

# Delivering superior throughput for EDA verification workloads

How Cadence Xcelium Parallel Logic Simulation and HPE Apollo systems featuring Marvell ThunderX2 Arm-based processors are changing the economics of semiconductor device verification

## Contents

|                                                              |    |
|--------------------------------------------------------------|----|
| Improving the economics of chip design .....                 | 2  |
| Verification challenges in EDA.....                          | 2  |
| IT landscape under pressure.....                             | 4  |
| Performance gains in verification have stalled .....         | 5  |
| Cadence Xcelium Parallel Logic Simulation.....               | 6  |
| Arm processor technology.....                                | 6  |
| Marvell ThunderX2 .....                                      | 6  |
| HPE and Cadence Xcelium solution .....                       | 7  |
| HPE Apollo 70 systems—Built for performance and density..... | 8  |
| Performance where it matters for EDA verification.....       | 9  |
| Evaluating HPE Apollo 70 systems at HPE .....                | 9  |
| A compelling solution for verification workloads.....        | 10 |
| Multicore architectures improve operational efficiency.....  | 10 |
| Evolving infrastructure for EDA data centers.....            | 11 |
| Delivering business value.....                               | 12 |
| Resources.....                                               | 12 |

## Improving the economics of chip design

Perhaps no industry is more competitive than modern electronics manufacturing and chip design. As consumers, we take it for granted that electronic devices continue to get faster, cheaper, and more capable with each generation. From smart watches to industrial controls to electronic heart-rate monitors, electronics manufacturers are challenged to build smarter, more complex devices leveraging system-on-a-chip (SoC) designs for an increasingly connected world. With the number of IoT devices forecast to grow to over 75 billion by 2025, consumers and manufacturers in the supply chain increasingly value durability, safety, battery life, and security including resistance to malware and hacking attempts.<sup>1</sup>

This level of change is unprecedented and poses a significant challenge for firms engaged in electronic design automation (EDA). Chip designers are faced with seemingly irreconcilable pressures such as shorter product design cycles, increasing complexity, increased requirements for quality, and continuous pressures on costs.

In this paper, we discuss the challenge of device verification, a key issue in EDA, and explain how new high-performance systems and software are promising to improve the economics of chip design enabling firms to innovate faster and create high-quality products.

## Verification challenges in EDA

Verification is about making sure that designs execute correctly and reliably. Most of us are familiar with Moore's Law that asserts device complexity as measured by transistor counts doubles, roughly every 18 to 24 months. With advanced SoC designs having tens of millions of gates, verification is perhaps the largest single challenge faced by device manufacturers. As engineers make changes to a design, they want to make sure that functionality doesn't regress.

Regression testing involves running and rerunning up to millions of tests developed along with the design to ensure that it functions as expected throughout the SoC design project. The process of taping out a design to create a photomask can cost millions of dollars, so designs need to be error-free before they are sent to a foundry. By some estimates, regression testing and verification account for roughly 80%<sup>2</sup> of the simulation workload in semiconductor design environments.

Verification becomes more difficult as designs become larger and more sophisticated, or incorporate new intellectual property (IP). As the number of registers and memory elements in a device increases (call this  $n$ ), the number of potential states to be modeled can increase exponentially ( $2^n$ ). For billion gate SoC in a mobile device,  $n$  may be 100,000,000 or larger.

Engineers at companies that design new chips perform extensive computer simulations on large-scale server farms, iterating and refining designs as flaws are discovered and corrected. Aside from the sheer size and complexity of the designs, verification engineers face business pressures as well:

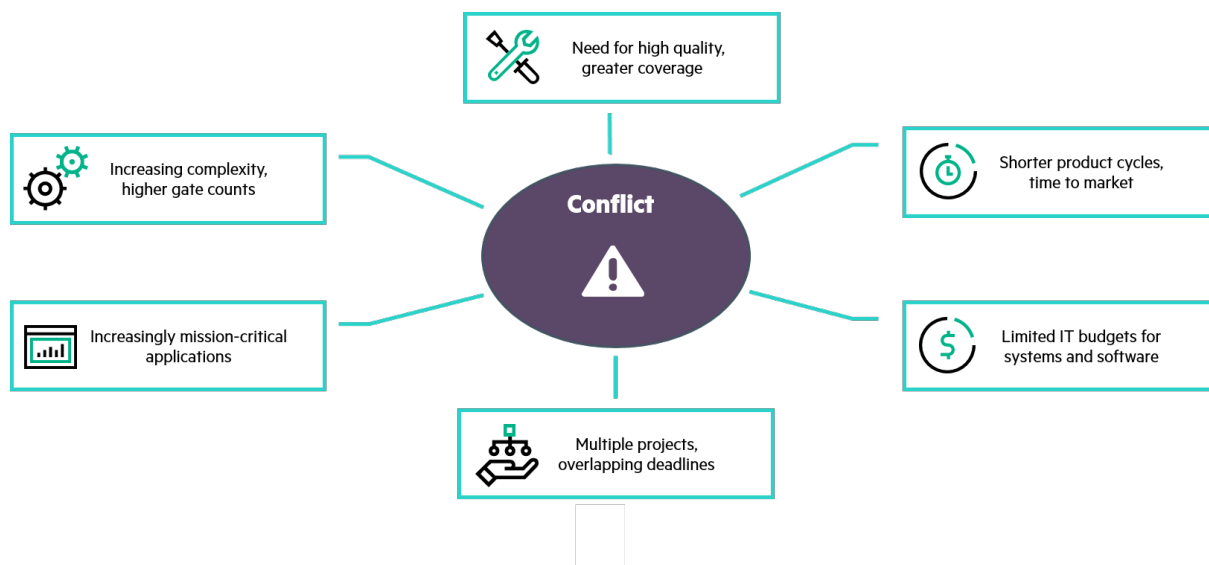
- Shorter product cycles and time to market
- More stringent reliability requirements as devices become mission critical
- Limited IT budgets for systems and EDA software

These pressures can conflict with one another as shown in Figure 1. For example, the need to get to market quickly may be inconsistent with ensuring high product quality and verification coverage.

<sup>1</sup> "Internet of Things (IoT) connected devices installed base worldwide from 2015 to 2025 (in billions)," Statista, 2018

<sup>2</sup> Based on a 2018 internal estimate from HPE VLSI design environment





**Figure 1.** Multiple challenges faced by design firms in getting electronic devices to market

Different types of devices carry different expectations regarding quality. For example, a mean time to failure (MTTF) of 30,000 hours may be acceptable for an inexpensive consumer device but unacceptable for a device such as an automotive engine controller, implantable cardioverter-defibrillator (ICD), or critical piece of avionics in an airliner. Mission- or life-critical devices, or devices that are hard to replace once deployed, will require more stringent verification.

Verification is performed across multiple disciplines. To make sure a device does what it's supposed to (functional verification), designers run extensive digital and analog simulations using different tools to verify various aspects of a design. Some of the more common Cadence tools used for different types of verification are listed in Table 1.

With digital abstractions, design engineers model the flow of data (functionality) from just a few clock cycles to billions of cycles. This depends on the functionality to be verified, for example, modeling how the internal state of a processor changes as each line of code is executed. For this type of verification, a key metric is the number of simulated cycles that can be performed per unit of elapsed time.

**Table 1.** Commonly used verification tools and run-time characteristics

| Category                    | Verification technology                                                                                                   | Verification type                                             | Description and verification run time                                                                                                                                                                                                                                                                                         |
|-----------------------------|---------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Digital abstractions</b> | Xcelium™ Parallel Logic Simulation, Palladium® Z1 Enterprise Emulation Platform, JasperGold® Formal Verification Platform | Gate-level simulations (GLS)                                  | Models consisting of 2 billion or more instances of logic gates; 10–100 simulated cycles per second (cps) for Xcelium and 500K–2M cps using Palladium with emulation; run times range from hours to weeks depending on model and simulator. JasperGold provides formal model checking, so performance is not measured in cps. |
|                             |                                                                                                                           | Register-transfer level (RTL)                                 | Model consists of up to 100 million lines of C-like code 100K–10K simulated cps; run times range from seconds to a week.                                                                                                                                                                                                      |
|                             | Xcelium Parallel Logic Simulation                                                                                         | Transaction-level model (TLM)                                 | Model consists of up to 1 million lines of C++-like code 10K–1M simulated cps; run times range from seconds to hours.                                                                                                                                                                                                         |
| <b>Analog abstractions</b>  | Spectre® Circuit Simulator and Spectre Accelerated Parallel Simulator                                                     | Transistor Level (SPICE)                                      | Model consists of analog primitives: resistors, capacitors, transistors, and others.                                                                                                                                                                                                                                          |
|                             |                                                                                                                           | Verilog-AMS/VHDL-AMS (VAMS)                                   | Model consists of behavioral code operating on voltage and current values in an analog simulator to solve the network.                                                                                                                                                                                                        |
|                             |                                                                                                                           | SystemVerilog Real Number Modeling (SVRNM)/Wired-Real (WREAL) | Model consists of C-like code executed on a digital simulator (Xcelium).                                                                                                                                                                                                                                                      |



Depending on the tool, the simulation type, and the complexity of the design, simulation rates can vary from 10 to 10,000 cycles per second (cps). For devices with clock speeds in the gigahertz range (sub-nanosecond cycle times), the compute requirements to simulate even a few hundred milliseconds of run time is enormous. Because of these variations, verification jobs can have run times ranging from seconds to weeks. Also, the resources required can vary depending on the simulator and the nature of the tests.

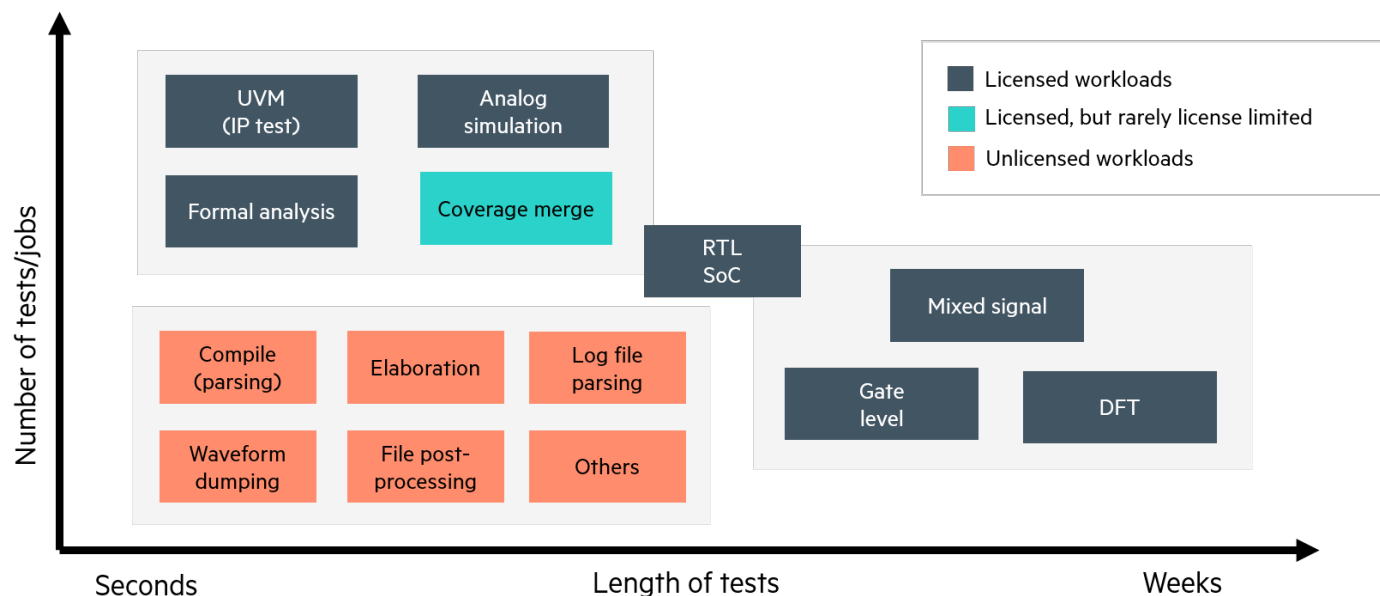


Figure 2. Wide variety of workloads in verification environments

Some simulations are smaller in size, run time, or event density (activity on each clock cycle) and run on a single core, whereas others can take advantage of multiple cores. This variety of tests is illustrated in Figure 2. Some tests require software licenses to run, whereas other jobs that are also critical to production, such as compilations, log file processing, and waveform dumping, do not.

Waveform dumping is a good example of the trade-offs faced by verification engineers. Engineers have the option of turning waveform dumping on or off when they run a simulation. For diagnosability, engineers would prefer to have them on because the data in waveform dumps is useful for troubleshooting. The resulting files can be large (tens of gigabytes) and the associated filesystem I/O activity reduces performance, leading to longer run times for regression tests. For this reason, engineers often turn the feature off to obtain better performance. If there is a problem, however, they'll need to rerun the regression with waveform dumping enabled before they can diagnose a problem. This process costs them more time.

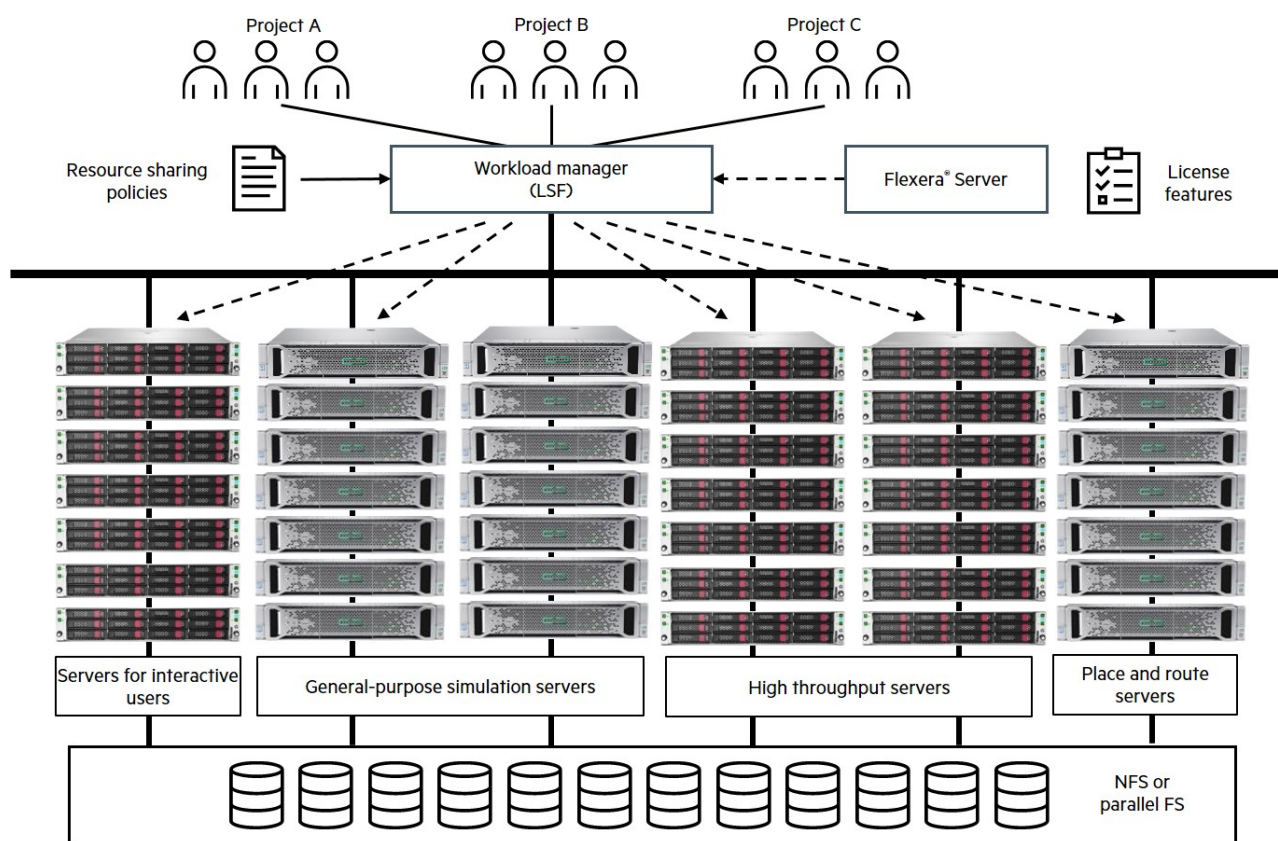
Verification engineers maximize coverage and quality using verification planning and management solutions such as Cadence vManager™ Metric-Driven Signoff Platform. vManager collects results from the Xcelium Parallel Logic Simulation and other simulation tools to track key metrics and analyze coverage in multiuser design environments.

## IT landscape under pressure

As is often the case, the business and technical pressures faced by chip designers translate into IT challenges. Verification is critical in getting quality devices to market, so device manufacturers compete based on the performance and cost-efficiency of their infrastructure.

Figure 3 illustrates a typical server farm environment. Design teams may work on multiple projects, with each design at different stages of development. Device simulation is essentially a high-performance computing (HPC) workload. Simulations run on large-scale server farms often comprised of hundreds of servers and thousands of cores. As with most HPC users, verification engineers have an insatiable need for capacity running hundreds of thousands to millions of jobs and keeping clusters as busy as possible.





**Figure 3.** Typical EDA environment with server farm supporting multiple project teams

A unique characteristic of EDA environments is that investments in software licenses and design talent is usually much greater than the cost of the computing infrastructure. While hardware costs are important, IC design teams are usually constrained by the number of software licenses. Before a simulation tool runs, it typically connects to a Flexera FlexNet Publisher™ (formerly FLEXlm™) server and requests a license for the software feature used. Some tools may checkout multiple licenses. Once a tool completes, license features are returned to a shared pool.

Because the number of shared license features is finite, efficient job scheduling is critical. Cluster administrators need to run scarce and expensive licenses on the fastest possible machine to minimize license checkout time, thereby boosting job throughput and overall efficiency.

Semiconductor firms typically run multiple types of machines optimized for different workloads. For example, place-and-route applications may benefit from machines with large amounts of memory. Multithreaded simulators such as Xcelium Parallel Logic Simulation will benefit from machines containing processors with many cores. Other simulators that run on a single core will benefit from processors that deliver fast single-threaded execution, and typically multiple jobs will be run simultaneously on a multicore processor.

Workload management plays an important role matching workloads to resources subject to various business priorities, policy constraints, and throttling jobs to avoid licenses being oversubscribed. IBM® Spectrum LSF is often used as a workload manager in EDA environments.

## Performance gains in verification have stalled

As designs get more complex, and workloads continue to increase, administrators are constantly looking for faster, high-capacity servers. Historically, server infrastructure was refreshed regularly to take advantage of performance improvements in new processors. A challenge for IT managers and verification engineers is that while workload demands continue to increase, the advances in processor performance traditionally predicted by Moore's Law gets stalled as clock speeds come up against hard limits dictated by physics.

Modern processors increasingly rely on higher core counts and on-chip parallelism for performance gains but exploiting these advantages requires changes to software. Despite the best efforts of technically savvy IT organizations experienced at wringing every ounce of performance from servers, performance gains from new generations of hardware have been marginal at best.

Fortunately, innovations in verification techniques, software, and high-throughput systems are promising to help bridge this gap and keep pace with increasingly compute and data-intensive verification workloads.



## Cadence® Xcelium™ Parallel Logic Simulation

A critical component in delivering fast and more thorough verification is the simulation engine itself. First-generation simulators were interpretive engines, parsing and executing a single simulation language. Second-generation simulators supported multiple verification languages (IEEE 1800 SystemVerilog, IEEE 1364 VHDL, and IEEE 1647 e as examples) and used compiled code to boost performance but were generally single threaded running on a single processor core.

Cadence Xcelium Parallel Logic Simulation is a new breed of third-generation simulator combining features of second-generation simulators with a new multithreaded engine that takes advantage of multicore architectures. The Xcelium Parallel Logic Simulation multicore engine is designed for fast SoC simulation and can deliver performance gains of up to 3X for RTL simulations, 5X for GLS, and 10X for design for test (DFT).<sup>3</sup> The single-core engine in Xcelium has also been redesigned and provides up to 1.5X speed up over the Cadence Incisive® Enterprise Simulator.<sup>3</sup>

Xcelium Parallel Logic Simulation together with the Cadence vManager Metric-Driven Signoff Platform provides design engineers with traceable verification quality, enabling the design of highly reliable systems suitable for deployment in harsh, mission-critical environments. Xcelium Parallel Logic Simulation provides additional improvements including incremental and parallel-build technology, process-based save/restart, and dynamic test reload allowing engineers to complete verification workloads faster, or run more extensive verifications within an allotted time.

## Arm® processor technology

With more than 125 billion Arm-based processors produced, and with partners shipping over 21 billion Arm-based processors in 2017 alone, Arm is the most widely used processor design by volume.<sup>4</sup> Historically, Arm has been most well-known for its leadership in low-power, performant designs for mobile, embedded systems, and consumer electronics. However, with changing compute requirements and the surge in data from billions of intelligent connected devices, Arm is well positioned to be the architecture of choice for next-generation cloud and networking infrastructure.

Arm licenses its processor designs to multiple semiconductor companies. Software built for Arm systems can run on any Arm processor regardless of who designed or manufactured the chip. For customers, this fosters a competitive, multivendor environment where companies can compete based on their implementation of the Arm instruction set. Presently, there are over 1590 Arm licensees.<sup>5</sup>

Over the past two years, Arm has emerged as a serious contender for high-end data center processors with designs from Fujitsu, Cavium® (now Marvell), and Ampere, who recently announced plans for Arm-based server processors.

Looking at the HPC space specifically, Armv8-A based server class chips are now shipping in volume and have garnered design wins at major supercomputing centers including the Catalyst UK program,<sup>6</sup> Sandia National Labs' Vanguard program,<sup>7</sup> and Japan's Exascale Post-K supercomputer. Arm-based servers are also being deployed in HPC clusters at multiple universities.

## Marvell ThunderX2®

Purpose-built for HPC applications, Marvell ThunderX2 second-generation Armv8-A server-level processor delivers socket-level performance in line with state-of-the-art high-end processors. It also delivers best-in-class memory-bandwidth and capacity, cache, and multithreading, all critical requirements for high-end simulators such as the Xcelium Parallel Logic Simulator.

EDA simulations typically require a lot of memory to accommodate large designs, and on traditional servers, cache, and memory bandwidth, they tend to be bottlenecks. If cache or memory bandwidth is not sufficient, performance can tail off rapidly when multiple jobs are run simultaneously on the same processor. This impacts performance, license utilization, and can impact the productivity of the entire design environment. These concerns are addressed in the Marvell ThunderX2 processor, whose specifications are shown in Figure 4.

<sup>3</sup> "Xcelium Parallel Logic Simulation," Cadence, 2018

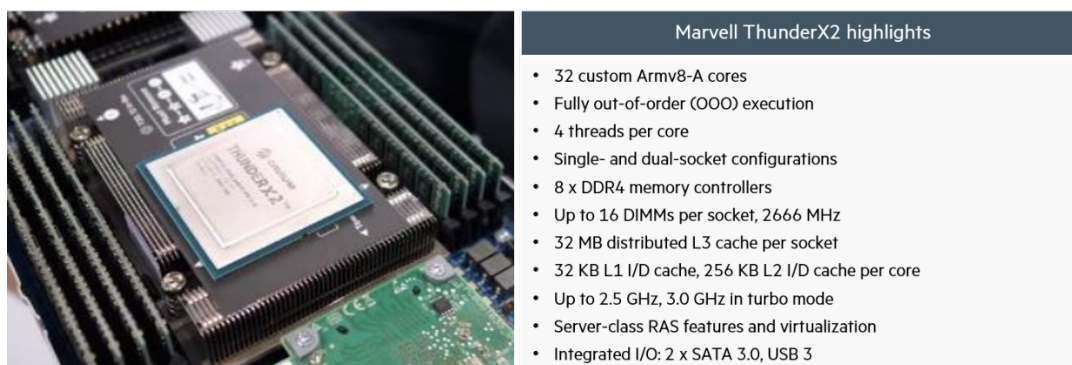
<sup>4</sup> "Arm Limited Q1 2018 Roadshow Slides," Arm, 2018

<sup>5</sup> Based on Arm Holdings FY2018 Q1 report

<sup>6</sup> Catalyst: Accelerating the Arm Ecosystem for HPC

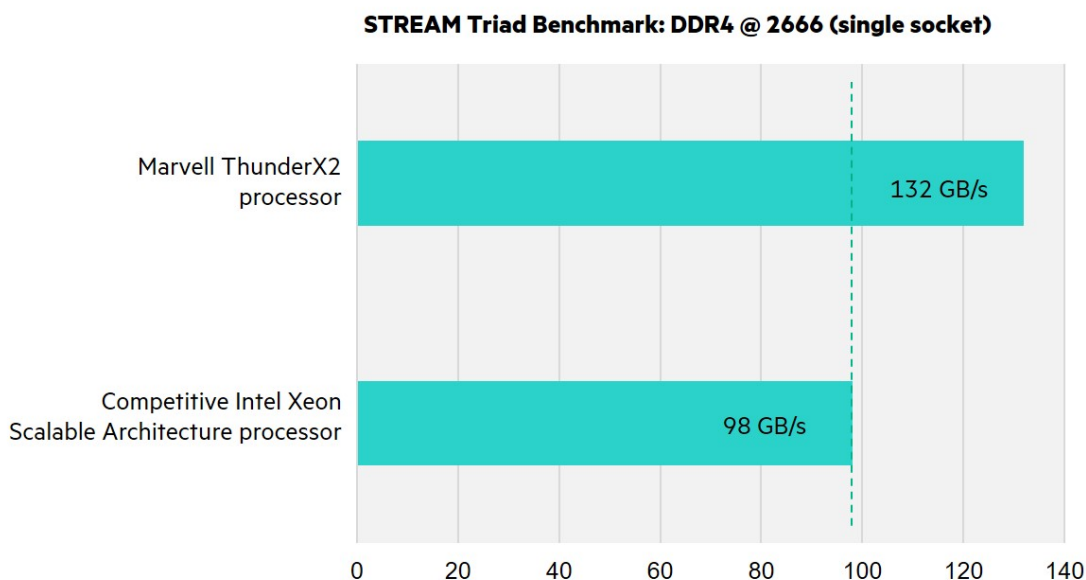
<sup>7</sup> Sandia National Labs' Vanguard program





**Figure 4.** Marvell ThunderX2 built for high-performance applications

With 32 MB of L3 cache and eight separate memory channels, ThunderX2 avoids these bottlenecks, allowing more simultaneous simulations to run on the same socket, thereby maximizing job throughput and reducing run times for regression tests. Figure 5 illustrates the ThunderX2 performance advantage for memory-intensive workloads showing roughly a 33% increase<sup>8</sup> over Intel® Xeon® Scalable Architecture processors on the STREAM Triad Benchmark.



**Figure 5.** ThunderX2—Delivers up to a 33% gain in memory bandwidth versus competitors

For simulators such as Xcelium Parallel Logic Simulation that can take advantage of multiple cores, the ThunderX2 provides enormous opportunities for parallelism with 32 cores and 4-way multithreading supporting up to 256 virtual cores in a 2-socket system.

## HPE and Cadence Xcelium solution

For engineers needing high-performance alternatives to existing verification solutions, HPE, Cadence, and Marvell have partnered to deliver a complete solution based on Marvell ThunderX2 processor. Cadence supports Xcelium Parallel Logic Simulation and other Cadence tools on HPE Apollo 70 Arm-based servers and is continuing to work with HPE, Marvell, and other Arm OEMs to enhance the solution. We shall discuss the HPE Apollo 70 system and how it delivers superior verification throughput and value with Xcelium Parallel Logic Simulation in the following.

<sup>8</sup> Based on comparison of dual-socket ThunderX2 based system with 32 cores per socket versus competitor: [nextplatform.com/2017/11/27/cavium-truly-contender-one-two-arm-server-punch/](https://nextplatform.com/2017/11/27/cavium-truly-contender-one-two-arm-server-punch/). Actual server level performance gains will depend on the speed of the memory subsystem

## HPE Apollo 70 systems—Built for performance and density

HPE Apollo 70 system, based on Marvell ThunderX2 processor, is built to deliver high levels of performance for a variety of HPC requirements. Designed by HPE, and available in 1U or 2U form factors, the dual-socket server delivers high-performance, high core counts, and high memory bandwidth required for compute- and memory-intensive workloads. HPE Apollo 70 system is a powerful, dense, cost-effective solution for EDA customers looking for more capable alternatives to existing server platforms.

HPE Apollo 70 system leverages the proven HPE Apollo 2000 architecture supporting up to four dual-socket servers in two rack units. This equates to up to 2 TB of RAM, eight Arm processors, and 256 multithreaded physical cores in just 2U of rack space making efficient use of data center real estate.

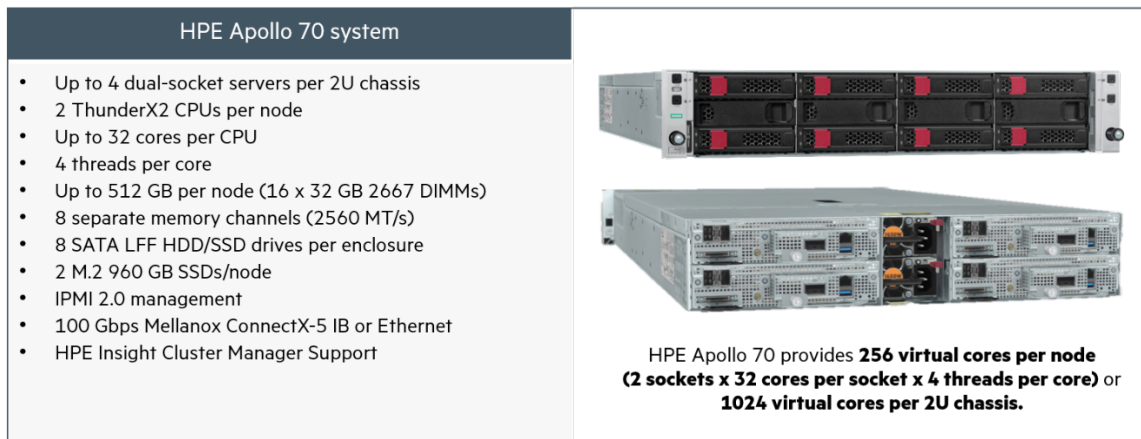


Figure 6. HPE Apollo 70 system that delivers superior performance and density

In EDA environments, large regression tests can run for weeks making system reliability, availability, and serviceability (RAS) important. Even where simulations are checkpointable, engineers don't want to take risks with long, running workloads. HPE Apollo 70 systems are designed for quality, reliability, and ease-of-management featuring redundant power, an Intelligent Platform Management Interface (IPMI) 2.0, and a variety of RAS features. Solutions such as HPE Performance Cluster Manager make server farm deployments easy to manage.

In addition to being a highly capable server, HPE Apollo 70 system has a broad portfolio that is software integrated, validated, performance optimized, and supported<sup>9</sup> by HPE and partners as shown in Figure 7.

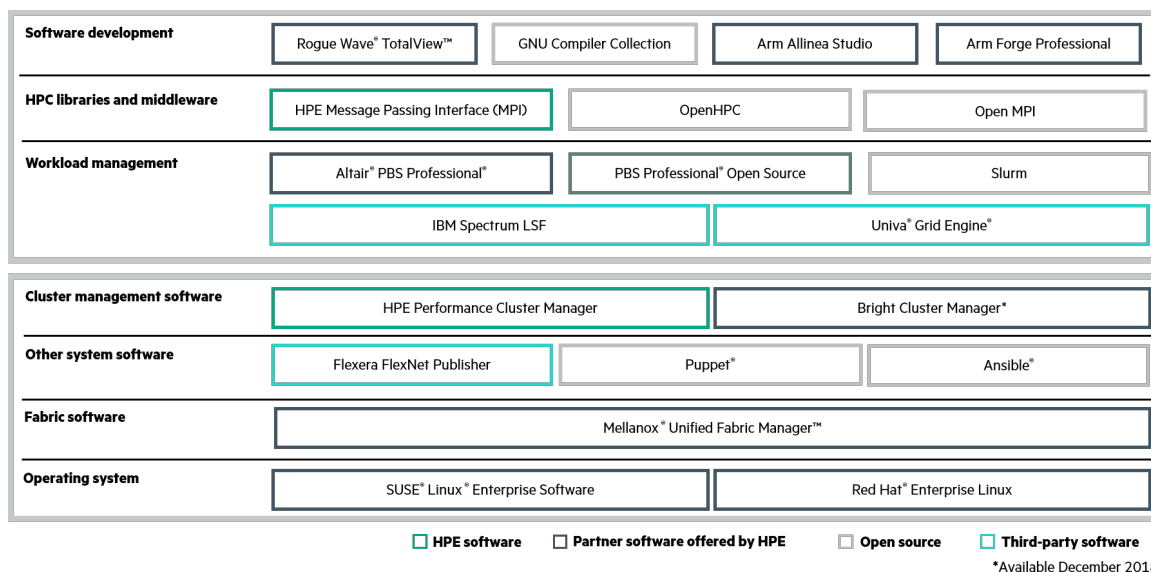


Figure 7. HPE Apollo 70 system software ecosystem

<sup>9</sup> Support excludes open source software components and third-party supported products

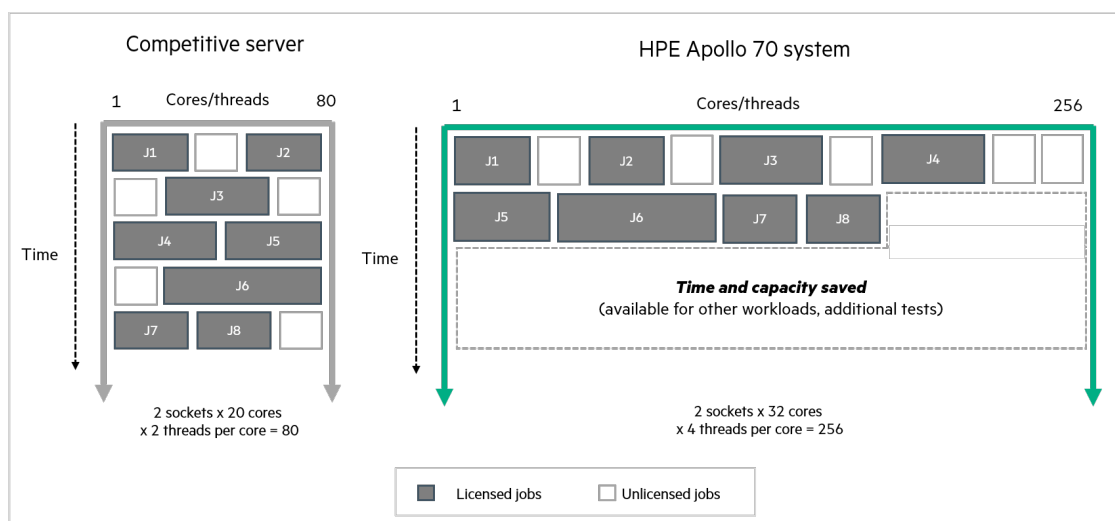


## Performance where it matters for EDA verification

For verification workloads, the ability to manage a variety of workloads and deliver high throughput is important. HPE Apollo 70 system offers a balanced design supporting multiple, simultaneous, single-threaded jobs as well as multithreaded jobs. HPE Apollo 70 system addresses common performance bottlenecks such as cache capacity, memory bandwidth, and multicore scalability.

A key advantage of HPE Apollo 70 system is its large number of cores and threads per system. From the operating system's perspective, HPE Apollo 70 system presents up to **3X the number of virtual cores** compared to other HPE industry-standard servers.<sup>10</sup> For workloads that are single threaded and able to take advantage of only a single core or thread, this means that more jobs can run simultaneously on the same server. For jobs that are multithreaded, the same is true—more multithreaded jobs can run at the same time leading to better throughput and efficiency.

The benefit of this multicore, multithreaded architecture is illustrated in Figure 8. By providing more virtual cores, administrators can either run more jobs in the same amount of time or complete larger regression tests faster even though individual job run times may be higher. The jobs pictured in Figure 8 can be parallel (multicore) jobs or job arrays comprised of many single-threaded jobs.



**Figure 8.** With support for more simultaneous jobs, HPE Apollo 70 system delivers higher throughput

For simulators that write intermediate data such as waveform dumps or log files to storage, HPE Apollo 70 system provides fast, integrated solid-state drives (SSDs) and supports up to eight additional large form factor (LFF) SSDs or hard disk drives per enclosure. This flexibility enables systems to be configured to support large amounts of local scratch storage minimizing run times and license checkout time. High-performance network interfaces provide fast access to a variety of storage subsystems including NAS filers, SAN storage, or fast parallel filesystems such as WekaIO Matrix<sup>11</sup> for maximum flexibility.

## Evaluating HPE Apollo 70 systems at HPE

To demonstrate the value of HPE Apollo 70 systems for EDA workloads, HPE has been conducting internal tests running Xcelium Parallel Logic Simulation in its internal VLSI semiconductor design environment. These tests involve running both multicore and multiple single-core simulations on dual-processor HPE Apollo 70 systems and comparing results to systems in existing server farm of HPE.

<sup>10</sup> Based on internal benchmark tests, 2018

<sup>11</sup> [WekaIO](#) is an HPE business partner

While public benchmarks are not yet available, early results are promising and show that high-throughput Arm-based HPE Apollo 70 systems can deliver higher regression throughput when running multiple simultaneous simulations. HPE Apollo 70 system has shown that it can complete more regression tests per day running multiple simultaneous instances of Xcelium Parallel Logic Simulation per server. Improved throughput results from the Arm-based ThunderX2's multicore design, efficient 4-way hardware threading, and its large cache and fast memory subsystem.

Based on internal price comparisons, HPE Apollo 70 system are up to **19% less expensive**<sup>12</sup> than similarly configured HPE industry-standard servers. HPE is working with partners including Cadence and Marvell to make HPE Apollo 70 systems available in its corporate benchmark centers for customer and partner benchmarks.

## A compelling solution for verification workloads

In production environments, server farm administrators constantly balance priorities, deadlines, and licenses scheduling licensed and unlicensed jobs on the most appropriate hardware. Investments in software can be considerable, so administrators strive to use license features as efficiently as possible.

A key benefit of the HPE Apollo 70 system is **administrators can potentially run more simultaneous jobs per physical server and reduce total regression time** per dollar invested in infrastructure.<sup>13</sup> At a cluster and workload management level, this translates into more **job slots per rack** enabling **higher job throughput and improved productivity**.

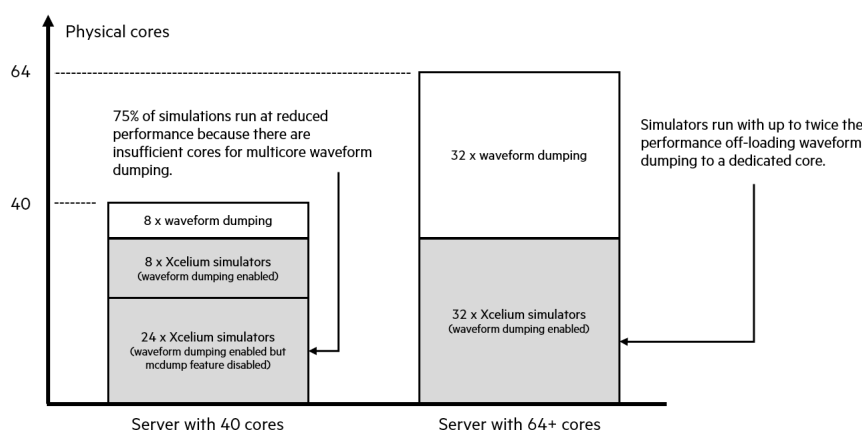
## Multicore architectures improve operational efficiency

Earlier we discussed the importance of waveform dumping when running simulations. With earlier single-threaded Cadence simulators, running with waveform dumping enabled could cut performance roughly in half because the processor core running the simulation spent much of its time waiting on filesystem I/O.

Engineers often turn off waveform dumping to improve performance, but this can be counterproductive and reduce overall efficiency. If there is a problem with a test, the regression may need to be run again to gather sufficient debugging information to identify and correct an error costing time and impacting schedules.

Xcelium Parallel Logic Simulation can exploit multicore parallelism for waveform dumping using an `-mcdump` setting. This setting can result in up to twice the performance when waveform dumping is enabled because filesystem writes are performed asynchronously by a separate thread (requiring a separate core) so that the single-core simulator is not blocked.<sup>14</sup>

Figure 9 illustrates why additional cores are beneficial for waveform dumping. Consider a case where two different servers, one with 40 physical cores and another with 64 physical cores are each running 32 simultaneous instances of the Xcelium single-core simulator with waveform dumping enabled.



**Figure 9.** Example of comparing systems running multiple Xcelium single-core simulations with waveform dumping enabled

<sup>12</sup> Based on internal USD list price comparisons

<sup>13</sup> Based on internal HPE benchmark tests and agreed to by Cadence who reviewed the benchmark tests, 2018

<sup>14</sup> Based on a Cadence internal estimate, 2018

On a system with 64 physical cores like HPE Apollo 70 system, 32 simulator instances with 32 companion waveform dump threads can run simultaneously without oversubscribing cores. This means that 32 Xcelium instances can run at full throughput without blocking on filesystem I/O. On the system with fewer cores, only eight simulations can run with multicore waveform dumps enabled. This means that 24 simulations (75% of the workload) will be up to twice as slow reducing overall regression throughput.

The dynamics are complicated of course, but for multicore workloads, unlicensed jobs, jobs that are not license constrained, and high-throughput HPE Apollo 70 systems running Xcelium Parallel Logic simulation can be important additions to EDA data centers. From a system and application software standpoint, there is minimal impact in adopting HPE Apollo 70 Arm-based systems since ISVs tend to price Arm ThunderX2 offerings similar to offerings on other high-end processors. Administrator and user skillsets should be fully portable to Arm systems since the software tools and interfaces are the same.

Organizations should probably not count on significant savings related to power and cooling, as the ThunderX2 power draw is comparable to other high-end processors. However, HPE Apollo 70 systems offer an exceptional degree of rack density (over 5000 cores per 42U rack or 20,000 threads at four threads per core) and excellent cluster management software, so depending on the systems you are comparing to, HPE Apollo 70 systems may provide additional benefits.

## Evolving infrastructure for EDA data centers

Keeping pace with increasing verification requirements for large designs will depend on parallel, multicore simulators, and infrastructure that can support higher levels of job level parallelism and throughput.

As a new entrant to the EDA infrastructure game, HPE Apollo 70 systems have made an impressive debut delivering **3X the virtual cores at a lower price, and superior job throughput and price-performance for large regression tests** involving multiple simultaneous simulations.<sup>15</sup>

While software licensing is always a factor, for verification workloads being able to take advantage of multicore parallelism, HPE Apollo 70 systems with Xcelium Parallel Logic Simulation provide greater flexibility to manage regression time and cost trade-offs. Results will vary by customer depending on the nature of workloads, mix of licensed and unlicensed workloads, whether they are using features such as multicore waveform dumps, and such.

As clock speeds and single-core performance plateau, parallel simulators and multicore architectures represent the best opportunity to improve verification performance, both areas where Xcelium Parallel Logic Simulation and HPE Apollo 70 systems excel.

Users can expect that application performance will only continue to improve as:

- Cadence continues to optimize its applications for the Arm architecture.
- The Arm architecture continues to evolve.
- Chip manufacturers compete on data-center-class Arm designs.
- Compilers and development tools continue to improve.
- ISVs optimize and tune their software for high-throughput Arm systems.
- Systems companies such as HPE continue to deliver more capable servers.

<sup>15</sup> Based on HPE and Cadence internal estimate benchmark tests, 2018



## Delivering business value

Hardware designers grapple with multiple challenges including increasing design complexity, time-to-market pressures, and the need for more thorough verification as new applications in healthcare, IoT sensors, and autonomous vehicles demand greater reliability and safety from SoC designs.

The Xcelium Parallel Logic Simulation and high-performance HPE Apollo 70 systems powered by Marvell ThunderX2 processors provide an important new tool and added flexibility for EDA firms needing to improve the productivity and efficiency of chip design environments.

By deploying HPE Apollo HPE 70 systems with Xcelium Parallel Logic Simulation, customers can:

- Speed verification and regression tests to meet time-to-market pressures
- Improve product quality and meet increasing reliability requirements with the capacity to run additional verification workloads within available time frames
- Maximize limited IT budgets by deploying more cost-effective, higher throughput systems delivering improved farm utilization and better engineering productivity

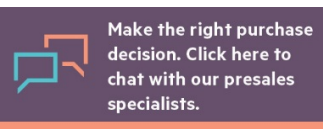
## Resources

To learn more about the Cadence Xcelium Parallel Logic Simulation for Arm systems, visit [cadence.com/go/verificationsuiteeco](https://cadence.com/go/verificationsuiteeco)

To learn more about the Marvell ThunderX2 processor, visit [cavium.com/product-thunderx2-arm-processors.html](https://cavium.com/product-thunderx2-arm-processors.html)

## Learn more at

[hpe.com/servers/apollo70](https://hpe.com/servers/apollo70)



 **Share now**

 **Get updates**

---

© Copyright 2018 Hewlett Packard Enterprise Development LP. The information contained herein is subject to change without notice. The only warranties for Hewlett Packard Enterprise products and services are set forth in the express warranty statements accompanying such products and services. Nothing herein should be construed as constituting an additional warranty. Hewlett Packard Enterprise shall not be liable for technical or editorial errors or omissions contained herein.

Arm is a registered trademark of Arm Limited. Intel Xeon is a trademark of Intel Corporation in the U.S. and other countries. Red Hat is a registered trademark of Red Hat, Inc. in the United States and other countries. Linux is the registered trademark of Linus Torvalds in the U.S. and other countries. All other third-party marks are property of their respective owners.

a00058962ENW, November 2018

