

## INTELLIGENCE Q&As

# IT sustainability: achieving more transactions per megawatt hour

This Q&A brings together contributor questions and expert answers from an Uptime webinar on the transactions per megawatt hour metric. Topics discussed include: power management settings, server refresh rates and sustainability, greenhouse gas emissions accounting, IT metrics and utilization rates.

These topics are explored in many Uptime Intelligence publications and webinars. Some of the questions arising from these are answered here. Questions have been altered and merged for clarity and consistency.

---

## Power management settings

**Q.** Is there data and / or graphs demonstrating the latency / performance effects of enabling power management settings?

**A.** The Uptime Institute report *Server energy efficiency: five key insights* discusses the performance impact of applying power management using SERT data. It varies according to workloads and utilization states for each machine. The impact on performance rarely exceeds 6% at any utilization point, while power savings range from 6% to 21% in active mode.

The Green Grid has published a publicly available white paper titled *Trade-offs of processor power management functions of servers (WP 84)*, which discusses the latency impacts of the power management mode:

- **P-states.** Active power reductions and transitions between P-states increase latency to between 20 and 30 microseconds for second-generation Intel Xeon-based servers. The P-state latency delays have decreased with each generation of CPU.
- **C-states.** When there is reduced power due to no workload being present, the latency impact of the C6 (highest energy savings in idle) setting is approximately 80 microseconds.

Intel and AMD publish data on the latency impacts of power management mode in both P-states and C-states.

---

Uptime Institute Intelligence is an independent unit of Uptime Institute dedicated to identifying, analyzing and explaining the trends, technologies, operational practices and changing business models of the mission-critical infrastructure industry. For more about Uptime Institute Intelligence, visit [uptimeinstitute.com/ui-intelligence](https://uptimeinstitute.com/ui-intelligence) or contact [research@uptimeinstitute.com](mailto:research@uptimeinstitute.com). Intelligence reports do not represent Uptime Institute's position and do not constitute investment or technical advice

---

**Q. Can you explain the risk associated with implementing power management settings?**

**A.** Operators are reluctant to implement power management settings for two reasons: an increase in latency (reduction in response time) and the small monetary savings that result when measured against the cost of breaching the service level agreement (SLA). A breach of the SLA can cost hundreds of thousands of dollars, tarnish an organization's reputation and, in extreme cases, result in the data center staff responsible being subject to disciplinary action. Enabling power management can increase latency by 20 microseconds to 80 microseconds, depending on the power management settings. This latency increase may appear small but can be unacceptable for latency-critical workloads (including financial trading). IT operators are concerned that an increase in latency and a reduced response time for latency-sensitive applications (or during periods of high workloads) will cause unacceptable degradation of system performance or user experience, potentially violating the SLA. Operators can manage these concerns if workloads are tested for compatibility with power management settings and assigned to servers with power management enabled. Operators can then place the remaining workloads on servers managed for the highest performance levels. Uptime Institute is aware of several IT operators that successfully deploy power management in this way.

**Q. Are SERT scores dependent on a server's BIOS setting when power management settings have been set to reduce the frequency of the CPU at lower utilization? For example, would a 64-core server with low utilization consume significantly more energy if BIOS settings are set to "performance" mode (disabled power management)?**

**A.** Typically, the BIOS and firmware functions should be set to support power management functions. Typically, the power management functions and overall profile are set in the operating system or hypervisor. If a hypervisor is deployed, power management functions should be assigned there.

A server with performance settings will have an idle-to-maximum power ratio of 60% to 80%. A server with full P-state and C-state settings (to C6) will have an idle to maximum power ratio of 20% to 40% and a 12.5% utilization rate to 100% utilization power ratio of 30% to 50%. A server operating consistently below 20% utilization can save significant energy and cost by enabling of power management. A better option, however, is to consolidate workloads from multiple servers with low utilization rates to run on a few servers and remove the excess servers from service.

---

## Server refresh cycle and sustainability

**Q. Extending or slowing down hardware refresh cycles reduces Scope 3 (embedded) emissions at the cost of potential reductions in Scope 2 emissions from more efficient servers. Has Uptime identified a break-even point for reducing Scope 2 versus Scope 3 emissions for refresh cycles?**

**A.** It is a complex equation and every case is different. Factors include energy savings and carbon emissions reductions (a function of the electricity grid emission factor) achieved with new and more efficient servers and the electricity grid emission factor (metric tons of CO<sub>2</sub> per megawatt-hour), and the embedded CO<sub>2</sub> emissions (carbon emitted to manufacture the servers) value of the new servers (Scope 3). **Table 1** shows the importance of the grid emission factor to the relative carbon benefits of a shorter or longer refresh rate.

**Table 1 Example showing benefits of 4- and 6-year refresh cycles**

|                                                          |                             | Embedded emissions<br>value of purchased server |                                      | Cost of servers     |                      | 12-year<br>server energy<br>use (MWh) | 12-year energy cost |                | 12-year CO <sub>2</sub><br>emissions (MT CO <sub>2</sub> ) |                                    |
|----------------------------------------------------------|-----------------------------|-------------------------------------------------|--------------------------------------|---------------------|----------------------|---------------------------------------|---------------------|----------------|------------------------------------------------------------|------------------------------------|
|                                                          |                             | 0.5 MT CO <sub>2</sub> e<br>/ server            | 1.2 MT CO <sub>2</sub> e<br>/ server | \$6,000 /<br>server | \$10,000 /<br>server |                                       | \$70 /<br>MWh       | \$150 /<br>MWh | 0.5 MT<br>CO <sub>2</sub> /<br>MWh                         | 0.2 MT<br>CO <sub>2</sub> /<br>MWh |
| 30% generation-to-generation server capacity improvement | Number of servers purchased |                                                 |                                      |                     |                      |                                       |                     |                |                                                            |                                    |
| 4-year refresh cycle                                     | 182                         | 91                                              | 218                                  | 1,092,000           | 1,820,000            | 2,400                                 | 168,000             | 360,000        | 1,200                                                      | 480                                |
| 6-year refresh cycle                                     | 136                         | 68                                              | 163                                  | 816,000             | 1,360,000            | 2,600                                 | 182,000             | 390,000        | 1,300                                                      | 520                                |
| Difference                                               | 46                          | 23                                              | 55                                   | 276,000             | 460,000              | -200                                  | -14,000             | -30,000        | -100                                                       | -40                                |
| 75% generation-to-generation server capacity improvement | Number of servers purchased |                                                 |                                      |                     |                      |                                       |                     |                |                                                            |                                    |
| 4-year refresh cycle                                     | 126                         | 63                                              | 151                                  | 756,000             | 1,260,000            | 1,700                                 | 119,000             | 255,000        | 850                                                        | 340                                |
| 6-year refresh cycle                                     | 101                         | 51                                              | 121                                  | 606,000             | 1,010,000            | 2,000                                 | 140,000             | 300,000        | 1,000                                                      | 400                                |
| Difference                                               | 25                          | 12.5                                            | 30                                   | 150,000             | 250,000              | -300                                  | -21,000             | -45,000        | -150                                                       | -60                                |

Calculations assume starting with 100 servers and utilization is constant at each refresh

Refresh rate benefits depend on assumptions of workload improvement, energy costs, embedded carbon levels, and MT CO<sub>2</sub>/MWh

Longer refresh rates reduce capital costs

Shorter refresh rates reduce energy use and associated CO<sub>2</sub> emissions, as long as the refreshed servers have the same or increased utilization rates

UPTIME INTELLIGENCE 2023

uptime  
INTELLIGENCE

If the grid has a low emissions factor, there are carbon benefits to a longer refresh cycle, by avoiding the embedded carbon (Scope 3) of purchasing new servers. If the emission factor is high, carbon benefits will result from a shorter refresh cycle.

The clear message of the analysis is that longer refresh cycles offer significant savings in capital costs (although not necessarily carbon savings). Ultimately, the refresh rate will be set based on a business assessment of the potential benefits of performance and efficiency improvements against the cost of capital or cloud expenses. Typically, a longer refresh cycle is better.

The *Uptime Institute Global Data Center Survey 2023* found that data center operators, including cloud providers, have been reducing the frequency of their server refreshes. The slower refresh rate is influenced by a need to reduce capital expenditures, supply chain challenges and the fact that higher-performance servers are not required for most workloads.

As a secondary observation, the full benefit of a server refresh can only be realized if workloads running on multiple servers are consolidated onto fewer, more efficient servers. Refreshing at a one-to-one ratio will, at best, improve data center efficiency by 5% and will waste the additional work capacity available in new servers.

**Q. When upgrading to more efficient servers, what is the environmentally responsible way to decommission old servers?**

**A.** Operators are expected to manage their end-of-life equipment responsibly, maximizing the refurbishment and reuse of server, storage and network products, reusing components for spare parts, and reusing and recycling critical materials. Efforts should be made to minimize the percentage of end-of-life equipment sent to landfills. There is an ecosystem of reputable, certified product recyclers / reclaimers to manage end-of-life products. Operators should periodically assess their recyclers / reclaimers to validate that they are handling equipment in accordance with their contract requirements.

Most countries, states and / or provinces have laws to encourage or enforce recycling and reuse. Many jurisdictions have specific regulations governing the waste management, transport and the cross-boundary shipment of IT. Data center operators should understand and comply with restrictions governing end-of-life products from their data center facilities.

## Greenhouse gas emissions accounting

**Q.** When Uptime refers to greenhouse gas Scope 3 emissions, are you referring to the Greenhouse Gas Protocol? Do you also use ISO 14064-1?

**A.** ISO 14064 is the Greenhouse Gas Protocol's corporate standard incorporated into an ISO standard.

## IT metrics

**Q.** In Uptime's view, is the industry any closer to replacing power usage effectiveness (PUE) with a more meaningful data center energy efficiency metric such as transactions per watt, bytes stored per watt and bits transmitted per watt?

**A.** Introduced in 2007, PUE serves as the de facto data center efficiency metric. PUE is limited: it only measures the percentage of data center energy consumption required to support the facility's infrastructure. It does not assess IT efficiency in a meaningful way.

The value of the PUE depends on IT equipment utilization, facility design and age, system redundancies and local climate conditions. It tracks year-to-year facility system improvements but is not a good metric for comparing data center performance. It will continue to serve as a valuable metric to assess the performance of the facility operations at a data center.

The industry is moving slowly toward adopting a work-per-energy metric for data center operations. This is a complex undertaking and will require five to 10 years to develop a practical methodology and a standardized metric to measure work per watt.

Several groups are working to develop and establish a work per watt metric. Initial efforts will focus on estimating the total and utilized capacity of the server and storage equipment. Those estimates will then be combined with data center energy use measurements to calculate the work per watt. Over time, the methodology will likely be refined to match server configuration information, such as the central processing unit part number and memory capacity, and to average out utilization and power measurements from the data center to generate a work per watt value.

A conclusive methodology and the installation of data center monitoring and management systems to collect real-time operating data will take several years to finalize. In the meantime, data center operators should focus on increasing the average capacity utilization of their servers and consolidate workloads during a refresh cycle to do more work with fewer, more efficient servers.

**Q.** What are Uptime's recommendations for bare metal servers with low utilization rates? Is power management the best (or only) option for this scenario?

**A.** The procurement of bare metal servers in the public cloud is subject to the same considerations as the procurement of individual servers for an enterprise data center. The IT operator needs to maximize the available workload capacity on the server, taking into account the average and maximum workload demands of the application(s) running on the server — and the performance and response time requirements of the applications.

Bare metal servers should be assessed for ways to increase the use of virtual machines and containers, to efficiently consolidate workloads and to optimize and / or maximize the use of the deployed hardware. Integrating a group of bare metal servers

into a cloud-type architecture may help increase utilization. Additionally, it may be advantageous to use platform or infrastructure as a service to host the applications and improve the utilization of processor, memory and storage capacity.

Once hardware utilization is optimized, consideration can be given to the enablement of power management, based on the ability of the applications to tolerate latency delays.

This discussion highlights that optimizing IT system performance demands knowledgeable system administrators, robust workload management and placement tools, constant diligence and the management's commitment to achieving optimal environmental performance and value from IT assets..

---

**Q. How many data centers still use diesel and how many are able to use geothermal or a different renewable energy source?**

**A.** Nearly all data center backup power systems rely on diesel generators. Some data centers have deployed natural gas-powered generators but these generators can be more finicky on start-up and fuel supplies cannot easily be stored on-site, creating a dependence on the natural gas distribution system.

Geothermal is not a feasible energy source for backup power. The initial cost of a geothermal power system demands that the system be a primary power source operating at 90% or better availability.

Renewable energy, primarily solar with batteries, can serve as the backup generation for edge deployments. For traditional facilities, on-site renewable energy plus storage are not economical, reliable or feasible to carry the entire facility's electrical load in the event of a power outage.

To make diesel generators more sustainable, some operators have embraced hydrotreated vegetable oil (HVO) fuel to replace diesel. HVO is a second-generation biofuel, chemically distinct from "traditional" biodiesel. It can be manufactured from waste food stocks, raw plant oils, used cooking oils or animal fats, and either blended with petroleum diesel or used in 100% concentrations. HVO reduces carbon dioxide emissions by up to 90%, particulate matter by 10% to 30% and nitrogen oxides by 6% to 15% when compared with petroleum diesel (see *Vegetable oil promises a sustainable alternative to diesel*).

---

**Q. Can tools capture and analyze power demand, capacity utilization and transaction per second data in real time?**

**A.** The data presented in Uptime Intelligence webinars and reports are primarily generated from published Server Efficiency Rating Tool (SERT) measurement data. Average central processing unit (CPU) utilization values are assumed and matched to CPU-specific SERT capacity and power measurements to estimate utilized workload capacity and average power demand.

For operational or regulatory metrics, CPU capacity utilization data and power data could be captured in real-time by data center management software, likely as averages in 15- or 60-minute increments. This is then matched with the SERT work capacity data to calculate utilized work capacity and work per watt.

There are Data Center Infrastructure Management (DCIM) and IT infrastructure or operations management software packages that can collect and aggregate the measured utilization capacity and power data and match it to IT equipment configuration data. Although DCIM deployment is creeping up, most data centers do not currently have software deployed to collect and aggregate average IT equipment utilization and power measurements.

**Q.** SERT tells us the potential number of transactions per second, but the actual number of transactions per second depends on how an application uses the underlying hardware. What is the best way to measure or estimate the exact number?

**A.** Measuring transactions per second on operating servers is feasible but not practical. A system administrator understands the work capacity of the servers in their data centers and uses utilization as a proxy for tracking capacity utilization. For example, if a server is regularly approaching 80% or higher utilizations (or is experiencing spikes at or above 100% capacity) an application will require additional hardware equipment to support the workload (or the workload will need to be redistributed).

Efficient operation of the IT infrastructure demands attention to the average and maximum utilization of individual servers and the entire server fleet. Rather than running all the work on one server with high utilization, work often runs on multiple servers at low utilization to ensure the resiliency and reliability of the system. A critical application may run on two (rather than one) servers in three availability zones in a public cloud or in mirrored data centers to provide a hot backup if the enterprise is to meet its business commitments. These infrastructure systems are not the most efficient, but they are sometimes needed to meet reliability and resiliency requirements.

To make a comparison of transactions per second and work per watt or kilowatt hour, it is necessary to:

- Collect the average utilization of the systems supporting a given workload or set of workloads.
- Determine the workload by multiplying the maximum SERT work capacity for the servers.
- Divide the workload value by the energy use to calculate the work per watt.

---

**Q.** What can you share about the upcoming SERT3.0?

**A.** The SPECpower Committee is evaluating an expansion of the Server Efficiency Rating Tool (SERT). Performance power tests and metrics are being considered for AI-focused servers, high-performance computing or graphics processing unit servers and heavy storage servers. Each server type is expected to have a specific efficiency rating tool that executes representative workloads, generating a relative work per watt efficiency score.

---

## Utilization rates

**Q.** What recommendations (e.g., adopting additional virtualization or the use of data center infrastructure management software) would Uptime give operators who want to increase server utilization rates? What are the challenges?

**A.** There are several considerations to increase utilization.

First, it is necessary to have the appropriate data center infrastructure, IT infrastructure or IT operations management software to collect and track, at a minimum, average CPU and memory utilization over time. The workload activity information for each application / VM (virtual machine) / container is essential to assist in combining VMs on a server if utilization, performance and reliability are to be optimized.

There are also workload placement packages that can collect this information and combine it with reliability and availability requirements for each VM. These packages will provide recommendations for consolidating workloads onto the minimum number of servers to ensure reliability and performance parameters are met. These software

packages can also help managers reduce the number of servers required to support a fixed workload by 10% to 35%, along with all the associated additional benefits of energy and space savings.

Analyzing how the workload is divided between batch jobs and office or enterprise applications is crucial. If batch jobs account for a low percentage of the workload, they can be scheduled to use available compute capacity and fill in periods of low utilization on individual machines. An important control metric for batch jobs is to track aborted jobs, which are wasted CPU cycles. These should be investigated, and adjustments made to the batch job to ensure it runs to completion.

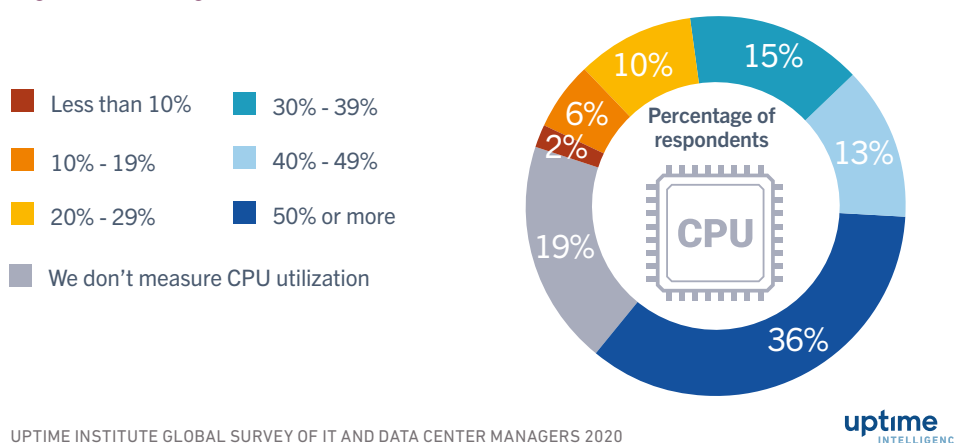
The journey to higher utilization levels takes time. The IT operations team should begin by finding one or two relatively "easy" projects. The initial projects will provide meaningful learning regarding the processes, techniques and tools needed to combine VMs. During the server refresh process, workloads can be incorporated into fewer, higher work-capacity servers.

### Q. Why are server utilization rates so low?

**A.** Operators and IT managers are highly risk averse and incentivized to prioritize system reliability, resiliency and performance over efficiency. Traditionally, a single application was placed on a single server, often resulting in low server utilization. The recent increases in containerization and virtualization, and improvements in hardware and software systems, have enabled the average server utilization rates to be increased over the last decade.

Despite these increases, average server utilization remains low (see **Figure 1**), providing ample opportunity for IT managers to make significant efficiency gains and reduce energy and water use. Increasing server utilization can save up to 50% through reduced equipment counts and capital costs, lower space requirements and energy use and lower software licensing costs. Servers remain underutilized, in part, because some applications would need to be rearchitected to run on virtual machines and / or containers. Cloud service operators face a particular utilization challenge, due to the need to keep spare dial-up capacity ready on demand for customers.

**Figure 1 Average server utilization levels**



Questions and answers collated by Jay Dietrich, Research Director of Sustainability at Uptime Intelligence and Lenny Simon, Research Associate at Uptime Intelligence.  
For further queries please contact: [research@uptimeinstitute.com](mailto:research@uptimeinstitute.com)

#### Relevant reports:

*IT efficiency: the critical core of digital sustainability*

*Three key elements: water, circularity and siting*