

INTELLIGENCE UPDATE

Integrated cold plates will help realize free cooling



Daniel Bizo

12 Nov 2025

The primary reason for adopting direct liquid cooling (DLC) in servers is the push to pack more processors and memory chips into each rack. The close integration of IT electronics and the compact chassis form factors of high-density server hardware often make air cooling impractical. Liquid coolants can remove concentrated heat from tight spaces more effectively. The need to cool high-density racks remains the primary driver for data center operators to support DLC systems in their facilities in 2025.

As cooling and IT designs evolve, one seemingly small technical detail will have an outsized effect on the overall business case: the interface between the cold plates (or immersion heat sinks) and the IT chips. Several manufacturers — including Accelsius, Boyd Corporation, Chilldyne, CoolIT Systems, Fabric8, JetCool, Motivair, ZutaCore — introduced improvements to their cold plates in 2025. Some have also expressed interest in closer collaboration with chipmakers to reduce thermal bottlenecks at the points where their respective products meet.

When deployed at scale, DLC promises significant business benefits: lower facility capital costs, improved overall infrastructure energy performance, and major sustainability gains through reduced consumption of energy, water and refrigerants. However, these promises depend on the data center's ability to reject heat year-round using only dry coolers — with no mechanical refrigeration and, ideally, no water consumption.

Several site engineering considerations determine whether this design objective is feasible or desirable. These include site climatic conditions, the accuracy of forecasting IT load requirements (specifically the ratio of DLC to air capacity), temperature set points, site layout and space limitations. The boundary between DLC equipment and CPUs, GPUs and other high-performance IT components represents an opportunity for further improvement of overall thermal performance, which in turn strengthens the business case for DLC.

Keeping up with the flow

Chip vendors are driving a continued escalation in silicon thermal design power (TDP), well

above today's already high levels. Next-generation server CPUs are expected to reach the 600-800 W band within the next few years, while high-end GPU-based accelerator modules will likely surpass 2 kW.

This creates thermal management challenges in three major areas, at least for single-phase coolant such as water cold plates:

- Preheating of coolant as it reaches some parts of the cold plate.
- Higher heat fluxes in silicon hot spots.
- Processor package temperature restrictions (known as T_{case}) to drive maximum performance and prevent any throttling, for select models.

The first issue is a result of several major design factors. Cold plates are becoming much larger, in step with rapidly increasing package sizes of CPUs and GPUs. Because coolant typically flows from a single inlet, by the time it reaches certain "remote" parts of the cold plate, it may not be able to meet temperature targets, leading to hot spots. In addition, due to space limitations, high-density IT hardware often serializes two (or more) cold plates. As a result, downstream cold plates receive preheated coolant, which can also exacerbate the issue.

High heat concentration on the silicon presents a distinct issue. As cold plates (and silicon packages) grow larger, it generally becomes easier for thermal solutions to keep up with TDP escalation — but silicon heat loads remain strongly non-uniform. With each generation, logic circuitry shrinks, and heat flux (watts per square millimeter) jumps in areas of high transistor switching activity, such as arithmetic units or control logic, compared with the relatively low-activity areas of cache memory arrays. Even when handling high TDPs itself is not a significant issue, maintaining stable temperatures in these high heat flux areas is increasingly difficult.

T_{case} restrictions ensure that there is sufficient temperature difference (drop) between the silicon and the case to ensure high heat flux over silicon hot spots, preventing overheating even under sustained extreme workloads. Although T_{case} is specified as a single maximum temperature for the entire chip package, it is typically dictated by the needs of the highest-heat areas on the silicon.

Cold plates have evolved considerably in recent years to address these challenges.

Manufacturers have invested to improve coolant distribution across the internal surface area of the cold plate and reduce impediments to flow. Increasingly, engineers are optimizing cold plate designs to match the heat map of silicon so that they have more directed flows and larger contact surfaces at high-heat flux areas — without further increasing the overall flow rate. All these techniques aim to reduce effective thermal resistance — a system's resistance to heat flow, measured in $^{\circ}\text{C}/\text{W}$.

Still, coolant flow rates have increased and are expected to rise further with future generation systems. Some sectors of the data center industry (including DLC equipment makers, IT hardware vendors and some large operators) have agreed, through coordination by the Open Compute Project (OCP) and ASHRAE workgroups, to target flow rates in the range of 1.2-1.5 L/min·kW (liters per minute per kilowatt).

These are reference points for performance comparison rather than hard targets, but some IT vendors and operators prefer this range as a balance between pumping energy and heat exchange capacity in the coolant distribution units (CDUs). Convective heat transfer improves with flow, but this also increases the energy required for pumping.

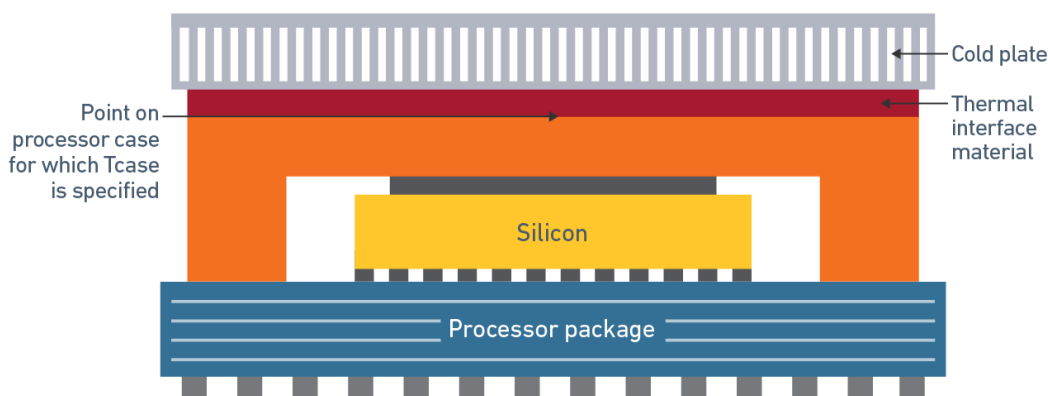
At these reference flow rates, modern cold plates can achieve very low average thermal resistance — around $0.02^{\circ}\text{C}/\text{W}$ including conductive and convective resistance. For example, a cold plate with this thermal resistance would require a temperature drop of only 10°C or less from its base plate to the coolant to transfer 500 W of thermal power.

Breaking the resistance

Despite all this, there are additional fixed components that contribute to the overall thermal resistance:

- The thermal interface material (TIM) between the cold plate base and the case of the chip.
- The case (heat spreader lid) itself.
- The internal interface between the case and the silicon, which also uses a TIM.

Figure 1 Cross section of a processor and cold plate stack



UPTIME INSTITUTE 2025

uptime
INTELLIGENCE

These components all add significant conductive thermal resistance, which limits how closely the coolant temperature can approach the silicon's operating temperature while still providing sufficient cooling. In practical terms, the total thermal resistance of the stack is why there is typically already a 40°C (104°F) limit on coolant supply temperature for GPU rack systems. Directionally, this will likely come down as chip TDPs and rack densities escalate. Some supercomputing systems, for example, allow coolant supply up to 32°C (89.6°F).

Typically, T_{case} specification (the maximum temperature allowed by the chip vendor to guarantee performance) is around or above 80°C (176°F). However, future high-performance processors and GPUs will likely drop maximum T_{case} to address the heat flux needs of silicon hot spots, potentially as low as 60°C (140°F) or below for some products.

So far, such systemic Tcase reductions have largely been avoided by increasing total silicon area (often doubling or tripling compared with just five years ago) in the package across which heat needs dissipating, effectively curbing or even reducing average heat flux. However, this approach does not solve long-term hot spot issues as circuits become smaller, and increasing total silicon area to drive performance will also reach its limits — package size limitations and economics will bring it to a stop within a few years.

Together with the limitations outlined above in this report, this will temper free cooling ambitions, even before broader facility infrastructure considerations, as previous Uptime Intelligence reports cautioned (see [*DLC will not come to the rescue of data center sustainability*](#)).

To address this, the data center industry, including chip vendors, will need to start cutting into the fixed components of the thermal resistance between the silicon and the cooling solution. The first step is improved precision mounting of the chip package through better alignment and the use of extremely thin, higher-performance TIMs. Cold plate manufacturers and IT vendors are already exploring new materials, including phase-change TIMs that require an initial “bake” to form the interface between the case and the base of the cold plate (or heat sink). This method can already enable several degrees of higher coolant temperature at the same flow rate.

Some approaches propose eliminating both the TIM and the cold plate base altogether, allowing the coolant to come into direct contact with the case (heat spreader lid). However, this may not be compatible with all cold plate designs or manufacturing techniques. The main benefit of this approach is that it does not depend on changes to processor and GPU packaging, and can be taken on by IT vendors or larger system integrators.

The next logical step is to eliminate not only the TIM but also the heat spreader from the stack (known as “delidding”) and mount the cold plate almost directly onto the silicon, with only a thin thermal interface in between. This approach has precedent in the IT industry but is no longer common. It requires careful mounting and calibration of pressure to avoid damaging the CPU or GPU silicon, which can cost several thousand dollars each. At scale, the direct-mount approach would likely require new automated assembly lines in the server supply chain to ensure quality control.

However, the benefit of delidding CPUs and GPUs to bring the thermal solution into closer contact with silicon is too big to ignore. While there is no single number that represents the potential benefit of such a development, in general it can enable coolant temperatures that are 5-10°C higher.

A dramatic reduction in effective thermal resistance from the silicon to the thermal solution will have a measurable effect on coolant temperature limits and the corresponding design assumptions for facility heat rejection. This opens the path toward designing facility systems that support ASHRAE W40, W45 and W+ category cooling for DLC loads, making them entirely dry-cooled without mechanical refrigeration in most climates.

Arguably, some DLC approaches — notably two-phase cold plates (which benefit from the high cooling capacity and thermal stability provided by latent heat of liquid evaporation) and single-phase immersion with forced convection — already make it possible to cool even dense compute systems using extreme TDP silicon (see [Immersion cooling evolves in response to IT power density](#)). However, water cold plates remain the most common choice today and will likely be so in the near future. As such, they need to be considered when engineering facility cooling infrastructure for future systems.

The Uptime Intelligence View

As cold plate and heat sink designs evolve — and flow rates increase — to improve their thermal performance in step with silicon advancements, engineers will turn their attention to other sources of thermal resistance within the cooling solution. Until now, chip and cold plate designs have effectively had a separation boundary between them. By chipping away at this boundary, opportunities are emerging to further improve overall cooling performance. Extreme silicon TDPs and highly efficient cooling need not be mutually exclusive — if the data center industry, IT systems vendors and chipmakers all work together.

ABOUT THE AUTHOR



Daniel Bizo

Over the past 15 years, Daniel has covered the business and technology of enterprise IT and infrastructure in various roles, including industry analyst and advisor. His research includes sustainability, operations, and energy efficiency within the data center, on topics like emerging battery technologies, thermal operation guidelines, and processor chip technology.

[**dbizo@uptimeinstitute.com**](mailto:dbizo@uptimeinstitute.com)

About Uptime Institute

Uptime Institute is the Global Digital Infrastructure Authority. Its Tier Standard is the IT industry's most trusted and adopted global standard for the proper design, construction, and operation of data centers – the backbone of the digital economy. For over 25 years, the company has served as the standard for data center reliability, sustainability, and efficiency, providing customers assurance that their digital infrastructure can perform at a level that is consistent with their business needs across a wide array of operating conditions.

With its data center Tier Standard & Certifications, Management & Operations reviews, broad range of related risk and performance assessments, and accredited educational curriculum completed by over 10,000 data center professionals, Uptime Institute has helped thousands of companies, in over 100 countries to optimize critical IT assets while managing costs, resources, and efficiency.