
Accelerating Everything

Jeremy Kepner, Chansup Byun, Piotr Luszczek, LaToya Anderson, William Arcand, David Bestor, William Bergeron, Alex Bonn, Daniel Burrill, Vijay Gadepally, Michael Houle, Matthew Hubbell, Hayden Jananthan, Michael Jones, Lauren Milechin, Julie Mullen, Andrew Prout, Albert Reuther, Antonio Rosa, Charles Yee, Peter Michaleas

2026 June



Research was sponsored by the Department of the Air Force Artificial Intelligence Accelerator and was accomplished under Cooperative Agreement Number FA8750-19-2-1000. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of the Department of the Air Force or the U.S. Government. The U.S. Government is authorized to reproduce and distribute reprints for Government purposes notwithstanding any copyright notation herein.



Levels of Communication

Words

“How do I ... ”

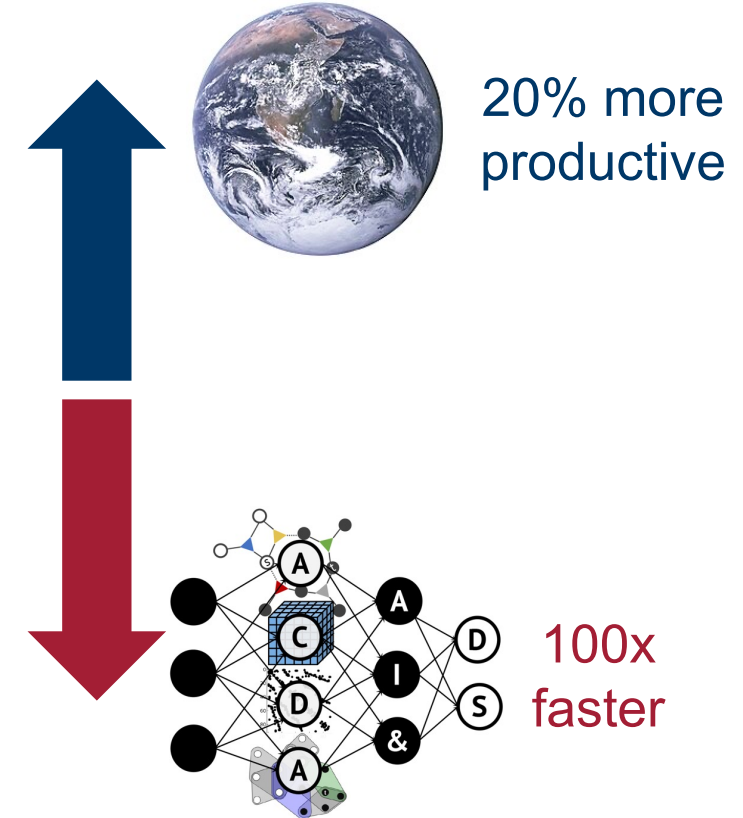
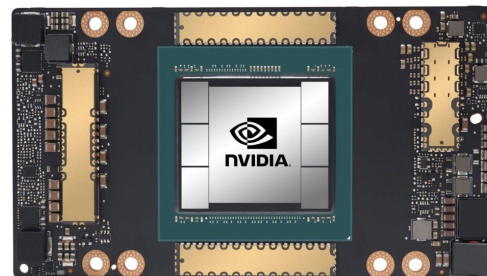
Code

```
for i=1:N
  for j=1:M
    for k=1:L
      C(i,j)=C(i,j)+A(i,k)*B(k,j);
```

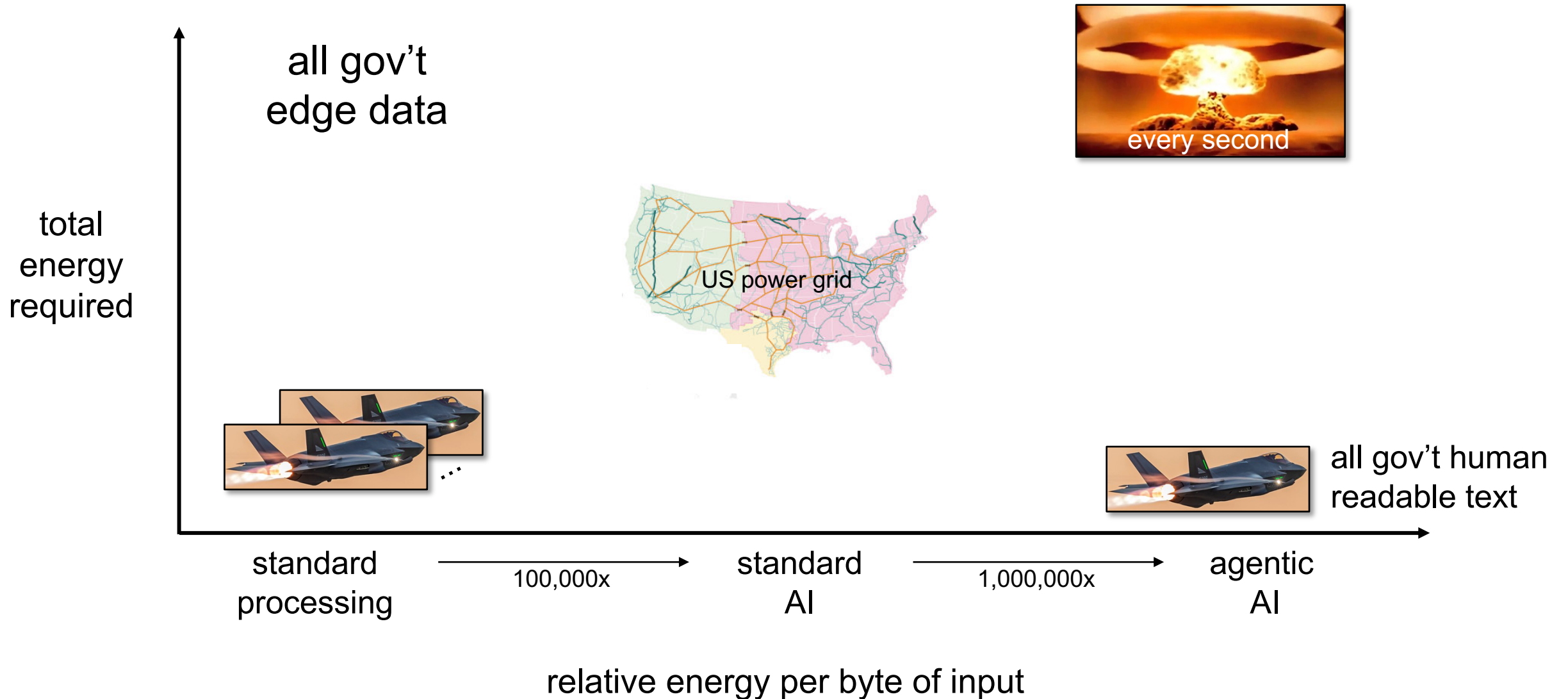
Math

$$C = A \times B$$

01101101

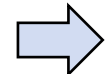


One Motivation: AI Compute Energy





Outline

- 
- **Hardware, Precision, Bandwidth**
 - **Data & Math**
 - **Benchmarks, Applications, Vectorization**
 - **Results**
 - **Summary**

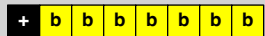


Numerical Precision Overview

Common Integer Representations

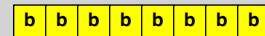
int8_t:

Range: -128 (-2^7) to 127 ($2^7 - 1$)



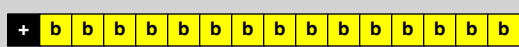
char, uint8_t:

Range: 0 to 255 ($2^8 - 1$)



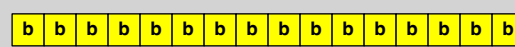
short, int16_t:

Range: -32,768 (-2^{15}) to 32,767 ($2^{15} - 1$)



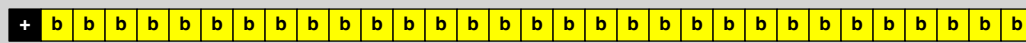
unsigned short, uint16_t:

Range: 0 to 65,535 ($2^{16} - 1$)



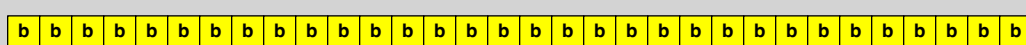
int, int32_t:

Range: -2,147,483,648 (-2^{31}) to 2,147,483,647 ($2^{31} - 1$)



unsigned int, uint32_t:

Range: 0 to 4,294,967,295 ($2^{32} - 1$)

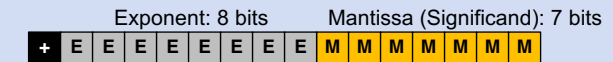


Common Floating-Point Representations

bfloat16: Brain Floating Point Format

Range: $\sim 1 \times 10^{-38}$ to $\sim 3 \times 10^{38}$

Number of Decimal Digits: ~ 2.3



fp16: Half-precision IEEE Floating Point Format

Range: $\sim 5.96 \times 10^{-8}$ to $\sim 6.50 \times 10^{-4}$

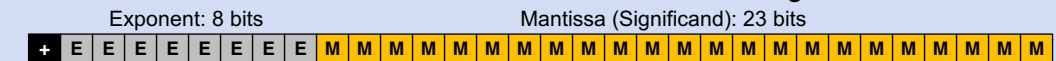
Number of Decimal Digits: ~ 3.3



fp32: Single-precision IEEE Floating Point Format

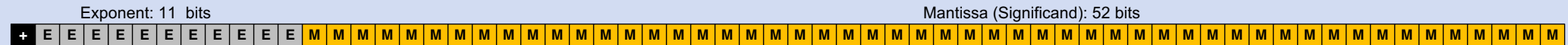
Range: $\sim 1 \times 10^{-38}$ to $\sim 3 \times 10^{38}$

Number of Decimal Digits: ~ 7.2



fp64: Double-precision IEEE Floating Point Format

Range: $\sim 1 \times 10^{-308}$ to $\sim 3 \times 10^{308}$



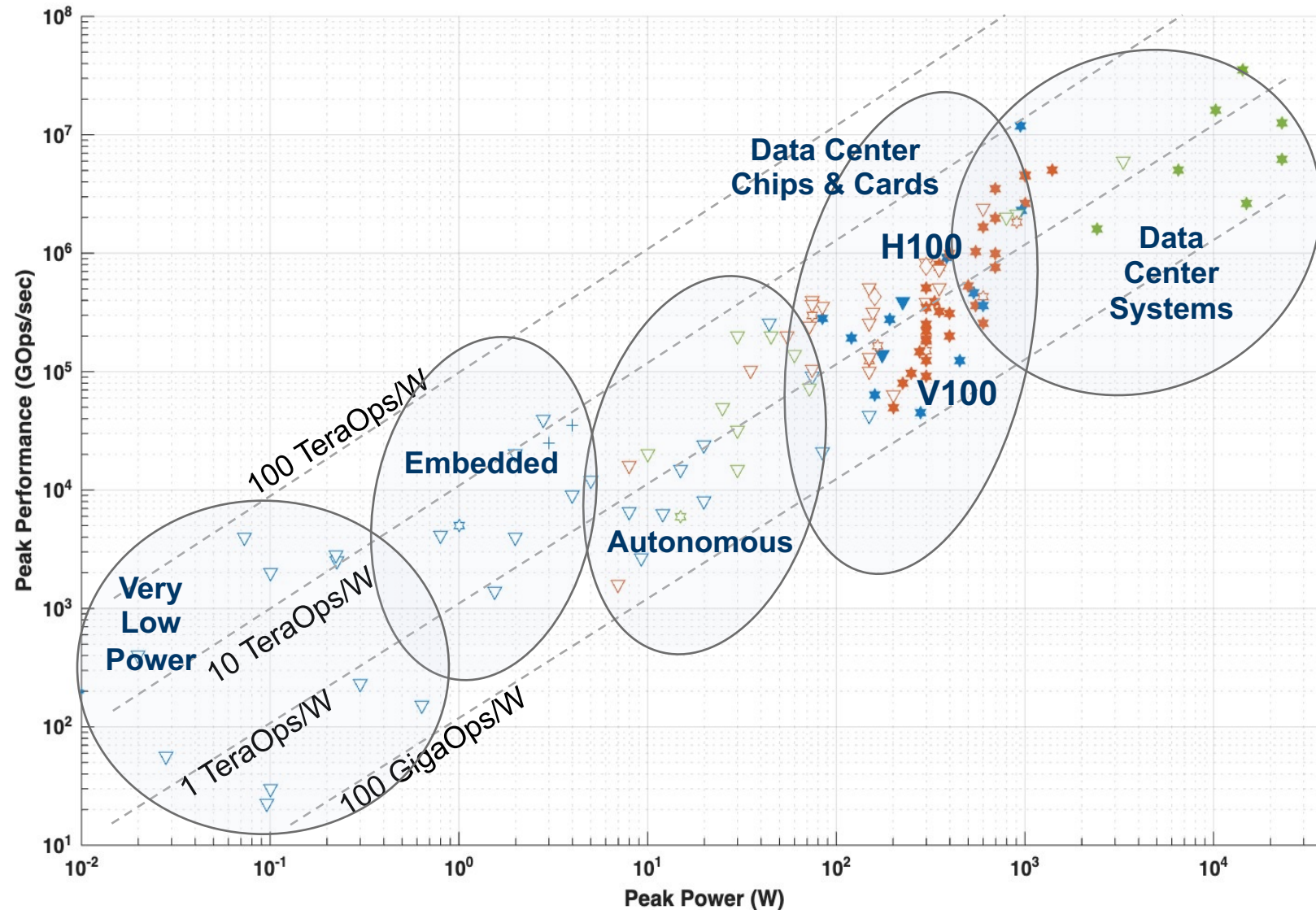
Number of Decimal Digits: ~ 15.9

Mantissa (Significand): 52 bits

Example: $123.45 = 1.2345 \times 10^2 = 0.12345 \times 10^3 = 0100\ 0010\ 1111\ 0110\ 1110\ 0110\ 0110\ 0110$ (fp32)



Lincoln AI Computing Survey (LAICS)



Legend

Computation Precision

+	analog	
◀	int1	■
▶	int2	■ b
•	int4	■ b b b
▼	int8	■ b b b b b b b
◆	fp8	■ E E E E M M M
▲	int16	■ b b b b b b b b b b b b b b b
×	int32	■ b b b b b b b b b b b b b b b b b
★	fp16	■ E E E E E M M M M M M M M M M
☆	bf16	■ E E E E E E E E M M M M M M M M
●	fp32	■ E E E E E E E E E M M M M M M M M M M
*	fp64	■ E E E E E E E E E E E M M M M M M M M M

Form Factor

- Chip
- Card
- System

Computation Type

- Inference
- Training



Mixed Precision Performance Across Industry



Peak specifications	Azure Maia 200	AWS Trainium3	Google TPU v7
Process Node	3nm	3nm	3nm
FP4 TFLOPS 	10,145	2,517	n/a
FP8 TFLOPS 	5,072	2,517	4,614
BF16 TFLOPS 	1,268	671	2,307



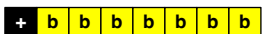
Genesis Mission Will Lean Heavily on **Ozaki Scheme** for FP64 Capability

While the latest generation of GPUs have emphasized lower precision performance that are favored for AI workloads,

- Many leading accelerators only have low-precision
- Rely on novel methods for emulating high-precision with low-precision



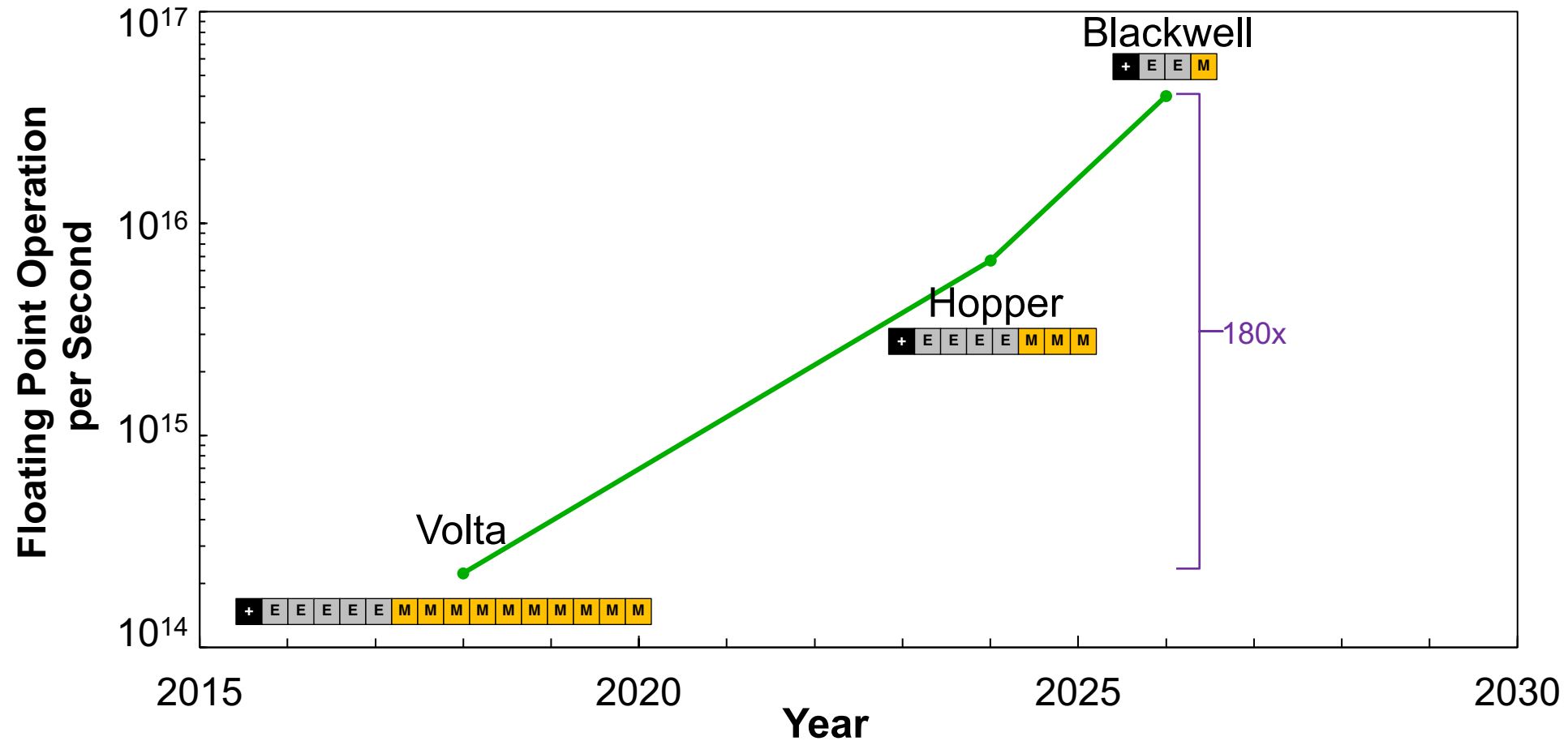
Example: Nvidia H100

Technical Specifications		
		H100 NVL
FP64	30 teraFLOPS	
FP64 Tensor Core	60 teraFLOPS	
FP32	60 teraFLOPS	
TF32 Tensor Core*	835 teraFLOPS	 =  X 
BFLOAT16 Tensor Core*	1,671 teraFLOPS	
FP16 Tensor Core*	1,671 teraFLOPS	
FP8 Tensor Core*	3,341 teraFLOPS	
INT8 Tensor Core*	3,341 TOPS	



Precision Performance Trend

- Nvidia Dual GPU -

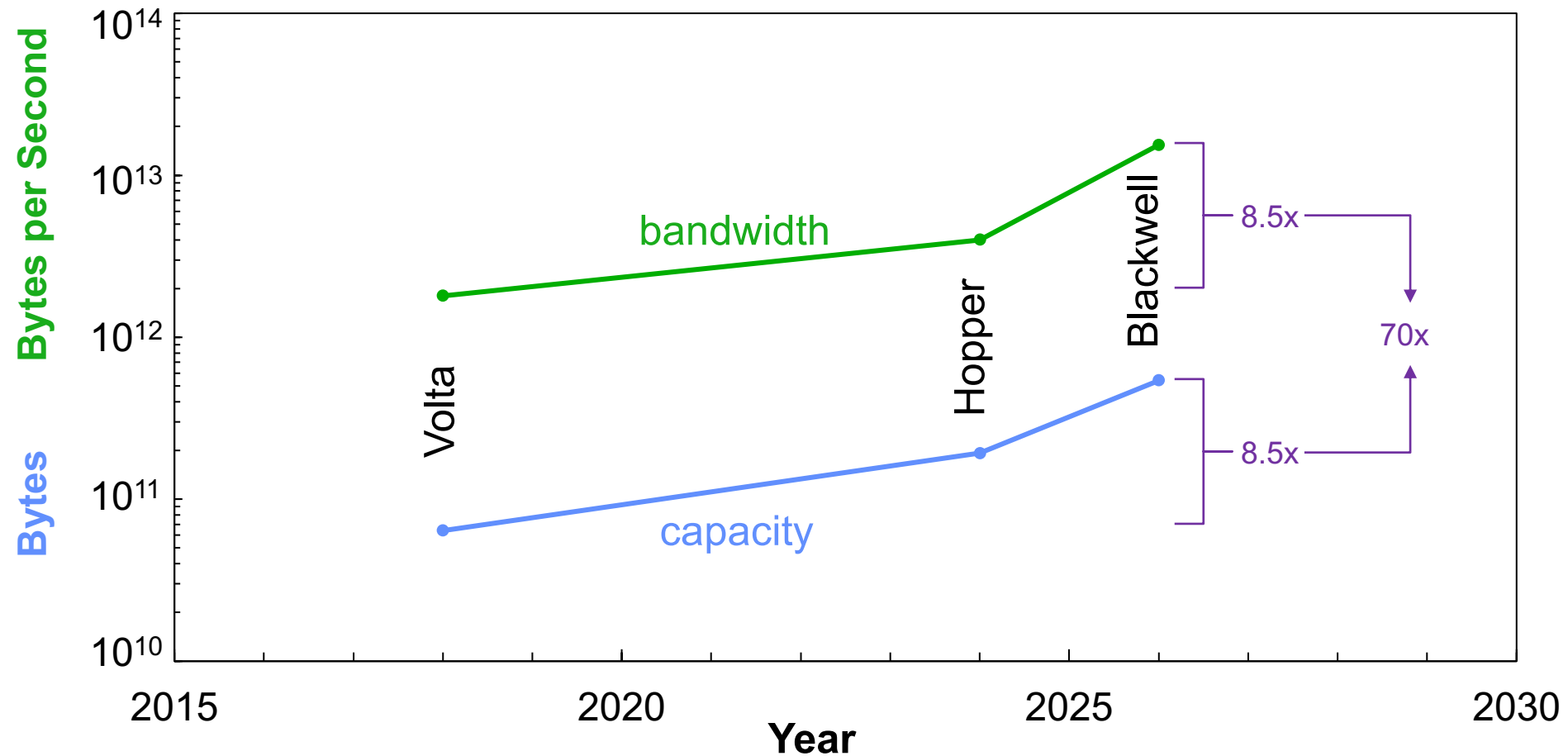


**Significant improvements for reduced precision algorithms
and emulation of high precision with low precision**



Memory Bandwidth Capacity Trend

- Nvidia Dual GPU -



Significant improvements for sparse, hypersparse, associative (token) array algorithms and spectral methods



Outline

- Hardware, Precision, Bandwidth
- ➡ • Data & Math
- Benchmarks, Applications, Vectorization
- Results
- Summary



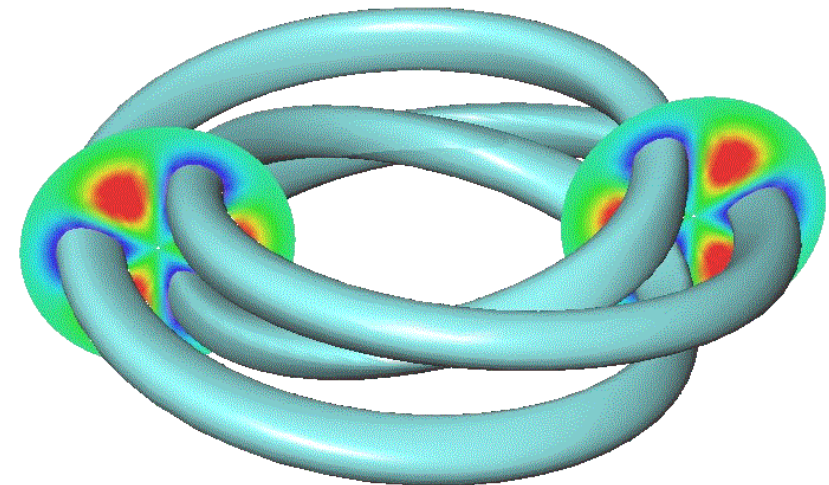
Accelerated Computing Data Foundations

Dense Double Precision (64 bits)

1 $\sim 2^{15}$

1	0.1234	0.1234	0.1234	0.1234	0.1234
	567890	567890	567890	567890	567890
	12345	12345	12345	12345	12345
...					
...	0.1234	0.1234	0.1234	0.1234	0.1234
	567890	567890	567890	567890	567890
	12345	12345	12345	12345	12345
...					
...	0.1234	0.1234	0.1234	0.1234	0.1234
	567890	567890	567890	567890	567890
	12345	12345	12345	12345	12345
...					
...	0.1234	0.1234	0.1234	0.1234	0.1234
	567890	567890	567890	567890	567890
	12345	12345	12345	12345	12345
...					
$\sim 2^{15}$	0.1234	0.1234	0.1234	0.1234	0.1234
	567890	567890	567890	567890	567890
	12345	12345	12345	12345	12345

Example: Physical Simulation





Accelerated Computing Data Foundations

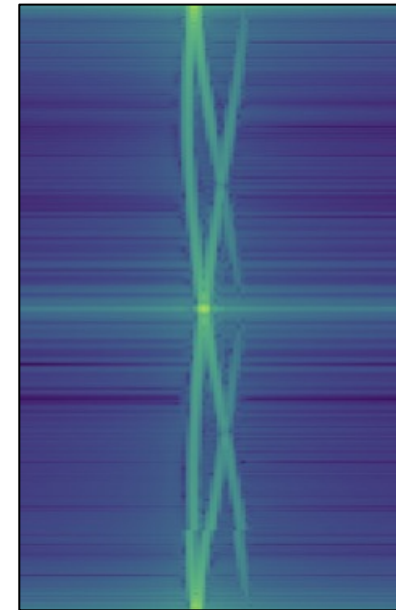
Dense Half Precision (16 bits)

1 $\sim 2^{15}$

1	0.1234	0.1234	0.1234	0.1234	0.1234
...					
...	0.1234	0.1234	0.1234	0.1234	0.1234
...					
...	0.1234	0.1234	0.1234	0.1234	0.1234
...					
...	0.1234	0.1234	0.1234	0.1234	0.1234
...					
...	0.1234	0.1234	0.1234	0.1234	0.1234
$\sim 2^{15}$					

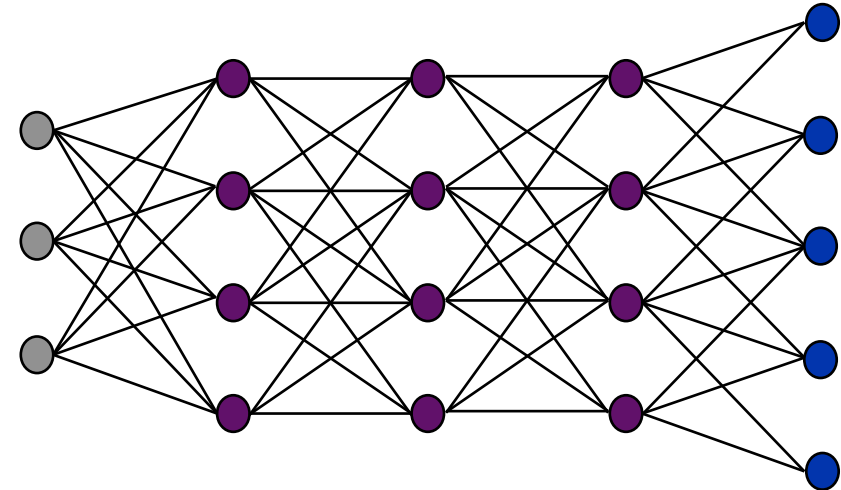
+ E E E E E M M M M M M M M M M

Example: Sensor Processing





Example: Networks





Example: Entity Relationships





Example: Arbitrary Data

```
abcdef
ghijkl
mnop...
```

Jeremy Kepner and Hayden Jananthan

$$2 \times (3 + 4) = (2 \times 3) + (2 \times 4)$$

$$2 \otimes (3 \oplus 4) = (2 \otimes 3) \oplus (2 \otimes 4)$$

$$\oplus = +$$

$$\otimes = \times$$

$$\oplus = \max$$

$$\otimes = +$$

$$\oplus = \cup$$

$$\otimes = \cap$$

$$A \otimes (B \oplus C) = (A \otimes B) \oplus (A \otimes C)$$

$$A = 2$$

$$B = 3$$

$$C = 4$$

Associative Arrays

A B C matrices

A B C graphs/networks

A B C database tables

A B C spreadsheets

A B C LLMs

$$A(i,j) = v$$

$$i \in I \quad j \in J \quad v \in V$$

matrices/graphs/networks: I, J integers V reals

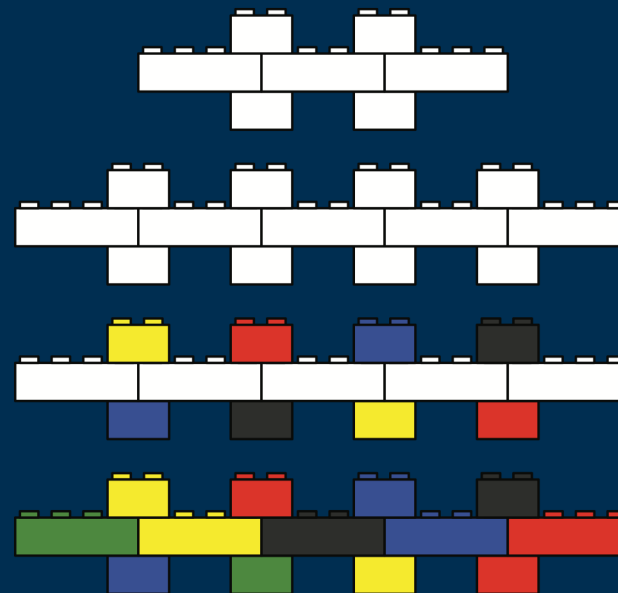
database tables/spreadsheets/LLMs: I, J, V tokens

Associative Arrays: I, J, V strict totally ordered sets

Mathematics of Big Data

Spreadsheets, Databases, Matrices, and Graphs

Jeremy Kepner and Hayden Jananthan



Foreword by Charles Leiserson

MIT LINCOLN LABORATORY SERIES



Outline

- **Hardware, Precision, Bandwidth**
- **Data & Math**
-  • **Benchmarks, Applications, Vectorization**
- **Results**
- **Summary**



Benchmark: STREAM Triad (memory bandwidth)

$$\mathbf{C} = \mathbf{A} + q \mathbf{B}$$

Python

$$C = A + q*B$$

Matlab

$$C = A + q*B;$$





The diagram illustrates the addition of two 8-bit numbers using a ripple-carry adder. The top row shows the addition of 10101010 (170) and 01010101 (85), resulting in 11111111 (255). The bottom row shows the addition of 11111111 (255) and 00000001 (1), resulting in 00000000 (0) with a carry-out of 1. The carry is shown as a '1' in a box to the left of the first column.

Diagram illustrating the iterative process of the bubble sort algorithm. It shows three rows of data: the initial array, the first pass, and the second pass. Each row has a '+' sign in a black box followed by 15 colored boxes representing elements. Row 1 (Initial): 14 yellow boxes, 1 black box. Row 2 (First Pass): 13 yellow boxes, 2 black boxes. Row 3 (Second Pass): 12 yellow boxes, 3 black boxes. The black boxes represent elements that are sorted and do not move in subsequent passes.

[illegible]

[illegible]



Associative (Token) Array: Construction, Addition, Multiplication

$$\mathbf{A} = \mathbb{A}(I, J, 1)$$

$$\mathbf{C} = \mathbf{A} + \mathbf{B}$$

$$\mathbf{C} = \mathbf{A} \times \mathbf{B}$$

Python

```
A = Assoc(I, J, 1)
```

```
C = A + B
```

```
C = A @ B
```

Matlab/Octave

```
A = Assoc(I, J, 1, @sum);
```

```
C = A + B;
```

```
C = A * B;
```

b b b b b b b b b b b b b b b b ...

b b b b b b b b b b b b b b b b ...

b b b b b b b b b b b b b b b b ...



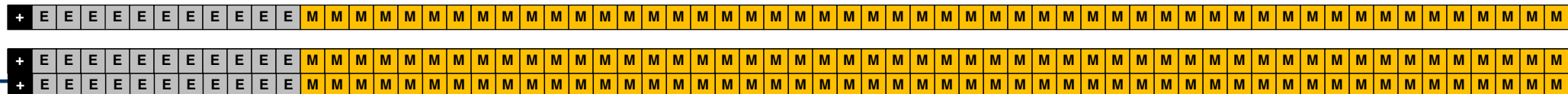
-
- The diagram illustrates the proposed beamforming architecture. It shows a sequence of operations: sensor response (from $\mathcal{X}_0$ to $\mathcal{X}_1$), beamform (from $\mathcal{X}_1$ to $\mathcal{X}_2$), and accumulate and sum (from $\mathcal{X}_2$ to $\mathbf{X}_3$). The input $\mathcal{X}_0$ is a 3D volume with dimensions N_t (time) and N_b (number of sensors). The intermediate $\mathcal{X}_1$ is a 3D volume with dimensions N_t and N_s (number of antennas). The output $\mathbf{X}_3$ is a 2D heatmap with dimensions N_t and N_b . The diagram also shows the physical layout of the sensors and antennas.

+	E	E	E	E	M	M	M	M	M	M	M	M	M
+	E	E	E	E	M	M	M	M	M	M	M	M	M

[illegible][illegible]

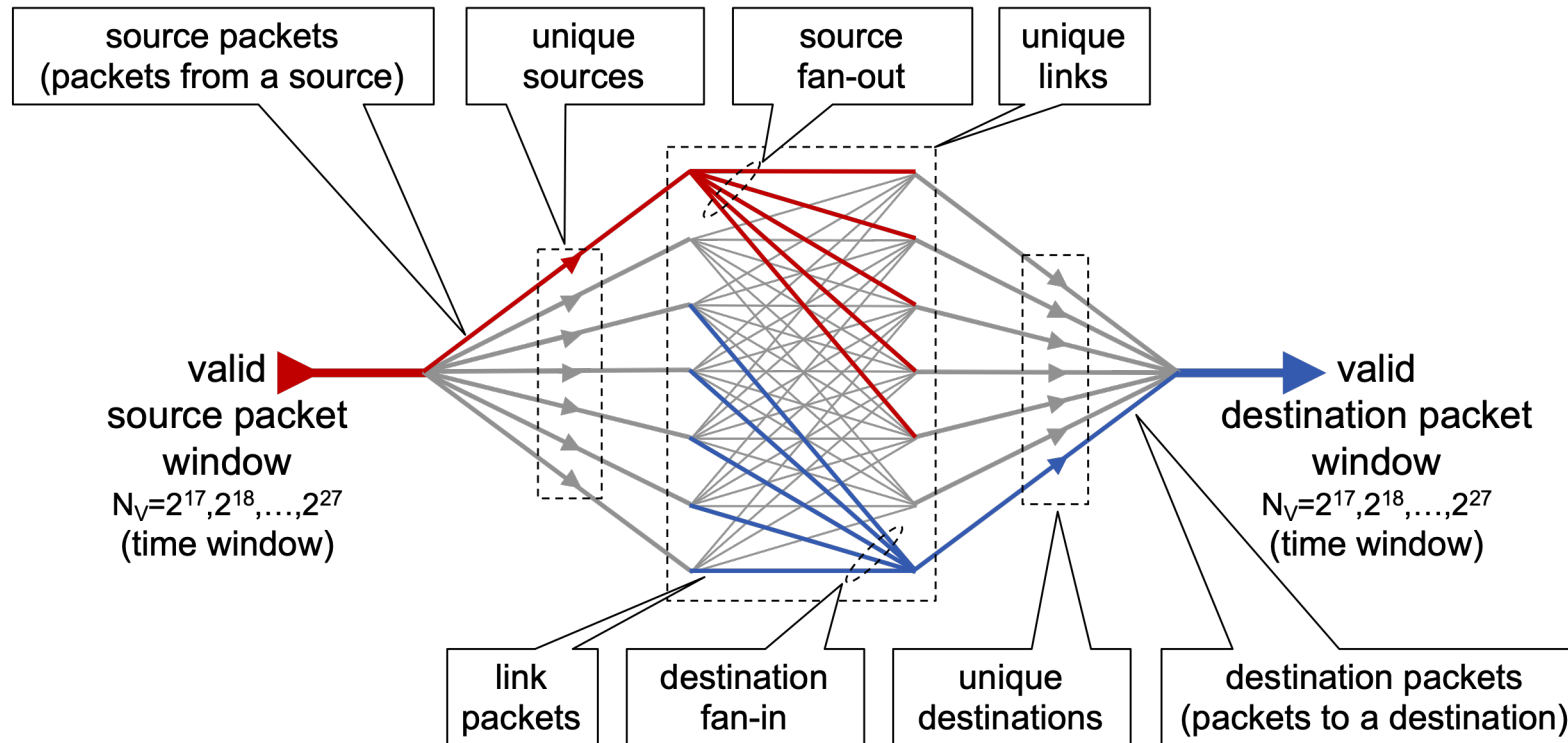
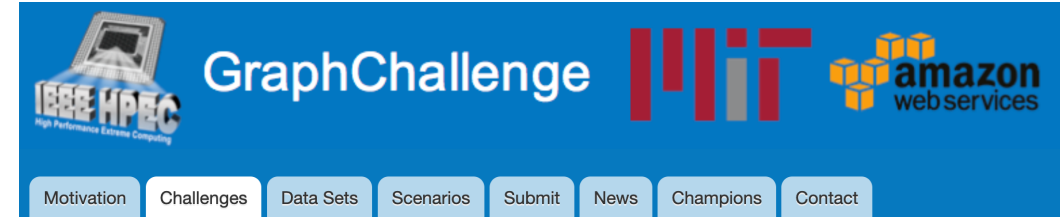
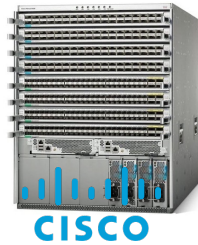
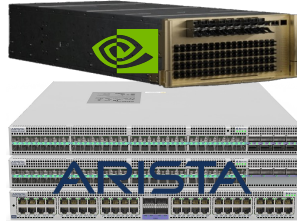
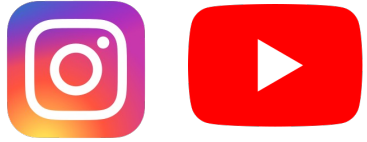


- 





Application: Network Sensing (associative arrays)



b b b b b b b b b b b b b b b b ...



Key Code Concept: Vectorized Computation

$$\mathbf{C} = \mathbf{A} \times \mathbf{B}$$

Standard Code

Python

```
for i in range(N):  
    for j in range(M):  
        for k in range(L):  
            C[i,j] += A[i,k] * B[k,j]
```

Matlab

```
for i=1:N  
    for j=1:M  
        for k=1:L  
            C(i,j)=C(i,j)+A(i,k)*B(k,j);
```

10,000x
Faster

Vectorized Code

Python

```
C = A @ B
```

Matlab

```
C = A * B;
```

Key Code Concept: Vectorized Row/Column Selection

$$\mathbf{A}_{\text{sub}} = \mathbf{A}(I,J) = \mathbf{I}_{\text{diag}} \times \mathbf{A} \times \mathbf{J}_{\text{diag}}$$

Standard Code

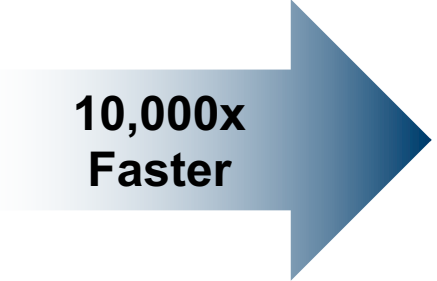
Python

```
for i in I:
    for j in J:
        Asub[i][j] = A[i][k]
```

Matlab

```
for i=I
    for j=J
        Asub(i,j)= A(i,k);
```

10,000x
Faster



Vectorized Code

Python

```
Idiag = csr_array(I,I,1)
Jdiag = csr_array(J,J,1)
Asub = Idiag @ A @ Jdiag
```

Matlab

```
Idiag = sparse(I,I,1);
Jdiag = sparse(J,J,1);
Asub = Idiag * A * Jdiag;
```



Outline

- **Hardware, Precision, Bandwidth**
- **Data & Math**
- **Benchmarks, Applications, Vectorization**
- ➡ • **Results**
- **Summary**



Hardware

Node Label	Processor				Memory	
	Era	Part	Clock	Cores	Part	Size
amd-e9	2024	Dual AMD EPYC 9254	2.9 GHz	48	DDR5	750 GB
- h100nvl	2024	Dual Nvidia H100 NVL	1.7 GHz		HBM3	188 GB
xeon-p8	2020	Dual Xeon Platinum 8260	2.4 GHz	48	DDR4	192 GB
xeon-g6	2018	Dual Xeon Gold 6248	2.5 GHz	40	DDR4	384 GB
- v100	2018	Dual Nvidia V100	1.2 GHz		HBM2	64 GB



Operation Counts

Operation	Standard	Alternative	Standard/Alternative
Matrix Addition	N^2	$k N$ (sparse)	N/k
Matrix Multiplication	N^3	$k^2 N$ (sparse)	$(N/k)^2$
2D Convolution	$N^2 N_{\text{kern}}^2$	$N^2 \log(N^2)$ (FFT)	$N_{\text{kern}}^2 / \log(N^2)$

N = matrix size; k = average non-zeros per row; N_{kern} = kernel size

Large fast memory enables significant benefits for alternative sparse algorithms and spectral methods



Benchmark: STREAM Triad (memory bandwidth)

Accessible
Performance
GB/s

Benchmark	Precision / HW	Year	2018	2020	2024	2018	2024	2018	2024
		Language	Matlab	Matlab	Matlab	Matlab	Matlab	Python	Python
		Type	CPU	CPU	CPU	GPU	GPU	GPU	GPU
			xeon-g6	xeon-p8	amd-e9254	v100	h100nvl	v100	h100nvl
Stream FP64			80	89	249	967	4,257	402	3,962

← 20x →

$$\mathbf{C} = \mathbf{A} + q \mathbf{B}$$

Memory bandwidth performance is readily accessible in high level programming environments via vectorized operations



Benchmark: Dense Matrix Multiplication

Accessible
Performance
GigaFlops

Benchmark	Year Language Type Precision / HW	2018	2020	2024	2018	2024	2018	2024
		Matlab CPU xeon-g6	Matlab CPU xeon-p8	Matlab CPU amd-e9254	Matlab GPU v100	Matlab GPU h100nvl	Python GPU v100	Python GPU h100nvl
MatMul FP64		2,204	2,713	2,218	14,041	103,479	13,852	85,984
MatMul FP32		4,684	5,808	4,002	27,728	97,200	27,556	86,900
MatMul FP16		E	E	E	← u	686,129	184,460	1,006,242
MatMul INT16		u	u	u	u	u	E	E
MatMul FP8		u	u	u	← u	u	u	1,752,992
MatMul INT8		u	u	u	u	u	u	INT8
MatMul FP64 (complex)		2,628	2,990	2,262	13,936	104,017	C	3,230
MatMul FP32 (complex)		5,216	6,182	4,642	27,239	97,699	C	6,462
MatMul FP16 (complex)		E	E	E	← u	295,656	S	S
MatMul INT16 (complex)		u	u	u	u	u	u	u
MatMul FP8 (complex)		u	u	u	u	u	u	u
MatMul INT8 (complex)		u	u	u	u	u	u	u

$$C = A \times B$$

Mixed precision performance is readily accessible in high level programming environments via vectorized operations



Associative (Token) Array: Construction, Plus, MatMul

Accessible
Performance
GigaFlops

Benchmark	Year	2018	2020	2024	2018	2024	2018	2024
	Language	Matlab	Matlab	Matlab	Matlab	Matlab	Python	Python
	Type	CPU	CPU	CPU	GPU	GPU	GPU	GPU
Precision / HW		xeon-g6	xeon-p8	amd-e9254	v100	h100nvl	v100	h100nvl
Assoc Constuct		532	498	1,240	3,876	12,916	u	u
Assoc Plus		3,129	3,300	25,803	8,980	104,318	u	u
Assoc Minus		2,415	2,506	22,609	8,672	100,775	u	u
Assoc Times		6,095	5,622	27,142	19,333	112,329	u	u
Assoc MatMul		1,589,137,900	1,510,271,563	68,849,757,411	10,391,243,876	366,068,652,573	u	u

>100,000x >100,000x >10,00,000x >100,000x >1,000,000x

Faster Than Dense MatMul

$$\mathbf{A} = \mathbb{A}(I,J,1)$$

$$\mathbf{C} = \mathbf{A} + \mathbf{B}$$

$$\mathbf{C} = \mathbf{A} \times \mathbf{B}$$

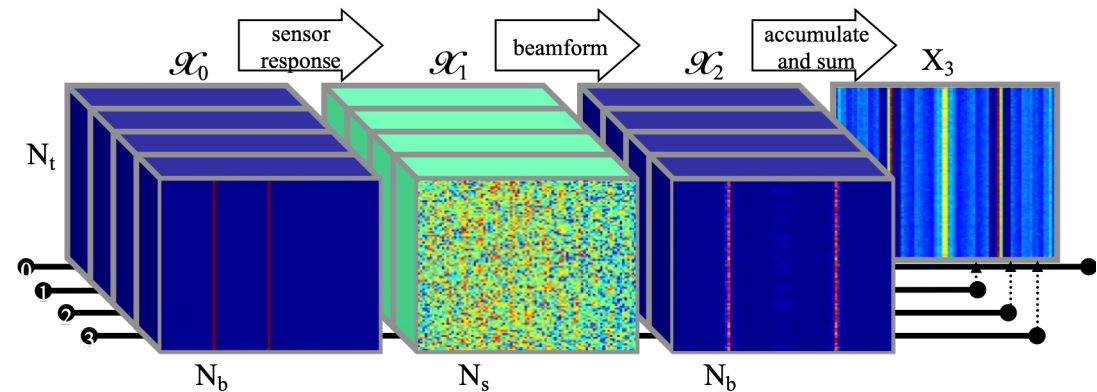
Sparse algorithm boost is readily accessible in high level programming environments via vectorized operations



Application: Beamforming (complex matrix multiply)

Accessible
Performance
GigaFlops

Benchmark	Language Type Precision / HW	2018	2020	2024	2018	2024	2018	2024
		Matlab CPU xeon-g6	Matlab CPU xeon-p8	Matlab CPU amd-e9254	Matlab GPU v100	Matlab GPU h100nvl	Python GPU v100	Python GPU h100nvl
Beamformer FP64 (complex)		2,091	2,516	2,200	13,817	77,149	12,149	69,892
Beamformer FP32 (complex)		4,546	5,432	4,643	27,125	86,048	26,447	218,530
Beamformer FP16 (complex)		E	E	E	← u 50x →	250,729	26,135	218,230
Beamformer INT16 (complex)		u	u	u		u	u	u
Beamformer FP8 (complex)		u	u	u		u	u	u
Beamformer INT8 (complex)		u	u	u		u	u	u



Complex mixed precision performance can be achieved in applications
via vectorized operations



Application: Blurimage (2D FFT Convolution)

Accessible
Performance
GigaFlops

Benchmark	Year	2018	2020	2024	2018	2024	2018	2024
	Language Type	Matlab CPU	Matlab CPU	Matlab CPU	Matlab GPU	Matlab GPU	Python GPU	Python GPU
Precision / HW		xeon-g6	xeon-p8	amd-e9254	v100	h100nvl	v100	h100nvl
Blurimage (FFT) FP64		17,992	21,996	79,093	216,172	538,203	356,450	1,869,907
Blurimage (FFT) FP32		30,173	36,541	129,510	345,396	705,916	889,540	3,220,154
Blurimage (FFT) FP16		27,502	32,894	113,965	u	u	851,716	3,018,172
Blurimage (FFT) INT16		18,261	23,194	78,058	216,017	504,129	352,583	1,803,563
Blurimage (FFT) FP8		u	u	u	u	u	u	u
Blurimage (FFT) INT8		18,625	23,020	79,222	229,887	509,873	353,763	1,801,938
Blurimage (FFT) FP64 (complex)		60,424	74,441	265,580	881,918	2,130,639	812,984	3,459,935
Blurimage (FFT) FP32 (complex)		118,680	138,015	481,290	1,350,745	2,812,656	1,623,058	6,400,063
Blurimage (FFT) FP16 (complex)		109,527	131,764	461,739	u	u	1,623,904	6,359,006
Blurimage (FFT) INT16 (complex)		60,891	73,908	263,376	864,854	2,020,847	813,222	3,489,245
Blurimage (FFT) FP8 (complex)		u	u	u	u	u	u	u
Blurimage (FFT) INT8 (complex)		62,341	76,883	271,335	878,376	2,032,205	813,089	3,488,037

>10x

>10x

>30x

>10x

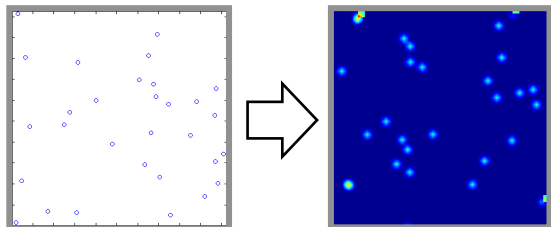
>10x

>30x

>30x

Faster Than Dense MatMul

Spectral method boost can be achieved in applications via vectorized operations





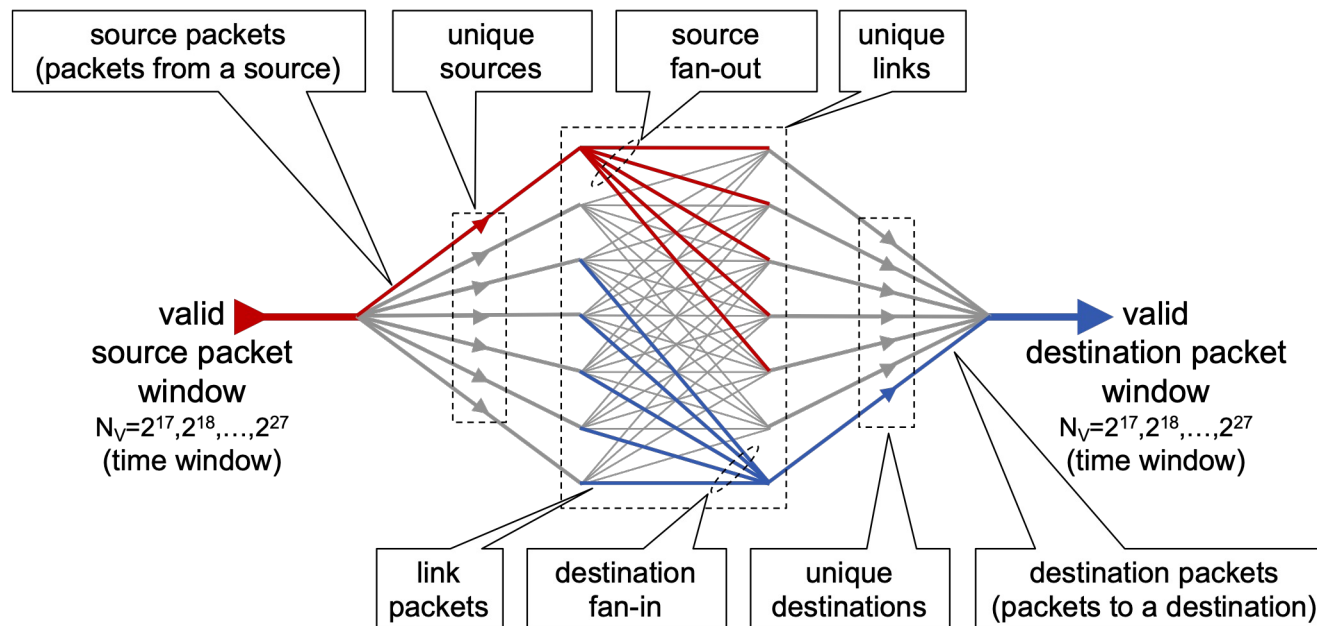
Application: Network Sensing (associative arrays)

Accessible
Performance
GigaFlops

Benchmark	Precision / HW	Year	2018	2020	2024	2018	2024	2018	2024
		Language	Matlab	Matlab	Matlab	Matlab	Matlab	Python	Python
		Type	CPU	CPU	CPU	GPU	GPU	GPU	GPU
			xeon-g6	xeon-p8	amd-e9254	v100	h100nvl	v100	h100nvl
AnonNetSense Assoc			24,181,370,389	12,145,394,560	17,909,144,186	281,593,137,894	862,723,041,517	u	u

>1,000,000x >1,000,000x >1,000,000x >10,00,000x >10,000,000x

Faster Than Dense MatMul



Sparse algorithm boost can be achieved in applications via vectorized operations

Summary

Words

“How do I ...”

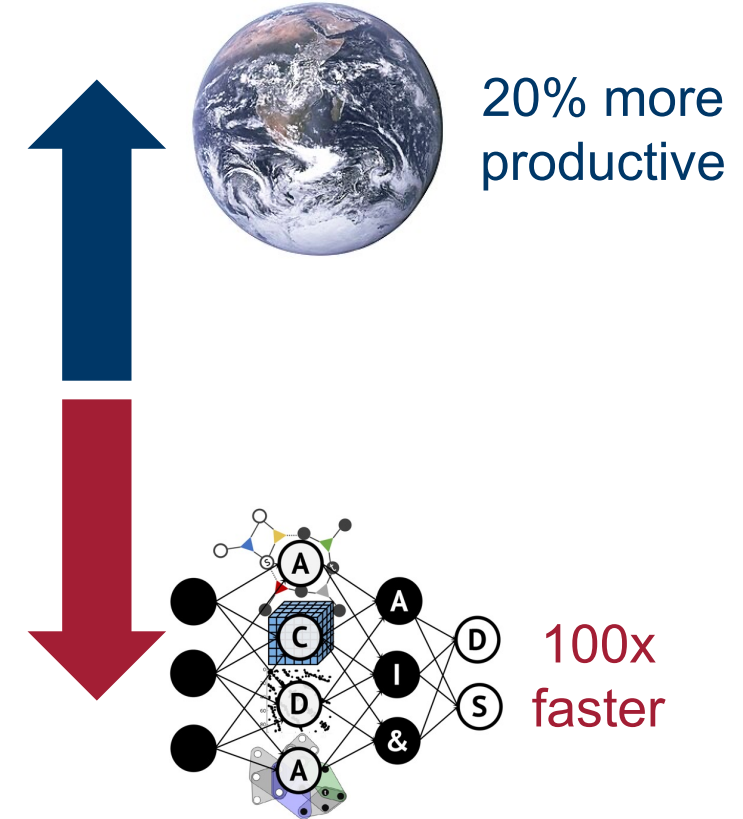
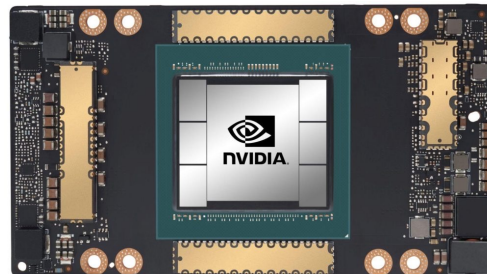
Code

```
for i=1:N
  for j=1:M
    for k=1:L
      C(i,j)=C(i,j)+A(i,k)*B(k,j);
```

Math

$$C = A \times B$$

01101101

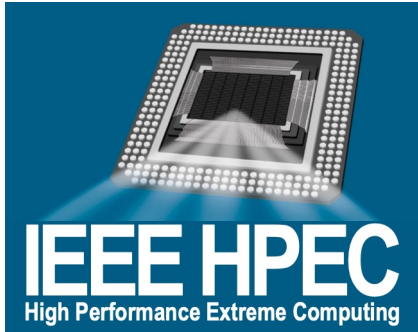




30th IEEE HPEC Virtual Conference September 14-18, 2026 (ieee-hpec.org)

Papers Due:
July 14

- Premiere conference on High Performance Extreme Computing (850+ participants in 2025)
- 150+ presentations; 2026 Distinguished Speakers:
 - Prof. Saman Amarasinghe (Perkins Prof EECS MIT; ACM Fellow)
 - Prof. Yaneer Bar-Yam (President NECSI)
 - Craig Connors (VP & CTO Cisco Infrastructure & Security)
 - Prof. Danna Freedman (Quantum@MIT Director; MacArthur Fellow)
 - Prof. Sigal Gottlieb (UMass Dartmouth; SIAM & AWM Fellow)
 - Gregory Kurtzer (CEO CIQ & Founder of CentOS & Singularity)
 - Dave Martinez (MIT Lincoln Lab Fellow; IEEE Fellow)
 - Dr. CJ Newburn (Nvidia Distinguished Engineer)
 - Prof. Katsuhisa Ozaki (Shibaura Institute of Technology)
 - Prof. Daniela Rus (CSAIL Director; AAAS, ACM, AAAI, IEEE Fellow)
- 2026 Special Sessions
 - Age of Mixed-Precision
 - Bridging Quantum and HPC
 - Large Language Models: Opportunities and Challenges
 - MIT/Amazon/IEEE Graph Challenge (also due **July 14**)
 - GraphBLAS Forum
 - BRAINS: Building Resilience through Artificial Intelligence for Networked Systems
 - Scaling Research Computing Education



IEEE HPEC
High Performance Extreme Computing

IEEE Boston Section

Sponsors	Thank You	Technical Organizer
 NVIDIA MITRE	New England Section of SIAM Boston Chapter of the IEEE Computer Society MIT Student Chapter of SIAM	

- **HPEC focuses on hardware, software, systems, and applications where performance matters**
- **HPEC welcomes experts and people who are new to the field**