAM + 48 GB HBM3, 1.5 TB/s , 256 TFLOPS BF16/512 TFLOPS FP8, 150W spec.	compiler/dataflow מנסה לשמור intermediates ב-SRAM ולהשתמש ב-HBM ל-weights. ⁶²	הוא RNGD ממשי silicon; Furiosa gpt-oss-120B על RNGD שני ms/ ב-5.8 output-token במדידת החברה. ⁶³	הוא benchmark vendor demo ולא standardized independent test.	Speculation: מתחרה אמיתי efficient inference, אך ecosystem וקנה מידה קטנים בהרבה מ-CUDA/NVIDIA.
Positron	Transformer-focused Archer accelerators; כולל Atlas 8 chips × 32 GB HBM = 256 GB HBM .	specialization מפחיתה general-GPU overhead; מערכת משלבת HBM עם host DDR.	מוגדר Atlas באתר "Shipping Today" ב-2026. ⁶⁴	>4× perf/W lower latency מול Hopper הם vendor benchmarks.	Speculation: יכול inference לנגוס ב- אבל עדיין לא cost, הוצגה architecture bandwidth/ בעלת capacity discontinuity ברמת CIM.
Qualcomm HBC	Near-Memory Compute / High Bandwidth Compute סביב LPDDR stack; AI250/AI300 roadmap.	computation data מתקרב ל- במקום להזרים central ל- tensor core. מפורסם AI250 GB/card עם 768 TB/s ו-133 effective bandwidth. ⁶⁵	HBC/AI250 הוצגו ב-2026 roadmap ויש רשמית.	18× effective bandwidth מול AI200 6× bandwidth/W הם HBM מול vendor claims; raw link BW אין HBM שקול ל- לצורך comparison. ⁶⁶	Speculation: אחת ההתפתחויות החשובות של 2026 effective-bandwidth יתורגמו ל- real LLM benchmarks, היא עשויה להוכיח LPDDR+near-memory יכול compute HBM לעקוף decode.

נקנתה על ידי NVIDIA; Celestial AI חתמה על רישוי טכנולוגיה ל- Groq: מעניין לראות שגם השוק עצמו התחיל להתכנס rack-scale system architecture ל- NVLink עצמה הופכת NVIDIA; UCIe מתיישרים עם Marvell; optical chiplets להיספג אל תוך הארכיטקטורה הטכנולוגיות החלופיות מתחילות. "alternative world" מול "GPU world" כלומר אין כרגע המרכזית. ⁶⁷

Compute-centric מול Memory-centric מול Dataflow-centric

אפשר לזקק כמעט את כל המאבק הארכיטקטוני לשלוש פילוסופיות

מאפיין	Compute-centric	Memory-centric	Dataflow-centric
דוגמאות	NVIDIA/AMD GPU, TPU בחקו	d-Matrix, Qualcomm HBC, EnCharge, PIM/ CIM	Cerebras, SambaNova, Groq, Tenstorrent
הנחת יסוד	הוא compute array memory; המרכז מזינה אותו hierarchy	הוא המרכז; צריך data אליו compute להביא	הם graph/data movement operations המרכז; מקמים spatially
Compute density	גבוהה מאוד	יכולה להיות גבוהה, אך memory cell/ADC area תופסים	spatial utilization בינונית-גבוהה; תלויה ב-
Raw HBM capacity	מצוינת	משתנה; SRAM CIM NMC+DRAM, נמוכה גבוהה	SRAM-only חלישה; tiered designs חזקים יותר
Local bandwidth	HBM גבוהה אך מוגבלת ב-	פוטנציאלית עצומה	on-chip עצומה
Latency	large batch, פחות, large batch streaming ב-	נמצא data מצוינת אם local array ב-	ממופה graph מצוינת כאשר היטב
Energy	יעיל arithmetic; movement יקר	פוטנציאלית הטובה ביותר	locality טובה מאוד בזכות
Programmability	בעיקר, הטובה ביותר CUDA ecosystem	הקשה ביותר, במיוחד analog/PIM	compiler-dependent; פחות flexible GPU
Scalability	rack/multi-rack מוכחת עד	עדיין לא מוכחת באותו קנה מידה	fabric/תלויה מאוד ב- topology
Yield	chiplets/known-good- die משפרים	3D/CIM process מוסיפים complexity	defect מגדיל wafer-scale exposure; redundancy יכולה לפצות
Precision flexibility	גבוהה	digital CIM בינונית; analog יותר נמוכה	design בינונית-גבוהה לפי
Cost	HBM + advanced package יקרים	אך דורש HBM עשוי לחסוך חדש integration	אך memory traffic יכול לחסוך silicon/fabric area גבוה
Best fit 2026 ב-	training + general AI + large-batch inference	memory-bound inference	deterministic / streaming / spatial inference

של Transformer של 2026 אינו flexibility. ניצח עד היום מסיבה שאסור לזלזל בה **Compute-centric** משתנים kernels, speculative decoding, משתנה MoE routing, משתנים attention variants, משתנים Formats. 2022.

3D stacked SRAM – logic-on-memory

Research Result → likely Industry Direction: SRAM layer יורד למיקרונים בודדים hybrid-bond pitch ככל ש- compute 3 מעל PJ/kq צרי-טווח עם links 0.88–0.66 מ-2026 מדווחים על D systems או מתחתיו נעשה מעניין במיוחד. מחקרי 3 מעל bit/datacenter, אך edii תוצאות מחקריות ולא מוצר. HBM-style interfaces גבוהה בהרבה מ- bandwidth-density ועל bit, הוא כיוון סביר מאוד “L3 cache on top of tensor die” הן מראות מדוע ⁷³

Speculation: SRAM/embedded cache של GB עד כמה MB שבהם מאות accelerators בשנים 2028–31 נצפה ליותר HBM בעוד, D נמצאים ב-3 area/leakage/cost, בגלל DRAM לא תחליף את SRAM. capacity tier. HBM נשאר כ- HBM לשנות משמעותית כמה פעמים צריך להגיע ל-.

Wafer-scale computing

Measured today: research. הוא בר-ייצור ובהפעלה, לא רק wafer-scale מוכיחה ש- Cerebras ⁵²

Speculation: packaged GPUs של mainstream replacement לא יהפוך בהכרח ל- wafer-scale — אבל הרעיון הבסיסי, Multi-reticle interposers, package/network boundaries — ולמנוע scale-up domain להגדיל את ה- הם למעשה וריאציות על אותו רעיון optical fabrics ו- giant packages, panel-scale interconnects

Compute-in-Memory

Speculation: Digital CIM לעומת full analog CIM, יותר מ- datacenter צפוי להגיע ל- Digital CIM deterministic משום שהוא שומר על, Analog CIM, כמו EnCharge, היא הדוגמה הבולטת לכך ב-2026 d-Matrix. פשוט יותר toolchain ועל digital precision BF16 training חשוב יותר מ- power budget שם, fixed/quantized inference, edge/client קודם להצליח ב- flexibility. ⁷⁴

logic צפוף מעל memory die: יתמזגו CIM ו- DRAM/SRAM ותהיה אם 3 datacenter הפריצה שתוכל לשנות באמת את מחקרי. boundary מפסיק להיות external HBM bus שכך שה, thousands/millions of vertical connections, עם die, 2026 monolithic-3D PIM כאלה **Research Result / early-stage.** ⁷⁵ כבר חוקרים בדיוק מבנים כאלה

Optical compute לעומת optical interconnect

Industry Projection: 2027–31. הסיבה אינה רק scale-up mainstream הוא מועמד חזק להפוך Optical I/O במהלך Lightmatter, Ayar, SerDes rates ככל שה- efficiency/reach מאבד copper bandwidth; rack dimensions גדלים ⁷⁶ chiplet/interposer/rack-validation stages. device demos עברו מ- Marvell/Celestial ו-

Speculation: specialized באותו חלון זמן. הוא עשוי להופיע כ- GPU/Tensor Cores עצמו לא יחליף optical compute ממשיכים להשאיר את precision ו- nonlinearity, memory, conversion, אבל, accelerator או linear operator electronics של optical communication של maturity תומך בהבחנה בין silicon photonics קריטית. המחקר הרחב על computing. ⁷⁷

יהפוך ליחידת התכנון Rack-scale

וז. לדעתי התחזית בעלת הביטחון הגבוה ביותר

NVIDIA Rubin NVL72 כבר מפרסמת GPUs לא כאוסף של **20.7 TB HBM4 capacity, 1,580 TB/s aggregate HBM bandwidth** ו- **260 TB/s NVLink Switch bandwidth.** Google מתכננת TPU pods עם זיכרון

collectives rack-scale; Cerebras CS-4 הוא rack-scale; Qualcomm Dragonfly משווקת rack; Kyrin מהארכיטקטורה collectives; Etched SambaNova fabric/memory מתכננות rack. ⁷⁸

יצטרך לחשוב על Architect **rack architecture** של 2030 תהיה במידה רבה "chip architecture" לכן

FLOPs + bytes/token + HBM placement + KV placement + expert placement + collective topology + J/l

transistor count ולא רק על MAC array size.

בשנת 2026 AI מהו צוואר הבקבוק האמיתי של מחשוב?

היכולת להציב את הנתון הנכון ליד יחידת החישוב הנכונה בזמן — **Data Supply** צוואר הבקבוק האמיתי הוא: **המסקנה** joules ובמינימום bytes הנכון, במינימום.

זה אינו רכיב יחיד.

Tensor FLOPS הבקבוק עדיין יכול להיות, Training של large GEMM.

interactive Decode, כלל, HBM bandwidth / weight streaming הוא בדרך כלל.

long-context serving, הוא KV-cache bandwidth + capacity.

MoE, expert-memory capacity + All-to-All network הוא לעיתים.

large-scale training, הוא יכול להיות, collective communication.

bits של הזזת כל אותם power budget הופך יותר ויותר ל-, datacenter, וברמת ⁷⁹

לכן הייתי מדרג את הבעיה ב-2026 כך:

עבור חלק הולך, **Data movement > local memory bandwidth/capacity > interconnect > raw arithmetic**, inference workloads וגדל של

מסוגל Rubin אבל העובדה שמאיץ. compute-bound גדולים עדיין יכולים להיות training GEMMs; הדירוג אינו אוניברסלי FLOPs היא אינדיקציה חזקה לכך שעוד HBM bandwidth של byte/s עבור כל NVFP4 FLOPs תאורטית לבצע ~2,273 FLOPs. לבדם כבר אינם התשובה ⁶

או שלב ביניים, AI הוא הארכיטקטורה ארוכת הטווח של GPU + HBM האם?

במובנו של 2020-24 הוא כמעט בוודאות שלב "GPU + HBM" לא עומד להיעלם, אך GPU + HBM: התשובה שלי ביניים.

יישארו עמוד HBM ו-GPU-class programmable compute בטווח של 2-5 שנים: **Speculation — confidence high** frontier training, השדרה של programmability, mature software, high precision, HBM capacity, manufacturing scale ו-multi-rack ecosystem.

בלבד "GPU + HBM" עצמו הולך להשתנות עד שקשה יהיה לקרוא לארכיטקטורה העתידית GPU אבל ה-

היא תהיה יותר קרובה ל:

Tensor compute + massive/3D SRAM + HBM4/HBMx + near-memory engines + chiplets + optical/electrical sc

Rubin משלב מאות MB TPU8 NVLink. dual reticle-scale dies, NV-HBI, HBM4, TMA 3.6 TB/s per-GPU. compute מזיזה HBC Qualcomm HBM/DDR. SambaNova SRAM/HBM/collective acceleration. memory neighborhood. d-Matrix MAC מזיזה Lightmatter/Ayar/Marvell. package/rack boundary אל optical domain. מהם תוקף. **bit מרחק אחר במסלול של אותו כל אחד** 80

אלא "GPU או PIM?" לכן השאלה הארכיטקטונית של סוף העשור כנראה לא תהיה:

state כמה עובר אל הזיכרון, כמה programmable tensor compute כמה מן החישוב נשאר rack ואיזה חלק ממערכת הזיכרון הופך למעשה לרשת אופטית של SRAM/HBM דנשמר ב-3 שלם?

אינו "קיר" שבסופו של דבר נפרוץ ונשכח ממנו. הוא כוח שמכריח את המחשב להשתנות. תחילה הוא הביא Memory Wall ה- optical scale- caches, CIM, dataflow, quantization D stacking, הוא מכישו הוא מביא 3 packaging; HBM 2.5 ID אחר כך, MAC ככל ש- up. **המיקום של הנתונים הופך בהדרגה לדבר היקר ביותר בארכיטקטורה**, עצמו נעשה זול יותר MAC ככל ש- 81

1 79 Efficient LLM Inference: Bandwidth, Compute ...

https://arxiv.org/html/2507.14397v1?utm_source=chatgpt.com

2 A Modern Primer on Processing-In-Memory

https://arxiv.org/html/2012.03112v5?utm_source=chatgpt.com

3 5 6 40 78 Rack-Scale Agentic AI Supercomputer | NVIDIA Vera Rubin NVL72

<https://www.nvidia.com/en-eu/data-center/vera-rubin-nvl72/>

4 9 32 Inside NVIDIA Blackwell Ultra: The Chip Powering the AI Factory Era | NVIDIA Technical Blog

<https://developer.nvidia.com/blog/inside-nvidia-blackwell-ultra-the-chip-powering-the-ai-factory-era/>

7 AMD Instinct™ MI350 Series GPUs

<https://www.amd.com/en/products/accelerators/instinct/mi350.html>

8 TPU 8t and TPU 8i technical deep dive | Google Cloud Blog

<https://cloud.google.com/blog/products/compute/tpu-8t-and-tpu-8i-technical-deep-dive>

10 1.2. Programming Model — CUDA Programming Guide

https://docs.nvidia.com/cuda/cuda-programming-guide/01-introduction/programming-model.html?utm_source=chatgpt.com

11 3.2. Advanced Kernel Programming — CUDA ...

https://docs.nvidia.com/cuda/cuda-programming-guide/03-advanced/advanced-kernel-programming.html?utm_source=chatgpt.com

12 2.3. Writing SIMT Kernels — CUDA Programming Guide

https://docs.nvidia.com/cuda/cuda-programming-guide/02-basics/writing-cuda-kernels.html?utm_source=chatgpt.com

13 CUDA C++ Programming Guide (Legacy)

https://docs.nvidia.com/cuda/cuda-c-programming-guide/?utm_source=chatgpt.com

14 81 Solving computing's memory problem

https://arxiv.org/html/2505.00458v2?utm_source=chatgpt.com

15 Co-Designing AI Model Attention for Fast, Interactive Long- ...

https://developer.nvidia.com/blog/co-designing-ai-model-attention-for-fast-interactive-long-context-inference/?utm_source=chatgpt.com

- 16 FlashAttention-3: Fast and Accurate Attention with Asynchrony and Low-precision
https://arxiv.org/abs/2407.08608?utm_source=chatgpt.com
- 17 Bringing HBM3 to RRGD: Challenges and Benefits
https://furiousa.ai/blog/bringing-hbm3-to-rrgd-key-challenges-and-benefits?utm_source=chatgpt.com
- 18 Q-Hitter: A Better Token Oracle for Efficient LLM Inference ...
https://proceedings.mlsys.org/paper_files/paper/2024/hash/bbb7506579431a85861a05fff048d3e1-Abstract-Conference.html?utm_source=chatgpt.com
- 19 Efficient Memory Management for Large Language Model ...
https://dl.acm.org/doi/10.1145/3600006.3613165?utm_source=chatgpt.com
- 20 46 DisagMoE: Computation-Communication overlapped MoE ...
https://arxiv.org/html/2605.11005?utm_source=chatgpt.com
- 21 The Industry is Striving for Custom Memory (Part 2)
https://www.eetimes.com/beyond-bandwidth-the-industry-is-striving-for-custom-memory-part-2/?utm_source=chatgpt.com
- 22 73 A Full-Stack Performance Evaluation Infrastructure for 3D- ...
https://arxiv.org/html/2604.08044v1?utm_source=chatgpt.com
- 23 NVIDIA Grace Hopper Superchip Architecture In-Depth
https://developer.nvidia.com/blog/nvidia-grace-hopper-superchip-architecture-in-depth/?utm_source=chatgpt.com
- 24 High-Bandwidth and Energy-Efficient Memory Interfaces for ...
https://www.researchgate.net/publication/383971948_High-Bandwidth_and_Energy-Efficient_Memory_Interfaces_for_the_Data-Centric_Era_Recent_Advances_Design_Challenges_and_Future_Prospects?utm_source=chatgpt.com
- 25 48 60 76 Passage M1000 EVK - 3D Photonic Interconnect for AI ...
https://lightmatter.co/products/passage-m1000-evk/?utm_source=chatgpt.com
- 26 51 77 Roadmapping the next generation of silicon photonics
https://www.nature.com/articles/s41467-024-44750-0?utm_source=chatgpt.com
- 27 The Memory Wall Has Three Exits - by Anni Sen
https://senanni.substack.com/p/the-memory-wall-has-three-exits?utm_source=chatgpt.com
- 28 Nvidia supplier SK Hynix begins mass production of next generation memory chip
https://www.reuters.com/technology/nvidia-supplier-sk-hynix-begins-mass-production-next-generation-memory-chip-2024-03-19/?utm_source=chatgpt.com
- 29 71 Micron in High-Volume Production of HBM4 Designed for ...
https://investors.micron.com/news/press-release/2026/Micron-in-High-Volume-Production-of-HBM4-Designed-for-NVIDIA-Vera-Rubin-PCIe-Gen6-SSD-and-SOCAMM2-03-16-2026/default.aspx?utm_source=chatgpt.com
- 30 Samsung Electronics Begins Shipment of Industry-First ...
https://news.samsung.com/global/samsung-electronics-begins-shipment-of-industry-first-hbm4e-samples?utm_source=chatgpt.com
- 31 From HBM to eSSD: How SK hynix is charting the future of ...
https://news.skhyunx.com/en/hbm-to-essd/?utm_source=chatgpt.com
- 33 CoWoS® - Taiwan Semiconductor Manufacturing ...
https://3dfabric.tsmc.com/english/dedicatedFoundry/technology/cowos.htm?utm_source=chatgpt.com

34 TSMC-SoIC® - Taiwan Semiconductor Manufacturing ...

https://3dfabric.tsmc.com/english/dedicatedFoundry/technology/SoIC.htm?utm_source=chatgpt.com

35 Data Center AI Inference Accelerators | Dragonfly

https://www.qualcomm.com/data-center/expertise/ai-accelerators?utm_source=chatgpt.com

36 Session 14 Overview: Compute-in-Memory

https://ieeexplore.ieee.org/iel8/10904417/10904496/10904651.pdf?utm_source=chatgpt.com

37 Special Topic on Energy-Efficient In-/Near-Memory ...

https://ieeexplore.ieee.org/iel8/6570653/10879154/11308150.pdf?utm_source=chatgpt.com

38 News

https://www.enchargeai.com/category/news?utm_source=chatgpt.com

39 68 74 d-Matrix Corsair AI Platform | In-Memory Computing for AI

https://www.d-matrix.ai/product/?utm_source=chatgpt.com

41 AWQ: Activation-aware Weight Quantization for On-Device ...

https://proceedings.mlsys.org/paper_files/paper/2024/file/42a452cbafa9dd64e9ba4aa95cc1ef21-Paper-Conference.pdf?utm_campaign=news&utm_medium=miragenews&utm_source=chatgpt.com

42 Low-Bit Quantization for Efficient and Accurate LLM Serving

https://proceedings.mlsys.org/paper_files/paper/2024/file/5edb57c05c81d04beb716ef1d542fe9e-Paper-Conference.pdf?utm_source=chatgpt.com

43 FlashAttention-3: fast and accurate attention with ...

https://dl.acm.org/doi/10.5555/3737916.3740109?utm_source=chatgpt.com

44 70 80 Inside NVIDIA Rubin GPU Architecture: Powering the Era of Agentic AI | NVIDIA Technical Blog

<https://developer.nvidia.com/blog/inside-nvidia-rubin-gpu-architecture-powering-the-era-of-agentic-ai/>

45 An Overview of the CXL 3.X Specification

https://computeexpresslink.org/webinars/an-overview-of-the-cxl-3-x-specification-3942/?utm_source=chatgpt.com

47 How NVIDIA GB200 NVL72 and NVIDIA Dynamo Boost ...

https://developer.nvidia.com/blog/how-nvidia-gb200-nvl72-and-nvidia-dynamo-boost-inference-performance-for-moe-models/?utm_source=chatgpt.com

49 Ayar Labs Unveils World's First UC1e Optical Chiplet for AI ...

https://ayarlabs.com/news/ayar-labs-unveils-worlds-first-uc1e-optical-chiplet-for-ai-scale-up-architectures/?utm_source=chatgpt.com

50 Marvell to Acquire Celestial AI, Accelerating Scale-up ...

https://investor.marvell.com/news-events/press-releases/detail/1000/marvell-to-acquire-celestial-ai-accelerating-scale-up-connectivity-for-next-generation-data-centers?utm_source=chatgpt.com

52 53 Product - System

https://www.cerebras.ai/cs4?utm_source=chatgpt.com

54 What is a Language Processing Unit?

https://groq.com/blog/the-groq-lpu-explained?utm_source=chatgpt.com

55 RDU | Next-Gen AI Chip for Inference at Scale

https://sambanova.ai/products/rdu-ai-chips?utm_source=chatgpt.com

56 **Introducing the SN50 RDU: Purpose-Built for Agentic ...**

https://sambanova.ai/blog/introducing-the-sn50-rdu-purpose-built-for-agentic-inference?utm_source=chatgpt.com

57 58 **Cards**

https://tenstorrent.com/hardware/cards?utm_source=chatgpt.com

59 **Etched**

https://www.etched.com/?utm_source=chatgpt.com

61 **May 28, 2026 - 10-Q: Quarterly report [Sections 13 or 15(d)]**

https://investor.marvell.com/sec-filings/all-sec-filings/content/0001835632-26-000019/mrvl-20260502.htm?utm_source=chatgpt.com

62 **FuriosaAI RNGD - Furiosa Docs**

https://developer.furiosa.ai/v2025.2.0/en/overview/rngd.html?utm_source=chatgpt.com

63 **Serving GPT-OSS 120B at 5.8 ms TPOT with RNGD**

https://furiosa.ai/blog/serving-gpt-oss-120b-at-5-8-ms-tpot-with-two-rngd-cards-compiler-optimizations-in-practice?utm_source=chatgpt.com

64 **Positron | Atlas**

https://www.positron.ai/atlas?utm_source=chatgpt.com

65 **AI inference at scale with AI250 rack-scale Platform**

https://www.qualcomm.com/data-center/products/qualcomm-dragonfly-ai250?utm_source=chatgpt.com

66 **Near-Memory Computing and the AI Memory Wall**

https://www.qualcomm.com/news/onq/2026/07/qualcomm-hbc-near-memory-computing?utm_source=chatgpt.com

67 **Groq and Nvidia Enter Non-Exclusive Inference ...**

https://groq.com/newsroom/groq-and-nvidia-enter-non-exclusive-inference-technology-licensing-agreement-to-accelerate-ai-inference-at-global-scale?utm_source=chatgpt.com

69 **The Decode Era of AI: Why Dataflow Matters More Than Ever**

https://sambanova.ai/blog/why-dataflow-matters-more-than-ever?utm_source=chatgpt.com

72 **[Tech Note] Hybrid Bonding: Evolving into a Foundational ...**

https://news.skynix.com/en/tech-note-series-ep2/?utm_source=chatgpt.com

75 **PALUTE: Processing-In-Memory Acceleration via Lookup ...**

https://arxiv.org/html/2606.08891v1?utm_source=chatgpt.com