

AMD FINANCIAL ANALYST DAY 2025

Dan McNamara

SVP & GM, Compute & Enterprise AI



together we advance_

Cautionary Statement

This presentation contains forward-looking statements concerning Advanced Micro Devices, Inc. (AMD) such as the features, functionality, performance, availability, timing and expected benefits of AMD products; AMD's ability to gain server market share; AMD EPYC™ CPU roadmap; AI driving additional CPU demand; AMD's data center CPU revenue growth and CAGR; and AMD's ability to accelerate server leadership, which are made pursuant to the Safe Harbor provisions of the Private Securities Litigation Reform Act of 1995. Forward-looking statements are commonly identified by words such as "would," "may," "expects," "believes," "plans," "intends," "projects" and other terms with similar meaning. Investors are cautioned that the forward-looking statements in this presentation are based on current beliefs, assumptions and expectations, speak only as of the date of this presentation and involve risks and uncertainties that could cause actual results to differ materially from current expectations. Such statements are subject to certain known and unknown risks and uncertainties, many of which are difficult to predict and generally beyond AMD's control, that could cause actual results and other future events to differ materially from those expressed in, or implied or projected by, the forward-looking information and statements. Investors are urged to review in detail the risks and uncertainties in AMD's Securities and Exchange Commission filings, including but not limited to AMD's most recent reports on Forms 10-K and 10-Q.

AMD does not assume, and hereby disclaims, any obligation to update forward-looking statements made in this presentation, except as may be required by law.

Accelerating Our Path To #1 Server Market Share

**Product
Leadership**

**Customer
Focus**

**Relentless
Execution**

EPYC™ CPU: The Foundation of Next Generation Compute



Cloud

Highest vCPU Density

Generational
TCO Gain



Enterprise

Highest Performance Per Core

Modernization
& Consolidation



HPC

Highest FLOPS Per Node

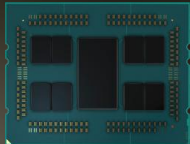
Faster Solutions to
Complex Problems

Optimized Performance & TCO for AI Use Cases

Leadership in Density, Performance & Efficiency

Customer Centricity Drives Innovation

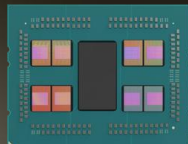
2nd Generation “Rome”



>2x Thread Density

Cloud Computing
HPC

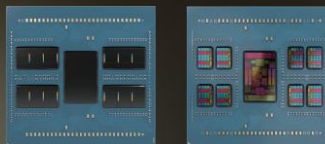
3rd Generation “Milan”



Core Performance
Leadership

Enterprise
Technical Compute

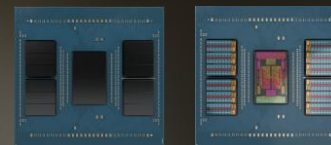
4th Generation “Genoa” | “Bergamo”



Optimized
Design Points

Cloud Native
Edge & Telco

5th Generation “Turin” Classic & Dense



AI Optimization

AI Inference &
Host Node

Next-Generation “Venice”



Next-Gen
Leadership

Broad AI Use Cases
Increased Efficiency

2019

2026

AMD
together we advance_

Cloud: Leading Across Applications

Enabling Secure Hybrid & Multi-Cloud Deployments



General IT Migration
Lowest TCO

>1.5x Performance/\$



Mission Critical
Leadership Performance

>1.7x VM Performance



Specialty Compute
Faster Time to Results

>2.7x Simulation Performance

EPYC™ CPUs Connect Billions of Users Daily

10/10

Top 10 Social Media Platforms



LinkedIn

10/10

Largest Cloud Native SaaS Services



NETFLIX

zoom



>3x Y/Y

Public Cloud VM Adoption



Uber

Enterprise Ecosystem Centers Around EPYC™ CPUs



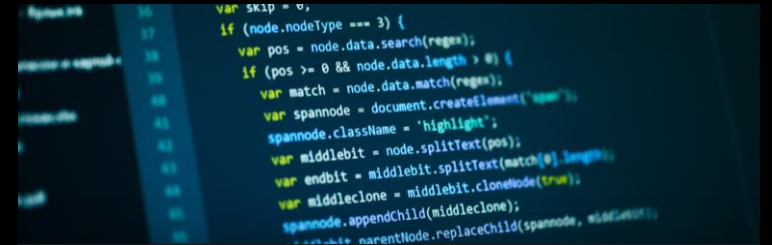
Optimized Platforms

>170 Platforms



Industry Workloads

Accelerate Time to Results by 1.8x



Core IT

~80% Lower TCO*



Enterprise Adoption Accelerating

>60%
Customers

Fortune® 100 Adoption

>2x

New Customers YTD

Landing New Customers

>2x

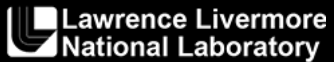
Consumption Growth

Expanding Footprint

Continued Exascale Leadership

Established Exascale HPC Leader
Double Precision Floating Point

AI to Accelerate Scientific Discovery
Optimized AI & HPC Data Types



**Molecular
Dynamics**

Genomics

Oil & Gas

**Training
LLMs**

**Extreme Weather
Prediction**

**National
Security**

AI Requires Leadership CPUs



Head Node

CPU Unlocks
Training & Inference
Cluster Performance

AI Ready Data Center

Modernization for
AI Driven Computing

Processing & Storage

Optimizing Data
Pipelines

CPU Inference

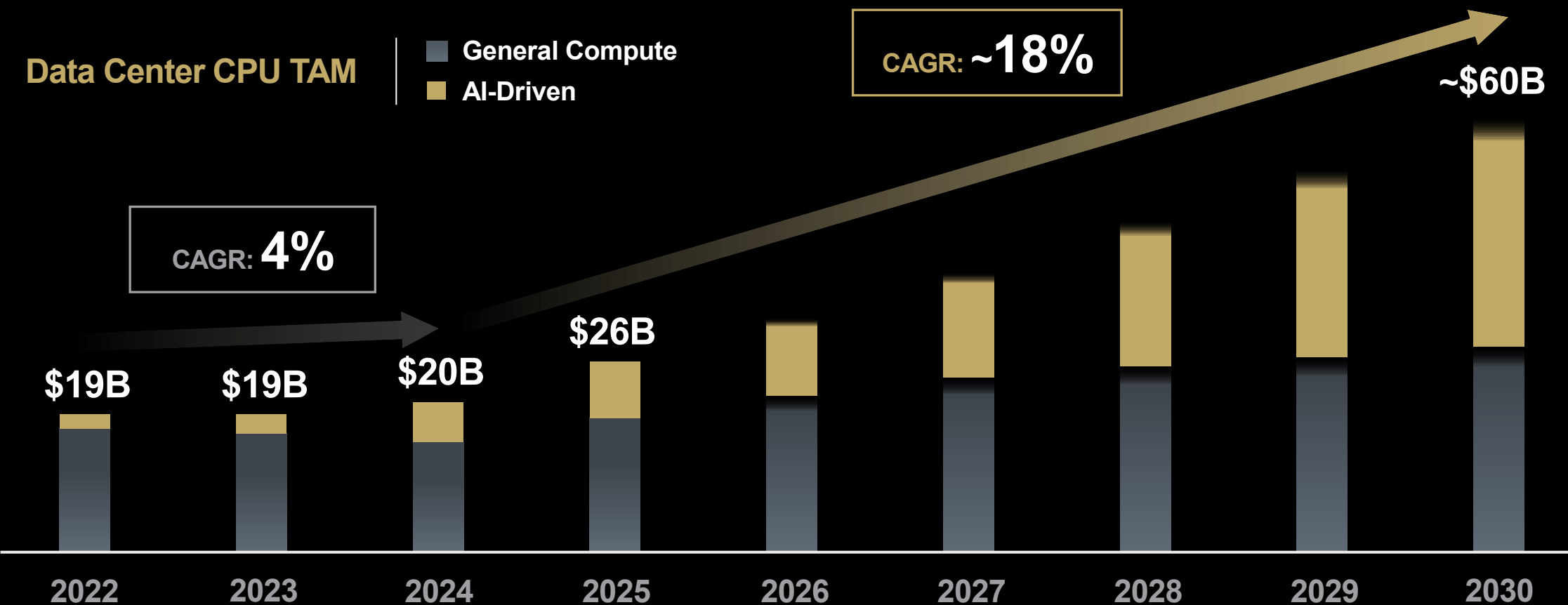
CPU Inference
Where Data Lives

Agentic Actions

Agentic Actions
Drive Data
Processing

AI Inflection Drives Server CPU Consumption

AI Will Create an Additional ~\$30B of CPU Demand by 2030



Delivering the Next Generation of Performance

With “Venice” The World’s Best Gets Even Better



Next Generation “Venice”

>1.3x

Thread Density

>1.7x

Performance
& Efficiency

Accelerating Server Leadership

**Sustained
Cloud Expansion**

**Accelerated
Enterprise Adoption**

**Leadership in
AI Growth**

Clear Path to >50% Market Share

Disclaimer & Attribution

DISCLAIMER: The information contained herein is for informational purposes only and is subject to change without notice. While every precaution has been taken in the preparation of this document, it may contain technical inaccuracies, omissions and typographical errors, and AMD is under no obligation to update or otherwise correct this information. Advanced Micro Devices, Inc. makes no representations or warranties with respect to the accuracy or completeness of the contents of this document, and assumes no liability of any kind, including the implied warranties of noninfringement, merchantability or fitness for particular purposes, with respect to the operation or use of AMD hardware, software or other products described herein. No license, including implied or arising by estoppel, to any intellectual property rights is granted by this document. Terms and limitations applicable to the purchase or use of AMD products are as set forth in a signed agreement between the parties or in AMD's Standard Terms and Conditions of Sale. GD-18u.

© 2025 Advanced Micro Devices, Inc. All rights reserved. AMD, the AMD Arrow logo, AMD Instinct, EPYC, Pensando, Radeon, ROCm, Ryzen, Versal, Xilinx, and combinations thereof are trademarks of Advanced Micro Devices, Inc. CXL is a registered trademark of Compute Express Link Consortium, Inc. OpenAI is a trademark of OpenAI, Inc. PCIe® is a registered trademark of PCI-SIG Corporation. UCIE is a trademark of Universal Chiplet Interconnect Express, Inc. Ultra Accelerator Link and UALink are trademarks of the UALink Consortium. Other product names used in this publication are for identification purposes only and may be trademarks of their respective owners. Certain AMD technologies may require third-party enablement or activation. Supported features may vary by operating system. Please confirm with the system manufacturer for specific features. No technology or product can be completely secure.

Fortune content is From Fortune Magazine. © 2025 Fortune Media IP Limited. All rights reserved. Used under license. Fortune is a registered trademark of Fortune Media IP Limited and is used under license.

Endnotes

- Endnote #1: Top 10 Social Media platforms (based on active users) and Top 10 Cloud Native SaaS Services (based on Market Cap)
- 9xx5c-012: Testing by AMD Performance Labs as of 3/12/2025. H4D highmem-192 vCPU instance with Local SSD provides an average ~2.70% performance uplift versus H3 standard 88 vCPU instance running industry standard benchmarks on single-instance including Ansys® Fluent®, StarCCM+™, OpenFOAM®, Altair® Radioss®, Ansys® LS-DYNA®, Ansys® Mechanical™, WRF®, and GROMACS. H4D and H3 are both SMT off VMs. Local SSD is not supported on H3. Average (mean) of 3 runs for each benchmark on each VM type: H3H4DUplift Fluent aircraft_wing_14mSolver Rating (Higher is Better)1752.303495.001.99 Fluent exhaust_system_33mSolver Rating (Higher is Better)644.871208.671.87 StarCCM+ hlmach10Avg Elapsed Time (Lower is Better)1.520.532.88 StarCCM+ lemans_seggregatedAvg Elapsed Time (Lower is Better)7.692.852.70 OpenFOAM 1004040Elapsed Time (Lower is better)997.99484.962.06 OpenFOAM 1084646Elapsed Time (Lower is better)1409.63615.062.29 OpenFOAM 1305252Elapsed Time (Lower is better)2163.55955.322.26 Radioss neonElapsed Time (Lower is better)1583.52437.813.62 Radioss 10mElapsed Time (Lower is better)8389.523082.642.72 LSDYNA 3carsElapsed Time (Lower is better)732.33193.333.79 LSDYNA car2carElapsed Time (Lower is better)6495.002541.672.56 LSDYNA odb10mElapsed Time (Lower is better)60246.6722692.332.65 Mechanical GearBoxSolver Rating (Higher is Better)361.92905.982.50 Mechanical SemiSubmersibleSolver Rating (Higher is Better)378.51887.682.35 WRF conus2.5kmElapsed Time (Lower is better)3.351.113.03 GROMACS benchPEPNS per day (Higher is Better)1.123.763.35 GROMACS lignocellulose-rfNS per day (Higher is Better)6.8922.323.24 Average (mean) of all workload benchmark uplifts: 2.70 Cloud performance results presented are based on the test date in the configuration. Results may vary due to changes to the underlying configuration, and other conditions such as the placement of the VM and its resources, optimizations by the cloud service provider, accessed cloud regions, co-tenants, and the types of other workloads exercised at the same time on the system.
- 9xx5c-041: Testing by AMD Performance Labs as of 2/21/2025. C4D-standard-16 score comparison to C3D-standard-16 estimated SPECrate®2017_int_base benchmark scores Performance differences (normalized to C3D): C4D 16 vCPU vs C3D 16 vCPU Perf (est) Perf/\$ (est) OpEx Savings SPECrate®2017 Integer 1.271.24~19.6% On-demand hourly pricing from <https://cloud.google.com/compute/vm-instance-pricing> (us-central1) as of 8/1/2025: C4D: \$0.740, C3D \$0.726 Cloud performance results presented are based on the test date in the configuration. Results may vary due to changes to the underlying configuration, and other conditions such as the placement of the VM and its resources, optimizations by the cloud service provider, accessed cloud regions, co-tenants, and the types of other workloads exercised at the same time on the system. SPEC®, SPECrate® and SPECint® are registered trademarks of Standard Performance Evaluation Corporation. Learn more at spec.org.

Endnotes

- 9xx5c-046: Testing by AMD Performance Labs as of 10/25/2025. M8a.4xlarge score comparison to M8i.4xlarge running the following benchmarks: - FFmpeg v6.1.1 (average of 2x encode and 2x transcode) - Redis 7.2 (average of get and set) - NGINX 1.1.9-2, WRK 4.2 - MySQL HammerDB TPROC-C 4.2 - SQL Server (TPC_E derivative nd TPC_H) Performance differences (normalized to M8a): M8A.4XL 16 vCPU vs M8i.4XL 16 vCPU Perf Perf/\$ OpEx Saving
FFmpeg Average 1.90 1.66 40.0% Redis Average 2.02 1.75 42.0% NGINX 2.14 1.86 46.0% MySQL TPROC-C 1.53 1.33 25.0% SQL Server TPC_E derivative 1.35 1.18 15% SQL Server TPC_H 1.47 1.28 22% Top 6 Workload Benchmarks 1.73 1.51 .31.7. 0% On-demand hourly pricing from <https://aws.amazon.com/ec2/pricing/on-demand/> (us-east) as of 10/25/2025: M8A.4XL: \$0.973, M8i.4XL \$0.846 Cloud performance results presented are based on the test date in the configuration. Results may vary due to changes to the underlying configuration, and other conditions such as the placement of the VM and its resources, optimizations by the cloud service provider, accessed cloud regions, co-tenants, and the types of other workloads exercised at the same time on the system. Cloud OpEx savings calculated from relative runtime and instance cost ratio. The HammerDB TPROC-C for MySQL workload is an open-source workload derived from TPC-Benchmark™ Standard, and as such is not comparable to published TPC-C™ results, as the results do not comply with the TPC-C Benchmark Standard .. TPC, TPC Benchmark, TPC-C are trademarks of the Transaction Processing Performance Council. The SQL Server TPC-E derivative is derived from TPC-E Benchmark standard and as such is not comparable to published TPC-E results as the results do not comply with the TPC-E Benchmark standard.
- 9xx5-142: AMD testing as of 04/02/2025. The detailed results show the average uplift of the performance metric (ns/day) of this benchmark for a 2P 192-Core AMD EPYC™ 9965 powered reference system compared to a 2P Intel® Xeon® 6980P powered production system running select tests on Open-Source GROMACS 2023.1. Uplifts for the performance metric normalized to the Intel® Xeon® 6980P follow for each benchmark: * gmx_water1536K_PME: ~1.76x * benchPEP: ~1.83x System Configurations CPU: 2P Intel® Xeon® 6980P (256 total cores) Memory: 24x 96 GB DDR5-6400 Storage: SOLIDIGM SBFPF2BU076T 7.68 TB NVMe Platform and BIOS: BIOS Options: SMT=OFF, SNC=3, HPC Workload Profile OS: rhel 9.4 5.14.0-427.16.1.el9_4.x86_64 Kernel Options: BOOT_IMAGE=(hd1,gpt2)/vmlinuz-5.14.0-427.16.1.el9_4.x86_64 crashkernel=1G-4G:192M,4G-64G:256M,64G-:512M rhgb mitigations=off tsc=nowatchdog nmi_watchdog=0 intel_pstate=disable processor.max_cstate=1 intel_idle.max_cstate=0 iommu=pt Runtime Options: cpupower idle-set -d 2 cpupower frequency-set -g performance echo 3 > /proc/sys/vm/drop_caches echo 0 > /proc/sys/kernel/nmi_watchdog echo 0 > /proc/sys/kernel/numa_balancing echo 0 > /proc/sys/kernel/randomize_va_space echo 'always' > /sys/kernel/mm/transparent_hugepage/enabled echo 'always' > /sys/kernel/mm/transparent_hugepage/defrag CPU: 2P 192-Core AMD EPYC™ 9965 (384 total cores) Memory: 24x 64 GB DDR5-6400 Storage: SAMSUNG MZWLO3T8HCLS-00A07 3.84 TB NVMe Platform and BIOS: RVOT1000C BIOS Options: SMT=Off NPS=4 Power Determinism Mode OS: rhel 9.4 5.14.0-427.16.1.el9_4.x86_64 Kernel Options: amd_iommu=on iommu=pt mitigations=off tsc=nowatchdog nmi_watchdog=0 Runtime Options: cpupower idle-set -d 2 cpupower frequency-set -g performance echo 3 > /proc/sys/vm/drop_caches echo 0 > /proc/sys/kernel/nmi_watchdog echo 0 > /proc/sys/kernel/numa_balancing echo 0 > /proc/sys/kernel/randomize_va_space echo 'always' > /sys/kernel/mm/transparent_hugepage/enabled echo 'always' > /sys/kernel/mm/transparent_hugepage/defrag Results may vary based on system configurations, software versions, and BIOS settings.

Endnotes

- 9xx5-261: The AMD EPYC 9965 has the highest theoretical Floating Operations per Node (FLOPS per node) in any 1P or 2P server. This is calculated by: $2.25 \text{ GHz Base Frequency} \times 192 \text{ Cores} \times 32 \text{ CPU instructions per cycle} \times 1 \text{ or } 2 \text{ CPUs per node}$. For: $2.25 \times 192 \times 32 \times 1 = 13824 \text{ GFLOPS}$ OR $2.25 \times 192 \times 32 \times 2 = 27648 \text{ GFLOPS}$. For reference, the Intel Xeon 6980P: $2.0 \text{ GHz Base Frequency} \times 128 \text{ Cores} \times 32 \text{ CPU instructions per cycle} \times 1 \text{ or } 2 \text{ CPUs per node}$. For: $2 \times 128 \times 32 \times 1 = 8192 \text{ GFLOPS}$ OR $2 \times 128 \times 32 \times 2 = 16384 \text{ GFLOPS}$. MLN-058B: Based on SPECrate®2017_fp_base on 07/06/2021, a server powered by two 8c AMD EPYC 72F3 CPU has scored 247, <http://spec.org/cpu2017/results/res2021q3/cpu2017-20210621-27506.html> with a per core score of 15.4 which is a higher per core floating point base performance score than any currently posted in any SPEC.org publication. SPEC®, SPECrate® and SPEC CPU® are registered trademarks of the Standard Performance Evaluation Corporation. See www.spec.org for more information.
- 9xx5-262: AMD EPYC 9175F has the highest 1P and 2P SPECrate®2017_int_base / Core as of November 04, 2025. 2P AMD EPYC 9175F (618 SPECrate®2017_int_base, 16 cores), 19.313 SPECrate®2017_int_base / Core, <https://www.spec.org/cpu2017/results/res2025q2/cpu2017-20250602-48622.html> 1P AMD EPYC 9175F (303 SPECrate®2017_int_base, 16 cores), 18.938 SPECrate®2017_int_base / Core, <https://www.spec.org/cpu2017/results/res2025q2/cpu2017-20250324-47116.html> Highest Intel Xeon 1P and 2P results for Reference: 2P Intel Xeon 6507P (216 SPECrate®2017_int_base, 8 cores), 13.500 SPECrate®2017_int_base C / Core, <https://www.spec.org/cpu2017/results/res2025q4/cpu2017-20250922-49721.html> 1P Intel Xeon 6349P (87.1 SPECrate®2017_int_base, 6 cores), 14.517 SPECrate®2017_int_base / Core, <https://www.spec.org/cpu2017/results/res2025q2/cpu2017-20250518-48122.html> SPEC®, SPEC CPU®, and SPECrate® are registered trademarks of the Standard Performance Evaluation Corporation. See www.spec.org for more information.
- 9xx5TCO-005: This scenario contains many assumptions and estimates and, while based on AMD internal research and best approximations, should be considered an example for information purposes only, and not used as a basis for decision making over actual testing. The AMD Server & Greenhouse Gas Emissions TCO (total cost of ownership) Estimator Tool - version 1.3, compares the selected AMD EPYC™ and Intel® Xeon® CPU based server solutions required to deliver a TOTAL_PERFORMANCE of 391000 units of SPECrate®2017_int_base performance as of November 21, 2024. This estimation compares upgrading from a legacy 2P Intel Xeon 28 core Platinum_8280 based server with a score of 391 (<https://spec.org/cpu2017/results/res2020q3/cpu2017-20200915-23984.pdf>) versus 2P EPYC 9965 (192C) powered server with a score of 3100 (<https://spec.org/cpu2017/results/res2024q4/cpu2017-20241004-44979.pdf>) and compared to a 2P Intel Xeon 64 core Platinum_8592+ (64C) based server with a SPECrate2017_int_base score of 1130, <https://spec.org/cpu2017/results/res2023q4/cpu2017-20231127-40064.pdf>. For additional details, see <https://www.amd.com/en/claims/epyc.html?q=9xx5TCO-005>. SPEC®, SPEC CPU®, and SPECpower® are registered trademarks of the Standard Performance Evaluation Corporation. See www.spec.org for more information. Intel CPU specifications at <https://ark.intel.com/>. FOR USE ONLY WHEN ENVIRONMENTAL DATA IS CITED: Environmental impact estimates made leveraging this data, using the Country / Region specific electricity factors from Country Specific Electricity Factors - 2024, and the United States Environmental Protection Agency Greenhouse Gas Equivalencies Calculator.

Endnotes

- VEN-001A: SPECrate®2017_int_base comparison based on AMD internal estimates for top of stack 2P 6th Gen EPYC CPU and 5th Gen EPYC measurements as of 10/30/2025. Preliminary performance estimates based on AMD engineering projections or measurements as of 10/30/2025 and subject to change.
- EPYC-058: Top 100 U.S. companies by revenue according to 2025 Fortune 500® list as of June 2, 2025. <https://fortune.com/ranking/fortune500/>. 'Fortune 100' refers to the top 20% ranked companies in the 2025 Fortune 500 list, published in June 2025. From Fortune Magazine. ©2025 Fortune Media IP Limited. All rights reserved. Used under license. Fortune and Fortune Media IP Limited are not affiliated with, and do not endorse products or services of Advanced Micro Devices, Inc.