

# The AI Disruption: Challenges and Guidance for Data Center Design

## White Paper 110

Version 1.1

### Energy Management Research Center

by Victor Avelar

Patrick Donovan

Paul Lin

Wendy Torell

Maria A. Torres Arango

#### Executive summary

From large training clusters to small edge inference servers, AI is becoming a larger percentage of data center workloads. This represents a shift to higher rack power densities. AI start-ups, enterprises, colocation providers, and internet giants must now consider the impact of these densities on the design and management of the data center physical infrastructure. This paper explains relevant attributes and trends of AI workloads, and describes the resulting data center challenges. Guidance to address these challenges is provided for each physical infrastructure category including power, cooling, racks, and software management.

RATE THIS PAPER



# Introduction

In recent years, we have witnessed an extraordinary acceleration in the growth of artificial intelligence (AI), transforming the way we live, work, and interact with technology. Generative AI (e.g., ChatGPT) is a catalyst for this growth. Predictive algorithms are having an impact on industry sectors ranging from healthcare<sup>1</sup> and finance to manufacturing<sup>2</sup>, transportation<sup>3</sup> and entertainment. The data requirements associated with AI are driving new chip and server technologies resulting in extreme rack power densities. At the same time, there is massive demand for AI. Together these present new challenges in designing and operating data centers to support this demand.

## AI growth projection

We estimate that AI represents 4.3 GW of power demand today and project this to grow at a CAGR of 26% to 36%, resulting in a total demand of 13.5 GW to 20 GW by 2028. This growth is two to three times that of overall data center power demand CAGR of 11%. See **Table 1** for more details. One key insight is that inference<sup>4</sup> loads will increase over time as more newly trained models are transitioned to production. The actual energy demand will heavily depend on technology factors including successive generations of servers, more efficient instruction sets, improved chip performance, and continued AI research.

**Table 1**  
*Overview of AI workloads in data centers.*

| Schneider Electric estimate         | 2023                        | 2028                        |
|-------------------------------------|-----------------------------|-----------------------------|
| Total data center workload          | 54 GW                       | 90 GW                       |
| AI workload                         | 4.3 GW                      | 13.5-20 GW                  |
| AI workload (% of total)            | 8%                          | 15-20%                      |
| AI workload (Training vs Inference) | 20% Training, 80% Inference | 15% Training, 85% Inference |
| AI workload (Central vs Edge)       | 95% Central, 5% Edge        | 50% Central, 50% Edge       |

This paper explains important AI attributes and trends that create challenges for each data center physical infrastructure category including power, cooling, racks, and software management. We then give guidance on how to address these challenges.<sup>5</sup> Finally, we provide a forward-looking view of what's to come in data center design. This paper is not about applying AI to physical infrastructure systems. **While next-generation physical infrastructure systems will eventually leverage more AI, this paper focuses on supporting AI workloads with *existing* systems available today.**

<sup>1</sup> Federico Cabitza, et al., [Rams, hounds and white boxes: Investigating human-AI collaboration protocols in medical diagnosis](#), Artificial Intelligence in Medicine, 2023, vol 138

<sup>2</sup> Jongsuk Lee, et al., [Key Enabling Technologies for Smart Factory in Automotive Industry: Status and Applications](#), International Journal of Precision Engineering and Manufacturing, 2023, vol 1

<sup>3</sup> Christian Birchler, et al., [Cost-effective simulation-based test selection in self-driving cars software](#), Science of Computer Programming, 2023, vol 226

<sup>4</sup> See section "AI attributes and trends" for definition.

<sup>5</sup> This guidance also applies to other high-density workloads such as high-performance computing (HPC). A major difference with HPC applications is that they tend to be one-off installations that may employ custom IT, power, cooling, and/or rack solutions. In contrast, the massive demand for AI applications requires standard gear (IT and supporting infrastructure) to scale.

## AI attributes and trends

Four AI attributes and trends underlie physical infrastructure challenges:

- AI workloads
- [Thermal design power](#) (TDP) of GPUs
- Network latency
- AI cluster size

### AI workloads

AI workloads fall into two general categories: training and inference.

**Training** workloads are used to train AI models like large language models (LLMs). The type of training workload we refer to in this paper is large-scale [distributed training](#) (large number of machines running in parallel<sup>6</sup>), because of the challenges it poses to data centers today. These workloads require massive amounts of data fed to specialized servers with processors known as accelerators. A graphics processing unit (GPU) is an example of an accelerator<sup>7</sup>. Accelerators are very efficient at performing parallel processing tasks like the ones used in training LLMs. In addition to servers, training also requires data storage and a network to connect it all together. These elements are assembled into an array of racks known as an AI cluster which essentially trains a model as a single computer. The accelerators in a *well-designed* AI cluster operate at near 100% utilization for most of its training duration, which ranges from hours to months. This means that the average power draw of a training cluster is nearly equal to its peak power draw (peak-to-average ratio  $\approx 1$ ).

The larger the model, the more accelerators are required. Rack densities in large AI clusters can range from 30 kW to 100 kW depending on the GPU model and quantity. Clusters can range from a few racks to hundreds of racks and are commonly described by quantity of accelerators used. For example, a [22,000 H100 GPU](#) cluster uses approximately 700 racks and requires about 31 MW to power, with an average rack density of 44 kW. Note that this power excludes physical infrastructure requirements like cooling. Finally, training workloads save the model as “[check-points](#)”. If the cluster fails or loses power, it can continue from where it left off.

**Inference** means that the previously trained model is put into production to predict the output of new queries (inputs). From the user's perspective, there's a trade-off between an output's accuracy and the inference time (i.e., latency). If I'm a scientist, I may be willing to pay a premium and wait longer in between queries in order to get highly accurate outputs. On the other hand, if I'm a copywriter looking for writing ideas, I want a free chatbot with instant answers. In short, the business need determines the size of the inference model, but very rarely is the full original trained model used. Instead, a light-weight version of the model is deployed to reduce inference time with an acceptable loss of accuracy.

Inference workloads tend to use accelerators for large models, and may also depend heavily on CPUs, depending on the application. Applications like autonomous vehicles, recommendation engines, and ChatGPT likely all have different IT stacks, “tuned” to their requirements. Depending on the size of the model, hardware requirements per instance can range from an edge device (e.g., smart phone) to several racks of servers. This means that the rack densities can range from a few

<sup>6</sup> The large number of [parameters](#) and [tokens](#) in a model require that the processing workload be [split up across many GPUs](#) to decrease the time it takes to train the model.

<sup>7</sup> Other examples of accelerators are tensor processing units (TPUs), field programmable gate arrays (FPGAs), and application-specific integrated circuits (ASICs).

hundred watts to over 10 kW. Unlike training, the number of inference servers increase with the number of users / queries. In fact, it's likely that a popular model (e.g., ChatGPT) requires many more times the quantity racks for inference as it did for training since their queries are now in the [millions per day](#). Finally, inference workloads are often business-critical which requires resiliency (e.g., a UPS and or geographic redundancy).

Thermal design power (TDP) of GPUs

While training or inference is impossible without storage and network, we focus on the GPU because it represents about half of an AI cluster's power consumption.<sup>8</sup> GPU power is trending higher with every new generation. A chip's power consumption, measured in watts, is commonly specified with [TDP](#). While we discuss the GPU specifically, this general trend of increasing TDP applies to other accelerators as well. The increasing TDP per GPU generation is a consequence of designing the GPU for an increased number of operations, in order to train models and infer in less time and with lower cost. **Table 2** compares three generations of Nvidia GPUs in terms of TDP and performance<sup>9</sup>.

**Table 2**  
*TDP and performance  
across different generations  
of Nvidia GPUs*

| GPU            | TDP (W) <sup>10</sup> | TFLOPS <sup>11</sup><br>(Training) | Performance<br>over V100 | TOPS <sup>12</sup><br>(Inference) | Performance<br>over V100 |
|----------------|-----------------------|------------------------------------|--------------------------|-----------------------------------|--------------------------|
| V100 SXM2 32GB | 300                   | 15.7                               | 1X                       | 62                                | 1X                       |
| A100 SXM 80GB  | 400                   | 156                                | 9.9X                     | 624                               | 10.1X                    |
| H100 SXM 80GB  | 700                   | 500                                | 31.8X                    | 2,000                             | 32.3X                    |

Network latency

With distributed training, [every GPU must have a network port](#) to establish the compute network fabric. For example, if an AI server has eight GPUs, that server will require eight compute network ports. This compute fabric allows all GPUs in a large AI cluster to communicate in concert at high speeds (e.g., 800 gigabit/second). As GPU processing speeds increase, so must the network speeds, in an effort to reduce the time and cost of training models. For example, using GPUs that process data from memory at 900 GB/s with a 100 GB/s compute fabric would decrease the average GPU utilization because it's waiting on the network to orchestrate what the GPUs do next. This is like buying a 500-horsepower autonomous vehicle with an array of fast sensors communicating over a slow network; the car's speed will be limited by the network speed, and therefore won't fully use the engine's power.

High-speed network cables are expensive. For example, InfiniBand optical connections are on the order of 10 times the cost of copper. Therefore, data scientists, in collaboration with IT teams, try to specify AI training clusters with copper such that the network cabling distances stay within acceptable latencies. Increasing the ports per rack decreases the cabling distances but increases the number of GPUs per rack and therefore rack density. Eventually, the rack cluster grows so large that latency forces designers to switch to fiber, increasing the cost. Note that it's more

<sup>8</sup> [At 400W the NVIDIA V100 GPU represents 55% in this cluster](#) and [at 700W the H100 represents 49%](#)

<sup>9</sup> While the GPU is key to these performance gains, other system improvements were made to take advantage of improved GPUs such as increasing memory and inter-GPU communication.

<sup>10</sup> [V110, A100, H100](#)

<sup>11</sup> TFLOPS - tera (trillion) floating-point operations per second - measure of matrix multiplication throughput at tensor float 32 ([TF32](#)) precision, generally used with training workloads. [V100, A100, H100](#)

<sup>12</sup> TOPS is tera (trillion) operations per second is a measure of integer math throughput at 8-bit integer ([INT8](#)) precision, generally used with inference workloads. [V100, A100, H100](#)

difficult to parallelize GPUs for inference workloads and therefore this rack density relationship doesn't typically apply to inference.<sup>13</sup>

## AI cluster size

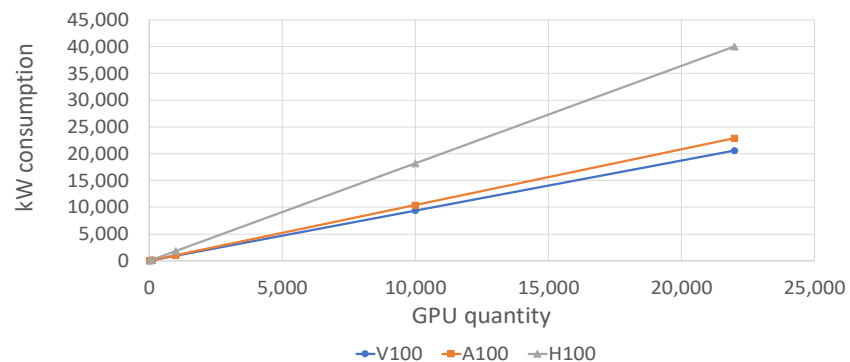
As discussed above, training large models can require thousands of GPUs acting in concert. Given that the GPU represents about half of a cluster's power consumption, GPU quantity becomes a helpful proxy for estimating data center power consumption. **Figure 1** estimates the data center power consumption as a function of GPU quantity in an AI training cluster across three GPU generations (from **Table 2**). To put these values into perspective, a 40,000 kW power plant is capable of powering about 31,000 [average US homes](#). Note that the three trend lines do not equate to the same productivity. In other words, while the power consumption of a data center with H100 GPUs exceeds one with V100 GPUs, the productivity gains of the H100 data center far exceeds its power consumption premium.

**Figure 1**

*Estimated data center power consumption as a function of GPU quantity*

*Data center PUE = 1.3*

*Note that productivity gains are not presented in this chart.*



The four attributes and trends described have a direct impact on rack power density. Most data centers today can support peak rack power densities of about 10 to 20 kW.<sup>14</sup> However, deploying tens or hundreds of racks all greater than 20 kW in an AI cluster will present physical infrastructure challenges to data center operators. They may be specific to power or may touch across two or more physical infrastructure categories. These challenges are not insurmountable, but operators should proceed with a full understanding of the requirements, not only with respect to IT, but to physical infrastructure, especially for existing data center facilities. The older the facility, the more challenging it will be to support AI training workloads. The main sections below explain these challenges in more detail for each physical infrastructure category and provide guidance to overcome these challenges. Note, some of the recommended design approaches only apply to new data center builds, where others are relevant to both new and brownfield (retrofit) buildings.

## Power

AI workloads present six key challenges that impact the power train, including switchgear, distribution, and rack power distribution units (rPDUs).

- 120/208 V distribution is impractical to deploy
- Small power distribution block sizes waste IT space
- Standard 60/63 A rack PDUs are impractical to deploy
- Increased risk of arc flash hazard complicates work practices
- Lack of load diversity increases risk of upstream breaker tripping
- High rack temperatures increase risk of failures & hazards

<sup>13</sup> [Accelerating Deep Learning Inference with Hardware and Software Parallelism](#), Aril 2020

<sup>14</sup> Uptime Institute, [Rack Density is Rising](#), 12/2022

## 120/208 V distribution is impractical to deploy

120/208 V, a voltage historically used in North American data centers, served its purpose when densities were relatively low (on the order of two to three kilowatts per rack) and servers were supplied with 120 V power cords. Today, with high-density loads such as AI clusters, this voltage is too low. While it's still possible to power these loads at 120/208 V, it creates challenges, which stem from the following relationship: power is equal to volts times amps ( $P = V \times A$ ). As this equation demonstrates, the lower the voltage, the more current you need for the same power. Consequently, the wire must be larger to safely provide greater current.

Now consider an AI training rack of (8) HPE Cray XD670 GPU-accelerated servers, totaling a rack density of 80 kW. At 120/208 V, it would take five 60-amp circuits to power the rack (each circuit equals  $120\text{ V} \times 3\text{ phases} \times 60\text{ A} \times 80\% \text{ derating} = 17,280\text{ W} = 17.3\text{ kW}$ ) at 1N redundancy. If 2N was required (although uncommon for AI training loads), this number would double to ten. With 5 to 10 circuits per rack, imagine the chaos of power cables distributed to an AI cluster of 100 racks. The outcome is likely a makeshift, hodgepodge installation of power cables dangling above/near the rack, which could lead to issues including human error and airflow constraints. This is impractical. In addition, there are cost implications of installing and managing the excessive number of circuits.

**GUIDANCE:** Since doubling the voltage means doubling the power, existing data centers with 120/208 V distribution should retrofit their distribution to 240/415 V. New data centers should already design with 240/415 V in mind. See White Paper 128, [High-Efficiency AC Power Distribution for Data Centers](#), for more on this topic. This leads to the next challenge which is related to constraints in *how* to distribute 240/415 V power.

Note, a large part of the globe does not have this same challenge, as many countries distribute power at higher voltage of 230/400 V, which is suitable to achieve the power demands in AI racks.

## Small power distribution block sizes waste IT space

There are three main types of data center power distribution: transformer-based power distribution units (PDU), remote power panels (RPP), and busway. The distribution block size represents the capacity (kW) of each distribution solution. Even with a higher distribution voltage of 240/415 V (230 V IEC countries), the traditional distribution block sizes are too small to support today's AI cluster capacities. Ten years ago, a 300 kW (833 A at 120/208 V) distribution block could support 100 racks (five 20-rack rows at an average rack density of 3 kW/rack). Today that same block couldn't even support the minimum configuration of an [NVIDIA DGX Super-POD](#) (a single 358 kW 10-rack row at 36 kW/rack). Using multiple blocks of distribution for a single row of racks is impractical for various reasons. For example, PDU and RPP footprints double at a minimum. Multiple blocks also increase cost compared to a single higher-capacity block.

**GUIDANCE:** Distribution block sizes must increase to meet the demands of high-density clusters. We recommend that you choose a distribution block size high enough to accommodate, at a minimum, a full row of the cluster. A block size of 800 amps is a standard capacity size available today for all three distribution types at 240/415 V distribution. This provides 576 kW (461 kW derated).

## Standard 60/63 A rack PDUs are impractical to deploy

Even at a higher voltage, it's still a challenge to provide sufficient capacity with a standard rPDU. Most decision makers prefer off-the-shelf rPDUs because they have

shorter lead times, are readily available, cost effective, and are sold in similar configurations by multiple vendors.

Today, the highest capacity standard off-the-shelf rPDU is rated at 60 A (NEMA) / 63 A (IEC). **Table 3** illustrates the usable capacity of rPDUs at various current ratings and voltages. Based on this, we see that the 60 and 63 A rating limits a single rPDU capacity to 34.6 kW and 43.5 kW respectively. This leads to the dilemma of how best to handle rack densities greater than this.

Table 3

Usable 3-phase power density per rPDU based on circuit breaker amp rating and voltage (line-to-neutral)

Top: NEMA (e.g., North America)

|           | Standard |         | Custom  |         |         |          |
|-----------|----------|---------|---------|---------|---------|----------|
| NEMA      | 40 A     | 60 A    | 100 A   | 125 A   | 150 A   | 175 A    |
| 120/208 V | 11.5 kW  | 17.3 kW | 28.8 kW | 36.0 kW | 43.2 kW | 50.4 kW  |
| 240/415 V | 23.0 kW  | 34.6 kW | 57.6 kW | 72.0 kW | 86.4 kW | 100.8 kW |

Note, these values are de-rated to 80% based on typical code requirements.

| IEC       | 32 A    | 63 A    | 100 A   | 125 A   | 150 A    | 160 A    |
|-----------|---------|---------|---------|---------|----------|----------|
| 230/400 V | 22.1 kW | 43.5 kW | 69.0 kW | 86.3 kW | 103.5 kW | 110.4 kW |

**GUIDANCE:** For rack densities greater than 34.6 kW (NEMA) and 43.5 kW (IEC), there are two approaches.

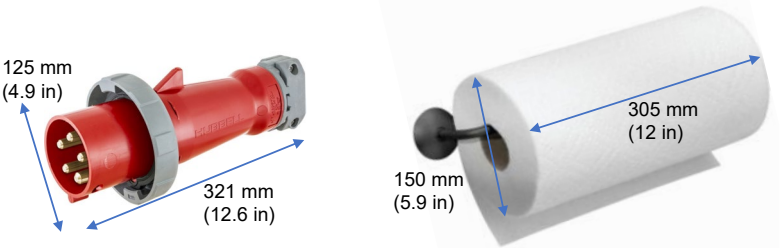
- 1. Multiple standard off-the-shelf rPDUs
- 2. Custom rPDU greater than 60 A and 63 A

Most zero U rPDUs today are roughly 2 meters (80 in) in height. With these standard offers, you are likely to fit, at most, 4 rPDUs in a single air-cooled rack (e.g., 4 x 60/63 A rPDUs is 138 kW / 174 kW). Or if a liquid cooling manifold is required, then 2 rPDUs in a single rack (e.g., 2 x 60/63 A rPDUs is 69 kW / 87 kW). These rPDUs can be combined for increased capacity, or applied for redundancy (e.g., 2N).

If there is a space constraint due to quantity of rPDUs, then we recommend custom rPDUs. For example, as shown in **Table 3**, powering a 100 kW rack is possible with a 175 A rPDU in North America or 150 A in Europe. Custom rPDUs can come with a pin and sleeve connector or be hardwired and give you the flexibility of quantity and type of receptacles. With higher current ratings, pin and sleeve connectors require more work to install and feed through a rack due to their physical size (see **Figure 2**). Note, at current ratings greater than 60A, installation and operation may require an electrician.

Figure 2

240/415 V 125 A pin and sleeve connector relative to the size of a paper towel roll. Mating a pair of such large connectors is challenging for a single person.



Increased risk of arc flash hazard complicates work practices

According to White Paper 194, [Arc Flash Considerations for Data Center IT Space](#), the term “arc flash” describes what happens when electrical short circuit current flows through the air. In an arc flash, the current literally travels through the air from one point to the other, releasing a large amount of energy, known as incident

energy<sup>15</sup>, in less than a second. This energy is released in the form of heat, sound, light, and explosive pressure all of which can cause injuries. Some specific injuries can include burns, blindness, electric shock, hearing loss, and fractures.

A consequence of increasing rPDU current ratings is that they have larger wire diameters which allow more fault current through the rPDU. If the available fault current at the rPDU results in incident energy of 1.2 calories/cm<sup>2</sup> or more, workers are not allowed in that area without proper training and personal protective equipment (PPE).<sup>16</sup> Risk increases as rPDU amp rating increases. The safety of data center personnel is a challenge that must be addressed.

**GUIDANCE:** With so many variables involved, we recommend starting with an arc flash risk assessment to analyze available fault current, as this informs the best solutions for a specific site. It is important this study is conducted from the medium voltage gear all the way down to the rack level. Examples of solutions include:

- Specifying upstream transformers with higher impedance
- Using line reactors (i.e., inductors) to impede the flow of short circuit current
- Using [current limiter blocks](#)
- Using [current limiting breakers](#)

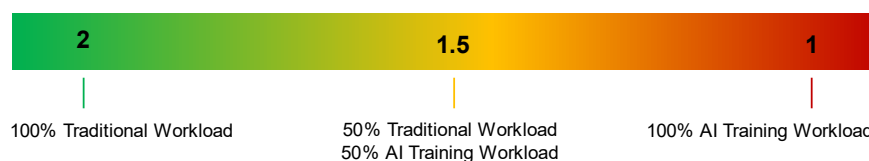
See White Paper, [Arc Flash Mitigation](#), and White Paper 253, [Benefits of Limiting MV Short-Circuit Current in Large Data Centers](#), for more details on addressing arc flash hazards.

### Lack of load variation increases risk of upstream breaker tripping

The power consumption of diverse data center workloads typically peaks at random times. Statistically speaking, there's a very low probability that all these peaks occur at the same time. Therefore, if you were to sum the peak of all individual workloads and divide this value by the total average power consumption, you would find a peak-to-average ratio of 1.5 to 2 or more for a typical large data center. This is what allows designers to “oversubscribe” power and cooling systems. But as discussed in the “AI attributes and trends” section, AI training loads lack diversity. These workloads can run for hours, days, or even weeks at peak power. The result is an increased likelihood of tripping a large upstream breaker. This is like what happens when many large loads run simultaneously in a home and the main panel breaker trips open. **Figure 3** illustrates the typical spectrum of peak-to-average ratio (also referred to as diversity factor) as the loads in a data center move to 100% AI loads.

**Figure 3**

*Spectrum of typical peak-to-average ratios from 100% traditional mixed loads to 100% AI training*



**GUIDANCE:** In the case of a new data center hall with greater than 60-70% of AI training workloads, we recommend sizing the main breaker based on the sum of the individual feeder breakers downstream. In other words, assume a peak-to-average ratio of 1, where the average power draw is equal to the peak power draw. The practice of oversubscribing and depending on diversity is not advised.

<sup>15</sup> According to NFPA 70E (2015), incident energy is “The amount of thermal energy impressed on a surface, at a certain distance from the source, generated during an electrical arc event.”

<sup>16</sup> For more information see White Papers 13, “[Mitigating Electrical Risk While Swapping Energized Equipment](#)” and 194, “[Arc Flash Considerations for Data Center IT Space](#)”.

For existing data centers, calculate the total AI load the upstream breaker can support. For example, if there is a 1,000 A main breaker upstream of the AI workload cluster(s), make sure the AI loads don't add up to more than 1,000 A.

High rack temperatures increase risk of failures & hazards

Between densities climbing and a focus on operating efficiency, IT environments are getting hotter. Higher operating temperatures improve cooling system efficiency, but they also cause added stress on components. When components are exposed to temperatures they are not rated for, the following can result:

- **Pre-mature component failures** – although systems operate as expected on day 1, the life expectancy of components can be significantly shortened when exposed to conditions outside the specified range.
- **Safety hazards** – using cords not rated for the operating range could lead to safety hazards like melting cords.

IEC 60320 is the recognized international standard used by most of the globe for the connection of power supply cords. There are specific IEC connectors that are rated for higher temperatures. **Table 4** compares the standard C19/C20 connectors to the high temperature C21/C22 connectors.

**Table 4**  
*Comparison of IEC 60320 standard and high temperature connectors for 250 V and 16/20 A*

|                  | Female                                                                                  | Male                                                                                     | Limit | Notes                                                                                               |
|------------------|-----------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------|-------|-----------------------------------------------------------------------------------------------------|
| Standard         |  C19  |  C20  | 65°C  | C20 is commonly used as a jumper cable, providing power from a rack PDU to high powered IT devices. |
| High temperature |  C21 |  C22 | 155°C | C21 mates with either C22 or C20 connectors and used when temperatures exceed the C19 rating.       |

**GUIDANCE:** We recommend analyzing all loads within the AI cluster to ensure the appropriate connectors and receptacles are used. C21/C22 connectors are becoming more common with higher density compute loads like AI servers. AI servers are often configured with these high temperature rated cords/receptacles, but other devices in your rack may not, like the top of rack switch. It's important to understand the environment your equipment will operate in, and ensure all devices are rated accordingly, including the rack PDU and all its subcomponents.

When specifying rack PDUs, it's important to not only look at the voltage, amperage, and number of outlets, but also its temperature rating. High-temperature rated rPDUs are available on the market for this type of application. Although they generally come at a cost premium, that added cost generally outweighs the cost of latent failures waiting to happen. Note, placing temperature sensors in the back of the rack (monitored by DCIM) are also recommended to validate the operating conditions are as expected.

## Cooling

The densification of AI training server clusters is forcing an evolution from air-cooled to liquid-cooled to address their increasing TDPs. While less dense clusters and inference servers will still use more conventional data center cooling, we see the following six key cooling challenges that data center operators need to address:

- Air cooling is not suitable for AI clusters above 20 kW/rack
- Lack of standardized designs and site constraints complicate liquid cooling retrofits
- Unknown future TDPs increase the risk of cooling design obsolescence
- Inexperience complicates installation, operation, & maintenance
- Liquid cooling increases the risk of leaks within IT racks
- Limited fluid options exist to operate liquid cooling sustainably

### Air cooling is not suitable for AI clusters above 20 kW/rack

Liquid cooling for IT has been around for over half a century for specialized high-performance computing. Air cooling has been the mainstream choice and can support average rack power densities of about 20 kW when properly designed with hot aisle containment. With a single 8-10U AI server consuming [12 kW](#), it's easy to exceed this 20 kW threshold. Adding to this challenge is that servers in large AI clusters cannot be spread out (to lower rack density) due to latency limitations. Liquid-cooled versions of AI training servers are increasingly available with some being exclusively liquid cooled, driven by increasing TDPs.

**GUIDANCE:** Smaller AI clusters and inference server racks that are configured at 20 kW per rack or less can be air cooled. For these racks, good airflow management practices (e.g., [blanking panels](#), [aisle containment](#)) should be followed to ensure more effective and efficient cooling. If an air-cooled system is still constrained, spreading out the AI servers across multiple racks is a strategy for reducing rack density. For example, if a cluster is 20 racks at 20 kW/rack, spreading the servers out to 40 racks would reduce the rack density to 10 kW/rack. Note that spreading racks may not be possible if the increased network cabling distances degrade the AI cluster's performance.

When AI rack densities go above 20 kW, strong consideration should be given to liquid-cooled servers. There are several liquid cooling technologies and architectures. Direct-to-chip (DTC), sometimes called conductive or cold plate, and immersion are the two main categories. Compared to immersion, direct-to-chip is currently the preferred choice, as it has better compatibility with existing air cooling and is also easier for retrofit applications. If given the choice, data center operators should select liquid-cooled servers to improve performance and reduce energy cost, which can offset the investment premium. For example, the HPE Cray XD670 GPU-accelerated server consumes 10 kW when air-cooled vs. 7.5 kW when liquid-cooled due to reduced fan power requirements and lower leakage currents in the silicon. For more information on liquid cooling, see White Paper 279, [Five Reasons to Adopt Liquid Cooling](#), and White Paper 265, [Liquid Cooling Technologies for Data Centers and Edge Applications](#).

Note that liquids have a much greater capacity to capture heat by unit volume which allows liquid cooling technologies to remove heat more efficiently than air cooling. However, if the fluid flow stops, the chip temperature will increase much faster than with air, leading to faster shutdown. Placing pumps on a UPS helps address this issue.

## Lack of standardized designs and site constraints complicate liquid cooling retrofits

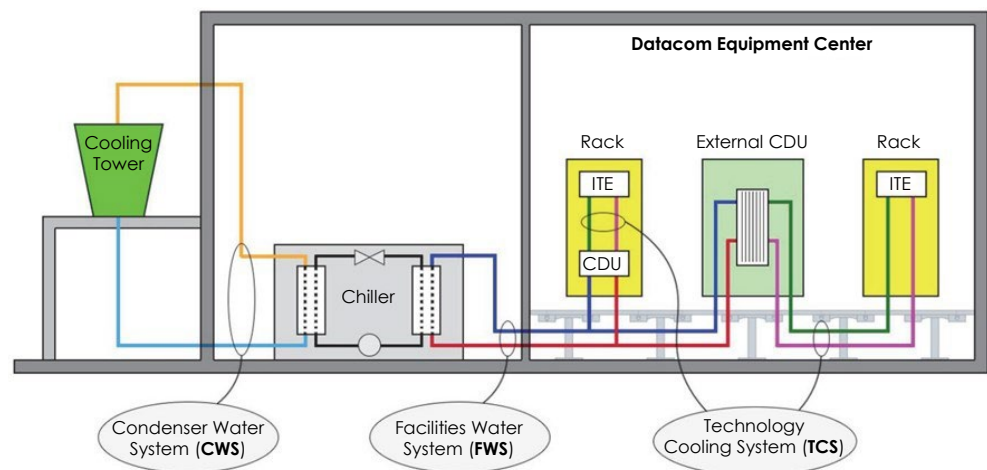
Compared to traditional chilled water systems, direct-to-chip liquid-cooled servers have more stringent requirements in water temperature, flow, and chemistry. This means that operators cannot run water directly from a chiller system through a chip's cold plate.<sup>17</sup> While water quality is certainly part of the challenge of retrofitting a data center to liquid cooling, the biggest issue is the lack of standardized designs for AI loads at this scale (i.e., hundreds of racks). Consider the fact that there are several mounting options and locations for coolant distribution units (CDUs)<sup>18</sup>. It could be floor-mounted on the room perimeter, end of row, or rack-mounted in each server rack. There are several methods for distributing piping to racks, many locations for cooling system equipment, several approaches to controlling temperatures, etc. To help visualize the components of a liquid-cooled system, **Figure 3** illustrates different water loops and CDUs.

Retrofitting for liquid cooling is also disruptive to a production data center and may encounter physical constraints such as limited floor space and insufficient raised floor height to run water piping. Even if 100% of the servers are direct-to-chip liquid-cooled, there is still a need for supplemental air-cooling to cool other equipment like network switches and heat conduction from liquid-cooled servers. In short, retrofitting is a challenge due to the high number of design permutations combined with limited analysis and few large-scale liquid-cooling deployments to learn from. Note that some data centers don't have a chilled water system which makes retrofitting even more challenging.

**Figure 3**

*Liquid cooling using CDUs in a data center*

Source: ASHRAE, *Water-Cooled Servers: Common Designs, Components, and Processes*, page 10



**GUIDANCE:** We recommend data center operators perform a design assessment of the proposed liquid-cooled loads and the facility's existing conditions in advance of deploying liquid cooling. Expert review is essential in order to evaluate possible designs and avoid the cost implications of unforeseen building constraints. For example, pipes may obstruct airflow under the raised floor or interfere with power or cable trays. For more information see White Paper 133, *Data Center Design Practices for Integrating Liquid-cooled AI Workloads*.

<sup>17</sup> Running untreated water through a server's cold plate can cause corrosion, biological growth, and fouling. All of these compromise heat transfer from the GPUs, eventually leading the GPUs to throttle or shutdown to prevent damage.

<sup>18</sup> A CDU physically isolates the chilled water loop from the "clean" water loop that supplies the servers.

## Unknown future TDPs increase the risk of cooling design obsolescence

AI technologies are evolving at such rapid pace that it is likely that next generation GPUs will have higher TDPs and increased cooling requirements. For example, a current server with eight GPUs may have sixteen in its next generation. As a result, a data center's cooling distribution designed for today's loads may become insufficient to support tomorrow's loads.

**GUIDANCE:** We recommend designing the cooling system to accommodate air cooling and liquid cooling, scale up as needed, and support different generations of accelerators. For example, using high-temperature chillers for air cooling today, can be easily switched to higher temperature liquid cooling tomorrow. Another recommended practice is to design the chilled water piping system with tap-offs for future CDUs. This will allow a transition to 100% direct-to-chip liquid-cooled loads combined with rear door heat exchangers for supplemental air cooling.

## Inexperience complicates installation, operation, & maintenance

Data center operators are quite familiar with air-cooled systems as it has been used for decades, however, liquid cooling is new for most operators. Liquid cooling uses components such as cold plates, manifolds, blind-mate valves, etc. These components come with additional installation, operation, and maintenance procedures unfamiliar to these operators. For example, the tiny water channels in direct-to-chip cold plate servers are more susceptible to fouling, which means operators may need to learn new operation and maintenance procedures to control water chemistry. Another example is moving water into servers which introduces the risk of leaks.

**GUIDANCE:** Liquid cooling design plays a critical role in minimizing installation, operation, and maintenance work. We recommend that data center operators, unfamiliar with supporting liquid-cooled servers, seek experts to provide a thorough review of their design and issue detailed standard operating procedures (SOP) and method of procedures (MOP) for day-to-day operations. This will minimize failures and human error, especially relating to leaks.

## Liquid cooling increases the risk of leaks within IT racks

Direct-to-chip technologies use water (e.g., deionized water, alcohol-based solutions) in cold plates within a server. Water leaks are a safety and reliability concern and must be considered in the design and procurement phase.

**GUIDANCE:** We recommend working with reputable providers to ensure their systems go through rigorous pressure testing to minimize the risk of leaks. Additionally, leak detection at the server and rack level can help catch leaks before they become more serious. Instead of traditional CDU pumping systems, consider CDUs with innovative leak-prevention-systems (LPS). LPSs maintain the water loop at a slight vacuum (negative pressure) to eliminate the risk of leaks within IT equipment. Immersion liquid cooling uses dielectric fluids which also eliminate the risk of water leaks within the servers. These may be an option from your AI server or integration vendor. Finally, emergency operating procedures (EOP) should be developed to address leaks if they do occur.

## Limited fluid options exist to operate liquid cooling sustainably

Compared to traditional air-cooled IT, liquid cooling offers some environmental sustainability benefits in that it reduces both energy consumption and water usage. This is due to higher energy efficiency for both IT servers and cooling systems as most or all of the server fans are removed, and higher water temperatures enable

increased economizer hours<sup>19</sup>. However, some liquid-cooled systems use engineered chemicals that are harmful to the environment. For example, fluorocarbon fluids are widely used as dielectric fluids in immersion liquid cooling technologies<sup>20</sup> due to their heat transfer performance. Unfortunately, some fluorocarbons have global warming potentials (GWPs) on the order of 8,000. For comparison, HFC-143a refrigerant, common in refrigerators, has a GWP of 1,430. In addition, societal pressures have prompted manufactures to eliminate PFAS (per-and polyfluoroalkyl substances) from products like refrigerants (to mitigate environmental impact) and to transition to refrigerants with lower GWP. Sustainability has become a top priority for most data center operators, leaving them fewer options.

**GUIDANCE:** We recommend avoiding fluorocarbon fluids. In the past these were used in both direct-to-chip and immersion cooling systems. Today, direct-to-chip uses water and therefore is not an issue. If deploying immersion liquid cooling, we recommend using oil-based dielectric fluids which have zero GWP (unlike 2-phase engineered fluids). However, because oil-based dielectric fluids are not as effective at transferring heat as direct-to-chip using water, direct-to-chip has become the preferred liquid cooling architecture today. Note that it's likely that vendors will develop sustainable alternative dielectrics to fluorocarbon fluids. This would significantly improve the heat removal efficiency of immersion liquid cooling and perhaps precipitate a change in cooling architecture. See White Paper 291, [Comparison of Dielectric Fluids for Immersive Liquid Cooling of IT Equipment](#) for more information.

## Racks

Some of the power and cooling challenges described in the previous sections also trickle down to the IT rack (i.e., IT cabinet or enclosure). We see the following four rack system challenges driven by AI workloads:

- Standard-width racks lack space for needed power & cooling gear
- Standard-depth racks lack space for deep AI servers & cabling
- Standard-height racks lack space for required quantity of servers
- Standard racks lack sufficient weight-bearing capacity for AI gear

### Standard-width racks lack space for needed power & cooling gear

Because AI servers are getting deeper, there is less space in the back of the rack to mount rack PDUs and liquid cooling manifolds. As server power densities continue to increase, it will become very difficult if not impossible to accommodate the necessary power and cooling distribution in the back of a standard-width rack (i.e., 600 mm / 24 in). Furthermore, narrow racks are likely to congest the exhaust airflow behind the rack due to power and network cables.

**GUIDANCE:** We recommend at least 750 mm (29.5 in) wide racks to accommodate rack PDUs, and in the case of liquid cooling, manifolds for liquid-cooled servers. Although these racks will not line up with 600 mm wide raised floor tiles like standard 600 mm racks do, this is no longer a relevant constraint. This is because air-cooled AI servers require high airflow rates and raised floors are not typically used for air distribution but rather for piping and cabling.

<sup>19</sup> Economization occurs when the outdoor temperature is lower than the water temperature. The return water temperatures from direct-to-chip servers are much higher than traditional chilled water return temperatures. At these higher temperatures there are more hours in the year to free-cool the water.

<sup>20</sup> Immersion cooling submerges the entire chip or even server in the dielectric fluid.

## Standard-depth racks lack space for deep AI servers & cabling

Servers optimized for AI workloads can reach depths that exceed the maximum mounting depth of some standard racks. Even if a deep server can mount into a shallow rack, sufficient rear clearance is needed to accommodate network cabling while still allowing for sufficient airflow.

**GUIDANCE:** IT racks have adjustable mounting rails to accommodate different IT equipment depths, however, maximum mounting depths vary. We recommend racks at least 1,200 mm (47.2 in) deep with maximum mounting depths greater than 1,000 mm (40 in).

## Standard-height racks lack space for required quantity of servers

Depending on the height of AI servers, common 42U high racks will likely be too short to accommodate all the servers, switches, and other equipment. For example, a 64-port network switch implies that the rack would have 8 servers, each with 8 GPUs. At this density, and assuming a 5U server height, the servers alone would consume 40U, leaving only 2U of remaining space to accommodate other devices.

**GUIDANCE:** We recommend deploying AI training clusters in 48U racks or higher with the presumption that data center doorways are tall enough to accommodate them. [1U](#) is equal to 44.45 mm (1.75 in)<sup>21</sup>.

## Standard racks lack sufficient weight-bearing capacity for AI gear

With heavy AI servers, a high-density rack can weigh over 900 kg (2000 lb). This places a significant load on IT racks and raised floors, both in terms of static and dynamic (rolling) load bearing capacity. Racks not rated for these weights may experience deformation to frames, leveling feet, and/or casters. Furthermore, raised floors may not support these heavy racks.

**GUIDANCE:** IT rack weight-bearing capacities are specified as static and dynamic. Static refers to the weight a rack can support while stationary. Dynamic refers to the weight a rack can support while moving. We recommend specifying racks with a static weight capacity greater than 1,800 kg (3,968 lb) and a dynamic weight capacity greater than 1,200 kg (2,646 lb). These rack capacities should be validated by an independent third party.<sup>22</sup> Even if your current AI deployment is small and doesn't yet require these capacities, racks tend to have a longer service life than IT equipment. It's likely that the next generation of your AI deployment will require some or all of these rack recommendations. Finally, in some cases IT racks are pre-configured offsite and then transported to the data center. These racks must be able to sustain the dynamic forces generated during transportation and the associated packaging must also protect the racks and the valuable IT gear they support.

Data center floors, and raised floors in particular, should be assessed to ensure they can support the weight of an AI cluster. This is especially important to raised floor dynamic capacity when moving heavy racks around the data center.

## Software tools

Physical infrastructure software tools support the design and operation of the data center and include [DCIM](#), [EPMS](#), [BMS](#), and digital [electrical design tools](#). Having clusters of high-power density and liquid-cooled IT alongside traditional air-cooled IT means that certain software functions become more critical. Even though some

<sup>21</sup> For example, 48U means that there is 2.13 m (84 in) of interior vertical space available for equipment.

<sup>22</sup> We recommend Underwriters Laboratory (UL) and International Safe Transit Association (ISTA). For more information see White Paper 201, "[How to Choose an IT Rack](#)".

AI training workloads may not require high availability, poor design and monitoring can lead to downtime risks for adjacent racks and tenants that are likely to be business critical. The following two challenges highlight important management software functions that become more relevant in the context of high-density AI training workloads:

- Extreme power density and demand of AI cluster leads to design uncertainty
- Less margin for error increases operational risk in a dynamic environment

### AI cluster extreme power density and power demand leads to design uncertainty

Before retrofitting an existing site to accommodate new AI clusters, a feasibility study is needed to confirm there is enough power and cooling capacity, as well as the infrastructure needed to distribute that capacity to the new loads. In typical cases with rack power densities well below 10kW and with excess bulk power and cooling capacity, adding standard IT might be relatively easy and not require as much scrutiny and verification. Point-in-time power and cooling measurements might be used along with common power distribution components and existing cooling units that you're familiar with. This more manual, "eye-ball" retrofit design approach will not suffice for large high density AI training clusters. An AI cluster drawing hundreds of kilowatts has much bigger consequences if you make a design mistake (i.e., not knowing actual peak to average power draws, being unsure of what loads are on which circuits, etc.). You can't afford to have unknowns and uncertainties with the design. Also, because AI cluster designs are so unique (e.g., non-standard high amperage rPDUs/busway, use of liquid cooling, etc.) there is greater uncertainty about how the cluster will perform on startup.

**GUIDANCE:** We recommend using **EPMS** and **DCIM** to provide an accurate view of the current power capacity and its trends, both at the bulk power level and the distribution level within the IT space. These tools will show what the actual peak power draw is over a long period of time. This is important to understand to make sure you don't inadvertently trip a breaker. This capacity assessment will help determine the capability of hosting AI loads. Note, this assumes that the necessary power meters are in place. Next, prior to any changes, we recommend performing safety and technical studies including capacity analysis, protection coordination, arc flash study, and short-circuit & device evaluation<sup>23</sup>. Using **electrical design (a.k.a., power system engineering) software tools** simplifies the amount of data collection and calculations.

After the assessment, changes to the electric network will likely be required in order to add the AI clusters. In this case, electrical design software tools ensure you have the right data to select optimal electrical equipment, prevent electrical faults, develop effective methods of procedure, and implement proper safety protocols when working on and servicing the electrical network in the IT space.

It's important to note that existing data centers with [digitalized single-line diagrams \(iSLDs\)](#)<sup>24</sup> could simplify the assessment process described above. When accurate, intelligent, iSLDs are used, the time and expertise needed to collect data and perform the calculations are greatly reduced. An iSLD is a more advanced single-line diagram stored and managed in specialized software that includes advanced functionality and awareness of the devices' characteristics and operating behavior. It creates a digital twin of the physical electrical network. In essence, this one

---

<sup>23</sup> i.e., evaluating capacity, kA ratings, and other specs for suitability for the given design

<sup>24</sup> Some vendors offer iSLD creation and maintenance as a service.

software platform can be used to design the electrical network, create and maintain the SLD, and perform all technical studies and safety assessments.

### Less margin for error increases operational risk in a dynamic environment

Assuming you implemented an optimal data center design using the guidance in the first challenge, your “day one” operation should run smoothly. However, compared to other types of facilities, data centers are dynamic environments where frequent moves, adds, and changes of IT equipment take place. As capacity safety margins shrink, as is likely with the addition of a large AI cluster, the risk of tripping a breaker, creating a hot spot, or stranding resources increases as loads change over time within the IT space. The underlying reasons for the increased risk are the high rack densities and low peak-to-average ratios (close to 1) of AI clusters, discussed earlier. Less margin for error means operators need to be increasingly situationally aware to prevent downtime and ensure efficient use of available resources throughout the life of the data center.

**GUIDANCE:** We recommend creating a digital twin of the entire IT space (including the equipment and VMs in the racks) which minimizes or prevents the challenges listed above. This layout must then be maintained over time. DCIM planning and modeling functions allow you to operate effective IT space floor layouts using a rules-based tool. By digitally adding or moving IT loads, you can validate that there is sufficient power, cooling, and floor weight capacities to support them. DCIM creates a digital twin of the IT space and documents all equipment dependencies on resources. This informs decisions to avoid stranding resources and to minimize human error that might lead to downtime. EPMS and DCIM together allow you to monitor power capacities across all PDUs, UPSs, rPDUs, etc. to receive early warning of exceeding power thresholds to avert downtime. DCIM software will advise the best place to locate new equipment based on power, cooling, redundancy level requirements, as well as available U space, network port, and weight capacity. This applies more to non-AI equipment and to AI inference servers. Unlike inference loads, AI training loads require a pre-designed configuration that seldom changes, if at all.

Many DCIM planning and modeling software tools include a computational fluid dynamics (CFD) tool to ensure adequate air flow given the physical layout of the equipment and heat load. DCIM can be used to help optimize cooling capacity by releasing stranded cooling capacity through the optimal placement and configuration of infrastructure and loads. In terms of moves, adds, and changes of AI loads, CFD applies more to AI inference loads since more servers are added to meet user demand (i.e., queries). Note that in some cases, the AI training or inference cluster is isolated on its own power segment and cooling architecture. In these cases, non-AI loads are less susceptible to the effects of the AI cluster. However, in both cases establishing a digital twin of these spaces is beneficial.

## Future outlook of physical infrastructure to support AI

The guidance thus far focuses on technologies and design approaches available today. This section provides a brief description of some *future* technologies and design approaches we think will further help with the challenges presented.

- **Standard AI-optimized rPDUs** – Form factors will change to accommodate more power-dense servers with fewer stranded receptacles. Eliminating unnecessary receptacles allows more rPDUs in each rack or a single higher-capacity rPDU (rated for up to 150 amps at 240V, 86 kW derated). These rPDUs would also provide receptacles for low-density equipment like switches.
- **Medium voltage to 415/240 V transformers in the technical / IT space** – Distributing power at medium voltage (e.g., 13 kV) reduces copper, requires

fewer conductors, and reduces installation time. For example, IT distribution would use a 2 MW transformer to feed a 3,000 A busway at 415/240 V thereby powering an entire AI cluster or a portion of one greater than 2 MW. This distribution architecture also eliminates the traditional 13kV to 480/277 V transformers and switchgear upstream of the IT distribution. This may also mitigate supply chain constraints of 480V distribution gear.

- **Solid state transformers** – These are essentially power electronics converters. They use semiconductor components to change the primary voltage to a secondary voltage. They use a medium frequency transformer (MFT) [galvanically isolate](#) the primary and secondary sides. While traditional transformers are heavy and work with only alternating current (AC), solid state transformers are small and light, and convert between AC and DC voltage.
- **Solid state circuit breakers** – These circuit breakers use semiconductors to turn the flow of current on or off. This is especially important when interrupting the flow of current into a fault. However, to be considered a circuit breaker, they must also use a mechanical switch in series with the semiconductors to provide [galvanic isolation](#). Solid state breakers would allow faster operation and the ability to more tightly control fault currents. This would be very beneficial to reduce arc flash energy at high-density AI racks.
- **Sustainable dielectric fluids** – These may replace water as today's choice for direct-to-chip cooling if they increase the heat transfer efficiency and allow for higher chip TDPs.
- **Ultra-deep IT racks** – As deeper accelerator-based servers are introduced, deeper racks would accommodate not only the server, but the network cabling, water piping, and rack PDUs.
- **Increased interaction/optimization with the grid** – Scheduling workloads based on utility and micro grid conditions helps with balancing the grid and saving on electricity. Migrating loads to different redundancy zones or placing a UPS on battery operation are examples of workload management.

## Conclusion

The rapid growth and application of AI is changing the design and operation of data centers. We estimate that AI workloads will represent 15% to 20% of total data center energy by 2028. **Inference** workloads, although expected to consume much more power than training clusters, operate at a wide range of rack densities. **AI training** workloads, on the other hand, consistently operate at very high densities, ranging from 20-100 kW per rack or more. Networking demands and cost drive these training racks to be clustered together. These clusters of extreme power density are fundamentally what challenges the power, cooling, racks, and software management design in data centers. In this paper, guidance is provided on how the challenges are addressed. They are summarized below:

**POWER:** The use of 120/208 V distribution (in NAM) is no longer sufficient, and instead, 240/415 V distribution is advised to limit the number of circuits within high-density racks. Even at the higher voltage, it is still a challenge to provide sufficient capacity with standard 60/63 amp rack PDUs. For example, liquid-cooled racks are limited to two rPDUs, providing 69 / 87 kW. For personnel safety, we recommend an arc flash risk assessment and a load analysis to ensure the appropriate connectors, receptacles, and rPDUs are used based on their exposed temperatures. Upstream distribution block sizes must be large enough to support a single row of an AI cluster.

**COOLING:** Although air cooling will still exist for the near future, we predict a transition from air cooling to liquid cooling as a preferred or necessary solution for data centers with AI clusters. Compared with air cooling, liquid cooling provides many benefits such as improved processor reliability and performance, space savings with higher rack densities, more thermal inertia with water in piping, increased energy efficiency, improved power utilization (more power goes to IT), and reduced water usage. Data center operators can use our proposed guidance to achieve a successful transition from air cooling to liquid cooling to support AI workloads.

**RACKS:** With AI clusters, servers are deeper, power demands are greater, and cooling is more complex. As a result, we recommend using racks of greater dimensions and weight capacity, specifically: at least 750 mm (29.5 in) wide, 1,200 mm (47.2 in) deep, 48U high, with 1,000 mm (40 in) mounting depths, static weight capacity greater than 1,800 kg (3,968 lb), and a dynamic weight capacity greater than 1,200 kg (2,646 lb).

**SOFTWARE MANAGEMENT:** When managing AI clusters, software tools such as, DCIM, EPMS, BMS, and digital electrical design tools, become critical. They decrease the risk of unexpected behavior with complex electrical networks. They also provide a digital twin of the data center to identify constrained power and cooling resources to inform layout decisions.



## About the authors

**Victor Avelar** is a Senior Research Analyst at Schneider Electric's Energy Management Research Center. He is responsible for data center design and operations research, and consults with clients on risk assessment and design practices to optimize the availability and efficiency of their data center environments. Victor holds a bachelor's degree in mechanical engineering from Rensselaer Polytechnic Institute and an MBA from Babson College. He is a member of AFCOM.

**Patrick Donovan** is a Senior Research Analyst for the Energy Management Research Center at Schneider Electric. He has over 27 years of experience developing and supporting critical power and cooling systems for Schneider Electric's Secure Power Business unit including several award-winning power protection, efficiency, and availability solutions. An author of numerous white papers, industry articles, and technology assessments, Patrick's research on data center physical infrastructure technologies and markets offers guidance and advice on best practices for planning, designing, and operation of data center facilities.

**Paul Lin** is the Research Director and Edison Expert at Schneider Electric's Energy Management Research Center. He is responsible for data center design and operation research and consults with clients on risk assessment and design practices to optimize the availability and sustainability of their data center environment. He is a recognized expert, and a frequent speaker and panelist at data center industry events. Before joining Schneider Electric, Paul worked as an R&D Project Leader in LG Electronics for several years. He is also a registered professional engineer and holds over 10 patents. Paul holds both a Bachelor's and Master's of Science degree in mechanical engineering from Jilin University. He also holds a certificate in Transforming Schneider Leadership Programme from INSEAD.

**Wendy Torell** is a Senior Research Analyst at Schneider Electric's Data Center Science Center. In this role, she researches best practices in data center design and operation, publishes white papers & articles, and develops TradeOff Tools to help clients optimize the availability, efficiency, and cost of their data center environments. She also consults with clients on availability science approaches and design practices to help them meet their data center performance objectives. She received her Bachelor's of Mechanical Engineering degree from Union College in Schenectady, NY and her MBA from University of Rhode Island. Wendy is an ASQ Certified Reliability Engineer.

**Maria A. Torres Arango** is a Research Analyst at Schneider Electric's Energy Management Research Center. In this role Maria investigates technical strategic topics to inform decision making, currently focusing on energy storage systems and sustainability. Maria holds a BS in Aeronautical Engineering from Universidad Pontificia Bolivariana, Colombia; and a MSc in Aerospace Engineering and PhD in Materials Science and Engineering from West Virginia University.

## RATE THIS PAPER





[High-Efficiency AC Power Distribution for Data Centers](#)  
**White Paper 128**



[Data Center Design Practices for Integrating Liquid-cooled AI Workloads](#)  
**White Paper 133**



[Arc Flash Considerations for Data Center IT Space](#)  
**White Paper 194**



[Benefits of Limiting MV Short-Circuit Current in Large Data Centers](#)  
**White Paper 253**



[Liquid Cooling Technologies for Data Centers and Edge Applications](#)  
**White Paper 265**



[Five Reasons to Adopt Liquid Cooling](#)  
**White Paper 279**



[Comparison of Dielectric Fluids for Immersive Liquid Cooling of IT Equipment](#)  
**White Paper 291**



[Arc Flash Mitigation](#)  
**White Paper**



[Browse all](#)  
[white papers](#)  
**[whitepapers.apc.com](https://whitepapers.apc.com)**



[Browse all](#)  
[TradeOff Tools™](#)  
**[tools.apc.com](https://tools.apc.com)**

**Note:** Internet links can become obsolete over time. The referenced links were available at the time this paper was written but may no longer be available now.

## Contact us

For feedback and comments about the content of this white paper:

Schneider Electric Energy Management Research Center  
[dcsc@schneider-electric.com](mailto:dcsc@schneider-electric.com)

If you are a customer and have questions specific to your data center project:

Contact your Schneider Electric representative at  
[www.apc.com/support/contact/index.cfm](https://www.apc.com/support/contact/index.cfm)