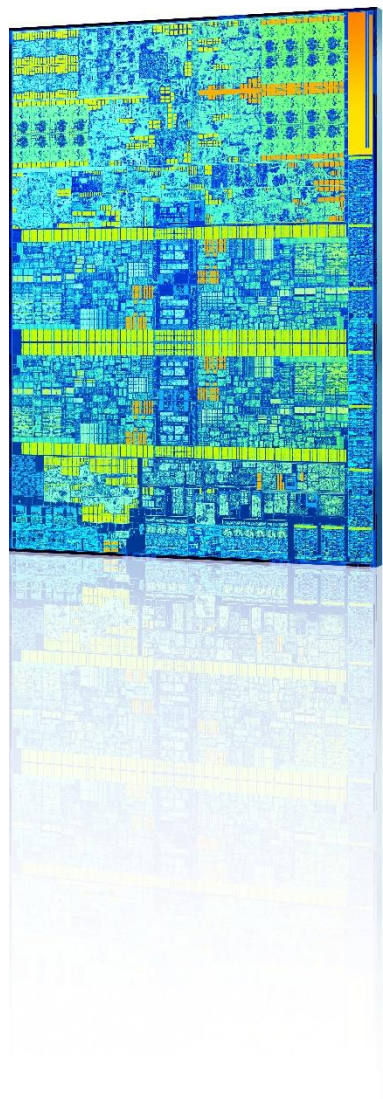


The Compute Architecture of Intel® Processor Graphics Gen9

Version 1.0



External Revision History

Ver	Date	Comment
v1.0	8/14/2015	IDF-2015 release of “The Compute Architecture of Intel Processor Graphics Gen9” – Stephen Junkins

1 CONTENTS

2	Audience	2
3	Introduction	2
3.1	What is Intel Processor Graphics?.....	2
4	SoC Architecture	3
4.1	SoC Architecture.....	4
4.2	Ring Interconnect.....	4
4.3	Shared LLC	4
4.4	Optional EDRAM.....	4
5	The Compute Architecture of Intel Processor Graphics Gen9.....	5
5.1	New Changes for Intel Processor Graphics Gen9.....	5
5.2	Modular Design for Product Scalability	5
5.3	Execution Unit (EUs) Architecture	6
5.3.1	Simultaneous Multi-Threading and Multiple Issue Execution	7
5.3.2	SIMD FPUs	7
5.3.3	Branch and Send Units.....	7
5.3.4	EU ISA and Flexible Width SIMD	7
5.3.5	SIMD Code Generation for SPMD Programming Models.....	8
5.4	Subslice Architecture	8
5.4.1	Sampler.....	9
5.4.2	Data Port.....	10
5.5	Slice Architecture	10
5.5.1	Level-3 Data Cache.....	11
5.5.2	Shared Local Memory.....	11
5.5.3	Barriers and Atomics	12
5.5.4	64-Byte Data Width	12
5.6	Product Architecture.....	12
5.6.1	Command Streamer and Global Thread Dispatcher	14
5.6.2	Graphics Technology Interface (GTI)	14
5.6.3	Unslice	15
5.6.4	Product EU Counts.....	15
5.7	Memory	15
5.7.1	Unified Memory Architecture.....	15
5.7.2	Shared Memory Coherency	15
5.8	Architecture Configurations, Speeds, and Feeds	17
6	Example Compute Applications	17
7	Acknowledgements	19
8	More Information	20
9	Notices.....	21

2 AUDIENCE

Software, hardware, and product engineers who seek to understand the architecture of Intel® processor graphics gen9. More specifically, those seeking to understand the architecture characteristics relevant to compute applications on Intel processor graphics.

This gen9 whitepaper updates the material found in “The Compute Architecture of Intel Processor Graphics Gen8” so that it can stand on its own. But where necessary, specific architecture changes for gen9 are noted.

3 INTRODUCTION

Intel’s on-die integrated processor graphics architecture offers outstanding real time 3D rendering and media performance. However, its underlying compute architecture also offers general purpose compute capabilities that approach teraFLOPS performance. The architecture of Intel processor graphics delivers a full complement of high-throughput floating-point and integer compute capabilities, a layered high bandwidth memory hierarchy, and deep integration with on-die CPUs and other on-die system-on-a-chip (**SoC**) devices. Moreover, it is a modular architecture that achieves scalability for a family of products that range from cellphones to tablets and laptops, to high end desktops and servers.

3.1 WHAT IS INTEL PROCESSOR GRAPHICS?

Intel processor graphics is the technology that provides graphics, compute, media, and display capabilities for many of Intel’s processor SoC products. At Intel, architects colloquially refer to Intel processor graphics architecture as simply “**Gen**”, shorthand for Generation. A specific generation of the Intel processor graphics architecture may be referred to as “Gen7” for generation 7, or “gen” for generation 8, etc. The branded products Intel HD Graphics 5600, Intel Iris™ Graphics 6100, and Intel Iris Pro Graphics 6200 are all derived from instances of Intel processor graphics gen8 architecture. Intel HD Graphics 530 is the first released product derived from an instance of Intel processor graphics gen9 architecture. (Note: graphics product naming conventions changed with gen9, from 4 digits to 3.)

This whitepaper focuses on just the compute architecture components of Intel processor graphics gen9. For shorthand, in this paper we may use the term **gen9 compute architecture** to refer to just those compute components. The whitepaper also briefly discusses the gen9 derived instantiation of Intel HD Graphics 530 in the recently released Intel Core™ i7 processor 6700K for desktop form factors. Additional processor products that include Intel processor graphics gen9 will be released in the near future.

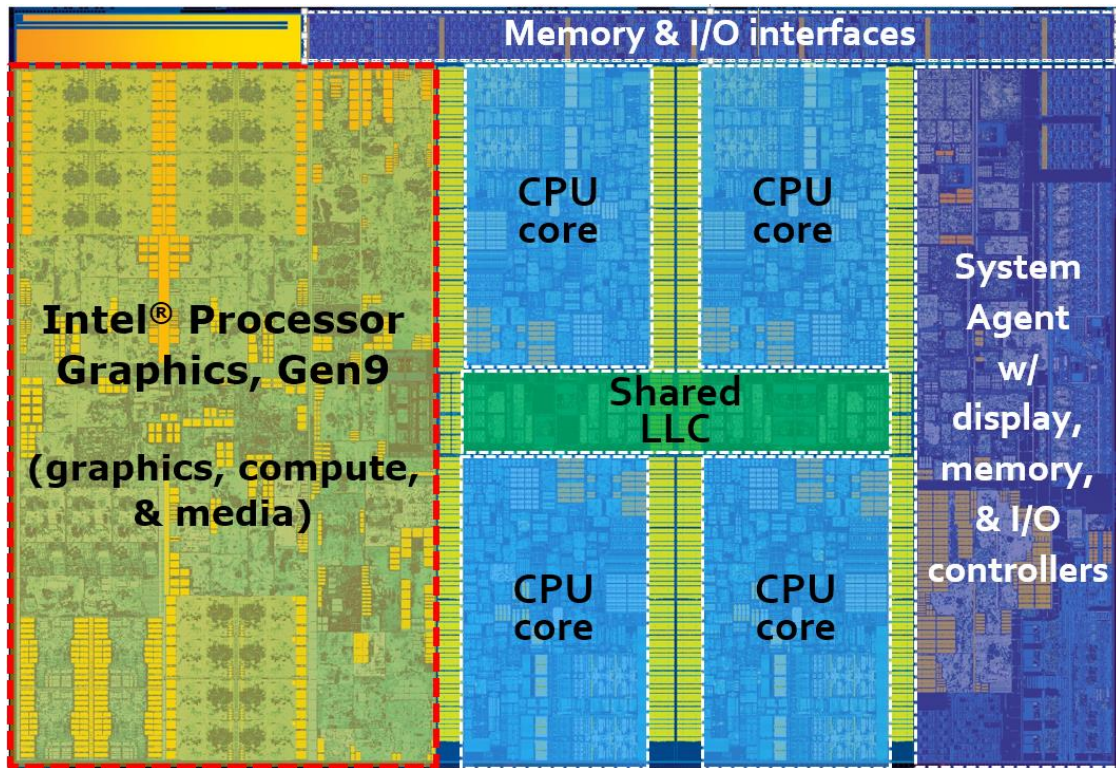


Figure 1: Architecture components layout for an Intel® Core™ i7 processor 6700K for desktop systems. This SoC contains 4 CPU cores, outlined in blue dashed boxes. Outlined in the red dashed box, is an Intel® HD Graphics 530. It is a one-slice instantiation of Intel processor graphics gen9 architecture.

4 SoC ARCHITECTURE

This section describes the SoC architecture within which Intel processor graphics is a component.

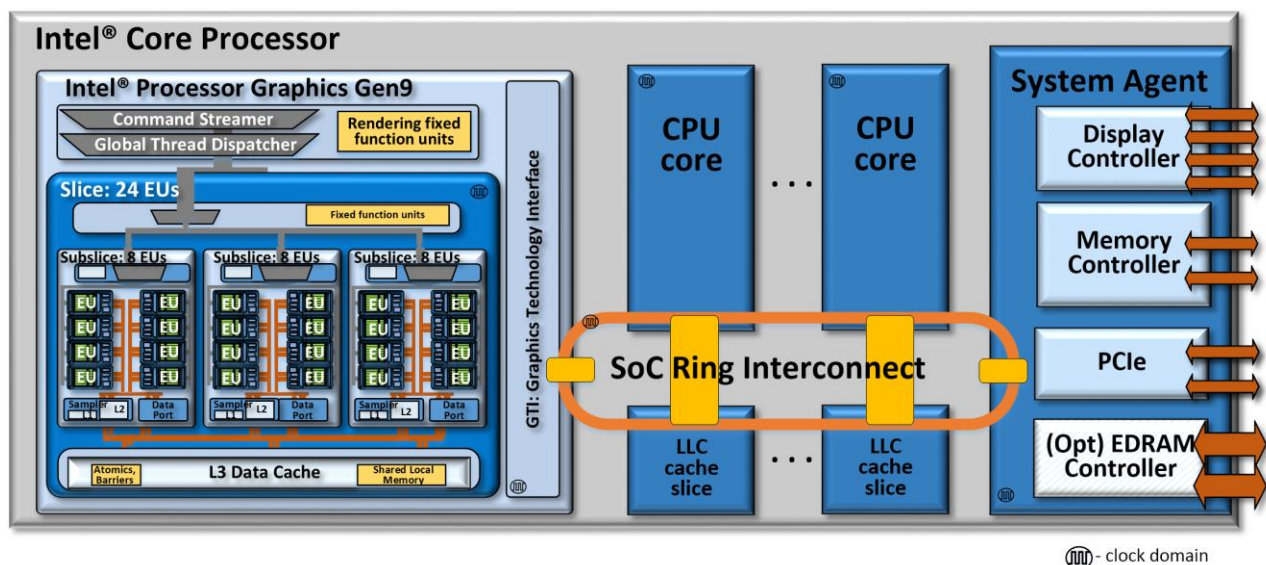


Figure 2: An Intel® Core™ i7 processor 6700K SoC and its ring interconnect architecture.

4.1 SoC ARCHITECTURE

Intel 6th generation Core processors are complex SoCs integrating multiple CPU cores, Intel processor graphics, and potentially other fixed functions all on a single shared silicon die. The architecture implements multiple unique clock domains, which have been partitioned as a per-CPU core clock domain, a processor graphics clock domain, and a ring interconnect clock domain. The SoC architecture is designed to be extensible for a range of products, and yet still enable efficient wire routing between components within the SoC.

4.2 RING INTERCONNECT

The on-die bus between CPU cores, caches, and Intel processor graphics is a ring based topology with dedicated local interfaces for each connected “agent”. This **SoC ring interconnect** is a bi-directional ring that has a 32-byte wide data bus, with separate lines for request, snoop, and acknowledge. Every on-die CPU core is regarded as a unique agent. Similarly, Intel processor graphics is treated as a unique agent on the interconnect ring. A **system agent** is also connected to the ring, which bundles the DRAM memory management unit, display controller, and other off chip I/O controllers such as PCI Express*. Importantly, all off-chip system memory transactions to/from CPU cores and to/from Intel processor graphics are facilitated by this interconnect, through the system agent, and the unified DRAM memory controller.

4.3 SHARED LLC

Some SoC products include a shared **Last Level Cache (LLC)** that is also connected to the ring. In such SoCs, each on-die core is allocated a slice of cache, and that cache slice is connected as a unique agent on the ring. However, all of the slices work together as a single cache, albeit a shared and distributed cache. An address hashing scheme routes data requests to the cache slice assigned to its address. This distributed LLC is also shared with Intel processor graphics. For both CPU cores and for Intel processor graphics, LLC seeks to reduce apparent latency to system DRAM and to provide higher effective bandwidth.

4.4 OPTIONAL EDRAM

Some SoC products may include 64-128 megabytes of embedded DRAM (**EDRAM**), bundled into the SoC’s chip packaging. For example, the Intel processor graphics gen8 based Intel Iris Pro 6200 products bundle a 128 megabyte EDRAM. The EDRAM operates in its own clock domain and can be clocked up to 1.6GHz. The EDRAM has separate buses for read and write, and each are capable of 32 byte/EDRAM-cycle. EDRAM supports many applications including low latency display surface refresh. For the compute architecture of Intel processor graphics gen9, EDRAM further supports the memory hierarchy by serving as a “memory-side” cache between LLC and DRAM. Like LLC, EDRAM caching is shared by both Intel processor graphics and by CPU cores. On an LLC or EDRAM cache miss, data from DRAM will be filled first into EDRAM. (An optional mode also allows bypass to LLC.) Conversely, as cachelines are evicted from LLC, they will be written back into EDRAM. If compute kernels wish to read or write cachelines currently stored in EDRAM, they are quickly re-loaded into LLC, and read/writing then proceeds as usual.

Look for more details about Intel processor graphics gen9-based products with EDRAM in future product announcements.

5 THE COMPUTE ARCHITECTURE OF INTEL PROCESSOR GRAPHICS GEN9

5.1 NEW CHANGES FOR INTEL PROCESSOR GRAPHICS GEN9

Intel processor graphics gen9 includes many refinements throughout the micro architecture and supporting software, over Intel processor graphics gen8. Generally, these changes are across the domains of memory hierarchy, compute capability, and product configuration. They are briefly summarized here, with more detail integrated throughout the paper.

Gen9 Memory Hierarchy Refinements:

- Coherent SVM write performance is significantly improved via new LLC cache management policies.
- The available L3 cache capacity has been increased to 768 Kbytes per slice (512 Kbytes for application data).
- The sizes of both L3 and LLC request queues have been increased. This improves latency hiding to achieve better effective bandwidth against the architecture peak theoretical.
- In Gen9 EDRAM now acts as a memory-side cache between LLC and DRAM. Also, the EDRAM memory controller has moved into the system agent, adjacent to the display controller, to support power efficient and low latency display refresh.
- Texture samplers now natively support an NV12 YUV format for improved surface sharing between compute APIs and media fixed function units.

Gen9 Compute Capability Refinements:

- Preemption of compute applications is now supported at a thread level, meaning that compute threads can be preempted (and later resumed) midway through their execution.
- Round robin scheduling of threads within an execution unit.
- Gen9 adds new native support for the 32-bit float atomics operations of min, max, and compare/exchange. Also the performance of all 32-bit atomics is improved for kernel scenarios that issued multiple atomics back to back.
- 16-bit floating point capability is improved with native support for denormals and gradual underflow.

Gen9 Product Configuration Flexibility:

- Gen9 has been designed to enable products with 1, 2 or 3 slices.
- Gen9 adds new power gating and clock domains for more efficient dynamic power management. This can particularly improve low power media playback modes.

5.2 MODULAR DESIGN FOR PRODUCT SCALABILITY

The gen9 compute architecture is designed for scalability across a wide range of target products. The architecture's modularity enables exact product targeting to a particular market segment or product power envelope. The architecture begins with compute components called execution units. Execution units are clustered into groups called subslices. Subslices are further

clustered into slices. Together, execution units, subslices, and slices are the modular building blocks that are composed to create many product variants based upon Intel processor graphics gen9 compute architecture. Some example variants are shown in Figure 6, Figure 7, and in Figure 8. The following sections describe the architecture components in detail, and show holistically how they may be composed into full products.

5.3 EXECUTION UNIT (EUs) ARCHITECTURE

The foundational building block of gen9 compute architecture is the **execution unit**, commonly abbreviated as **EU**. The architecture of an EU is a combination of simultaneous multi-threading (**SMT**) and fine-grained interleaved multi-threading (**IMT**). These are compute processors that drive multiple issue, single instruction, multiple data arithmetic logic units (**SIMD ALUs**) pipelined across multiple threads, for high-throughput floating-point and integer compute. The fine-grain threaded nature of the EUs ensures continuous streams of ready to execute instructions, while also enabling latency hiding of longer operations such as memory scatter/gather, sampler requests, or other system communication.

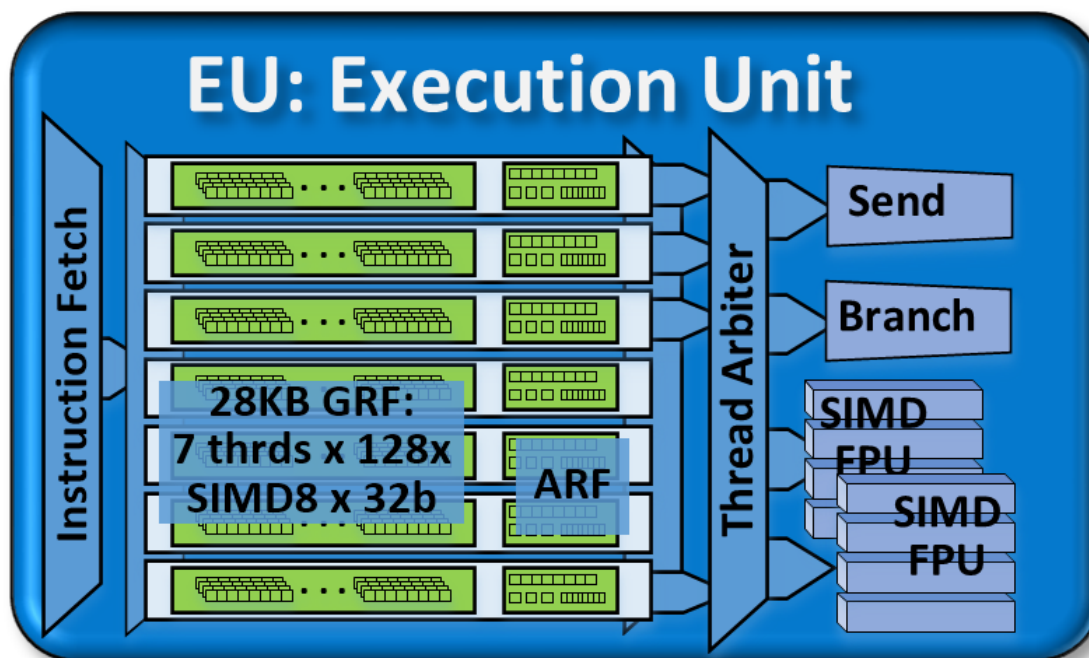


Figure 3: The Execution Unit (EU). Each gen9 EU has seven threads. Each thread has 128 SIMD-8 32-bit registers (GRF) and supporting architecture specific registers (ARF). The EU can co-issue to four instruction processing units including two FPUs, a branch unit, and a message send unit.

Product architects may fine-tune the number of threads and number of registers per EU to match scalability and specific product design requirements. For gen9-based products, each EU thread has 128 general purpose registers. Each register stores 32 bytes, accessible as a SIMD 8-element vector of 32-bit data elements. Thus each gen9 thread has 4 Kbytes of general purpose register file (**GRF**). In the gen9 architecture, each EU has seven threads for a total of 28 Kbytes of GRF per EU. Flexible addressing modes permit registers to be addressed together to build effectively wider registers, or even to represent strided rectangular block data

structures. Per-thread architectural state is maintained in a separate dedicated architecture register file (**ARF**).

5.3.1 Simultaneous Multi-Threading and Multiple Issue Execution

Depending on the software workload, the hardware threads within an EU may all be executing the same compute kernel code, or each EU thread could be executing code from a completely different compute kernel. The execution state of each thread, including its own instruction pointers, are held in thread-specific ARF registers.

On every cycle, an EU can co-issue up to four different instructions, which must be sourced from four different threads. The EU's **thread arbiter** dispatches these instructions to one of four functional units for execution. Although the issue slots for the functional units pose some instruction co-issue constraints, the four instructions are independent, since they are dispatched from four different threads. It is theoretically possible for just two non-stalling threads to fully saturate the floating-point compute throughput of the machine. More typically all seven threads are loaded to deliver more ready-to-run instructions from which the thread arbiter may choose, and thereby promote the EU's instruction-level parallelism.

5.3.2 SIMD FPUs

In each EU, the primary computation units are a pair of SIMD floating-point units (**FPUs**). Although called FPUs, they support both floating-point and integer computation. These units can SIMD execute up to four 32-bit floating-point (or integer) operations, or SIMD-execute up to eight 16-bit integer or 16-bit floating-point operations. The 16-bit float (half-float) support is new for gen9 compute architecture. Each SIMD FPU can complete simultaneous add and multiply (MAD) floating-point instructions every cycle. Thus each EU is capable of 16 32-bit floating-point operations per cycle: (add + mul) x 2 FPUs x SIMD-4. In gen9, both FPUs support native 32-bit integer operations. Finally, one of the FPUs provides extended math capability to support high-throughput transcendental math functions and double precision 64-bit floating-point.

In each EU, gen9 compute architecture offers significant local bandwidth between GRF registers and the FPUs. For example, MAD instructions with three source operands and one destination operand are capable of driving 96 bytes/cycle read bandwidth, and 32 bytes/cycle write bandwidth locally within every EU. Aggregated across the whole architecture, this bandwidth can scale linearly with the number of EUs. For gen9 products with multiple slices of EUs and higher clock rates, the aggregated theoretical peak bandwidth that is local between FPUs and GRF can approach multiple terabytes of read bandwidth.

5.3.3 Branch and Send Units

Within the EUs, branch instructions are dispatched to a dedicated **branch unit** to facilitate SIMD divergence and eventual convergence. Finally, memory operations, sampler operations, and other longer-latency system communications are all dispatched via “send” instructions that are executed by the message passing **send unit**.

5.3.4 EU ISA and Flexible Width SIMD

The EU Instruction Set Architecture (**ISA**) and associated general purpose register file are all designed to support a flexible SIMD width. Thus for 32-bit data types, the gen9 FPUs can be viewed as *physically* 4-wide. But the FPUs may be targeted with SIMD instructions and registers that are *logically* 1-wide, 2-wide, 4-wide, 8-wide, 16-wide, or 32-wide.

For example, a single operand to a SIMD-16 wide instruction pairs two adjacent SIMD-8 wide registers, logically addressing the pair as a single SIMD-16 wide register containing a contiguous 64 bytes. This logically SIMD-16 wide instruction is transparently broken down by the microarchitecture into physically SIMD-4 wide FPU operations, which are iteratively executed. From the viewpoint of a single thread, wider SIMD instructions do take more cycles to complete execution. But because the EUs and EU functional units are fully pipelined across multiple threads, SIMD-8, SIMD-16, and SIMD-32 instructions are all capable of maximizing compute throughput in a fully loaded system.

The instruction SIMD width choice is left to the compiler or low level programmer. Differing SIMD width instructions can be issued back to back with no performance penalty. This flexible design allows compiler heuristics and programmers to choose specific SIMD widths that precisely optimize the register allocation footprint for individual programs, balanced against the amount of work assigned to each thread.

5.3.5 SIMD Code Generation for SPMD Programming Models

Compilers for single program multiple data (**SPMD**) programming models, such as RenderScript, OpenCL™¹, Microsoft DirectX® Compute Shader, OpenGL® Compute, and C++AMP, generate SIMD code to map multiple **kernel instances**² to be executed simultaneously within a given hardware thread. The exact number of kernel instances per-thread is a heuristic driven compiler choice. We refer to this compiler choice as the dominant **SIMD-width** of the kernel. In OpenCL and DirectX Compute Shader, SIMD-8, SIMD-16, and SIMD-32 are the most common SIMD-width targets.

On gen9 compute architecture, most SPMD programming models employ this style of code generation and EU processor execution. Effectively, each SPMD kernel instance appears to execute serially and independently within its own SIMD lane. In actuality, each thread executes a SIMD-width number of kernel instances concurrently. Thus for a SIMD-16 compile of a compute kernel, it is possible for SIMD-16 x 7 threads = 112 kernel instances to be executing concurrently on a single EU. Similarly, for a SIMD-32 compile of a compute kernel, 32 x 7 threads = 224 kernel instances could be executing concurrently on a single EU.

For a given SIMD-width, if all kernel instances within a thread are executing the same instruction, then the SIMD lanes can be maximally utilized. If one or more of the kernel instances chooses a divergent branch, then the thread will execute the two paths of the branch separately in serial. The EU branch unit keeps track of such branch divergence and branch nesting. The branch unit also generates a “live-ness” mask to indicate which kernel instances in the current SIMD-width need to execute (or not execute) the branch.

5.4 SUBSLICE ARCHITECTURE

In gen9 compute architecture, arrays of EUs are instantiated in a group called a **subslice**. For scalability, product architects can choose the number of EUs per subslice. For most gen9-based products, each subslice contains 8 EUs. Each subslice contains its own **local thread**

¹ OpenCL and the OpenCL logo are trademarks of Apple Inc. used by permission by Khronos.

² We use the generic term *kernel instance* as equivalent to OpenCL *work-item*, or DirectX Compute Shader *thread*.

dispatcher unit and its own supporting instruction caches. Given these 8 EUs with 7 threads each, a single subslice has dedicated hardware resources and register files for a total of 56 simultaneous threads. Each subslice also includes a sampler unit and a data port memory management unit. Compared to the Gen7.5 design which had 10 EUs per subslice, the gen8 and gen9 designs reduce the number EUs sharing each subslice's sampler and data port. From the viewpoint of each EU, this has the effect of improving effective bandwidth local to the subslice.

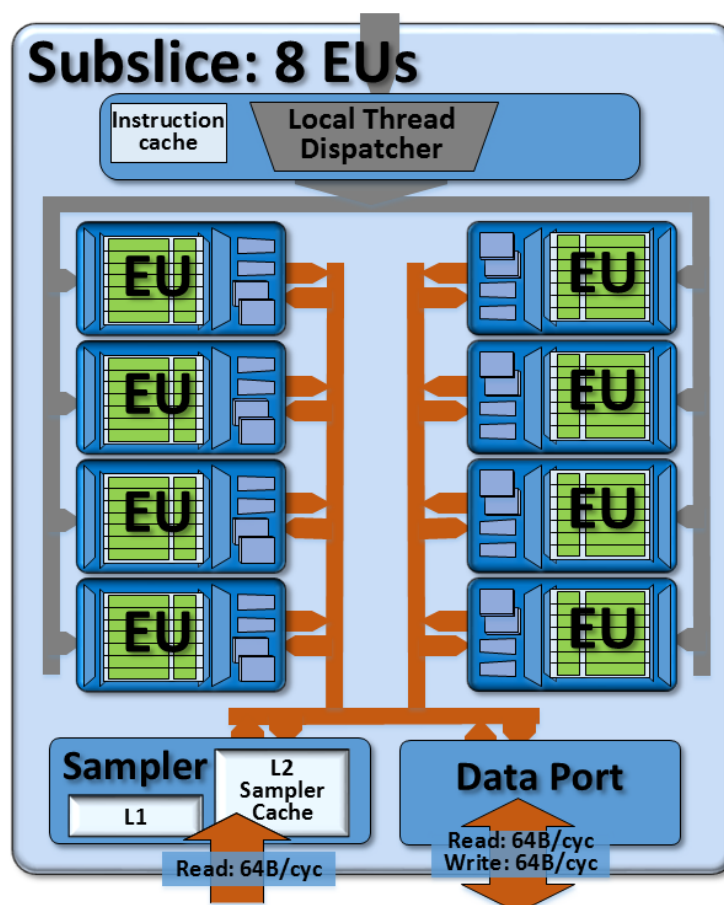


Figure 4: The Intel processor graphics gen9 subslice, containing 8 EUs each. The subslice also instantiates sampler and data port units per subslice.

5.4.1 Sampler

The **sampler** is a read-only memory fetch unit that may be used for sampling of tiled (or not tiled) texture and image surfaces. The sampler includes a level-1 sampler cache (**L1**) and a level-2 sampler cache (**L2**). Between the two caches is dedicated logic to support dynamic decompression of block compression texture formats such as DirectX BC1-BC7, DXT, and OpenGL compressed texture formats. The sampler also includes fixed-function logic that enables address conversion on image (u,v) coordinates, and address clamping modes such as mirror, wrap, border, and clamp. Finally, the sampler supports a variety of sampling filtering modes such as point, bilinear, tri-linear, and anisotropic.

5.4.2 Data Port

Each subslice also contains a memory load/store unit called the **data port**. The data port supports efficient read/write operations for a variety of general purpose buffer accesses, flexible SIMD scatter/gather operations, as well as shared local memory access. To maximize memory bandwidth, the unit dynamically coalesces scattered memory operations into fewer operations over non-duplicated 64-byte cacheline requests. For example, a SIMD-16 gather operation against 16 unique offset addresses for 16 32-bit floating-point values, might be coalesced to a single 64-byte read operation if all the addresses fall within a single cacheline.

5.5 SLICE ARCHITECTURE

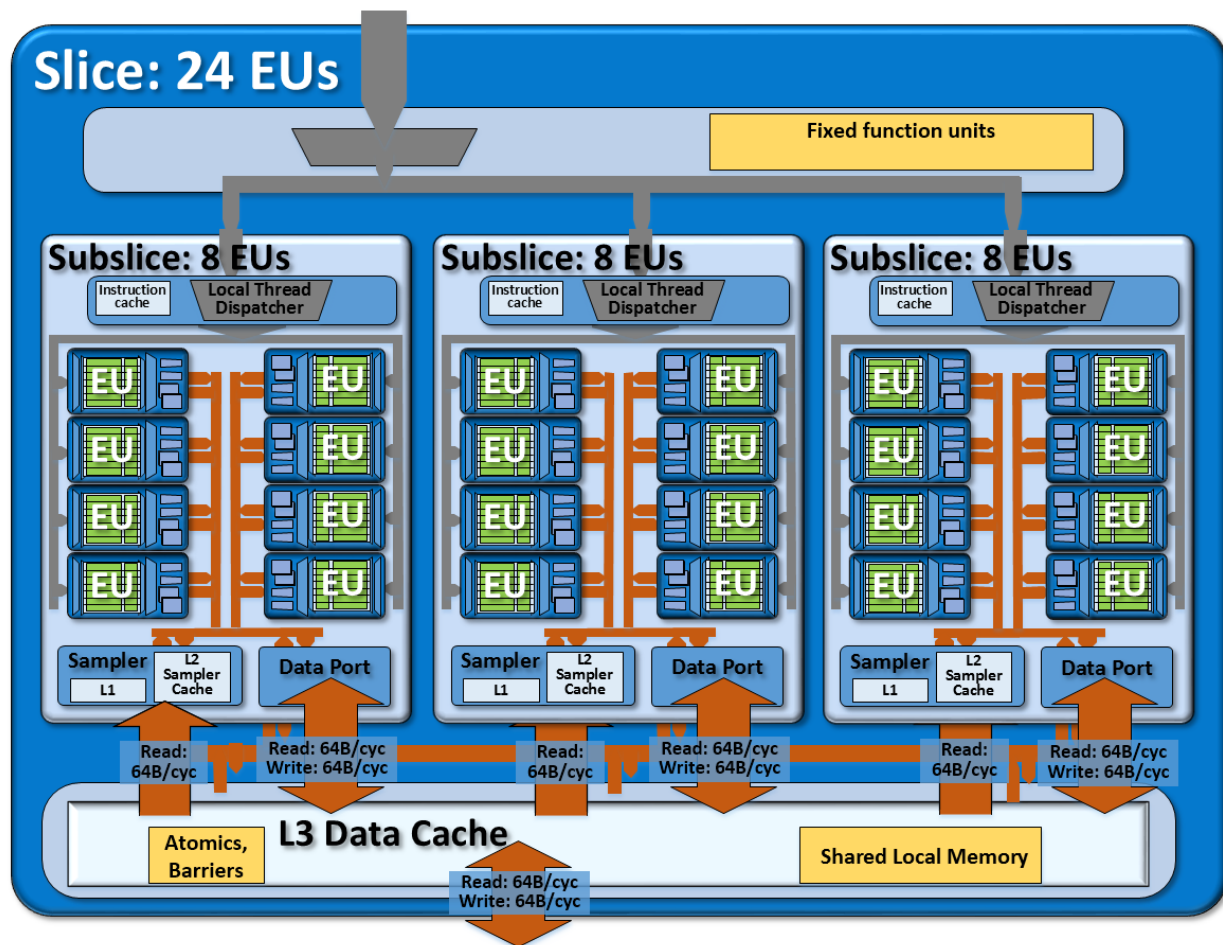


Figure 5: The Intel processor graphics gen9 slice, containing three subslices for a total of 24 EUs. The slice adds supporting L3 cache, shared local memory, atomics, barriers, and other fixed function units.

Subslices are clustered into **slices**. For most gen9-based products, 3 subslices are aggregated into 1 slice. Thus a single slice aggregates a total of 24 EUs³. Aside from grouping subslices, the slice integrates additional logic for thread dispatch routing, a banked level-3 cache, a smaller but highly banked shared local memory structure, and fixed function logic for atomics

³ Note some gen9-based products may enable fewer than 24 EUs in a slice.

and barriers. Additional fixed function units support the media and graphics capability, but are not discussed here.

5.5.1 Level-3 Data Cache

For gen9-based products, the level-3 (**L3**) data cache capacity has been increased to 768Kbytes total per slice. Each application context has flexibility as to how much of the L3 memory structure is allocated: 1) as application L3 data cache, 2) as system buffers for fixed-function pipelines, and 3) as shared local memory. For example, 3D rendering contexts often allocate more L3 as system buffers to support their fixed-function pipelines, instead of as shared local memory. For compute application contexts on gen9 compute architecture, the typical allocation is 512 Kbytes per slice as application data cache.

As with previous generations, gen9 products with multiple subslices will instantiate multiple L3 cache partitions. These cache partitions aggregate together via an L3 fabric and act as a single larger capacity monolithic cache. This L3 fabric is also expandable across multiple slices for even larger cache aggregations. Cachelines are 64 bytes each, and they are uniformly distributed across the entire aggregate cache.

All samplers and data ports are given their own separate memory interface to the L3. The interface between each data port and the L3 data cache enables both read and write of 64 bytes per cycle. Thus a slice containing three subslices, each with a unique data port, will have an aggregate L3 bandwidth of 192 bytes per cycle. For accesses that miss the L3 cache, the L3 fill logic can read and write system memory data at 64 bytes per cycle.

All data in and out of the samplers and data ports flows through the L3 data cache in units of 64-byte-wide cachelines. This includes read and write actions on general purpose buffers. It also includes sampler read transactions that miss the level-1 (L1) and level-2 (L2) sampler caches. L3 cache bandwidth efficiency is highest for read/write accesses that are cacheline-aligned and adjacent within a cacheline. Compute kernel instructions that miss the subslice instruction caches flow through the L3 cache.

5.5.2 Shared Local Memory

Shared local memory⁴ is a structure within the L3 complex that supports programmer-managed data for sharing among EU hardware threads within the same subslice. The read/write bus interface between each subslice and shared local memory is again 64-bytes wide. Latency wise, access to shared local memory is similar to accessing the L3 data cache. However, the shared local memory itself is more highly banked than the L3 data cache. The shared local memory banking can yield full shared local memory bandwidth for access patterns that may not be 64-byte aligned or that may not be contiguously adjacent in memory. For gen9-based products, 64 Kbytes of shared local memory are dedicated and available per subslice. Note that shared local memory is not coherent with other memory structures.

SPMD programming model constructs such as OpenCL's *local* memory space or DirectX Compute Shader's *shared* memory space are shared across a single work-group (thread-group). For software kernel instances that use shared local memory, driver runtimes typically map all instances within a given OpenCL work-group (or a DirectX 11 threadgroup) to EU

⁴ We use the term *shared local memory* to indicate the hardware memory structure that supports the software address space (OpenCL refers to it as *work-group local memory*, and which DirectX Compute Shader refers to as *thread-group shared memory*).

threads within a single subslice. Thus all kernel instances within a work-group will share access to the same 64 Kbyte shared local memory partition. Because of this property, an application's accesses to shared local memory should scale with the number of subslices.

5.5.3 Barriers and Atomics

Each slice within gen9 compute architecture bundles dedicated logic to support implementation of barriers across groups of threads. This barrier logic is available as a hardware alternative to pure compiler-based barrier implementation approaches. The gen9 logic can support barriers simultaneously in up to 16 active thread-groups per subslice.

Each slice also provides a rich suite of atomic read-modify-write memory operations. These operations support both operations to L3 cached global memory or to shared local memory. Gen9-based products support 32-bit atomic operations.

5.5.4 64-Byte Data Width

A foundational element of gen9 compute architecture is the **64-byte data width**. Recall that each EU thread's register file is composed of 128 32-byte registers (SIMD-8 x 32-bit). But recall also that operands to SIMD-16 instructions typically pair two such registers, treating the pair as a single 64-byte SIMD-16 register. Observe:

- A SIMD-16 instruction can source 64-byte wide operands from 64-byte wide registers.
- The data for such 64-byte wide registers are read and written from L3 over a 64-byte wide data bus.
- Within the L3 data cache, each cacheline is again 64-bytes wide.
- Finally the L3 cache's bus interface to the SoC-shared LLC is also 64-bytes wide.

5.6 PRODUCT ARCHITECTURE

Finally, SoC product architects can create product families or a specific product within a family by instantiating a single slice or groups of slices. Members of a product family might differ primarily in the number of slices. These slices are combined with additional front end logic to manage command submission, as well as fixed-function logic to support 3D, rendering, and media pipelines. Additionally the entire gen9 compute architecture interfaces to the rest of the SoC components via a dedicated unit called the graphics technology interface (GTI).

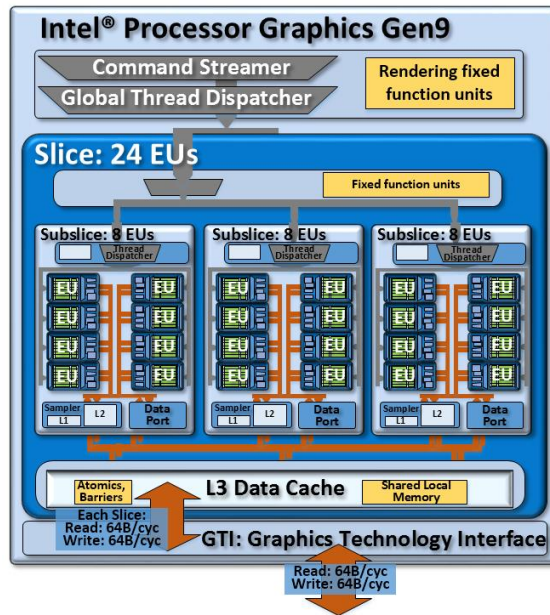


Figure 6: A potential product design that instantiates the compute architecture of Intel® processor graphics gen9. It is composed of a single slice with three subslices, for a total of 24 EUs. The Intel® Core™ i7 processor 6700K with Intel® HD Graphics 530 instantiates such a design.

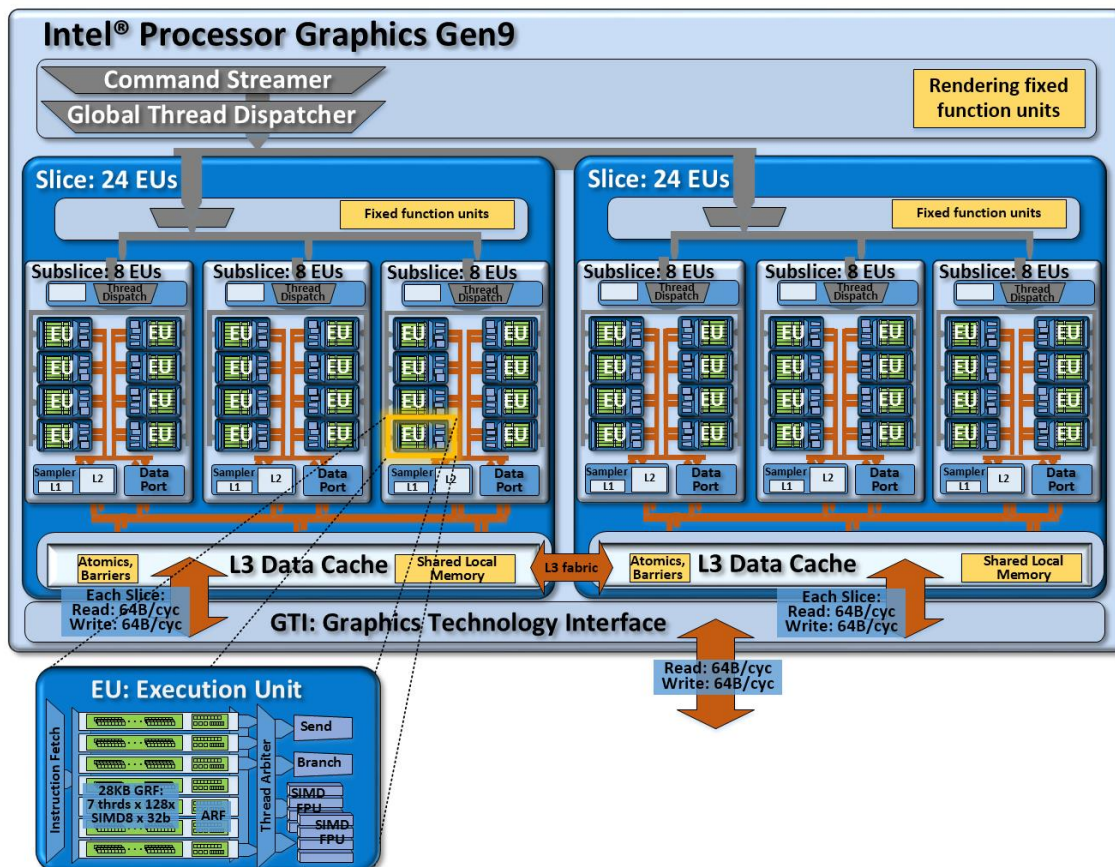


Figure 7: Another potential product design that instantiates the compute architecture of Intel® processor graphics gen9. This design is composed of two slices, of three subslices each for a total of 48 EUs.

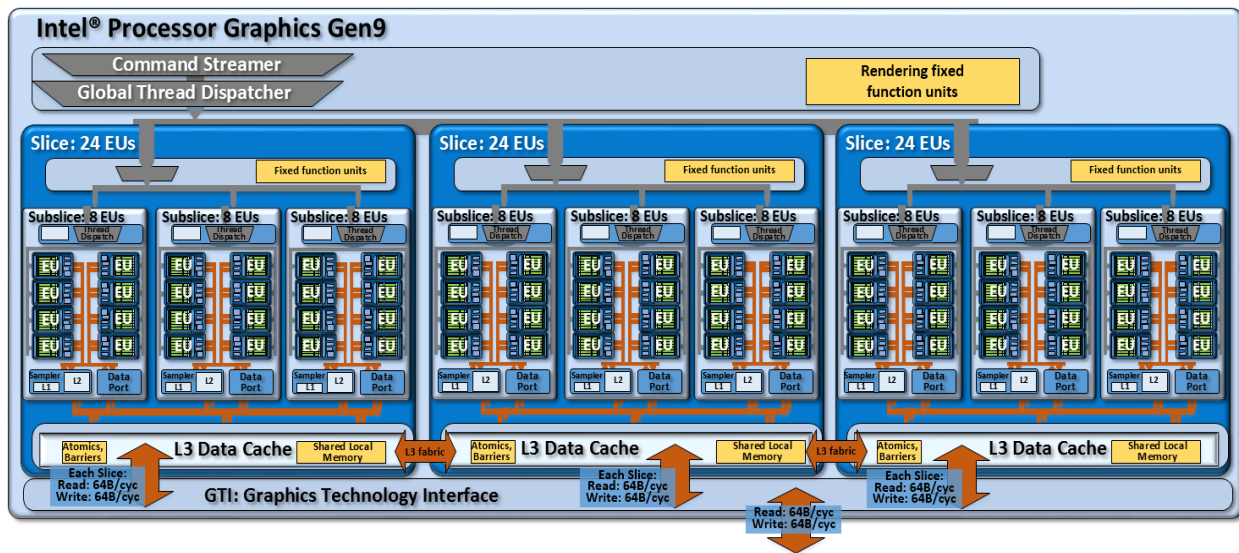


Figure 8: Another potential product design that instantiates the compute architecture of Intel® processor graphics gen9. This design is composed of three slices, of three subslices each for a total of 72 EUs.

5.6.1 Command Streamer and Global Thread Dispatcher

As its name implies, the **command streamer** efficiently parses command streams submitted from driver stacks and routes individual commands to their representative units. For compute workloads, the **global thread dispatcher** is responsible for load balancing thread distribution across the entire device. The global thread dispatcher works in concert with local thread dispatchers in each subslice.

The global thread dispatcher operates in two modes. For compute workloads that *do not* depend on hardware barriers or on shared local memory, the global thread dispatcher may distribute the workload over all available subslices to maximize throughput and utilization. Given the unit's global visibility, it is able to load balance across all the execution resources. For compute workloads that *do depend* upon hardware barriers or shared local memory, the global thread dispatcher will assign thread-group sized portions of the workload to specific subslices. Such an assignment ensures localized access to the barrier logic and shared local memory storage dedicated to each subslice.

5.6.2 Graphics Technology Interface (GTI)

The graphics technology interface, or **GTI**, is the gateway between gen9 compute architecture with the rest of the SoC. The rest of the SoC includes memory hierarchy elements such as the shared LLC memory, the system DRAM, and possibly embedded DRAM. GTI facilitates communication with the CPU cores and possibly with other fixed function devices such as camera imaging pipelines. GTI also implements global memory atomics that may be shared between Intel processor graphics gen9 and CPU cores. Finally GTI implements power management controls for Intel processor graphics gen9 and interfaces between the GTI clock domain and the (usually different) SoC clock domains.

The bus between each slice that is interfaced to GTI is capable of 64-bytes per cycle read and 64-bytes per cycle write. Internally, GTI has memory request queues to maintain memory order

and manage differing latencies on individual memory requests. For instances of Intel processor graphics gen9, the bus between GTI and LLC has two possible configurations. Like gen8, higher-performance gen9 configurations are capable of 64-bytes per cycle read and 64-bytes per cycle write. A second configuration for lower power is capable of 64-bytes per cycle read and 32-bytes per cycle write. Like other aspects of the architecture, this bus width is also scalable, and SoC designers may configure it for specific products.

5.6.3 Unslice

The command streamer, global thread dispatcher, and graphics technology interface all exist independent of the slice instantiations, in a domain typically called the “unslice.” New to gen9, this domain is given its own power gating and clocking that can run at the same or faster than the slice clock. This can enable intelligent power savings by dynamically diverting more power to GTI’s memory bandwidth, versus the EU slices’s compute capability. This can be particularly effective for low power media playback modes.

5.6.4 Product EU Counts

Although gen9 subslices generally contain 8 EUs each, complete gen9-based products can disable an EU within a subslice to optimize product yields from silicon manufacturing. For example, a three subslice-based product can have a total of 23 EUs by disabling an EU in one subslice.

5.7 MEMORY

5.7.1 Unified Memory Architecture

Intel processor graphics architecture has long pioneered sharing DRAM physical memory with the CPU. This unified memory architecture offers a number of system design, power efficiency, and programmability advantages over PCI Express-hosted discrete memory systems.

The obvious advantage is that **shared physical memory** enables **zero copy** buffer transfers between CPUs and gen9 compute architecture. By zero copy, we mean that no buffer copy is necessary since the physical memory is shared. Moreover, the architecture further augments the performance of such memory sharing with a shared LLC cache. The net effect of this architecture benefits performance, conserves memory footprint, and indirectly conserves system power not spent needlessly copying data. Shared physical memory and zero copy buffer transfers are programmable through the buffer allocation mechanisms in APIs such as OpenCL 1.0+ and DirectX11.2+.

5.7.2 Shared Memory Coherency

Gen9 compute architecture supports global memory coherency between Intel processor graphics and the CPU cores. SoC products with Intel processor graphics gen9 integrate new hardware components to support the recently updated Intel® Virtualization Technology (Intel® VT) for Directed I/O (Intel® VT-d) specification. This specification extends Intel VT, which generally addresses virtual machine to physical machine usage models and enables virtual machine monitor implementation. In particular, the recent Intel VT-d specification extensions define new page table entry formats, cache protocols, and hardware snooping mechanisms for shared memory between CPU cores and devices such as Intel processor graphics.

These new mechanisms can be used to maintain memory coherency and consistency for fine grained sharing throughout the memory hierarchy between CPU cores and devices. Moreover, the same virtual addresses can be shared seamlessly across devices. Such memory sharing is application-programmable through emerging heterogeneous compute APIs such as the shared virtual memory (SVM) features specified in OpenCL 2.0. The net effect is that pointer-rich data-structures can be shared directly between application code running on CPU cores with application code running on Intel processor graphics, without programmer data structure marshalling or cumbersome software translation techniques.

Within Intel processor graphics, the data port unit, L3 data cache, and GTI have all been upgraded to support a new globally coherent memory type. Reads and writes originating from Intel processor graphics to memory typed as globally coherent and drive Intel VT-d-specified snooping protocols to ensure data integrity with any CPU core cached versions of that memory. Conversely, the same is true for GPU cached memory and reads and writes that originate from the CPU cores. The coherent memory hierarchy is shown in Figure 9. Note that the sampler's L1 and L2 caches as well as the shared local memory structures are not coherent.

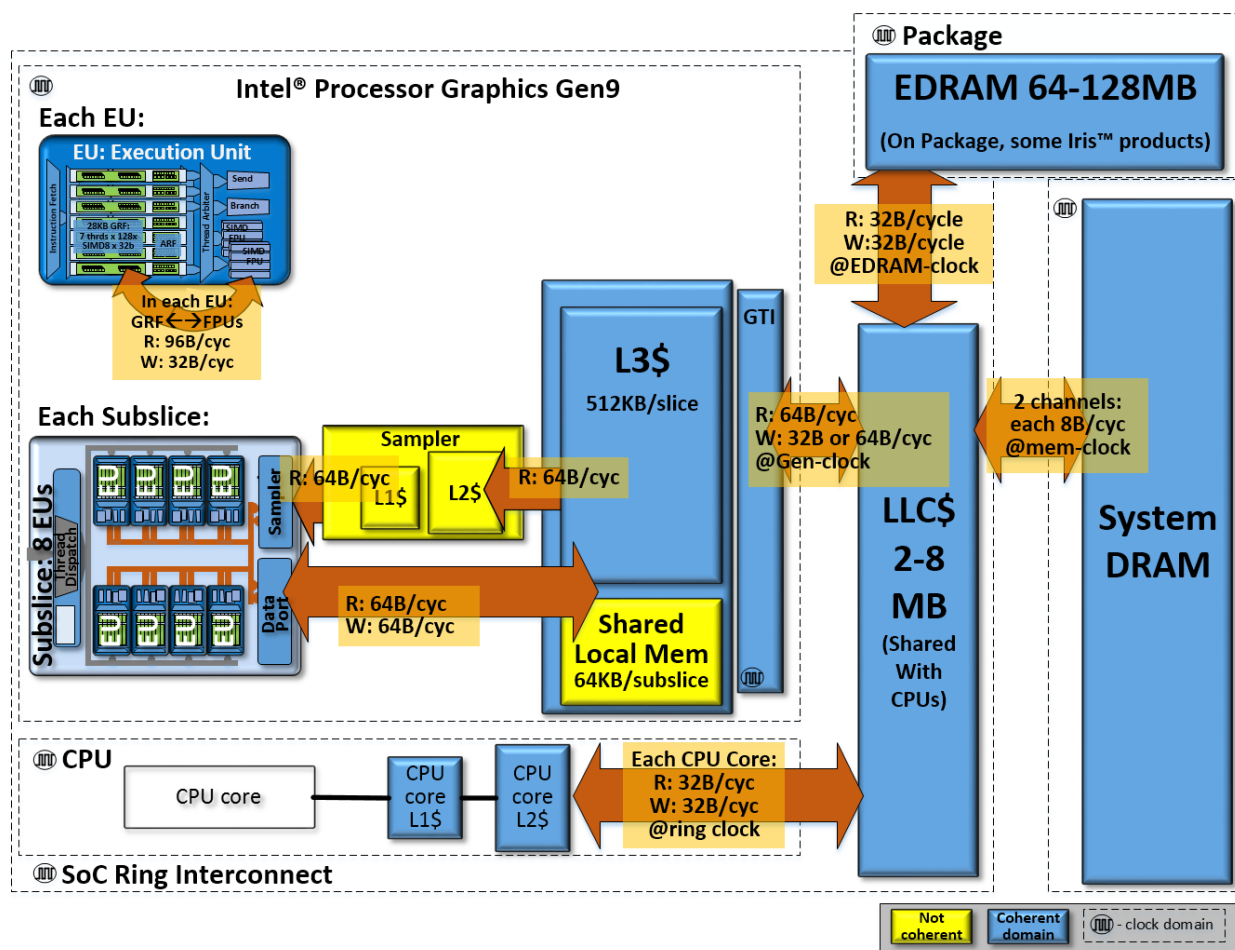


Figure 9: A view of the SoC chip level memory hierarchy and its theoretical peak bandwidths for the compute architecture of Intel processor graphics gen9.

5.8 ARCHITECTURE CONFIGURATIONS, SPEEDS, AND FEEDS

The following table presents the *theoretical peak throughput* of the compute architecture of Intel processor graphics, *aggregated* across the entire graphics product architecture. Values are stated as “per cycle”, as final product clock rates were not available at time of this writing.

	Intel® HD Graphics 530	Derivation, notes
Configurations:		
Execution units (EUs)	24 EUs	8 EUs x 3 subslices x 1 slices
Hardware threads	168 threads	24 EUs x 7 threads
Concurrent kernel instances (e.g. OpenCL™ work-items or DirectX® Compute Shader “threads”)	5376 instances	168 threads * SIMD-32 compile
Level-3 data cache (L3\$) size	512 Kbytes	1 slice x 512 Kbytes /slice
Max shared local memory size	192 Kbytes	3 subslices x 64 Kbytes /subslice
Last level cache (LLC\$) size	2-8 Mbytes	depending on product configuration
Package embedded DRAM size	n/a	
Peak Compute Throughput		
32b float FLOPS	384 FLOP/cycle	24 EUs x (2 x SIMD-4 FPU) x (MUL + ADD)
64b double float FLOPS	96 FLOP/cycle	24 EUs x SIMD-4 FPU x (MUL + ADD) x ½ throughput
32b integer IOPS	192 IOP/cycle	24 EUs x (2 x SIMD-4 FPU) x (ADD)

6 EXAMPLE COMPUTE APPLICATIONS

The following images provide a few visual examples of the kinds of compute applications and algorithms that have been accelerated on Intel processor graphics.

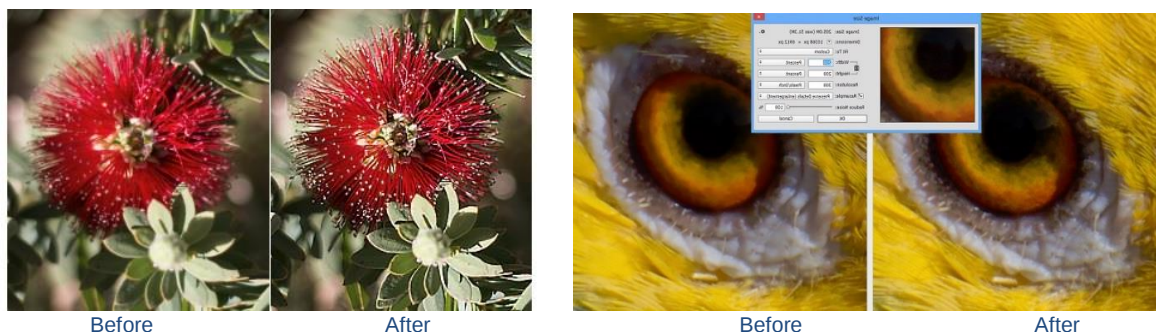


Figure 10: Intel® processor graphics acceleration in Adobe Photoshop®, via OpenCL™. In the left pair of images, Adobe Photoshop’s “Smart Sharpen” feature uses Intel processor graphics to efficiently analyze images to maximize

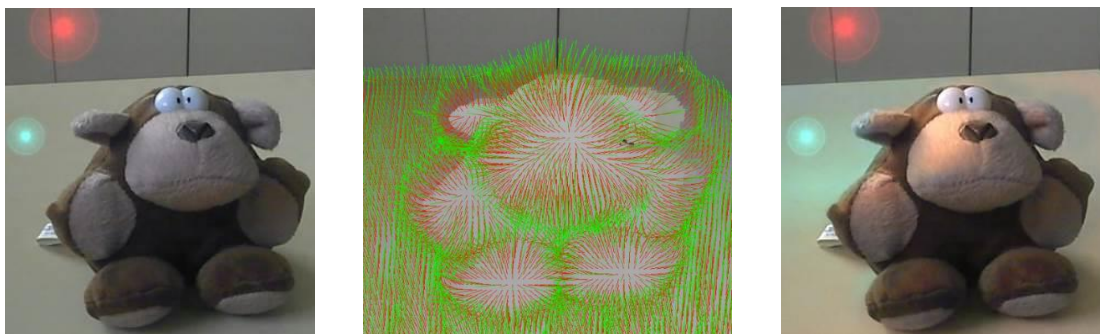
clarity and minimize visual noise and halos. In the right pair of images, Adobe Photoshop's "Intelligent upsampling" feature use Intel processor graphics to accelerate upsampling operations that preserve detail and sharpness without introducing visual noise. Images courtesy of Anita Banerjee and Ilya Albrekht.



Figure 11: A before(left) and after(right) image generated using CyberLink's PhotoDirector* Clarify effect, whose OpenCL™ implementation is optimized for Intel® processor graphics. CPU and GPU work concurrently on the same image using SVM to efficiently share the image data. Images courtesy of Ilya Albrekht.



Figure 12: Interactive Real-time volumetric rendered 3D smoke effect implemented using DirectX® 11 Compute Shader on Intel® processor graphics. Image courtesy Doug McNabb.



Raw video frame & virtual lights

Computed surface normals

Video frame with relighting

Figure 13: Interactive dynamic relighting of a real-time video feed. A localized surface polynomial approximation solver are implemented in OpenCL™ and applied to real-time depth captures. Derived surface normals and lighting computations can then synthetically re-light the scene based on virtual light sources. Images courtesy of Konstantin Rodyushkin. [See an interactive example video.](#)

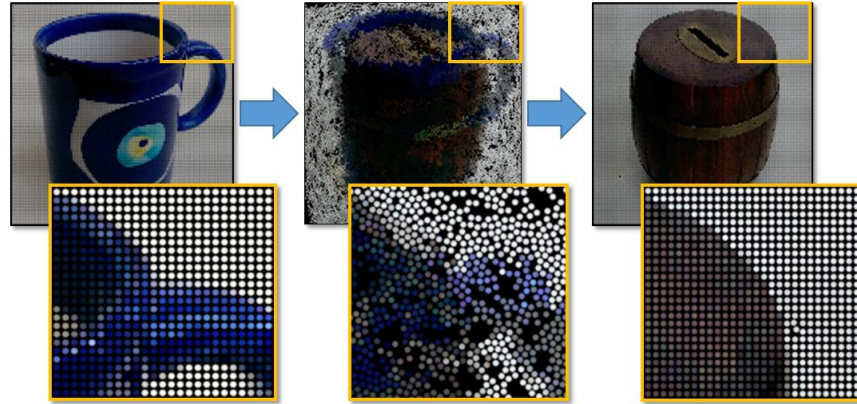


Figure 14: Crowd simulation-based transition effect between two photos (or videos). Particles carry colors of the source image and then change the colors while moving to form the destination image. Dynamic particle collision detection and response are calculated with UNC's RVO2 library ported to OpenCL™ 2.0 and running on Intel® processor graphics. Intel processor graphics and OpenCL 2.0's Shared Virtual Memory enable passing the original pointer-rich data structures of the RVO2 library directly to Intel processor graphics "as is". Neither data structure redesign nor fragile software data marshaling is necessary. Images courtesy of Sergey Lyalin and UNC. More info: <http://gamma.cs.unc.edu/RVO2/>.

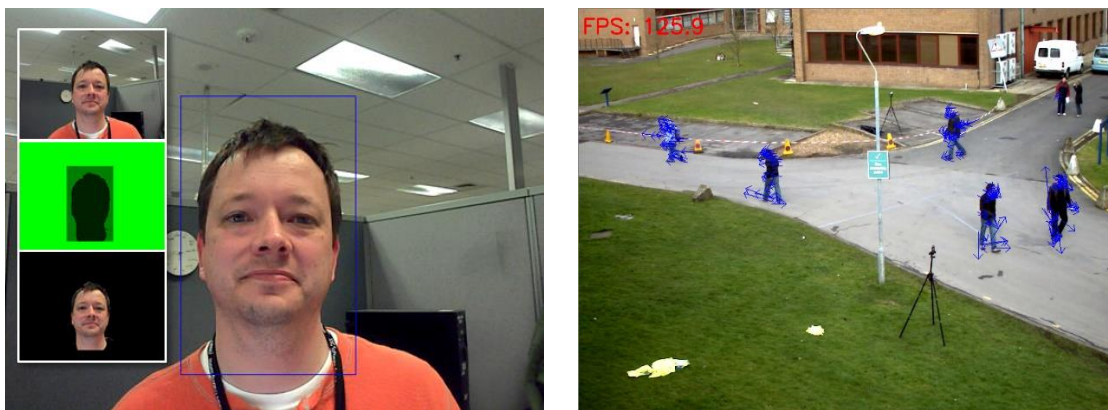


Figure 15: Two examples from OpenCV*. The left image is an OpenCV Face detection algorithm, accelerated via OpenCL™ and optimized for Intel® processor graphics. The right image shows optical flow vectors in blue, dynamically calculated using OpenCV's Lucas-Kanade implementation also accelerated via OpenCL and optimized for Intel processor graphics. (Images courtesy of Aaron Kunze and Maxim Shetsov.) OpenCV 3.0, is now a feature of Intel® INDE with Intel-optimized, Windows* and Android* pre-built binaries for academic and commercial use. Details at <http://software.intel.com/opencv>

7 ACKNOWLEDGEMENTS

Intel processor graphics architecture, products, supporting software, and optimized applications are the results of many years and the efforts of many engineers and architects, too many to list here. Also many reviewers contributed to this document. Thank you all. A particular thank-you to Jim Valerio, Murali Sundaresan, David Blythe, and Tom Piazza for their technical leadership and support in writing this document.

8 MORE INFORMATION

- [The Compute Architecture of Intel Processor Graphics Gen7.5](#)
- [Intel® Iris™ Graphics Powers Built-in Beautiful](#)
- [About Intel® Processor Graphics Technology](#)
- [Open source Linux documentation of Gen Graphics and Compute Architecture](#)
- [Intel® SDK for OpenCL](#)
- [Optimizing Heterogeneous Computing for Intel® Processor Graphics, IDF 2014 Shenzhen](#)
- [Intel® 64 and IA-32 Architectures Software Developers Manual](#)
- [Intel® Virtualization Technology \(Intel® VT\) for Directed I/O \(Intel® VT-d\): Enhancing Intel platforms for efficient virtualization of I/O devices](#)
- [Intel® Virtualization Technology for Directed I/O - Architecture Specification](#)

9 NOTICES

Copyright © 2015 Intel Corporation. All rights reserved

By using this document, in addition to any agreements you have with Intel, you accept the terms set forth below.

No license (express or implied, by estoppel or otherwise) to any intellectual property rights is granted by this document.

Intel disclaims all express and implied warranties, including without limitation, the implied warranties of merchantability, fitness for a particular purpose, and non-infringement, as well as any warranty arising from course of performance, course of dealing, or usage in trade.

This document contains information on products, services and/or processes in development. All information provided here is subject to change without notice. Contact your Intel representative to obtain the latest forecast, schedule, specifications and roadmaps.

The products and services described may contain defects or errors known as errata which may cause deviations from published specifications. Current characterized errata are available on request.

Copies of documents which have an order number and are referenced in this document may be obtained by calling 1-800-548-4725 or by visiting www.intel.com/design/literature.htm.

Intel, the Intel logo, {List the Intel trademarks in your document} are trademarks of Intel Corporation in the U.S. and/or other countries.

Software and workloads used in performance tests may have been optimized for performance only on Intel microprocessors. Performance tests, such as SYSmark* and MobileMark*, are measured using specific computer systems, components, software, operations and functions. Any change to any of those factors may cause the results to vary. You should consult other information and performance tests to assist you in fully evaluating your contemplated purchases, including the performance of that product when combined with other products. Configurations: [describe config + what test used + who did testing]. For more information go to <http://www.intel.com/performance>.

Intel, the Intel logo, Iris™, Core™ are trademarks of Intel Corporation in the U.S. and/or other countries.

*Other names and brands may be claimed as the property of others.

Intel® Graphics 4600, Iris™ Graphics, and Iris™ Pro Graphics are available on select systems. Consult your system manufacturer. visit <http://www.intel.com/content/www/us/en/architecture-and-technology/microarchitecture/latest-microarchitecture.html>