

WHITE PAPER



SGI[®] ICE[™] X

Ultimate Flexibility for the World's Fastest Supercomputer

TABLE OF CONTENTS

1.0 SGI ICE X: Ultimate Flexibility to Serve Any Workload	1
2.0 Architected for Density and Power Flexibility	1
2.1 Modular Enclosure with a Unified Backplane	1
2.2 Independently-Scalable Power Shelf	2
2.3 A Choice of Powerful Compute Blades	4
3.0 FDR InfiniBand Interconnect Flexibility at the Node, Switch, and Topology Levels	6
3.1 A Choice of FDR InfiniBand Mezzanine Cards	6
3.2 A Choice of FDR InfiniBand Switch Blades	7
3.3 A Broad Choice of InfiniBand Topologies	9
3.4 Scalable Out-of-Band Management	9
4.0 Cooling Flexibility at Node and Rack Level	10
4.1 SGI ICE X D-Rack for Standard Hot Aisle / Cold Aisle Installations	10
4.2 SGI ICE X Cells for Large-Scale HPC Installations	11
4.3 Innovative SGI Cold Sink Technology	12
5.0 Conclusion	13

1.0 SGI ICE X: Ultimate Flexibility to Serve Any Workload

Modern technical computing challenges drive an insatiable appetite for computing power in the quest to solve the world's most complex scientific, technical, and business-related problems. To meet these demands, horizontally-scalable clusters of dense and powerful compute nodes are now widely deployed in many industries, with supercomputing clusters now easily dominating the TOP500 list of supercomputing sites (www.top500.org). In order to be effective, however, these interconnected systems must be able to scale without arbitrary limitations, and overcome the very real constraints that organizations face in terms of real estate, power, and cooling resources.

Despite the trend towards large supercomputing clusters based on industry-standard components, there is no one-size-fits-all solution in high-performance computing (HPC). As HPC technology is applied to solve new problems, or extended to solve old problems at greater scale, system flexibility has become perhaps as important as overall performance. Different application needs drive requirements for specific computational configurations and interconnect topologies. Different physical environments drive the need for specific deployment solutions in terms of physical footprint, power, and cooling. Effective supercomputing architectures must simultaneously enable the latest technologies while providing maximum flexibility in terms of interconnect, power, and cooling.

SGI has considerable expertise and decades of experience deploying some of the largest InfiniBand clusters in existence. While some vendors push solutions based on their own limitations, SGI understands that the best infrastructure offers maximum performance along with the flexibility to match the needs of the application and the organization. Next-generation SGI ICE X systems embody this approach, providing the largest and fastest pure compute InfiniBand cluster in the world, along with ultimate scalability and flexibility to serve any workload.

- The SGI ICE platform has been the world's fastest distributed memory supercomputer for over four years running. This performance leadership is proven time and again, not just in the lab, but at customer sites including the largest and fastest pure compute InfiniBand cluster in the world.
- Petascale-class, and with a clear roadmap to exascale performance, these systems offer seamless scalability from tens of teraflops to hundreds of petaflops, all within a system based on industry-standard components.
- Unlike many competitors, SGI ICE X systems power up and go at scale, enabling deployments in hours or days, not weeks or months — allowing upgrades and expansion on demand without cluster downtime.

This paper describes the capabilities of the SGI ICE X system, with a particular focus on architectural innovations that drive flexibility in terms of power, cooling, FDR InfiniBand interconnect, and topology.

2.0 Architected for Density and Power Flexibility

SGI ICE X systems are designed for expandability, both within and across technology generations. With SGI ICE X, engineers paid careful attention to the chassis, specifying a modular chassis with a unified backplane that supports both power flexibility and a choice of powerful compute blades.

2.1 Modular Enclosure with a Unified Backplane

As shown in Figure 1, SGI ICE X systems are based on a 9.5 rack unit (9.5U) modular enclosure. Unlike previous-generation systems, a single unified backplane is provided across all system configurations,

supplying power to individual compute blade slots, and connecting the blade slots to the FDR InfiniBand switch blade slots. Enclosures are provided in pairs as a modular building block. Each dual-enclosure pair provides:

- 36 blade slots for either single- or dual-node SGI ICE X compute blades based on Intel® Xeon® processor E5-2600 family
- Four slots for Fourteen Data Rate (FDR) InfiniBand switch blades
- Four chassis management controller slots

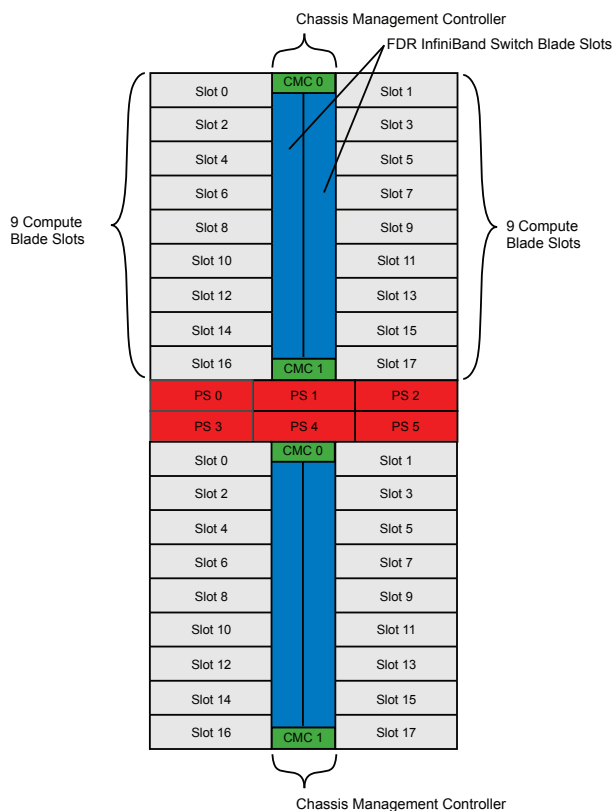


Figure 1. Each SGI ICE X enclosure links 36 blade slots (up to 72 compute blades) via FDR InfiniBand.

2.2 Independently-Scalable Power Shelf

Blade chassis power supplies are typically tightly integrated and designed to meet the power needs of specific components, along with necessary capacity and redundancy to accommodate continued operation in the face of potential individual power supply failures. While common, this approach can present several disadvantages and limitations, particularly when systems are deployed redundantly in large HPC deployments.

- Available chassis power ultimately limits the components and processor technology that can be deployed. For instance, as new generations of processors become available, blades with higher-power processors might not be supportable, forcing chassis or rack-level upgrades.
- Power supplies are typically configured in redundant fashion. While necessary to supply high availability, redundancy on an enclosure level can quickly result in substantial numbers of extra power supplies being deployed. Providing multiple blade chassis per rack, with each configured for N+1 redundancy, can quickly result in a significant number of extra (unused) power supplies.
- Extraneous power supplies continue to draw some power, even though they are not in active use powering components, wasting energy and generating excess heat.

The SGI ICE X independently-scalable power shelf addresses these issues by increasing the number of power supplies that can be combined to power the enclosures, while at the same time reducing the number of spare power supplies necessary for redundancy. The independently-scalable power shelf is shared across multiple enclosures, and all connected power supplies energize a shared 12V power bus. The shelf enables SGI to deliver considerable flexibility depending on rack configuration and the power requirements of installed components. As a result, enclosures can gain more usable power with fewer power supplies and less waste. For example, a 5+1 power supply configuration can provide more *usable* power than dual 2+1 configurations, while reducing the spare power supplies by half. These economies of scale only improve with larger configurations.

This innovative approach enables flexible scaling of available power to meet specific per-node power level needs, and accommodates future technology generations that may require greater power. The system is designed to scale from 400W per node to as much as 1,440W per node, and power shelves can be upgraded independently as new power technologies become available. When less power is required, power shelves can be populated only as required to provide the power needed to energize installed components.

Figure 2 illustrates differing power shelf configurations for the SGI ICE X D-Rack and M-Rack, respectively. In the D-Rack, two 9.5 rack-unit (9.5U) rack-mount enclosures are coupled with two shared and independently-scalable power shelves that provide power to both enclosures through the shared 12V power bus. In contrast, three power shelves (and up to 9 power supplies) power a 12V bus bar on each side of the enclosures in the M-Rack, yielding considerable power supply capacity and flexibility.

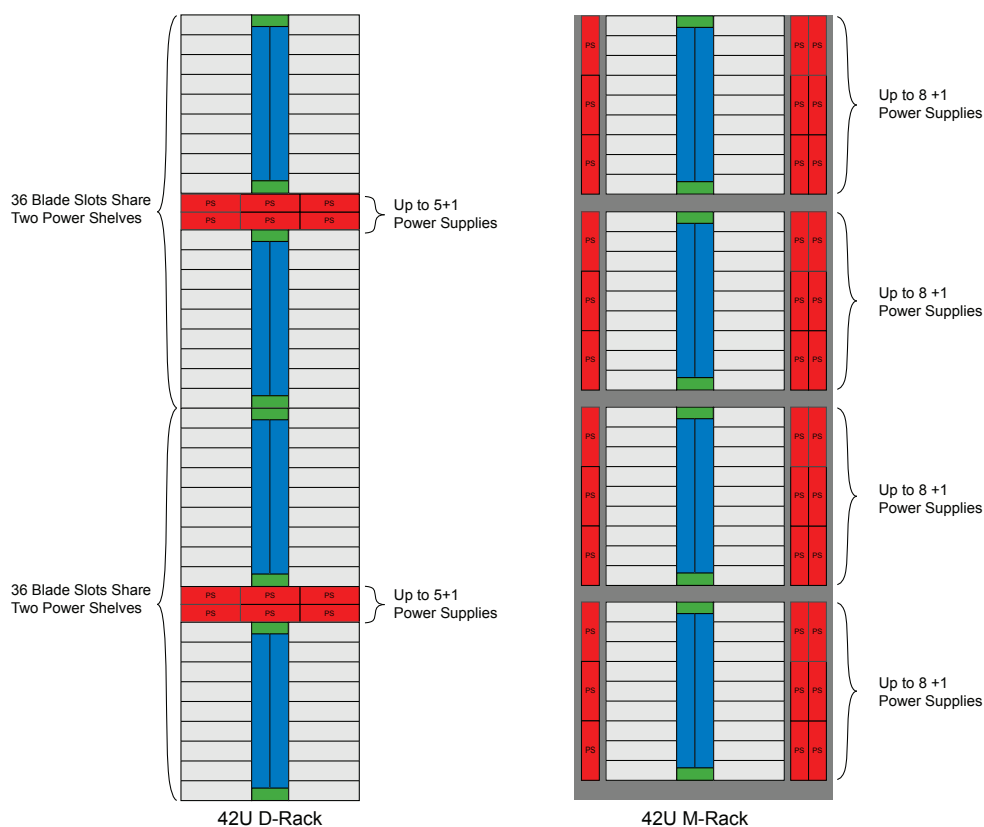


Figure 2. The independently-scalable power shelf can be used to provide power capacity and rack flexibility.

2.3 A Choice of Powerful Compute Blades

To accommodate a wide range of application needs, SGI ICE X systems provide a choice of compute blades based on the Intel® Xeon® processor E5-2600 family, with support for Intel® Advanced Vector Extensions, Intel® Trusted Execution Technology, and Intel® AES New instructions. These powerful compute blades will offer two sockets per node, with each socket supporting 12 cores and 24 threads. To provide maximum interconnect bandwidth, SGI has worked closely with Intel® and others to offer industry-leading FDR InfiniBand capabilities that are designed into the product. However, unlike previous generation SGI ICE systems, InfiniBand capabilities are provided via a mezzanine card, allowing a choice of InfiniBand interface options. InfiniBand and cooling options for supported compute blades are discussed later in this document.

SGI ICE X IP-113 Single-Node Compute Blade

The SGI ICE X IP-113 compute blade is designed for standard 19-inch (48.3cm) rack-mount deployments and special-purpose HPC deployments in the M-Rack and M-Cell. Shown in Figure 3, the compute blade features the latest high-performance Intel® Xeon® processor technology, and includes:

- Two sockets for Intel® Xeon® processor E5-2600 family
- Eight DIMM sockets per processor socket, offering up to 1866MT/s memory speed
- Two 2.5-inch SATA HDDs or solid state devices (SSDs) supported on each node

An interface on the compute node offers a choice of FDR InfiniBand mezzanine cards, with two x8 PCI Express connections, one leading to each of the two processor sockets. Information on available FDR InfiniBand options is provided later in this document.

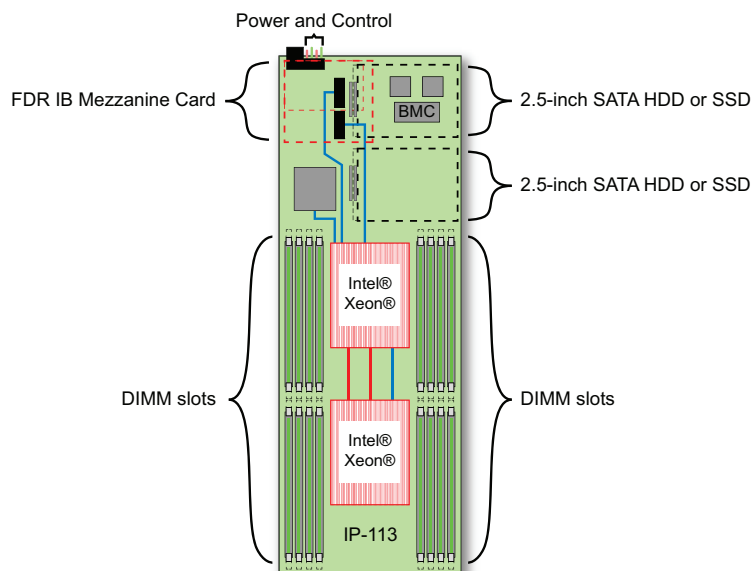


Figure 3. High-level perspective of the SGI ICE X IP-113 compute blade.

With up to 18 SGI ICE X IP-113 compute nodes installed in each enclosure, each SGI ICE D-Rack provides up to:

- 72 compute nodes
- 144 Intel® Xeon® processor E5-2600 family
- 1,728 cores and 3,456 threads

SGI ICE X IP-115 Dual-Node Compute Blade

The SGI ICE X IP-115 dual-node compute blade is designed for high-density and special-purpose HPC deployments in the M-Rack and M-Cell. Shown in Figure 4, the compute blade is actually two logically separate compute nodes that are physically joined together to share a single slot in the SGI ICE X enclosure. This dual-node configuration effectively doubles the achievable compute density over the IP-113 blade in the same physical space.

Each SGI ICE X IP-115 compute blade features:

- A dual-node configuration with physically joined, but logically separate, nodes
- A total of four sockets (two per node) for Intel® Xeon® processor E5-2600 family
- Dual single-port FDR InfiniBand connections to the unified backplane via mezzanine card
- Eight DDR3-1866 DIMMs per compute node (4 per socket, 1 per bus)

An interface on each compute node connects to a dual single-port InfiniBand mezzanine card, with each FDR InfiniBand HCA connected to one of the processors on each node with an x8 PCI Express connection.

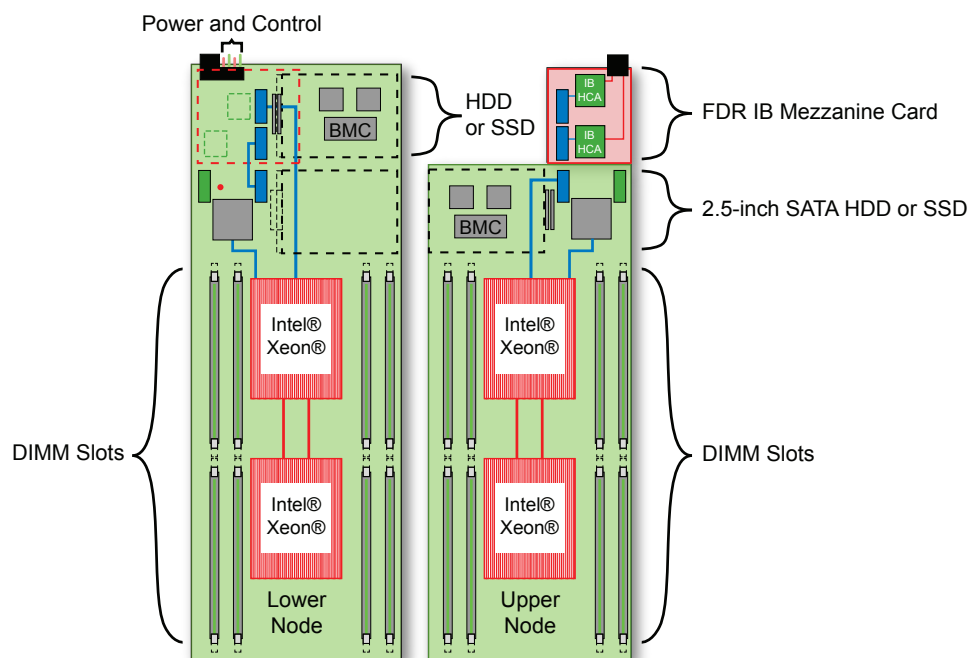


Figure 4. High-level perspective of the SGI ICE X IP-115 compute blade.

SGI ICE X IP-115 compute blades can be configured with traditional heat sinks or optional SGI cold sink technology, depending on processor power requirements, heat output, and deployment considerations. SGI cold sink technology replaces the heat sinks with liquid-cooled cold plates located between the processor cores on the upper and lower blades, as shown in Figure 5. Depending on the configuration, SGI cold sink technology will directly absorb approximately 55-70% of the heat dissipated by an individual dual-socket node. The remainder of the heat is then air-cooled by the closed-loop method provided by the SGI M-Rack, as described in the section on cooling flexibility later in this document.



Figure 5. SGI cold sink technology replaces traditional heat sinks with liquid-cooled cold plates for processor technology that requires additional cooling.

3.0 FDR InfiniBand Interconnect Flexibility at the Node, Switch, and Topology Levels

Organizations need the opportunity to make flexible topology choices rather than having to accept less than ideal configurations brought on by vendor limitations. Interconnect flexibility is particularly important to providing the right match between applications and deployed infrastructure. Different applications have different bandwidth needs and imply different topology choices. Those topology choices, in turn, can have an effect on the infrastructure requirements to meet performance needs at the best possible cost.

- Smaller clusters or lower bandwidth applications require fewer fabric connections and should be provided with affordable solutions that meet their needs.
- Larger topologies or more demanding applications may require redundant connections to spread out message passing load, or to separate message passing from I/O.

The SGI ICE X system was designed to provide a range of choices at the node, switch, and topology levels to allow maximum flexibility in interconnect design.

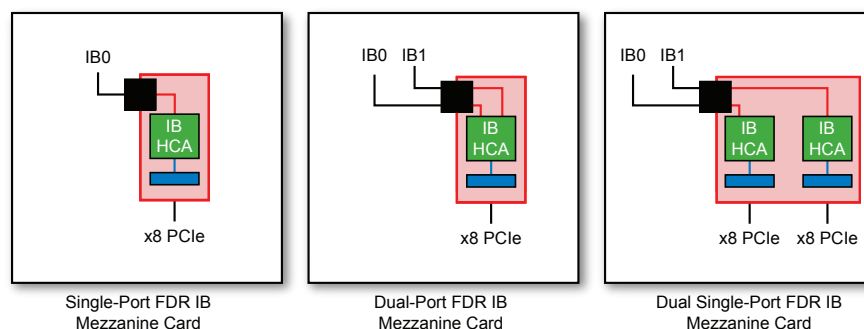
3.1 A Choice of FDR InfiniBand Mezzanine Cards

Integral to the SGI ICE X system, Fourteen Data Rate (FDR) InfiniBand is the next generation of InfiniBand technology developed and specified by the InfiniBand Trade Association (IBTA). FDR InfiniBand offers a 14 Gbps data rate per lane, as compared to only a 10 Gbps data rate per lane for Quad Data Rate (QDR) InfiniBand supported by most other vendors. Unlike previous generation SGI ICE systems that included InfiniBand on the compute node itself, SGI ICE X provides a choice of FDR InfiniBand mezzanine cards to accommodate the bandwidth needs of a wide range of applications. Supported FDR InfiniBand mezzanine cards are summarized in Table 1 and depicted in Figure 6.

- The Single-Port FDR InfiniBand mezzanine card offers a single FDR InfiniBand HCA connected to an x8 PCI Express interface provided by one of the SGI ICE X IP-113 compute blade processors.
- The Dual-Port FDR InfiniBand mezzanine card offers a dual-port FDR InfiniBand HCA connected to an x8 PCI Express interface provided by one of the SGI ICE X IP-113 compute blade processors.
- The Dual Single-Port FDR InfiniBand mezzanine card provides two single-port FDR InfiniBand HCAs, each connected to a separate x8 PCI Express interface. On the SGI ICE X IP-113 compute blade, the two single-port HCAs each connect to one Intel® Xeon® processor socket. On the SGI ICE X IP-115 compute blade, each of the single-node HCAs connects to one of the two compute nodes.

Table 1. Supported FDR InfiniBand mezzanine cards in SGI ICE X blades.

SGI ICE X Compute Blade	Single-Port FDR IB Card	Dual-Port FDR IB Card	Dual Single-Port FDR IB card
SGI ICE X IP-113 Single-Node Compute Blade	X	X	X
SGI ICE X IP-115 Dual-Node Compute Blade			X

**Figure 6.** SGI ICE X systems provide a choice of FDR InfiniBand mezzanine cards to suit application and topology needs.

3.2 A Choice of FDR InfiniBand Switch Blades

To enable a wide range of topologies and deployment options, SGI ICE X systems provide a choice of FDR InfiniBand switch blades. Two switch blades insert into the front of each enclosure, connecting the 18 compute blade slots together via the unified backplane. The switch blades provide different configurations of FDR InfiniBand switch chips as well as external Quad Small Form Factor Pluggable (QSFP) ports that can be used to connect to other chassis.

SGI ICE X Standard FDR IB Switch Blade

The SGI ICE X IB Standard Blade is optimal for building fat tree and all-to-all topologies, and can also be used for building hypercube or small enhanced hypercube systems. Shown in Figure 7, the standard blade provides a single 36-port Mellanox® FDR InfiniBand switch ASIC with 18 ports connected to the unified backplane, and 18 ports connected to the external bulkhead that is accessible when the switch blade is inserted into the enclosure. Up to two switch blades can be installed, depending on whether the deployment is single-plane or dual-plane. When two switch blades are installed, each connects to one of two connections provided in each compute blade slot in the enclosure.

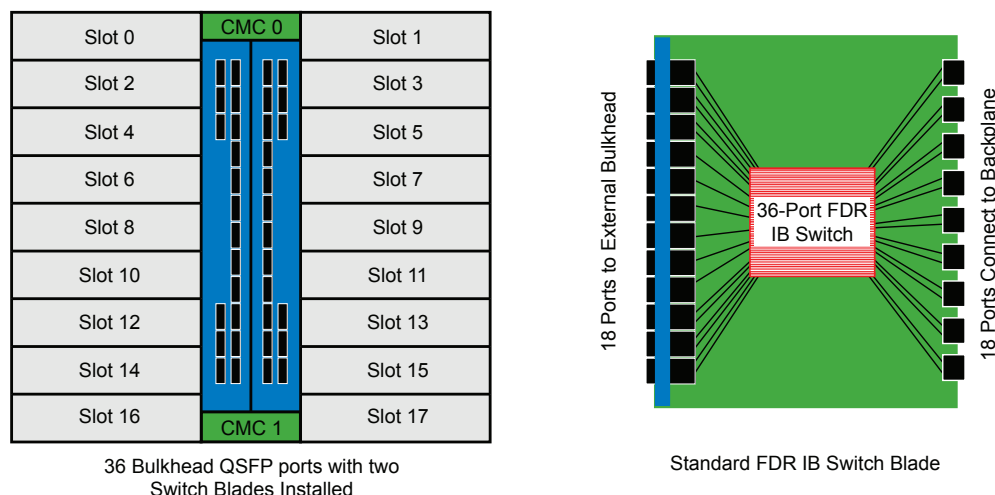


Figure 7. The SGI ICE X IB Standard Blade connects the 18 enclosure slots together with a single 36-port Mellanox FDR InfiniBand switch chip, providing 18 QSFP ports on the bulkhead for external connections.

SGI ICE X Premium FDR IB Switch Blade

The SGI ICE X IB Premium Blade is optimal for building higher-bandwidth all-to-all topologies as well as larger scale enhanced hypercube topologies. Shown in Figure 8, the premium blade provides the first level of interconnect via dual 36-port Mellanox FDR InfiniBand switch ASICs with connections as follows:

- 9 ports from each switch chip connect to the unified backplane, to connect the 18 compute node slots
- 3 ports on each chip provide connectivity between the chips
- 24 ports from each switch chip connect to the external bulkhead, for a total of 48

Up to two switch blades can be installed, depending on the nature of the deployment. When two switch blades are installed, each connects to one of two ports provided in each compute blade slot in the enclosure.

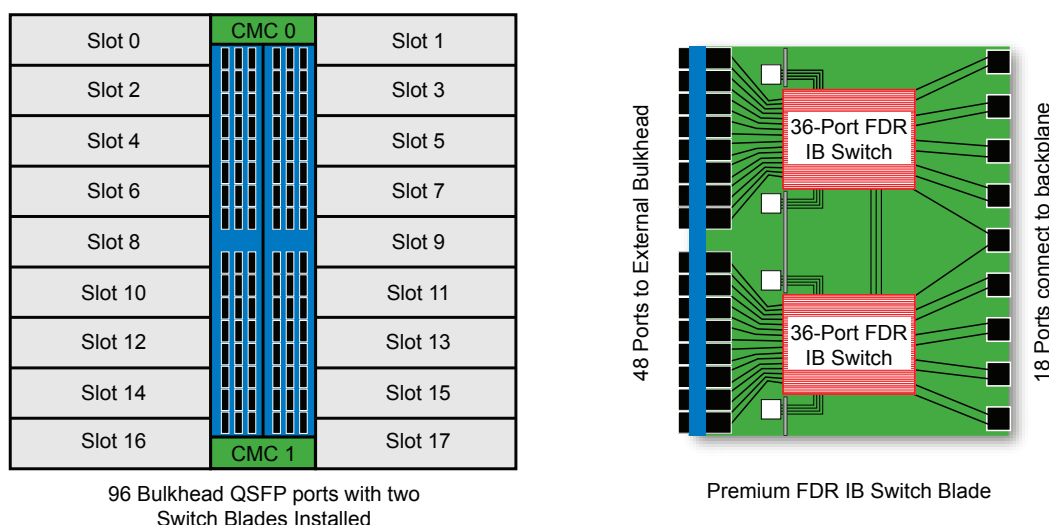


Figure 8. The SGI ICE X IB Premium Blade connects the 18 enclosure slots together with two 36-port Mellanox FDR InfiniBand switch chips, providing 48 QSFP ports on the bulkhead for external connections.

3.3 A Broad Choice of InfiniBand Topologies

SGI has considerable experience in the design and deployment of some of the largest InfiniBand clusters in existence. SGI utilizes that expertise to deliver a system that is designed for flexible and optimized InfiniBand topology configuration. To accommodate the widest possible range of application requirements and needs for both high bandwidth and low latency, SGI ICE X supports multiple InfiniBand topology choices, including:

- **All-to-all.** All-to-all topologies are ideal for applications that are highly sensitive to Message Passing Interface (MPI) latency since they provide minimal latency in terms of hop-count. Though all-to-all topologies can provide non-blocking fabrics and high bisection bandwidth, they are restricted to relatively small cluster deployments because of limited switch port counts.
- **Fat tree.** Fat tree or CLOS topologies are well suited for smaller node-count MPI jobs. Fat tree topologies can provide non-blocking fabrics and consistent hop-counts, resulting in predictable latency for MPI jobs. At the same time, fat tree topologies do not scale linearly with cluster size. Cabling and switching become increasingly difficult and expensive as cluster size grows, with very large-core switches required for larger clusters.
- **Standard hypercube.** Standard hypercube topologies are ideal for large node-count MPI jobs, provide rich bandwidth capabilities, and scale easily from small to extremely large clusters. Hypercubes add orthogonal dimensions of interconnect as they grow, and are easily optimized for both local and global communication within the cluster. Standard hypercube topology provides the lightest weight fabric at the lowest cost, with a single cable typically used for each dimensional link.
- **SGI enhanced hypercube.** Adding to the benefits of standard hypercube topologies, SGI enhanced hypercube topology makes use of additional available switch ports by adding redundant links at the lower dimensions of the hypercube to improve the overall bandwidth of the interconnect.

Torus topologies are not supported by SGI as they present scalability problems as well as introduce hop-count and latency with increasing cluster size. Hypercube topologies offer many of the advantages of torus topologies, while providing scalability and linear latency (hop-count) scalability. Hypercubes also offer higher connection capabilities and resources at the node level where they can be most effectively utilized.

3.4 Scalable Out-of-Band Management

Analyzing large amounts of environmental, compute, operating system, and general ecosystem data is critical in maintaining and managing large scale-out resources such as supercomputing clusters. Careful coordination of scheduling, power efficiency, failures, and monetary cost are essential to effective operation. With SGI ICE X, these monitoring activities are facilitated through the Failure Analysis and Power Management subsystems of SGI Management Suite, which provides the administrator with concise and pertinent data around error conditions and the state of their SGI ICE X system. This data, represented graphically in failure analysis and instrumentation panels, can then be utilized to monitor, correct, contain, or replace hardware.

SGI ICE X systems provide support for these administrative activities through a thinly-threaded hierarchical Gigabit Ethernet network. This out-of-band network leverages traditional TCP/IP and Ethernet networking models, without imposing on the InfiniBand network. A hierarchical structure ensures that management capabilities scale effectively to even the largest supercomputing clusters.

- **Blade Management Controllers (BMCs).** BMCs on each compute node control board-level hardware and monitor the compute node environment.
- **Chassis Management Controllers (CMCs).** Each BMC reports up to two or four CMCs per blade enclosure pair. The CMCs control master power to all compute nodes, and monitor power and the blade enclosure environment.
- **Rack Level Controllers (RLCs).** The CMCs aggregate up to an RLC, which is provided for each pair of blade enclosures. The RLC holds blade boot images, runs fabric management software, and collects cluster management data for the rack.
- **System Administration Controller (SAC).** An SAC is provided for each ICE system. The SAC provisions software to the RLC and pulls aggregated cluster management data from the RLC.

4.0 Cooling Flexibility at Node and Rack Level

Just as application and topology requirements vary widely, physical constraints can drive many aspects of supercomputing cluster deployments. Improved density is clearly a goal, but physical characteristics including effective and efficient cooling are essential in large HPC data centers. To address these needs, SGI ICE X provides considerable cooling flexibility:

- SGI ICE X D-Racks for traditional hot/cold aisle air-cooled installations
- SGI ICE Cells to support closed-loop cooling solutions and hot-aisle containment
- SGI cold sink technology to support increased node density per blade slot and to accommodate new technology

4.1 SGI ICE X D-Rack for Standard Hot Aisle/Cold Aisle Installations

Many organizations want to deploy systems in standard 19-inch (48.3cm) racks optimized for 24-inch (61.0cm) floor tile systems and hot-aisle/cold-aisle air handling. To support this need, SGI ICE X enclosures can be deployed in the SGI D-Rack and equipped with SGI ICE X IP-113 compute blades. In this configuration, each 42U rack can house two enclosure pairs, with each pair joined with two independent power shelves. Each D-Rack has capacity for 72 dual-processor blades, yielding a 1.4x density increase over previous generation SGI ICE systems.

Illustrated in Figure 9, the D-Rack supports both open-loop air cooling as well as optional chilled water coils with open-loop airflow. Rather than enclosure-level fans, the design provides larger and more efficient fans at the rack level. The optional chilled water coils are supplied as a door that affixes to the rear of the chassis to cool hot air as it exits the rack and before it escapes into the data center.

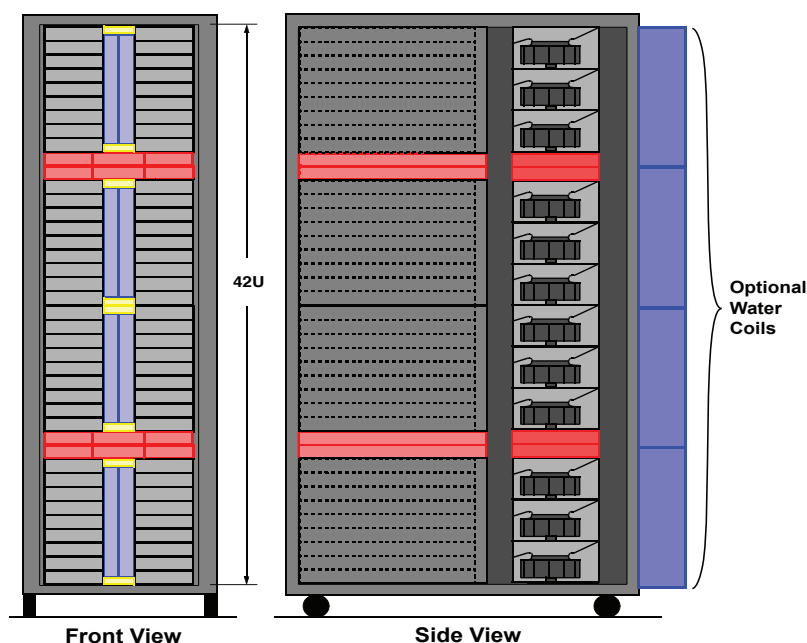


Figure 9. SGI ICE X IP-113 compute blades can be installed in the SGI D-Rack with open-loop air cooling, which can be combined with optional water coils.

4.2 SGI ICE X Cells for Large-Scale HPC Installations

Large-scale HPC installations often have scale and density requirements that are less well-served by standard 19-inch (48.3cm) rack deployments. In addition, modular data center deployments such as 40-foot ISO containers are increasingly popular for focused deployments of compute power in a range of environments. SGI ICE X Cells are designed to accommodate these demanding installations, offering highly-efficient closed-loop cooling solutions as well as improved density. As shown in Figure 10, each SGI ICE X Cell contains four compute racks and two cooling racks that combine air and water cooling. Access is provided through doors on both sides of the Cell.

SGI ICE X Cells provide distinct advantages, including:

- **Closed-loop airflow.** Each Cell features integrated hot aisle containment, and no air from within the Cell is mixed with data center air. All air within the Cell is water-cooled. This closed-loop airflow arrangement has the benefit of reducing acoustic emissions relative to open-loop systems, and essentially creates embedded hot aisle containment within the Cell.
- **Warm-water cooling.** The Cells support a broad range of facilities-supplied cooling water in the range of 45-85 degrees Fahrenheit. The higher temperature limits help enable more annual hours of free cooling, where the water can be supplied without the need for operating a chiller plant. This efficiency can result in cost savings, even from very dense supercomputing cluster deployments. An air-to-water heat exchanger is provided with all Cells, and a liquid-to-water heat exchanger is deployed with SGI cold sink technology.
- **Unified cooling racks.** Unlike standalone racks, compute racks deployed within the SGI ICE X Cell do not have their own rack-level cooling, instead relying on integral cooling racks within the Cell. The cooling racks draw hot air across water-cooled heat exchangers and recirculate it to cool the compute racks within the Cell. This approach provides greater efficiency over rack-level cooling, decreasing power costs associated with cooling and utilizing a single water source.

Shown in plan view in Figure 10 and detailed in Table 2, the SGI ICE X Cell is available to complement the D-Rack, and serve a variety of deployment requirements:

- **SGI ICE X M-Cell.** M-Cells offer the ability to double compute density over the D Rack with the addition of dual-node SGI ICE X IP-115 compute blades. M-Cells are also designed to support future higher-wattage technology upgrades. The M-Rack within the M-Cell provides more available power from more shared power supplies, since up to 9 power supplies share a single 12W power bus bar (Figure 2). In addition, SGI cold sink technology cools hotter components more effectively, allowing support for hotter components.

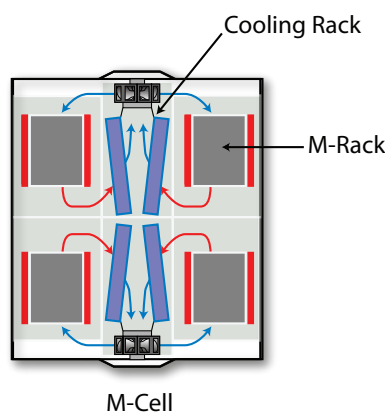


Figure 10. M-Cells utilize closed-loop airflow and warm-water cooling to create embedded hot-aisle containment within the Cell (Cell shown depicted from above).

Table 2. SGI ICE X D-Rack, D-Cell, and M-Cell capabilities.

Deployment Option	Blade Slots	Square Footage	Compute Nodes	Cooling	SGI ICE X Blades
SGI ICE X D-Rack	72	6.67	72	Open-loop air	IP-113 compute blade
SGI ICE X M-Cell	288	80	up to 576	Closed-loop air and warm water, with optional SGI cold sink technology	IP-113 compute blade IP-115 compute blade

4.3 Innovative SGI Cold Sink Technology

With over three decades of building and deploying sophisticated HPC systems, SGI has considerable depth of expertise in designing effective cooling systems. SGI cold sink technology is an example of this technology prowess, providing an elegant, yet efficient, cooling system that can take advantage of warm-water cooling infrastructure to facilitate increased computational density while minimizing the need for special-purpose chillers.

In addition to the closed-loop airflow, SGI ICE X M-Rack deployments using IP-115 compute nodes can utilize SGI cold sink technology. Depending on the wattage of deployed processors, SGI cold sink technology replaces the traditional heat sinks on the compute nodes with liquid-cooled cold plates located between the sockets on the upper and lower nodes. The cooling liquid flows through the pair of cold plates in series, and on to liquid distribution manifolds integrated with blade enclosures (Figure 11). Depending on the configuration, SGI cold sink technology can directly absorb approximately 50-65 percent of the heat dissipated by an individual dual-socket node. The remainder of the heat is air-cooled by the closed-loop airflow present in the M-Cell. The cooling rack within the M-Cell is also configured differently when using SGI cold sink technology rather than air cooling, and features a liquid-to-liquid plate heat exchanger and secondary pumping loop.

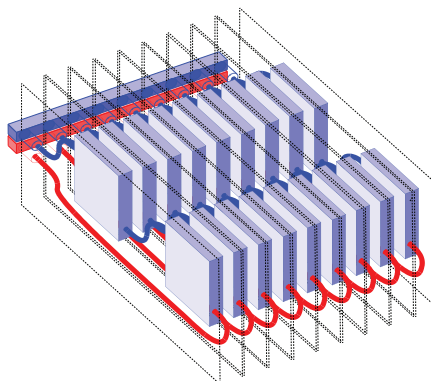


Figure 11. SGI cold sink technology provides water cooled cold plates between processors for higher-wattage versions of the SGI ICE X IP-115 compute blade in an M-Cell.

5.0 Conclusion

SGI ICE X was designed from the outset with ultimate flexibility to accommodate an extraordinarily wide range of HPC computing requirements. Independently-scalable power shelf provides more available power to the chassis and power flexibility for future technology advancements, while reducing redundant power supply deployment and costly power drain. Innovative FDR InfiniBand mezzanine cards, along with a choice of switches and InfiniBand topologies, provide for unparalleled interconnect flexibility. Cooling solutions include air and warm water cooling options to fit both system needs and physical plant requirements for today, tomorrow, and into the future.

With SGI ICE X, a choice of dual-processor compute nodes based on Intel® Xeon® processor E5-2600 family deliver compute power, memory capacity, and throughput, while rack-mount choices accommodate both traditional 19-inch (48.3cm) racks and the SGI high-density M-Cell. Innovations such as SGI cold sink technology leverage a legacy of successful designs and HPC deployments, resulting in a system that can be configured to address the broadest range of HPC computing needs. Best of all, even large SGI ICE X systems can be deployed rapidly, and can even be enhanced and upgraded while the cluster remains in production.

Global Sales and Support: sgi.com/global

©2013 Silicon Graphics International Corp. All rights reserved. SGI, the SGI logo and ICE are registered trademarks or trademarks of Silicon Graphics International Corp. or its subsidiaries in the United States and/or other countries. Intel, Xeon and the Intel Xeon logo are registered trademarks of Intel Corporation. All other trademarks are property of their respective holders. 07102012 4329

