

(43) International Publication Date
15 December 2016 (15.12.2016)

(51) International Patent Classification:

G06F 9/38 (2006.01) G06T 1/20 (2006.01)
G06F 9/54 (2006.01)

(21) International Application Number:

PCT/US2016/031872

(22) International Filing Date:

11 May 2016 (11.05.2016)

(25) Filing Language:

English

(26) Publication Language:

English

(30) Priority Data:

14/738,703 12 June 2015 (12.06.2015) US

(71) Applicant: INTEL CORPORATION [US/US]; 2200 Mission College Boulevard, Santa Clara, California 95054 (US).

(72) Inventors: APODACA, Michael; 1563 Ballou Circle, Folsom, California 95630 (US). CIMINI, David M.; 5611 Bench Court, Orangevale, California 95662 (US).

(74) Agents: MALLIE, Michael J. et al.; Blakely Sokoloff Taylor & Zafman LLP, 1279 Oakmead Parkway, Sunnyvale, California 94085 (US).

(81) Designated States (unless otherwise indicated, for every kind of national protection available): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JP, KE, KG, KN, KP, KR, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL, SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, ZA, ZM, ZW.

(84) Designated States (unless otherwise indicated, for every kind of regional protection available): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), Eurasian (AM, AZ, BY, KG, KZ, RU, TJ, TM), European (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

Published:

— with international search report (Art. 21(3))

(54) Title: FACILITATING EFFICIENT GRAPHICS COMMAND GENERATION AND EXECUTION FOR IMPROVED GRAPHICS PERFORMANCE AT COMPUTING DEVICES

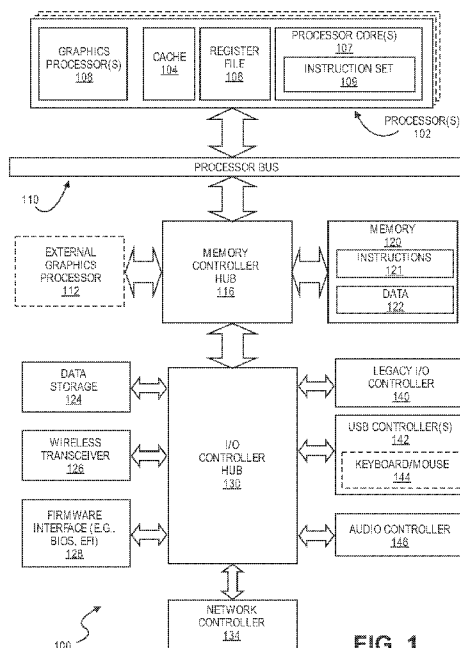


FIG. 1

(57) Abstract: A mechanism is described for facilitating efficient graphics command generation and execution for improved graphics performance at computing devices. A method of embodiments, as described herein, includes detecting an application programming interface (API) call to perform a plurality of transactions, where the API call is issued by an application at a first command buffer, where the plurality of transactions include a first set of transactions and a second set of transactions. The method may further include creating a second command buffer and appending the second command buffer to the first command buffer, where creating further includes separating the first set of transactions from the second set of transactions. The method may further include executing, via the second command buffer, the first set of transactions, prior to executing the first set of transactions.

FACILITATING EFFICIENT GRAPHICS COMMAND GENERATION AND EXECUTION FOR IMPROVED GRAPHICS PERFORMANCE AT COMPUTING DEVICES

5 FIELD

Embodiments described herein generally relate to computers. More particularly, embodiments for facilitating efficient graphics command generation and execution for improved graphics performance at computing devices.

BACKGROUND

10 Several attempts have been made to lower graphics driver overhead to promote efficient rendering of graphics contents. For example, a modern three-dimensional (3D) graphics application programming interface (API) may be capable of allowing an application to specify a method for a graphics processing interface (GPU) to generate work for itself; nevertheless, the technique is regarded as inefficient as it causes the GPU to suffer severe performance penalties,
15 such as when switching between 3D rendering and general-purpose GPU (GPGPU) compute.

BRIEF DESCRIPTION OF THE DRAWINGS

Embodiments are illustrated by way of example, and not by way of limitation, in the figures of the accompanying drawings in which like reference numerals refer to similar elements.

Figure 1 is a block diagram of a processing system, according to an embodiment.

20 **Figure 2** is a block diagram of an embodiment of a processor having one or more processor cores, an integrated memory controller, and an integrated graphics processor.

Figure 3 is a block diagram of a graphics processor, which may be a discrete graphics processing unit, or may be a graphics processor integrated with a plurality of processing cores.

Figure 4 is a block diagram of a graphics processing engine of a graphics processor in
25 accordance with some embodiments.

Figure 5 is a block diagram of another embodiment of a graphics processor.

Figure 6 illustrates thread execution logic including an array of processing elements employed in some embodiments of a graphics processing engine.

Figure 7 is a block diagram illustrating a graphics processor instruction formats according
30 to some embodiments.

Figure 8 is a block diagram of another embodiment of a graphics processor.

Figure 9A is a block diagram illustrating a graphics processor command format according to an embodiment and **Figure 9B** is a block diagram illustrating a graphics processor command sequence according to an embodiment.

35 **Figure 10** illustrates exemplary graphics software architecture for a data processing system

according to some embodiments.

Figure 11 is a block diagram illustrating an IP core development system that may be used to manufacture an integrated circuit to perform operations according to an embodiment.

Figure 12 is a block diagram illustrating an exemplary system on a chip integrated circuit that may be fabricated using one or more IP cores, according to an embodiment.

Figure 13 illustrates a conventional technique.

Figure 14 illustrates a conventional technique.

Figure 15 illustrates a computing device employing an efficient graphics commands generation/execution mechanism according to one embodiment.

Figure 16 illustrates an efficient graphics commands generation/execution mechanism according to one embodiment.

Figure 17A illustrates a transaction sequence for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment.

Figure 17B illustrates a transaction sequence for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment.

Figure 17C illustrates a transaction sequence for processing barrier APIs for efficient utilization of GPU resources according to one embodiment.

Figure 18A illustrates a method for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment.

Figure 18B illustrates a method for processing barrier APIs for efficient utilization of GPU resources according to one embodiment.

Figure 18C illustrates a method for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment.

DETAILED DESCRIPTION

In the following description, numerous specific details are set forth. However, embodiments, as described herein, may be practiced without these specific details. In other instances, well-known circuits, structures and techniques have not been shown in details in order not to obscure the understanding of this description.

Embodiments provide for a novel technique for efficient graphics command generation and execution for improved graphics performance at computing devices. In one embodiment, multiple back-to-back API calls (such as DirectX® 12 ExecuteIndirect API calls) may be combined to minimize the switching between GPGPU compute shader execution and 3D rendering. For example and in one embodiment, a graphics driver at a GPU may be facilitated to first execute a first phase of calls, such as all the “patch” call, followed by a second phase of calls, such as all the “submit” calls, as will be further described in this document.

GPGPU APIs may allow for arbitrary programming of a GPU, such as in games and high performance computing. For example, graphics driver may be used to perform various tasks to reduce the effort spent in optimizing applications in every GPU and thus improving the user (e.g., developer, end-user) experience, such as in case of 3D games and high-performance computing where the graphics performance of various algorithms maybe automatically improved through the graphics driver. For example, an application may include a game, a workstation application, etc., offered thorough an API, such as a free rendering API, such as Open Graphics Library (OpenGL®), DirectX® 11, DirectX® 12, etc.

It is contemplated and to be noted that “dispatch” may be interchangeably referred to as “work unit” or “draw” and similarly, “application” may be interchangeably referred to as “workflow”. For example, a workload, such as a 3D game, may include and issue any number and type of “frames” where each frame may represent an image (e.g., sailboat, human face). Further, each frame may include and offer any number and type of work units, where each work unit may represent a part (e.g., mast of sailboat, forehead of human face) of the image (e.g., sailboat, human face) represented by its corresponding frame. However, for the sake of consistency along with brevity, clarity, and ease of understanding, any single term, such as “dispatch”, “application”, etc., may be used throughout this document.

In some embodiments, terms like “display screen” and “display surface” may be used interchangeably referring to the visible portion of a display device while the rest of the display device may be embedded into a computing device, such as a smartphone, a wearable device, etc. It is contemplated and to be noted that embodiments are not limited to any particular computing device, software application, hardware component, display device, display screen or surface, protocol, standard, etc. For example, embodiments may be applied to and used with any number and type of real-time applications on any number and type of computers, such as desktops, laptops, tablet computers, smartphones, head-mounted displays and other wearable devices, and/or the like. Further, for example, rendering scenarios for efficient performance using this novel technique may range from simple scenarios, such as desktop compositing, to complex scenarios, such as 3D games, augmented reality applications, etc.

System Overview

Figure 1 is a block diagram of a processing system 100, according to an embodiment. In various embodiments the system 100 includes one or more processors 102 and one or more graphics processors 108, and may be a single processor desktop system, a multiprocessor workstation system, or a server system having a large number of processors 102 or processor cores 107. In on embodiment, the system 100 is a processing platform incorporated within a system-on-a-chip (SoC) integrated circuit for use in mobile, handheld, or embedded devices.

An embodiment of system 100 can include, or be incorporated within a server-based gaming platform, a game console, including a game and media console, a mobile gaming console, a handheld game console, or an online game console. In some embodiments system 100 is a mobile phone, smart phone, tablet computing device or mobile Internet device. Data processing system 100 can also include, couple with, or be integrated within a wearable device, such as a smart watch wearable device, smart eyewear device, augmented reality device, or virtual reality device. In some embodiments, data processing system 100 is a television or set top box device having one or more processors 102 and a graphical interface generated by one or more graphics processors 108.

In some embodiments, the one or more processors 102 each include one or more processor cores 107 to process instructions which, when executed, perform operations for system and user software. In some embodiments, each of the one or more processor cores 107 is configured to process a specific instruction set 109. In some embodiments, instruction set 109 may facilitate Complex Instruction Set Computing (CISC), Reduced Instruction Set Computing (RISC), or computing via a Very Long Instruction Word (VLIW). Multiple processor cores 107 may each process a different instruction set 109, which may include instructions to facilitate the emulation of other instruction sets. Processor core 107 may also include other processing devices, such a Digital Signal Processor (DSP).

In some embodiments, the processor 102 includes cache memory 104. Depending on the architecture, the processor 102 can have a single internal cache or multiple levels of internal cache. In some embodiments, the cache memory is shared among various components of the processor 102. In some embodiments, the processor 102 also uses an external cache (e.g., a Level-3 (L3) cache or Last Level Cache (LLC)) (not shown), which may be shared among processor cores 107 using known cache coherency techniques. A register file 106 is additionally included in processor 102 which may include different types of registers for storing different types of data (e.g., integer registers, floating point registers, status registers, and an instruction pointer register). Some registers may be general-purpose registers, while other registers may be specific to the design of the processor 102.

In some embodiments, processor 102 is coupled to a processor bus 110 to transmit communication signals such as address, data, or control signals between processor 102 and other components in system 100. In one embodiment the system 100 uses an exemplary 'hub' system architecture, including a memory controller hub 116 and an Input Output (I/O) controller hub 130. A memory controller hub 116 facilitates communication between a memory device and other components of system 100, while an I/O Controller Hub (ICH) 130 provides connections to I/O devices via a local I/O bus. In one embodiment, the logic of the memory controller hub 116

is integrated within the processor.

Memory device 120 can be a dynamic random access memory (DRAM) device, a static random access memory (SRAM) device, flash memory device, phase-change memory device, or some other memory device having suitable performance to serve as process memory. In one
5 embodiment the memory device 120 can operate as system memory for the system 100, to store data 122 and instructions 121 for use when the one or more processors 102 executes an application or process. Memory controller hub 116 also couples with an optional external graphics processor 112, which may communicate with the one or more graphics processors 108 in processors 102 to perform graphics and media operations.

10 In some embodiments, ICH 130 enables peripherals to connect to memory device 120 and processor 102 via a high-speed I/O bus. The I/O peripherals include, but are not limited to, an audio controller 146, a firmware interface 128, a wireless transceiver 126 (e.g., Wi-Fi, Bluetooth), a data storage device 124 (e.g., hard disk drive, flash memory, etc.), and a legacy I/O controller 140 for coupling legacy (e.g., Personal System 2 (PS/2)) devices to the system. One
15 or more Universal Serial Bus (USB) controllers 142 connect input devices, such as keyboard and mouse 144 combinations. A network controller 134 may also couple to ICH 130. In some embodiments, a high-performance network controller (not shown) couples to processor bus 110. It will be appreciated that the system 100 shown is exemplary and not limiting, as other types of data processing systems that are differently configured may also be used. For example, the I/O
20 controller hub 130 may be integrated within the one or more processor 102, or the memory controller hub 116 and I/O controller hub 130 may be integrated into a discreet external graphics processor, such as the external graphics processor 112.

Figure 2 is a block diagram of an embodiment of a processor 200 having one or more processor cores 202A-202N, an integrated memory controller 214, and an integrated graphics
25 processor 208. Those elements of **Fig. 2** having the same reference numbers (or names) as the elements of any other figure herein can operate or function in any manner similar to that described elsewhere herein, but are not limited to such. Processor 200 can include additional cores up to and including additional core 202N represented by the dashed lined boxes. Each of processor cores 202A-202N includes one or more internal cache units 204A-204N. In some
30 embodiments each processor core also has access to one or more shared cached units 206.

The internal cache units 204A-204N and shared cache units 206 represent a cache memory hierarchy within the processor 200. The cache memory hierarchy may include at least one level of instruction and data cache within each processor core and one or more levels of shared mid-level cache, such as a Level 2 (L2), Level 3 (L3), Level 4 (L4), or other levels of cache, where
35 the highest level of cache before external memory is classified as the LLC. In some

embodiments, cache coherency logic maintains coherency between the various cache units 206 and 204A-204N.

In some embodiments, processor 200 may also include a set of one or more bus controller units 216 and a system agent core 210. The one or more bus controller units 216 manage a set of peripheral buses, such as one or more Peripheral Component Interconnect buses (e.g., PCI, PCI Express). System agent core 210 provides management functionality for the various processor components. In some embodiments, system agent core 210 includes one or more integrated memory controllers 214 to manage access to various external memory devices (not shown).

In some embodiments, one or more of the processor cores 202A-202N include support for simultaneous multi-threading. In such embodiment, the system agent core 210 includes components for coordinating and operating cores 202A-202N during multi-threaded processing. System agent core 210 may additionally include a power control unit (PCU), which includes logic and components to regulate the power state of processor cores 202A-202N and graphics processor 208.

In some embodiments, processor 200 additionally includes graphics processor 208 to execute graphics processing operations. In some embodiments, the graphics processor 208 couples with the set of shared cache units 206, and the system agent core 210, including the one or more integrated memory controllers 214. In some embodiments, a display controller 211 is coupled with the graphics processor 208 to drive graphics processor output to one or more coupled displays. In some embodiments, display controller 211 may be a separate module coupled with the graphics processor via at least one interconnect, or may be integrated within the graphics processor 208 or system agent core 210.

In some embodiments, a ring based interconnect unit 212 is used to couple the internal components of the processor 200. However, an alternative interconnect unit may be used, such as a point-to-point interconnect, a switched interconnect, or other techniques, including techniques well known in the art. In some embodiments, graphics processor 208 couples with the ring interconnect 212 via an I/O link 213.

The exemplary I/O link 213 represents at least one of multiple varieties of I/O interconnects, including an on package I/O interconnect which facilitates communication between various processor components and a high-performance embedded memory module 218, such as an eDRAM module. In some embodiments, each of the processor cores 202-202N and graphics processor 208 use embedded memory modules 218 as a shared Last Level Cache.

In some embodiments, processor cores 202A-202N are homogenous cores executing the same instruction set architecture. In another embodiment, processor cores 202A-202N are heterogeneous in terms of instruction set architecture (ISA), where one or more of processor

cores 202A-N execute a first instruction set, while at least one of the other cores executes a subset of the first instruction set or a different instruction set. In one embodiment processor cores 202A-202N are heterogeneous in terms of microarchitecture, where one or more cores having a relatively higher power consumption couple with one or more power cores having a lower power consumption. Additionally, processor 200 can be implemented on one or more chips or as an SoC integrated circuit having the illustrated components, in addition to other components.

Figure 3 is a block diagram of a graphics processor 300, which may be a discrete graphics processing unit, or may be a graphics processor integrated with a plurality of processing cores.

In some embodiments, the graphics processor communicates via a memory mapped I/O interface to registers on the graphics processor and with commands placed into the processor memory. In some embodiments, graphics processor 300 includes a memory interface 314 to access memory. Memory interface 314 can be an interface to local memory, one or more internal caches, one or more shared external caches, and/or to system memory.

In some embodiments, graphics processor 300 also includes a display controller 302 to drive display output data to a display device 320. Display controller 302 includes hardware for one or more overlay planes for the display and composition of multiple layers of video or user interface elements. In some embodiments, graphics processor 300 includes a video codec engine 306 to encode, decode, or transcode media to, from, or between one or more media encoding formats, including, but not limited to Moving Picture Experts Group (MPEG) formats such as MPEG-2, Advanced Video Coding (AVC) formats such as H.264/MPEG-4 AVC, as well as the Society of Motion Picture & Television Engineers (SMPTE) 421M/VC-1, and Joint Photographic Experts Group (JPEG) formats such as JPEG, and Motion JPEG (MJPEG) formats.

In some embodiments, graphics processor 300 includes a block image transfer (BLIT) engine 304 to perform two-dimensional (2D) rasterizer operations including, for example, bit-boundary block transfers. However, in one embodiment, 2D graphics operations are performed using one or more components of graphics processing engine (GPE) 310. In some embodiments, graphics processing engine 310 is a compute engine for performing graphics operations, including three-dimensional (3D) graphics operations and media operations.

In some embodiments, GPE 310 includes a 3D pipeline 312 for performing 3D operations, such as rendering three-dimensional images and scenes using processing functions that act upon 3D primitive shapes (e.g., rectangle, triangle, etc.). The 3D pipeline 312 includes programmable and fixed function elements that perform various tasks within the element and/or spawn execution threads to a 3D/Media sub-system 315. While 3D pipeline 312 can be used to perform media operations, an embodiment of GPE 310 also includes a media pipeline 316 that is

specifically used to perform media operations, such as video post-processing and image enhancement.

In some embodiments, media pipeline 316 includes fixed function or programmable logic units to perform one or more specialized media operations, such as video decode acceleration, video de-interlacing, and video encode acceleration in place of, or on behalf of video codec engine 306. In some embodiments, media pipeline 316 additionally includes a thread spawning unit to spawn threads for execution on 3D/Media sub-system 315. The spawned threads perform computations for the media operations on one or more graphics execution units included in 3D/Media sub-system 315.

In some embodiments, 3D/Media subsystem 315 includes logic for executing threads spawned by 3D pipeline 312 and media pipeline 316. In one embodiment, the pipelines send thread execution requests to 3D/Media subsystem 315, which includes thread dispatch logic for arbitrating and dispatching the various requests to available thread execution resources. The execution resources include an array of graphics execution units to process the 3D and media threads. In some embodiments, 3D/Media subsystem 315 includes one or more internal caches for thread instructions and data. In some embodiments, the subsystem also includes shared memory, including registers and addressable memory, to share data between threads and to store output data.

3D/Media Processing

Figure 4 is a block diagram of a graphics processing engine 410 of a graphics processor in accordance with some embodiments. In one embodiment, the GPE 410 is a version of the GPE 310 shown in **Fig. 3**. Elements of **Fig. 4** having the same reference numbers (or names) as the elements of any other figure herein can operate or function in any manner similar to that described elsewhere herein, but are not limited to such.

In some embodiments, GPE 410 couples with a command streamer 403, which provides a command stream to the GPE 3D and media pipelines 412, 416. In some embodiments, command streamer 403 is coupled to memory, which can be system memory, or one or more of internal cache memory and shared cache memory. In some embodiments, command streamer 403 receives commands from the memory and sends the commands to 3D pipeline 412 and/or media pipeline 416. The commands are directives fetched from a ring buffer, which stores commands for the 3D and media pipelines 412, 416. In one embodiment, the ring buffer can additionally include batch command buffers storing batches of multiple commands. The 3D and media pipelines 412, 416 process the commands by performing operations via logic within the respective pipelines or by dispatching one or more execution threads to an execution unit array 414. In some embodiments, execution unit array 414 is scalable, such that the array includes a

variable number of execution units based on the target power and performance level of GPE 410.

In some embodiments, a sampling engine 430 couples with memory (e.g., cache memory or system memory) and execution unit array 414. In some embodiments, sampling engine 430 provides a memory access mechanism for execution unit array 414 that allows execution array
5 414 to read graphics and media data from memory. In some embodiments, sampling engine 430 includes logic to perform specialized image sampling operations for media.

In some embodiments, the specialized media sampling logic in sampling engine 430 includes a de-noise/de-interlace module 432, a motion estimation module 434, and an image scaling and filtering module 436. In some embodiments, de-noise/de-interlace module 432
10 includes logic to perform one or more of a de-noise or a de-interlace algorithm on decoded video data. The de-interlace logic combines alternating fields of interlaced video content into a single frame of video. The de-noise logic reduces or removes data noise from video and image data. In some embodiments, the de-noise logic and de-interlace logic are motion adaptive and use spatial or temporal filtering based on the amount of motion detected in the video data. In some
15 embodiments, the de-noise/de-interlace module 432 includes dedicated motion detection logic (e.g., within the motion estimation engine 434).

In some embodiments, motion estimation engine 434 provides hardware acceleration for video operations by performing video acceleration functions such as motion vector estimation and prediction on video data. The motion estimation engine determines motion vectors that
20 describe the transformation of image data between successive video frames. In some embodiments, a graphics processor media codec uses video motion estimation engine 434 to perform operations on video at the macro-block level that may otherwise be too computationally intensive to perform with a general-purpose processor. In some embodiments, motion estimation engine 434 is generally available to graphics processor components to assist with video decode
25 and processing functions that are sensitive or adaptive to the direction or magnitude of the motion within video data.

In some embodiments, image scaling and filtering module 436 performs image-processing operations to enhance the visual quality of generated images and video. In some embodiments, scaling and filtering module 436 processes image and video data during the sampling operation
30 before providing the data to execution unit array 414.

In some embodiments, the GPE 410 includes a data port 444, which provides an additional mechanism for graphics subsystems to access memory. In some embodiments, data port 444 facilitates memory access for operations including render target writes, constant buffer reads, scratch memory space reads/writes, and media surface accesses. In some embodiments, data
35 port 444 includes cache memory space to cache accesses to memory. The cache memory can be

a single data cache or separated into multiple caches for the multiple subsystems that access memory via the data port (e.g., a render buffer cache, a constant buffer cache, etc.). In some embodiments, threads executing on an execution unit in execution unit array 414 communicate with the data port by exchanging messages via a data distribution interconnect that couples each of the sub-systems of GPE 410.

Execution Units

Figure 5 is a block diagram of another embodiment of a graphics processor 500. Elements of **Fig. 5** having the same reference numbers (or names) as the elements of any other figure herein can operate or function in any manner similar to that described elsewhere herein, but are not limited to such.

In some embodiments, graphics processor 500 includes a ring interconnect 502, a pipeline front-end 504, a media engine 537, and graphics cores 580A-580N. In some embodiments, ring interconnect 502 couples the graphics processor to other processing units, including other graphics processors or one or more general-purpose processor cores. In some embodiments, the graphics processor is one of many processors integrated within a multi-core processing system.

In some embodiments, graphics processor 500 receives batches of commands via ring interconnect 502. The incoming commands are interpreted by a command streamer 503 in the pipeline front-end 504. In some embodiments, graphics processor 500 includes scalable execution logic to perform 3D geometry processing and media processing via the graphics core(s) 580A-580N. For 3D geometry processing commands, command streamer 503 supplies commands to geometry pipeline 536. For at least some media processing commands, command streamer 503 supplies the commands to a video front end 534, which couples with a media engine 537. In some embodiments, media engine 537 includes a Video Quality Engine (VQE) 530 for video and image post-processing and a multi-format encode/decode (MFX) 533 engine to provide hardware-accelerated media data encode and decode. In some embodiments, geometry pipeline 536 and media engine 537 each generate execution threads for the thread execution resources provided by at least one graphics core 580A.

In some embodiments, graphics processor 500 includes scalable thread execution resources featuring modular cores 580A-580N (sometimes referred to as core slices), each having multiple sub-cores 550A-550N, 560A-560N (sometimes referred to as core sub-slices). In some embodiments, graphics processor 500 can have any number of graphics cores 580A through 580N. In some embodiments, graphics processor 500 includes a graphics core 580A having at least a first sub-core 550A and a second core sub-core 560A. In other embodiments, the graphics processor is a low power processor with a single sub-core (e.g., 550A). In some embodiments, graphics processor 500 includes multiple graphics cores 580A-580N, each

including a set of first sub-cores 550A-550N and a set of second sub-cores 560A-560N. Each sub-core in the set of first sub-cores 550A-550N includes at least a first set of execution units 552A-552N and media/texture samplers 554A-554N. Each sub-core in the set of second sub-cores 560A-560N includes at least a second set of execution units 562A-562N and samplers 564A-564N. In some embodiments, each sub-core 550A-550N, 560A-560N shares a set of shared resources 570A-570N. In some embodiments, the shared resources include shared cache memory and pixel operation logic. Other shared resources may also be included in the various embodiments of the graphics processor.

Figure 6 illustrates thread execution logic 600 including an array of processing elements employed in some embodiments of a GPE. Elements of **Fig. 6** having the same reference numbers (or names) as the elements of any other figure herein can operate or function in any manner similar to that described elsewhere herein, but are not limited to such.

In some embodiments, thread execution logic 600 includes a pixel shader 602, a thread dispatcher 604, instruction cache 606, a scalable execution unit array including a plurality of execution units 608A-608N, a sampler 610, a data cache 612, and a data port 614. In one embodiment the included components are interconnected via an interconnect fabric that links to each of the components. In some embodiments, thread execution logic 600 includes one or more connections to memory, such as system memory or cache memory, through one or more of instruction cache 606, data port 614, sampler 610, and execution unit array 608A-608N. In some embodiments, each execution unit (e.g. 608A) is an individual vector processor capable of executing multiple simultaneous threads and processing multiple data elements in parallel for each thread. In some embodiments, execution unit array 608A-608N includes any number individual execution units.

In some embodiments, execution unit array 608A-608N is primarily used to execute “shader” programs. In some embodiments, the execution units in array 608A-608N execute an instruction set that includes native support for many standard 3D graphics shader instructions, such that shader programs from graphics libraries (e.g., Direct 3D and OpenGL) are executed with a minimal translation. The execution units support vertex and geometry processing (e.g., vertex programs, geometry programs, vertex shaders), pixel processing (e.g., pixel shaders, fragment shaders) and general-purpose processing (e.g., compute and media shaders).

Each execution unit in execution unit array 608A-608N operates on arrays of data elements. The number of data elements is the “execution size,” or the number of channels for the instruction. An execution channel is a logical unit of execution for data element access, masking, and flow control within instructions. The number of channels may be independent of the number of physical Arithmetic Logic Units (ALUs) or Floating Point Units (FPUs) for a

particular graphics processor. In some embodiments, execution units 608A-608N support integer and floating-point data types.

The execution unit instruction set includes single instruction multiple data (SIMD) instructions. The various data elements can be stored as a packed data type in a register and the execution unit will process the various elements based on the data size of the elements. For example, when operating on a 256-bit wide vector, the 256 bits of the vector are stored in a register and the execution unit operates on the vector as four separate 64-bit packed data elements (Quad-Word (QW) size data elements), eight separate 32-bit packed data elements (Double Word (DW) size data elements), sixteen separate 16-bit packed data elements (Word (W) size data elements), or thirty-two separate 8-bit data elements (byte (B) size data elements). However, different vector widths and register sizes are possible.

One or more internal instruction caches (e.g., 606) are included in the thread execution logic 600 to cache thread instructions for the execution units. In some embodiments, one or more data caches (e.g., 612) are included to cache thread data during thread execution. In some embodiments, sampler 610 is included to provide texture sampling for 3D operations and media sampling for media operations. In some embodiments, sampler 610 includes specialized texture or media sampling functionality to process texture or media data during the sampling process before providing the sampled data to an execution unit.

During execution, the graphics and media pipelines send thread initiation requests to thread execution logic 600 via thread spawning and dispatch logic. In some embodiments, thread execution logic 600 includes a local thread dispatcher 604 that arbitrates thread initiation requests from the graphics and media pipelines and instantiates the requested threads on one or more execution units 608A-608N. For example, the geometry pipeline (e.g., 536 of **Fig. 5**) dispatches vertex processing, tessellation, or geometry processing threads to thread execution logic 600 (**Fig. 6**). In some embodiments, thread dispatcher 604 can also process runtime thread spawning requests from the executing shader programs.

Once a group of geometric objects has been processed and rasterized into pixel data, pixel shader 602 is invoked to further compute output information and cause results to be written to output surfaces (e.g., color buffers, depth buffers, stencil buffers, etc.). In some embodiments, pixel shader 602 calculates the values of the various vertex attributes that are to be interpolated across the rasterized object. In some embodiments, pixel shader 602 then executes an application programming interface (API)-supplied pixel shader program. To execute the pixel shader program, pixel shader 602 dispatches threads to an execution unit (e.g., 608A) via thread dispatcher 604. In some embodiments, pixel shader 602 uses texture sampling logic in sampler 610 to access texture data in texture maps stored in memory. Arithmetic operations on the

texture data and the input geometry data compute pixel color data for each geometric fragment, or discards one or more pixels from further processing.

In some embodiments, the data port 614 provides a memory access mechanism for the thread execution logic 600 output processed data to memory for processing on a graphics processor output pipeline. In some embodiments, the data port 614 includes or couples to one or more cache memories (e.g., data cache 612) to cache data for memory access via the data port.

Figure 7 is a block diagram illustrating a graphics processor instruction formats 700 according to some embodiments. In one or more embodiment, the graphics processor execution units support an instruction set having instructions in multiple formats. The solid lined boxes illustrate the components that are generally included in an execution unit instruction, while the dashed lines include components that are optional or that are only included in a sub-set of the instructions. In some embodiments, instruction format 700 described and illustrated are macro-instructions, in that they are instructions supplied to the execution unit, as opposed to micro-operations resulting from instruction decode once the instruction is processed.

In some embodiments, the graphics processor execution units natively support instructions in a 128-bit format 710. A 64-bit compacted instruction format 730 is available for some instructions based on the selected instruction, instruction options, and number of operands. The native 128-bit format 710 provides access to all instruction options, while some options and operations are restricted in the 64-bit format 730. The native instructions available in the 64-bit format 730 vary by embodiment. In some embodiments, the instruction is compacted in part using a set of index values in an index field 713. The execution unit hardware references a set of compaction tables based on the index values and uses the compaction table outputs to reconstruct a native instruction in the 128-bit format 710.

For each format, instruction opcode 712 defines the operation that the execution unit is to perform. The execution units execute each instruction in parallel across the multiple data elements of each operand. For example, in response to an add instruction the execution unit performs a simultaneous add operation across each color channel representing a texture element or picture element. By default, the execution unit performs each instruction across all data channels of the operands. In some embodiments, instruction control field 714 enables control over certain execution options, such as channels selection (e.g., predication) and data channel order (e.g., swizzle). For 128-bit instructions 710 an exec-size field 716 limits the number of data channels that will be executed in parallel. In some embodiments, exec-size field 716 is not available for use in the 64-bit compact instruction format 730.

Some execution unit instructions have up to three operands including two source operands, src0 722, src1 722, and one destination 718. In some embodiments, the execution units support

dual destination instructions, where one of the destinations is implied. Data manipulation instructions can have a third source operand (e.g., SRC2 724), where the instruction opcode 712 determines the number of source operands. An instruction's last source operand can be an immediate (e.g., hard-coded) value passed with the instruction.

5 In some embodiments, the 128-bit instruction format 710 includes an access/address mode information 726 specifying, for example, whether direct register addressing mode or indirect register addressing mode is used. When direct register addressing mode is used, the register address of one or more operands is directly provided by bits in the instruction 710.

10 In some embodiments, the 128-bit instruction format 710 includes an access/address mode field 726, which specifies an address mode and/or an access mode for the instruction. In one embodiment the access mode to define a data access alignment for the instruction. Some embodiments support access modes including a 16-byte aligned access mode and a 1-byte aligned access mode, where the byte alignment of the access mode determines the access alignment of the instruction operands. For example, when in a first mode, the instruction 710
15 may use byte-aligned addressing for source and destination operands and when in a second mode, the instruction 710 may use 16-byte-aligned addressing for all source and destination operands.

In one embodiment, the address mode portion of the access/address mode field 726 determines whether the instruction is to use direct or indirect addressing. When direct register
20 addressing mode is used bits in the instruction 710 directly provide the register address of one or more operands. When indirect register addressing mode is used, the register address of one or more operands may be computed based on an address register value and an address immediate field in the instruction.

In some embodiments instructions are grouped based on opcode 712 bit-fields to simplify
25 Opcode decode 740. For an 8-bit opcode, bits 4, 5, and 6 allow the execution unit to determine the type of opcode. The precise opcode grouping shown is merely an example. In some embodiments, a move and logic opcode group 742 includes data movement and logic instructions (e.g., move (mov), compare (cmp)). In some embodiments, move and logic group 742 shares the five most significant bits (MSB), where move (mov) instructions are in the form
30 of 0000xxxxb and logic instructions are in the form of 0001xxxxb. A flow control instruction group 744 (e.g., call, jump (jmp)) includes instructions in the form of 0010xxxxb (e.g., 0x20). A miscellaneous instruction group 746 includes a mix of instructions, including synchronization instructions (e.g., wait, send) in the form of 0011xxxxb (e.g., 0x30). A parallel math instruction group 748 includes component-wise arithmetic instructions (e.g., add, multiply (mul)) in the
35 form of 0100xxxxb (e.g., 0x40). The parallel math group 748 performs the arithmetic operations

in parallel across data channels. The vector math group 750 includes arithmetic instructions (e.g., dp4) in the form of 0101xxxxb (e.g., 0x50). The vector math group performs arithmetic such as dot product calculations on vector operands.

Graphics Pipeline

5 **Figure 8** is a block diagram of another embodiment of a graphics processor 800. Elements of **Fig. 8** having the same reference numbers (or names) as the elements of any other figure herein can operate or function in any manner similar to that described elsewhere herein, but are not limited to such.

10 In some embodiments, graphics processor 800 includes a graphics pipeline 820, a media pipeline 830, a display engine 840, thread execution logic 850, and a render output pipeline 870. In some embodiments, graphics processor 800 is a graphics processor within a multi-core processing system that includes one or more general purpose processing cores. The graphics processor is controlled by register writes to one or more control registers (not shown) or via commands issued to graphics processor 800 via a ring interconnect 802. In some embodiments,
15 ring interconnect 802 couples graphics processor 800 to other processing components, such as other graphics processors or general-purpose processors. Commands from ring interconnect 802 are interpreted by a command streamer 803, which supplies instructions to individual components of graphics pipeline 820 or media pipeline 830.

20 In some embodiments, command streamer 803 directs the operation of a vertex fetcher 805 that reads vertex data from memory and executes vertex-processing commands provided by command streamer 803. In some embodiments, vertex fetcher 805 provides vertex data to a vertex shader 807, which performs coordinate space transformation and lighting operations to each vertex. In some embodiments, vertex fetcher 805 and vertex shader 807 execute vertex-processing instructions by dispatching execution threads to execution units 852A, 852B via a
25 thread dispatcher 831.

30 In some embodiments, execution units 852A, 852B are an array of vector processors having an instruction set for performing graphics and media operations. In some embodiments, execution units 852A, 852B have an attached L1 cache 851 that is specific for each array or shared between the arrays. The cache can be configured as a data cache, an instruction cache, or a single cache that is partitioned to contain data and instructions in different partitions.

35 In some embodiments, graphics pipeline 820 includes tessellation components to perform hardware-accelerated tessellation of 3D objects. In some embodiments, a programmable hull shader 811 configures the tessellation operations. A programmable domain shader 817 provides back-end evaluation of tessellation output. A tessellator 813 operates at the direction of hull
shader 811 and contains special purpose logic to generate a set of detailed geometric objects

based on a coarse geometric model that is provided as input to graphics pipeline 820. In some embodiments, if tessellation is not used, tessellation components 811, 813, 817 can be bypassed.

In some embodiments, complete geometric objects can be processed by a geometry shader 819 via one or more threads dispatched to execution units 852A, 852B, or can proceed directly to the clipper 829. In some embodiments, the geometry shader operates on entire geometric objects, rather than vertices or patches of vertices as in previous stages of the graphics pipeline. If the tessellation is disabled the geometry shader 819 receives input from the vertex shader 807. In some embodiments, geometry shader 819 is programmable by a geometry shader program to perform geometry tessellation if the tessellation units are disabled.

Before rasterization, a clipper 829 processes vertex data. The clipper 829 may be a fixed function clipper or a programmable clipper having clipping and geometry shader functions. In some embodiments, a rasterizer and depth test component 873 in the render output pipeline 870 dispatches pixel shaders to convert the geometric objects into their per pixel representations. In some embodiments, pixel shader logic is included in thread execution logic 850. In some embodiments, an application can bypass the rasterizer 873 and access un-rasterized vertex data via a stream out unit 823.

The graphics processor 800 has an interconnect bus, interconnect fabric, or some other interconnect mechanism that allows data and message passing amongst the major components of the processor. In some embodiments, execution units 852A, 852B and associated cache(s) 851, texture and media sampler 854, and texture/sampler cache 858 interconnect via a data port 856 to perform memory access and communicate with render output pipeline components of the processor. In some embodiments, sampler 854, caches 851, 858 and execution units 852A, 852B each have separate memory access paths.

In some embodiments, render output pipeline 870 contains a rasterizer and depth test component 873 that converts vertex-based objects into an associated pixel-based representation. In some embodiments, the rasterizer logic includes a windower/masker unit to perform fixed function triangle and line rasterization. An associated render cache 878 and depth cache 879 are also available in some embodiments. A pixel operations component 877 performs pixel-based operations on the data, though in some instances, pixel operations associated with 2D operations (e.g. bit block image transfers with blending) are performed by the 2D engine 841, or substituted at display time by the display controller 843 using overlay display planes. In some embodiments, a shared L3 cache 875 is available to all graphics components, allowing the sharing of data without the use of main system memory.

In some embodiments, graphics processor media pipeline 830 includes a media engine 837 and a video front end 834. In some embodiments, video front end 834 receives pipeline

commands from the command streamer 803. In some embodiments, media pipeline 830 includes a separate command streamer. In some embodiments, video front-end 834 processes media commands before sending the command to the media engine 837. In some embodiments, media engine 337 includes thread spawning functionality to spawn threads for dispatch to thread
5 execution logic 850 via thread dispatcher 831.

In some embodiments, graphics processor 800 includes a display engine 840. In some embodiments, display engine 840 is external to processor 800 and couples with the graphics processor via the ring interconnect 802, or some other interconnect bus or fabric. In some
10 embodiments, display engine 840 includes a 2D engine 841 and a display controller 843. In some embodiments, display engine 840 contains special purpose logic capable of operating independently of the 3D pipeline. In some embodiments, display controller 843 couples with a display device (not shown), which may be a system integrated display device, as in a laptop computer, or an external display device attached via a display device connector.

In some embodiments, graphics pipeline 820 and media pipeline 830 are configurable to
15 perform operations based on multiple graphics and media programming interfaces and are not specific to any one application programming interface (API). In some embodiments, driver software for the graphics processor translates API calls that are specific to a particular graphics or media library into commands that can be processed by the graphics processor. In some
20 embodiments, support is provided for the Open Graphics Library (OpenGL) and Open Computing Language (OpenCL) from the Khronos Group, the Direct3D library from the Microsoft Corporation, or support may be provided to both OpenGL and D3D. Support may also be provided for the Open Source Computer Vision Library (OpenCV). A future API with a compatible 3D pipeline would also be supported if a mapping can be made from the pipeline of the future API to the pipeline of the graphics processor.

Graphics Pipeline Programming

Figure 9A is a block diagram illustrating a graphics processor command format 900 according to some embodiments. **Figure 9B** is a block diagram illustrating a graphics processor command sequence 910 according to an embodiment. The solid lined boxes in **Fig. 9A** illustrate the components that are generally included in a graphics command while the dashed lines
30 include components that are optional or that are only included in a sub-set of the graphics commands. The exemplary graphics processor command format 900 of **Fig. 9A** includes data fields to identify a target client 902 of the command, a command operation code (opcode) 904, and the relevant data 906 for the command. A sub-opcode 905 and a command size 908 are also included in some commands.

35 In some embodiments, client 902 specifies the client unit of the graphics device that

processes the command data. In some embodiments, a graphics processor command parser examines the client field of each command to condition the further processing of the command and route the command data to the appropriate client unit. In some embodiments, the graphics processor client units include a memory interface unit, a render unit, a 2D unit, a 3D unit, and a media unit. Each client unit has a corresponding processing pipeline that processes the commands. Once the command is received by the client unit, the client unit reads the opcode and, if present, sub-opcode to determine the operation to perform. The client unit performs the command using information in data field. For some commands an explicit command size is expected to specify the size of the command. In some embodiments, the command parser automatically determines the size of at least some of the commands based on the command opcode. In some embodiments commands are aligned via multiples of a double word.

The flow diagram in **Fig. 9B** shows an exemplary graphics processor command sequence. In some embodiments, software or firmware of a data processing system that features an embodiment of a graphics processor uses a version of the command sequence shown to set up, execute, and terminate a set of graphics operations. A sample command sequence is shown and described for purposes of example only as embodiments are not limited to these specific commands or to this command sequence. Moreover, the commands may be issued as batch of commands in a command sequence, such that the graphics processor will process the sequence of commands in at least partially concurrence.

In some embodiments, the graphics processor command sequence may begin with a pipeline flush command to cause any active graphics pipeline to complete the currently pending commands for the pipeline. In some embodiments, the 3D pipeline and the media pipeline do not operate concurrently. The pipeline flush is performed to cause the active graphics pipeline to complete any pending commands. In response to a pipeline flush, the command parser for the graphics processor will pause command processing until the active drawing engines complete pending operations and the relevant read caches are invalidated. Optionally, any data in the render cache that is marked 'dirty' can be flushed to memory. In some embodiments, pipeline flush command can be used for pipeline synchronization or before placing the graphics processor into a low power state.

In some embodiments, a pipeline select command is used when a command sequence requires the graphics processor to explicitly switch between pipelines. In some embodiments, a pipeline select command is required only once within an execution context before issuing pipeline commands unless the context is to issue commands for both pipelines. In some embodiments, a pipeline flush command is required immediately before a pipeline switch

via the pipeline select command 913.

In some embodiments, a pipeline control command 914 configures a graphics pipeline for operation and is used to program the 3D pipeline 922 and the media pipeline 924. In some embodiments, pipeline control command 914 configures the pipeline state for the active pipeline.

5 In one embodiment, the pipeline control command 914 is used for pipeline synchronization and to clear data from one or more cache memories within the active pipeline before processing a batch of commands.

In some embodiments, return buffer state commands 916 are used to configure a set of return buffers for the respective pipelines to write data. Some pipeline operations require the
10 allocation, selection, or configuration of one or more return buffers into which the operations write intermediate data during processing. In some embodiments, the graphics processor also uses one or more return buffers to store output data and to perform cross thread communication. In some embodiments, the return buffer state 916 includes selecting the size and number of return buffers to use for a set of pipeline operations.

15 The remaining commands in the command sequence differ based on the active pipeline for operations. Based on a pipeline determination 920, the command sequence is tailored to the 3D pipeline 922 beginning with the 3D pipeline state 930, or the media pipeline 924 beginning at the media pipeline state 940.

The commands for the 3D pipeline state 930 include 3D state setting commands for vertex
20 buffer state, vertex element state, constant color state, depth buffer state, and other state variables that are to be configured before 3D primitive commands are processed. The values of these commands are determined at least in part based the particular 3D API in use. In some embodiments, 3D pipeline state 930 commands are also able to selectively disable or bypass certain pipeline elements if those elements will not be used.

25 In some embodiments, 3D primitive 932 command is used to submit 3D primitives to be processed by the 3D pipeline. Commands and associated parameters that are passed to the graphics processor via the 3D primitive 932 command are forwarded to the vertex fetch function in the graphics pipeline. The vertex fetch function uses the 3D primitive 932 command data to generate vertex data structures. The vertex data structures are stored in one or more return
30 buffers. In some embodiments, 3D primitive 932 command is used to perform vertex operations on 3D primitives via vertex shaders. To process vertex shaders, 3D pipeline 922 dispatches shader execution threads to graphics processor execution units.

In some embodiments, 3D pipeline 922 is triggered via an execute 934 command or event. In some embodiments, a register write triggers command execution. In some embodiments
35 execution is triggered via a 'go' or 'kick' command in the command sequence. In one

embodiment command execution is triggered using a pipeline synchronization command to flush the command sequence through the graphics pipeline. The 3D pipeline will perform geometry processing for the 3D primitives. Once operations are complete, the resulting geometric objects are rasterized and the pixel engine colors the resulting pixels. Additional commands to control pixel shading and pixel back end operations may also be included for those operations.

In some embodiments, the graphics processor command sequence 910 follows the media pipeline 924 path when performing media operations. In general, the specific use and manner of programming for the media pipeline 924 depends on the media or compute operations to be performed. Specific media decode operations may be offloaded to the media pipeline during media decode. In some embodiments, the media pipeline can also be bypassed and media decode can be performed in whole or in part using resources provided by one or more general purpose processing cores. In one embodiment, the media pipeline also includes elements for general-purpose graphics processor unit (GPGPU) operations, where the graphics processor is used to perform SIMD vector operations using computational shader programs that are not explicitly related to the rendering of graphics primitives.

In some embodiments, media pipeline 924 is configured in a similar manner as the 3D pipeline 922. A set of media pipeline state commands 940 are dispatched or placed into in a command queue before the media object commands 942. In some embodiments, media pipeline state commands 940 include data to configure the media pipeline elements that will be used to process the media objects. This includes data to configure the video decode and video encode logic within the media pipeline, such as encode or decode format. In some embodiments, media pipeline state commands 940 also support the use one or more pointers to “indirect” state elements that contain a batch of state settings.

In some embodiments, media object commands 942 supply pointers to media objects for processing by the media pipeline. The media objects include memory buffers containing video data to be processed. In some embodiments, all media pipeline states must be valid before issuing a media object command 942. Once the pipeline state is configured and media object commands 942 are queued, the media pipeline 924 is triggered via an execute command 944 or an equivalent execute event (e.g., register write). Output from media pipeline 924 may then be post processed by operations provided by the 3D pipeline 922 or the media pipeline 924. In some embodiments, GPGPU operations are configured and executed in a similar manner as media operations.

Graphics Software Architecture

Figure 10 illustrates exemplary graphics software architecture for a data processing system according to some embodiments. In some embodiments, software architecture includes a

3D graphics application 1010, an operating system 1020, and at least one processor 1030. In some embodiments, processor 1030 includes a graphics processor 1032 and one or more general-purpose processor core(s) 1034. The graphics application 1010 and operating system 1020 each execute in the system memory 1050 of the data processing system.

5 In some embodiments, 3D graphics application 1010 contains one or more shader programs including shader instructions 1012. The shader language instructions may be in a high-level shader language, such as the High Level Shader Language (HLSL) or the OpenGL Shader Language (GLSL). The application also includes executable instructions 1014 in a machine language suitable for execution by the general-purpose processor core 1034. The
10 application also includes graphics objects 1016 defined by vertex data.

In some embodiments, operating system 1020 is a Microsoft® Windows® operating system from the Microsoft Corporation, a proprietary UNIX-like operating system, or an open source UNIX-like operating system using a variant of the Linux kernel. When the Direct3D API is in use, the operating system 1020 uses a front-end shader compiler 1024 to compile any shader
15 instructions 1012 in HLSL into a lower-level shader language. The compilation may be a just-in-time (JIT) compilation or the application can perform shader pre-compilation. In some embodiments, high-level shaders are compiled into low-level shaders during the compilation of the 3D graphics application 1010.

In some embodiments, user mode graphics driver 1026 contains a back-end shader
20 compiler 1027 to convert the shader instructions 1012 into a hardware specific representation. When the OpenGL API is in use, shader instructions 1012 in the GLSL high-level language are passed to a user mode graphics driver 1026 for compilation. In some embodiments, user mode graphics driver 1026 uses operating system kernel mode functions 1028 to communicate with a kernel mode graphics driver 1029. In some embodiments, kernel mode graphics driver 1029
25 communicates with graphics processor 1032 to dispatch commands and instructions.

IP Core Implementations

One or more aspects of at least one embodiment may be implemented by representative code stored on a machine-readable medium which represents and/or defines logic within an integrated circuit such as a processor. For example, the machine-readable medium may include
30 instructions which represent various logic within the processor. When read by a machine, the instructions may cause the machine to fabricate the logic to perform the techniques described herein. Such representations, known as “IP cores,” are reusable units of logic for an integrated circuit that may be stored on a tangible, machine-readable medium as a hardware model that describes the structure of the integrated circuit. The hardware model may be supplied to various
35 customers or manufacturing facilities, which load the hardware model on fabrication machines

that manufacture the integrated circuit. The integrated circuit may be fabricated such that the circuit performs operations described in association with any of the embodiments described herein.

Figure 11 is a block diagram illustrating an IP core development system 1100 that may be used to manufacture an integrated circuit to perform operations according to an embodiment. The IP core development system 1100 may be used to generate modular, re-usable designs that can be incorporated into a larger design or used to construct an entire integrated circuit (e.g., an SOC integrated circuit). A design facility 1130 can generate a software simulation 1110 of an IP core design in a high level programming language (e.g., C/C++). The software simulation 1110 can be used to design, test, and verify the behavior of the IP core. A register transfer level (RTL) design can then be created or synthesized from the simulation model 1100. The RTL design 1115 is an abstraction of the behavior of the integrated circuit that models the flow of digital signals between hardware registers, including the associated logic performed using the modeled digital signals. In addition to an RTL design 1115, lower-level designs at the logic level or transistor level may also be created, designed, or synthesized. Thus, the particular details of the initial design and simulation may vary.

The RTL design 1115 or equivalent may be further synthesized by the design facility into a hardware model 1120, which may be in a hardware description language (HDL), or some other representation of physical design data. The HDL may be further simulated or tested to verify the IP core design. The IP core design can be stored for delivery to a 3rd party fabrication facility 1165 using non-volatile memory 1140 (e.g., hard disk, flash memory, or any non-volatile storage medium). Alternatively, the IP core design may be transmitted (e.g., via the Internet) over a wired connection 1150 or wireless connection 1160. The fabrication facility 1165 may then fabricate an integrated circuit that is based at least in part on the IP core design. The fabricated integrated circuit can be configured to perform operations in accordance with at least one embodiment described herein.

Figure 12 is a block diagram illustrating an exemplary system on a chip integrated circuit 1200 that may be fabricated using one or more IP cores, according to an embodiment. The exemplary integrated circuit includes one or more application processors 1205 (e.g., CPUs), at least one graphics processor 1210, and may additionally include an image processor 1215 and/or a video processor 1220, any of which may be a modular IP core from the same or multiple different design facilities. The integrated circuit includes peripheral or bus logic including a USB controller 1225, UART controller 1230, an SPI/SDIO controller 1235, and an I²S/I²C controller 1240. Additionally, the integrated circuit can include a display device 1245 coupled to one or more of a high-definition multimedia interface (HDMI) controller 1250 and a mobile

industry processor interface (MIPI) display interface 1255. Storage may be provided by a flash memory subsystem 1260 including flash memory and a flash memory controller. Memory interface may be provided via a memory controller 1265 for access to SDRAM or SRAM memory devices. Some integrated circuits additionally include an embedded security engine

5 1270.

Additionally, other logic and circuits may be included in the processor of integrated circuit 1200, including additional graphics processors/cores, peripheral interface controllers, or general purpose processor cores.

Figure 13 illustrates a conventional technique 1300. As illustrated, a graphics driver

10 involve a single command buffer, such as command buffer 1301, to handle various GPU commands 1305A, 1305B, 1305C in response to their corresponding API calls 1303A, 1303B, 1303C.

Figure 14 illustrates a conventional technique 1400. As illustrated, API call 1401 facilitates execution of a number of GPU commands 1405A, 1405B, 1405C, 1405D at GPU

15 command buffer 1403A, leading to GPU computer shader 1407 reading from parameters 1409 and writing to GPU command buffer 1403B.

Figure 15 illustrates a computing device 1500 employing an efficient graphics commands generation/execution mechanism 1510 according to one embodiment. Computing device 1300 (e.g., mobile computer, laptop computer, desktop computer, etc.) may be the same as data

20 processing system 100 of **Figure 1** and accordingly, for brevity, clarity, and ease of understanding, many of the details stated above with reference to **Figures 1-12** are not further discussed or repeated hereafter. For example, computing device 1500 may include a mobile computer (e.g., smartphone, tablet computer, laptops, game consoles, portable workstations, smart glasses and other smart wearable devices, etc.) serving as a host machine for hosting

25 graphics commands generation/execution mechanism (“commands mechanism”) 1510.

Commands mechanism 1510 may include any number and type of components for facilitating runtime transformation of graphics processing commands for efficient performance of graphics contents, such as images, videos, 3D applications, games, etc., according to one embodiment. Throughout the document, the term “user” may be interchangeably referred to as

30 “viewer”, “observer”, “person”, “individual”, “end-user”, and/or the like. It is to be noted that throughout this document, terms like “graphics domain” may be referenced interchangeably with “graphics processing unit” or simply “GPU” and similarly, “CPU domain” or “host domain” may be referenced interchangeably with “computer processing unit” or simply “CPU”.

Computing device 1500 may include any number and type of communication devices, such

35 as large computing systems, such as server computers, desktop computers, etc., and may further

include set-top boxes (e.g., Internet-based cable television set-top boxes, etc.), global positioning system (GPS)-based devices, etc. Computing device 1500 may include mobile computing devices serving as communication devices, such as cellular phones including smartphones, personal digital assistants (PDAs), tablet computers, laptop computers, e-readers, smart
5 televisions, television platforms, wearable devices (e.g., glasses, watches, bracelets, smartcards, jewelry, clothing items, etc.), media players, etc. For example, in one embodiment, computing device 1500 may include a mobile computing device employing an integrated circuit (“IC”), such as system on a chip (“SoC” or “SOC”), integrating various hardware and/or software components of computing device 1500 on a single chip.

10 As illustrated, in one embodiment, in addition to employing transformation mechanism 1510, computing device 1500 may further include any number and type of hardware components and/or software components, such as (but not limited to) GPU 1514 (having driver logic 1516), CPU 1512, memory 1508, network devices, drivers, or the like, as well as input/output (I/O)
15 sources 1504, such as touchscreens, touch panels, touch pads, virtual or regular keyboards, virtual or regular mice, ports, connectors, etc. Computing device 1500 may include operating system (OS) 1506 serving as an interface between hardware and/or physical resources of the computer device 1500 and a user. It is contemplated that CPU 1512 may include one or processors, such as processor(s) 102 of **Figure 1**, while GPU 1514 may include one or more graphics processors, such as graphics processor(s) 108 of **Figure 1**. In one embodiment and as
20 will be further described with reference to the subsequent figures, transformation mechanism 1510 may be in communication with its host driver logic 1516 which cooperates with GPU 1514 to facilitate any number and type of tasks facilitating generation and rendering of virtual 3D images as is described through this document.

It is to be noted that terms like “node”, “computing node”, “server”, “server device”,
25 “cloud computer”, “cloud server”, “cloud server computer”, “machine”, “host machine”, “device”, “computing device”, “computer”, “computing system”, and the like, may be used interchangeably throughout this document. It is to be further noted that terms like “application”, “software application”, “program”, “software program”, “package”, “software package”, and the like, may be used interchangeably throughout this document. Also, terms like “job”, “input”,
30 “request”, “message”, and the like, may be used interchangeably throughout this document.

It is contemplated and as further described with reference to **Figures 1-12**, some processes of the graphics pipeline as described above are implemented in software, while the rest are implemented in hardware. A graphics pipeline may be implemented in a graphics coprocessor design, where CPU 1512 is designed to work with GPU 1514 which may be included in or co-
35 located with CPU 1512. In one embodiment, GPU 1514 may employ any number and type of

conventional software and hardware logic to perform the conventional functions relating to graphics rendering as well as novel software and hardware logic to execute any number and type of instructions, such as instructions 121 of **Figure 1**, to perform the various novel functions of transformation mechanism 1510 as disclosed throughout this document.

5 As aforementioned, memory 1508 may include a random access memory (RAM) comprising application database having object information. A memory controller hub, such as memory controller hub 116 of **Figure 1**, may access data in the RAM and forward it to GPU 1514 for graphics pipeline processing. RAM may include double data rate RAM (DDR RAM), extended data output RAM (EDO RAM), etc. CPU 1512 interacts with a hardware graphics
10 pipeline, as illustrated with reference to **Figure 3**, to share graphics pipelining functionality. Processed data is stored in a buffer in the hardware graphics pipeline, and state information is stored in memory 1508. The resulting image is then transferred to I/O sources 1504, such as a display component, such as display device 320 of **Figure 3**, for displaying of the image. It is contemplated that the display device may be of various types, such as Cathode Ray Tube (CRT),
15 Thin Film Transistor (TFT), Liquid Crystal Display (LCD), Organic Light Emitting Diode (OLED) array, etc., to display information to a user.

Memory 1508 may comprise a pre-allocated region of a buffer (e.g., frame buffer); however, it should be understood by one of ordinary skill in the art that the embodiments are not so limited, and that any memory accessible to the lower graphics pipeline may be used.

20 Computing device 1500 may further include input/output (I/O) control hub (ICH) 130 as referenced in **Figure 1**, one or more I/O sources 1504, etc.

CPU 1512 may include one or more processors to execute instructions in order to perform whatever software routines the computing system implements. The instructions frequently involve some sort of operation performed upon data. Both data and instructions may be stored in
25 system memory 1508 and any associated cache. Cache is typically designed to have shorter latency times than system memory 1508; for example, cache might be integrated onto the same silicon chip(s) as the processor(s) and/or constructed with faster static RAM (SRAM) cells whilst the system memory 1508 might be constructed with slower dynamic RAM (DRAM) cells. By tending to store more frequently used instructions and data in the cache as opposed to the system
30 memory 1508, the overall performance efficiency of computing device 1500 improves. It is contemplated that in some embodiments, GPU 1514 may exist as part of CPU 1512 (such as part of a physical CPU package) in which case, memory 1508 may be shared by CPU 1512 and GPU 1514 or kept separated.

System memory 1508 may be made available to other components within the computing
35 device 1500. For example, any data (e.g., input graphics data) received from various interfaces

to the computing device 1500 (e.g., keyboard and mouse, printer port, Local Area Network (LAN) port, modem port, etc.) or retrieved from an internal storage element of the computer device 1500 (e.g., hard disk drive) are often temporarily queued into system memory 1508 prior to their being operated upon by the one or more processor(s) in the implementation of a software program. Similarly, data that a software program determines should be sent from the computing device 1500 to an outside entity through one of the computing system interfaces, or stored into an internal storage element, is often temporarily queued in system memory 1508 prior to its being transmitted or stored.

Further, for example, an ICH, such as ICH 130 of **Figure 1**, may be used for ensuring that such data is properly passed between the system memory 1508 and its appropriate corresponding computing system interface (and internal storage device if the computing system is so designed) and may have bi-directional point-to-point links between itself and the observed I/O sources/devices 1504. Similarly, an MCH, such as MCH 116 of **Figure 1**, may be used for managing the various contending requests for system memory 1508 accesses amongst CPU 1512 and GPU 1514, interfaces and internal storage elements that may proximately arise in time with respect to one another.

I/O sources 1504 may include one or more I/O devices that are implemented for transferring data to and/or from computing device 1500 (e.g., a networking adapter); or, for a large scale non-volatile storage within computing device 1500 (e.g., hard disk drive). User input device, including alphanumeric and other keys, may be used to communicate information and command selections to GPU 1514. Another type of user input device is cursor control, such as a mouse, a trackball, a touchscreen, a touchpad, or cursor direction keys to communicate direction information and command selections to GPU 1514 and to control cursor movement on the display device. Camera and microphone arrays of computer device 1500 may be employed to observe gestures, record audio and video and to receive and transmit visual and audio commands.

Computing device 1500 may further include network interface(s) to provide access to a network, such as a LAN, a wide area network (WAN), a metropolitan area network (MAN), a personal area network (PAN), Bluetooth, a cloud network, a mobile network (e.g., 3rd Generation (3G), etc.), an intranet, the Internet, etc. Network interface(s) may include, for example, a wireless network interface having antenna, which may represent one or more antenna(e). Network interface(s) may also include, for example, a wired network interface to communicate with remote devices via network cable, which may be, for example, an Ethernet cable, a coaxial cable, a fiber optic cable, a serial cable, or a parallel cable.

Network interface(s) may provide access to a LAN, for example, by conforming to IEEE

802.11b and/or IEEE 802.11g standards, and/or the wireless network interface may provide access to a personal area network, for example, by conforming to Bluetooth standards. Other wireless network interfaces and/or protocols, including previous and subsequent versions of the standards, may also be supported. In addition to, or instead of, communication via the wireless LAN standards, network interface(s) may provide wireless communication using, for example, Time Division, Multiple Access (TDMA) protocols, Global Systems for Mobile Communications (GSM) protocols, Code Division, Multiple Access (CDMA) protocols, and/or any other type of wireless communications protocols.

Network interface(s) may include one or more communication interfaces, such as a modem, a network interface card, or other well-known interface devices, such as those used for coupling to the Ethernet, token ring, or other types of physical wired or wireless attachments for purposes of providing a communication link to support a LAN or a WAN, for example. In this manner, the computer system may also be coupled to a number of peripheral devices, clients, control surfaces, consoles, or servers via a conventional network infrastructure, including an Intranet or the Internet, for example.

It is to be appreciated that a lesser or more equipped system than the example described above may be preferred for certain implementations. Therefore, the configuration of computing device 1500 may vary from implementation to implementation depending upon numerous factors, such as price constraints, performance requirements, technological improvements, or other circumstances. Examples of the electronic device or computer system 1500 may include (without limitation) a mobile device, a personal digital assistant, a mobile computing device, a smartphone, a cellular telephone, a handset, a one-way pager, a two-way pager, a messaging device, a computer, a personal computer (PC), a desktop computer, a laptop computer, a notebook computer, a handheld computer, a tablet computer, a server, a server array or server farm, a web server, a network server, an Internet server, a work station, a mini-computer, a main frame computer, a supercomputer, a network appliance, a web appliance, a distributed computing system, multiprocessor systems, processor-based systems, consumer electronics, programmable consumer electronics, television, digital television, set top box, wireless access point, base station, subscriber station, mobile subscriber center, radio network controller, router, hub, gateway, bridge, switch, machine, or combinations thereof.

Embodiments may be implemented as any or a combination of: one or more microchips or integrated circuits interconnected using a parentboard, hardwired logic, software stored by a memory device and executed by a microprocessor, firmware, an application specific integrated circuit (ASIC), and/or a field programmable gate array (FPGA). The term "logic" may include, by way of example, software or hardware and/or combinations of software and hardware.

Embodiments may be provided, for example, as a computer program product which may include one or more machine-readable media having stored thereon machine-executable instructions that, when executed by one or more machines such as a computer, network of computers, or other electronic devices, may result in the one or more machines carrying out operations in accordance with embodiments described herein. A machine-readable medium may include, but is not limited to, floppy diskettes, optical disks, CD-ROMs (Compact Disc-Read Only Memories), and magneto-optical disks, ROMs, RAMs, EPROMs (Erasable Programmable Read Only Memories), EEPROMs (Electrically Erasable Programmable Read Only Memories), magnetic or optical cards, flash memory, or other type of media/machine-readable medium suitable for storing machine-executable instructions.

Moreover, embodiments may be downloaded as a computer program product, wherein the program may be transferred from a remote computer (e.g., a server) to a requesting computer (e.g., a client) by way of one or more data signals embodied in and/or modulated by a carrier wave or other propagation medium via a communication link (e.g., a modem and/or network connection).

Figure 16 illustrates an efficient graphics commands generation/execution mechanism 1510 according to one embodiment. In one embodiment, commands mechanism 1510 may include any number and type of components to perform various tasks for efficient utilization of GPU 1514 for better rendering of display contents at computing device 1500. For example and in one embodiment, commands mechanism 1510 may include any number and type of components, such as (without limitation): reception/detection logic 1601; patch phase logic 1603; submit phase logic 1605; application/execution logic 1607; and communication/compatibility logic 1609. Computing device 1500 is further in communication with database 1630 (e.g., Cloud storage, non-Cloud storage, etc.) over one or more networks (e.g., Internet, Cloud-based network, proximity network, etc.).

As an initial matter, it is contemplated that in one embodiment, as illustrated here, commands mechanism 1510 may be hosted by driver logic 1516 of a graphics driver of GPU 1514 while, in another embodiment, commands mechanism 1510 may be hosted by another component of GPU 1514. In yet another embodiment, commands mechanism 1510 may not be hosted by GPU 1514 and that it may be hosted by an operating system, such as operating system 1506, as illustrated in **Figure 15**. Similarly, in yet another embodiment, various functionalities and/or components of commands mechanism 1510 may be provided as one or more hardware components of computing device 1500. Accordingly, embodiments are not limited to any particular form of implementation regarding hosting of commands mechanism 1510 at computing device 1500.

It is contemplated that an application (e.g., game application) may include any number and types of frames having any number and type of dispatches or draws, where each dispatch may represent a portion (e.g., desk in classroom, cloud in sky, nose on face, etc.) of an image (e.g., classroom, sky, face, etc.) represented by a frame. In one embodiment, any size and content of a dispatch, like its corresponding frame, may be predetermined by the application. It is contemplated that embodiments are not limited to any particular number or type of applications, frames, dispatches, etc., and similarly, embodiments are not limited to any particular number and type of image or size of image portions. It is contemplated that an application can have any number of frames and each frame may have any number of dispatches and each dispatch may have any type of image content. For example, in some embodiments, a dispatch may indicate a number of threads per group to be used and/or a number of groups are to be executed in a command sequence, etc.

As an initial matter, as will be further described throughout this document and illustrated with reference to **Figure 15**, embodiments are not limited a particular use case or optimization scenario and that any number and type of optimization scenarios may be exercised using commands mechanism 1510. For example, commands mechanism 1510 may enable or facilitate graphics driver logic (also referred to as “graphics logic”, “graphics driver”, or simply “driver”) 1516 to perform both high and low level optimizations for GPGPU applications (e.g., game applications) for improved performance on a corresponding GPU 1514. For example, such performance enhancement opportunities for GPU/GPGPU workloads may be achieved through any number and type of utilization scenarios for GPU optimization as facilitated by commands mechanism 1510, such optimization scenarios may include batching of various GPU-based commands for efficient generation and execution of such commands for improved and efficient rendering of graphics contents. For example, various utilization scenarios as facilitated by commands mechanism 1510 may allow for reducing the effort spent in optimizing applications for every GPU which, in turn, improves the viewing experience for the user, such as a software developer, a system administrator, an end-user, etc., when using 3D games and other high performing computing, etc.

It is contemplated that calls placed through 3D graphics APIs (e.g., DirectX® ExecuteIndirect API) include features that allow for an application to specify a method for GPU 1514 to dynamically generate work for itself, such as by allowing indirect dispatching to perform multiple draws with a single API call, giving the ability to both the CPU and GPU 1514 to control draw calls. Having GPU 1514 generate work for itself further allows for GPU 1514 to continue to run and perform its tasks without having to wait for the CPU to tell GPU 1514 about what changes to make. These features are implemented using GPGPU compute shaders

(“compute shaders” or simply “shaders”) to generate batch buffers that are then executed by a 3D render engine (“rendering engine”); however, serious performance penalties are accrued when the processes switch between compute shading and 3D rendering. For example, a single API call results in a single compute shader execution followed by a single 3D rendering and thus, multiple back-to-back API calls can result in multiple switching, resulting in a significant waste of GPU resources.

It is contemplated that embodiments are not limited to merely the processes of GPGPU computer shader and 3D renderer or any other number and type of GPU processes; however, for the sake of brevity, clarity, and ease of understanding, most of the discussion throughout this document pertains to GPGPU compute shading and 3D rendering. Further, throughout this document, GPGPU compute shader, which refers to a shader stage, may be interchangeably referenced as “patch phase”; similarly, 3D rendering, which refers to a 3D computer rendering process of converting models into images, may be referenced as “submit phase”.

Embodiments provide for combining or batching of multiple back-to-back API calls to minimize the switching between compute shading and 3D rendering such that, in one embodiment, patch phases relating to compute shading are batched together to be performed first using a second-level command buffer, followed by the execution of submit phase relating to 3D rendering. This novel technique provides for a much better utilization of GPU resources, such as enabling GPU 1514 to execute multiple patch phases in parallel, fully utilizing the total number of threads available for use while minimizing the number of transitions between patch phase processes, such as compute shading, and submit phase processes, such as 3D rendering.

Referring now to commands mechanism 1510, in one embodiment, reception/detection logic 1601 may be used to detect or receive an API call (e.g., API ExecuteIndirect) to then allow patch phase logic 1603 to batch together various patch phase processes, such as compute shader processes, and facilitate generation of a second-level command buffer that is separate from the first-level general command buffer. This second-level command buffer may then be used to execute patch phase processes as facilitated by application/execution logic 1607. Upon completion of the patch phase processes, any submit phase processes may then be batched together by submit phase logic 1605 and executed at the first-level command buffer as facilitated by application/execution logic 1607. In one embodiment, any subsequent processing of the results of the patch and submit phases processes may be performed using a third-level command buffer.

Unlike the conventional techniques of **Figures 13-14**, in one embodiment, as illustrated with reference to **Figures 17A, 17B, 17C**, patch phase logic 1603 may be used to propose and generate a separate second-level command buffer (also referred to as “patch command buffer” or

simply “patch buffer”), such as patch buffer 1733 of **Figures 17A-17C**, to process patch phase processes as opposed to processing them at first-level general command buffer 1731. The patch buffer may be kept open for any subsequent API calls and continues to process the corresponding patch phase processes (e.g., compute shading) in parallel. Upon completing patch phase processes for each API call, any submit phase processes (e.g., 3D rendering) are then executed by issuing a relevant submit command, such as execute submit batch” at the first-level command buffer. Upon execution of the phase and submit processes, any results obtained from these processes may then be processed at a third-level command buffer, such as command buffers 1735A-B of **Figures 17A-17C**.

For example and in one embodiment, as will be further described with reference to **Figures 17A-17C**, upon detecting a first API call by reception/detection logic 1601, patch phase logic 1603 creates and appends a second-level patch buffer which may then be use for processing of patch phase commands as facilitated by application/execution logic 1605, but do not close the second-level patch buffer with the GPU’s batch buffer end command. In other words, in one embodiment, this second-level patch buffer remains open for any subsequent API calls, as illustrated with respect to **Figure 17B**, to continue to process, in parallel, any subsequent patch phase commands. Then, as aforementioned, upon completion of the patch phase processing, submit phase processing is triggered by submit phase logic 1605 and executed at the first-level command buffer as facilitated by application/execution logic 1605. Any further processing then commences and continues at third-level command buffers, such as command buffers 1735A-B of **Figures 17A-17C**.

This novel technique of batching up of patch phase processes to be executed using a second-level patch buffer, separate from a first-level general command buffer, and the subsequent execution of submit phase process without having to switch between patch and submit processes, such as between compute shading and 3D rendering, respectively. Further, it is contemplated that in some embodiments, the GPU resources may need to be relinquished so they may be used for other relatively more important and/or urgent tasks by GPU 1514. For example, a Barrier API call may be issued to request the GPU resources for other tasks in which case, in one embodiment, any processing relating to the patch and/or submit phases may be terminated and the GPU resources may be relinquished as will be further described with reference to **Figures 17C and 18B**. A Barrier API call may be referred to as a call that indicates the transition of GPU memory, such as between read and write; typically, used for GPU cache coherency.

This allows for optimization of applications and efficient utilization of GPU resources as provided in the following API usage-models (provided as examples and without limitation): 1)

Barrier A, B, C to UAV; Dispatch to Buffer A; Dispatch to Buffer B; Dispatch to Buffer C; Barrier A, B, C to ArgBuffer; ExecuteIndirect from Buffer A; ExecuteIndirect from Buffer B; ExecuteIndirect from Buffer C; and 2) Barrier A, B, C to UAV; Dispatch to Buffer A; Dispatch to Buffer B; Dispatch to Buffer C; Barrier A, B, C to ArgBuffer; ExecuteIndirect from Buffer A; ExecuteIndirect from Buffer B; and ExecuteIndirect from Buffer C. As aforementioned, “ExecuteIndirect” is merely used as an example for brevity, clarity, and ease of understand and that it is contemplated and to be noted that embodiments are not limited to ExecuteIndirect and/or the like.

In one embodiment, any data relating to patch phases, submit phases, etc., may be stored at database 1630 for future reference and processing. As illustrated, computing device 1500 may be in communication with one or more repositories or databases, such as database 1630 having Cloud-based database, non-Cloud-based databases, etc., to store and maintain any amount and type of data (e.g., real-time sensory input data, historical contents, metadata, resources, policies, criteria, rules and regulations, upgrades, etc.). Similarly, as aforementioned, computing device 1300 may be in communication with any number and type of other computing devices over a communication medium, such as one or more networks including (without limitation) Cloud network, the Internet, intranet, Internet of Things (“IoT”), proximity network, and Bluetooth, etc. It is contemplated that embodiments are not limited to any particular number or type of communication medium or networks.

Communication/compatibility logic 1609 may be used to facilitate dynamic communication and compatibility between computing device 1500 and any number and type of other computing devices (such as mobile computing device, desktop computer, server computing device, etc.), processing devices (such as CPUs, GPUs, etc.), capturing/sensing/detecting devices (such as capturing/sensing components including cameras, depth sensing cameras, camera sensors, RGB sensors, microphones, etc.), display devices (such as output components including display screens, display areas, display projectors, etc.), user/context-awareness components and/or identification/verification sensors/devices (such as biometric sensors/detectors, scanners, etc.), memory or storage devices, databases, and/or data sources (such as data storage devices, hard drives, solid-state drives, hard disks, memory cards or devices, memory circuits, etc.), communication channels or networks (e.g., Cloud network, the Internet, intranet, cellular network, proximity networks, such as Bluetooth, Bluetooth low energy (BLE), Bluetooth Smart, Wi-Fi proximity, Radio Frequency Identification (RFID), Near Field Communication (NFC), Body Area Network (BAN), etc.), wireless or wired communications and relevant protocols (e.g., Wi-Fi®, WiMAX, Ethernet, etc.), connectivity and location management techniques, software applications/websites, (e.g., social and/or business networking websites, etc., business

applications, games and other entertainment applications, etc.), programming languages, etc., while ensuring compatibility with changing technologies, parameters, protocols, standards, etc.

Throughout this document, terms like "logic", "component", "module", "framework", "engine", and the like, may be referenced interchangeably and include, by way of example,

5 software, hardware, and/or any combination of software and hardware, such as firmware.

Further, any use of a particular brand, word, term, phrase, name, and/or acronym, such as

"GPU", "GPU domain", "GPGPU", "CPU", "CPU domain", "workload", "application", "work unit", "draw", "dispatch", "API call", "ExecuteIndirect", "command stream", "patch phase",

"submit phase", "command buffer", "GPGPU", "compute shader", "API barrier", "3D

10 rendering", etc., should not be read to limit embodiments to software or devices that carry that label in products or in literature external to this document.

It is contemplated that any number and type of components may be added to and/or

removed from commands mechanism 1510 to facilitate various embodiments including adding, removing, and/or enhancing certain features. For brevity, clarity, and ease of understanding of

15 commands mechanism 1510, many of the standard and/or known components, such as those of a computing device, are not shown or discussed here. It is contemplated that embodiments, as described herein, are not limited to any particular technology, topology, system, architecture, and/or standard and are dynamic enough to adopt and adapt to any future changes.

Figure 17A illustrates a transactional sequence 1700 for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment.

Transactional sequence 1700 may be performed by processing logic that may comprise hardware (e.g., circuitry, dedicated logic, programmable logic, etc.), software (such as instructions run on a processing device), or a combination thereof. In one embodiment, transactional sequence 1700 may be performed by commands mechanism 1510 of **Figures 15-17**. The processes of

25 transactional sequence 1700 are illustrated in linear sequences for brevity and clarity in presentation; however, it is contemplated that any number of them can be performed in parallel, asynchronously, or in different orders. For brevity, many of the details discussed with reference to the preceding **Figures 1-16** may not be discussed or repeated hereafter.

As previously described with reference to **Figures 15-16** and illustrated here, in one

30 embodiment, upon detecting or receiving API call A 1701A (e.g., ExecuteIndirect API) issued by an application running at computing device, a patch buffer creation command, such as

Execute Patch Batch 1703, is initiated at first-level general GPU command buffer 1731 to created a new second-level GPU command buffer, such as patch buffer 1733, to batch and

execute, in parallel, various patch phase processes (such as compute shading), at patch buffer

35 1733 by executing a relevant command, such as Enable GPGPU Mode 1705. This second-level

patch buffer 1733 may be appended to first level command buffer 1731 used for execution of various patch phase processes, leading to initiation and execution of a computer shader command, such as Execute GPGPU Compute Shader A 1707A, to generate GPU compute shader 1711A corresponding to API call A 1701A and its relevant patch phase processes.

5 In one embodiment, upon execution of various patch phase processes using second-level patch buffer 1733, the phase of processing, such as submit phase, is initiated by API call A 1701A by initiating and executing a relevant command, such as Enable 3D Mode 1721, at first-level command buffer 1731. This enablement of 3D mode leads to API call A 1701A issuing a submit command, such as Execute Submit Batch A 1723A, to batch and execute any submit
10 phase processes relating to the previously executed patch phase processes relating to API call A 1701A. In one embodiment, Execute Submit Batch A 1723A further facilitates third-level command buffer 1735A, where, as illustrated, compute shader 1711A reads from parameters 1709A (e.g., graphics parameters) and writes to third-level command buffer 1735A.

As previously described, in one embodiment, this novel technique prevents continuously
15 switching between patch phase processes and submit phase processes, such as between compute shading and 3D rendering, respectively, by allowing the novel second-level patch buffer 1733 for facilitating executing of any batched patch phase processes before executing any subsequent submit phase processes. Further, in one embodiment and as illustrated with reference to **Figure 17B**, patch buffer 1733 does not close with the end of processing relating to API call A 1701A
20 and that it remains open for any further processing of patch and submit phase transactions relating to any subsequent API calls, such as API call B 1701B of **Figure 17B**.

Figure 17B illustrates a transactional sequence 1750 for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment.

Transactional sequence 1750 may be performed by processing logic that may comprise hardware
25 (e.g., circuitry, dedicated logic, programmable logic, etc.), software (such as instructions run on a processing device), or a combination thereof. In one embodiment, transactional sequence 1750 may be performed by commands mechanism 1510 of **Figures 15-17**. The processes of transactional sequence 1750 are illustrated in linear sequences for brevity and clarity in presentation; however, it is contemplated that any number of them can be performed in parallel,
30 asynchronously, or in different orders. For brevity, many of the details discussed with reference to the preceding **Figures 1-17A** may not be discussed or repeated hereafter.

Continuing with transactional sequence 1700 of **Figure 17A**, transaction sequence 1750 of **Figure 17B** begins with another API call B 1701B (e.g., ExecuteIndirect API) being detected or received at first-level general command buffer 1731. As previously mentioned, upon completion
35 of phase and submit phase transactions relating to API call A 1701 of **Figure 17A**, in one

embodiment, the previously-created second-level patch buffer 1733 remains open for any subsequent API calls, such as API call B 1701B. Accordingly, since second-level patch buffer 1733 remains open, in one embodiment, upon issuance of API call B 1701B by the application running at the computing device, a corresponding compute shader command, such as Execute
 5 GPGPU Compute Shader B 1707B, is issued and execute at patch buffer 1733 to process and execute various batched patch phase transactions relating to API call B 1701B. Upon execution of Execute GPGPU Compute Shader B 1707B, another GPU computer shader 1711B is generated for processing of compute shader-relating transactions.

Further, upon completion of all patch phase transactions, API call B 1701B may then issue
 10 a submit-related command, such as Execute Submit Batch B 1723B, at first-level command buffer 1731, leading to processing of batched submit phase transactions. This Execute Submit Batch B 1723B further leads to generating of third-level command buffer 1735B that is appended to first-level command buffer 1731 for further processing, where, as illustrated, compute shader 1711B reads from parameters 1709B (e.g., graphics parameters) and writes to
 15 third-level command buffer 1735B.

It is contemplated that parameters 1709A-B may include any number and type of parameters relating to compute shading as it relates to their corresponding API calls 1701A-B. For example, with relating to ExecuteIndirect API, such parameters may include any number and type of standard parameters, such as (without limitation) *pCommandSignature*,
 20 *MaxCommandCount*, *pArgumentBuffer*, etc., and/or other custom parameters specific to, for example, patch buffer 1733, patch phase transactions, submit phase transactions, etc. Similarly, GPU/GPGPU compute shading may refer to a shader stage, such as for computing arbitrary information, while 3D rendering may refer to 3D computer graphics processing for converting 3D wire frame models into 2D images with 3D photorealistic effects, etc., for rendering on
 25 computing devices. However, it is to be noted that embodiments are not limited to any particular amount or type of parameters or any particular number or type of patch and/or submit phase transactions, and/or the like.

Figure 17C illustrates a transactional sequence 1770 for processing barrier APIs for efficient utilization of GPU resources according to one embodiment. Transactional sequence
 30 1770 may be performed by processing logic that may comprise hardware (e.g., circuitry, dedicated logic, programmable logic, etc.), software (such as instructions run on a processing device), or a combination thereof. In one embodiment, transactional sequence 1770 may be performed by commands mechanism 1510 of **Figures 15-17**. The processes of transactional sequence 1770 are illustrated in linear sequences for brevity and clarity in presentation; however,
 35 it is contemplated that any number of them can be performed in parallel, asynchronously, or in

different orders. For brevity, many of the details discussed with reference to the preceding **Figures 1-17B** may not be discussed or repeated hereafter.

It is contemplated that GPU resources may be used for any number and type of processes including certain occasions, when other more important and/or urgent processes need to be performed right away, requiring the resources (e.g., threads) to be relinquished for those other important and/or urgent processes. For example and in one embodiment, if any of the resources are to be changed from the indirect argument, such as **Figures 17A-17B**, to any other usage for the computing device, API barrier call 1741 may be issued by the application and detected by or received at first-level command buffer 1731 which then, in response, flushes the cache by issuing a relevant command, such as GPU cache flush 1743, as facilitated by application/execution logic 1607 of **Figure 16**. In response, in one embodiment, patch phase logic 1603 recommends to application/execution logic 1607 of **Figure 16** to signal the GPU driver to close out any existing second-level patch buffers, such as patch buffer 1733, by issuing a relevant command, such as batch buffer end 1745. Upon shutting down of patch buffer 1733, the GPU resources may be relinquished and used for other purposes.

Figure 18A illustrates a method 1800 for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment. Method 1800 may be performed by processing logic that may comprise hardware (e.g., circuitry, dedicated logic, programmable logic, etc.), software (such as instructions run on a processing device), or a combination thereof. In one embodiment, method 1800 may be performed by commands mechanism 1510 of **Figures 15-17**. The processes of method 1800 are illustrated in linear sequences for brevity and clarity in presentation; however, it is contemplated that any number of them can be performed in parallel, asynchronously, or in different orders. For brevity, many of the details discussed with reference to the preceding **Figures 1-17C** may not be discussed or repeated hereafter.

Method 1800 begins with detecting or receiving of an API call (e.g., ExecuteIndirect API) at a first-level command buffer at block 1801. At block 1803, a determination is made as to whether a second-level patch buffer is to be created, such as whether count = 0, to batch and process patch phase transactions. If yes, such as if count = 0, as illustrated with reference to **Figures 17A-17C**, the second-level patch buffer is created and appended with the first level command buffer at block 1805. At block 1807, a command, such as Enable GPGPU Mode, is issued to append the second-level patch buffer with the first-level command buffer and to initiate batching and processing of any patch phase transactions (e.g., GPGPU compute shading) using the second-level patch buffer.

In one embodiment, upon completion of the processing of patch phase transactions, a

second phase, such as submit phase, is initiated with another command, such as Enable 3D Mode, issued by the API call and executed at the first-level command buffer. Subsequently, or, referring back to block 1803, if count does not equal 0, patch phase transactions are appended and submit phase mode is initiated at block 1811. At block 1813, submit phase transactions (e.g., 3D rendering) are processed and subsequently, method 1800 ends at block 1815 with resetting of the count.

Figure 18B illustrates a method 1830 for processing barrier APIs for efficient utilization of GPU resources according to one embodiment. Method 1830 may be performed by processing logic that may comprise hardware (e.g., circuitry, dedicated logic, programmable logic, etc.), software (such as instructions run on a processing device), or a combination thereof. In one embodiment, method 1830 may be performed by commands mechanism 1510 of **Figures 15-17**. The processes of method 1830 are illustrated in linear sequences for brevity and clarity in presentation; however, it is contemplated that any number of them can be performed in parallel, asynchronously, or in different orders. For brevity, many of the details discussed with reference to the preceding **Figures 1-18A** may not be discussed or repeated hereafter.

Method 1830 begins at block 1831 with issuing of a barrier API by an application running at a computing device. As previously described with reference to **Figure 17C**, GPU resources may be used for other purposes as determined by the application and in that case, a barrier API may be issued to relinquish the resources being currently used by patch and submit phases so that they may be used and processed for other more immediate and/or important transactions. Subsequently, at block 1833, in one embodiment, GPU cache is flushed and, at block 1835, a command, such as Batch Buffer End, to terminate or close out the second-level patch buffer is issued and executed. At block 1837, the second-level patch buffer along with any patch and/or submit phase transactions is terminated and the count is reset.

Figure 18C illustrates a method 1850 for executing patch and submit phase transactions for efficient utilization of GPU resources according to one embodiment. Method 1850 may be performed by processing logic that may comprise hardware (e.g., circuitry, dedicated logic, programmable logic, etc.), software (such as instructions run on a processing device), or a combination thereof. In one embodiment, method 1850 may be performed by commands mechanism 1510 of **Figures 15-17**. The processes of method 1850 are illustrated in linear sequences for brevity and clarity in presentation; however, it is contemplated that any number of them can be performed in parallel, asynchronously, or in different orders. For brevity, many of the details discussed with reference to the preceding **Figures 1-18B** may not be discussed or repeated hereafter.

Method 1850 begins at block 1801 with receiving, at a first-level command buffer, an API

call (e.g., ExecuteIndirect API) from an application running at a computing device. At block 1803, a second-level patch buffer is created and appended to the first-level command buffer. At block 1805, various patch phase transactions (e.g., compute shading) are batched together and processed via the second-level patch buffer and any other components, such as GPU compute
5 shader, etc. At block 1807, upon completion of the patch phase transactions, a submit phase initiation command is executed to initiate processing of submit phase transactions.

At block 1809, submit phase transactions are processed and executed. At block 1811, a determination is made as to whether another API call or a barrier API call is received. If another API call is received, method 1850 continues with the process of block 1805, where the same
10 patch buffer may be used to process any new patch phase transactions relating to the new API call. If, however, the call received is a barrier API call, at block 1813, all processing of patch and/or submit phase transactions is terminated, the second-level patch buffer is closed out, GPU cache is flushed, etc., so that the resources may be relinquished and used for other more immediate and/or important transactions as determined by the application.

References to “one embodiment”, “an embodiment”, “example embodiment”, “various
15 embodiments”, etc., indicate that the embodiment(s) so described may include particular features, structures, or characteristics, but not every embodiment necessarily includes the particular features, structures, or characteristics. Further, some embodiments may have some, all, or none of the features described for other embodiments.

In the foregoing specification, embodiments have been described with reference to specific
20 exemplary embodiments thereof. It will, however, be evident that various modifications and changes may be made thereto without departing from the broader spirit and scope of embodiments as set forth in the appended claims. The Specification and drawings are, accordingly, to be regarded in an illustrative rather than a restrictive sense.

In the following description and claims, the term “coupled” along with its derivatives, may
25 be used. “Coupled” is used to indicate that two or more elements co-operate or interact with each other, but they may or may not have intervening physical or electrical components between them.

As used in the claims, unless otherwise specified the use of the ordinal adjectives “first”,
30 “second”, “third”, etc., to describe a common element, merely indicate that different instances of like elements are being referred to, and are not intended to imply that the elements so described must be in a given sequence, either temporally, spatially, in ranking, or in any other manner.

The following clauses and/or examples pertain to further embodiments or examples.
Specifics in the examples may be used anywhere in one or more embodiments. The various
35 features of the different embodiments or examples may be variously combined with some

features included and others excluded to suit a variety of different applications. Examples may include subject matter such as a method, means for performing acts of the method, at least one machine-readable medium including instructions that, when performed by a machine cause the machine to performs acts of the method, or of an apparatus or system for facilitating hybrid

5 communication according to embodiments and examples described herein.

Some embodiments pertain to Example 1 that includes an apparatus to facilitate efficient graphics command generation and execution for improved graphics performance at computing devices, comprising: reception/detection logic to detect an application programming interface (API) call to perform a plurality of transactions, wherein the API call is issued by an application
10 at a first command buffer, wherein the plurality of transactions include a first set of transactions and a second set of transactions; patch phase logic to create a second command buffer and append the second command buffer to the first command buffer, wherein the patch phase logic is further to separate the first set transactions from the second set of transactions; and application/execution logic to execute, via the second command buffer, the first set of
15 transactions, prior to executing the first set of transactions.

Example 2 includes the subject matter of Example 1, wherein the patch phase logic is further to issue a patch execution command to facilitate the application/execution logic to execute the first set of transactions, prior to executing the first set of transactions.

Example 3 includes the subject matter of Example 1, further comprising submit phase logic
20 to issue a submit execution command to facilitate the application/execution logic to execute the second set of transactions upon completion of the execution of the first set of transactions.

Example 4 includes the subject matter of Example 1 or 3, wherein the application/execution logic to execute the second set of transactions, upon completion of the execution of the first set of transactions.

Example 5 includes the subject matter of Example 1, wherein the first set of transactions
25 comprise patch transactions relating to one or more patch processes including compute shading of graphics contents, wherein the second set of transactions comprise submit transactions relating to one or more submit processes including three-dimensional (3D) rendering of graphics contents.

Example 6 includes the subject matter of Example 1 or 2, wherein the patch phase logic is further to maintain the second command buffer to remain open for one or more API calls to be subsequently issued by the application, wherein the second command buffer is used for
30 executing additional transactions relating to the one or more API calls.

Example 7 includes the subject matter of Example 1, wherein the reception/detection logic
35 is further to detect a barrier API call, wherein the barrier API call to requisite graphics

processing unit (GPU) resources for other processes.

Example 8 includes the subject matter of Example 7, wherein in response to the barrier API call, the application/execution logic is further to terminate the second command buffer to relinquish the GPU resources for the other processes.

5 Some embodiments pertain to Example 9 that includes a method for facilitating efficient graphics command generation and execution for improved graphics performance at computing devices: detecting an application programming interface (API) call to perform a plurality of transactions, wherein the API call is issued by an application at a first command buffer, wherein the plurality of transactions include a first set of transactions and a second set of transactions;
10 creating a second command buffer and appending the second command buffer to the first command buffer, wherein creating further includes separating the first set transactions from the second set of transactions; and executing, via the second command buffer, the first set of transactions, prior to executing the first set of transactions.

15 Example 10 includes the subject matter of Example 9, further comprising issuing a patch execution command to execute the first set of transactions, prior to executing the first set of transactions.

Example 11 includes the subject matter of Example 9, further comprising issuing a submit execution command to execute the second set of transactions upon completion of the execution of the first set of transactions.

20 Example 12 includes the subject matter of Example 9 or 11, further comprising executing the second set of transactions, upon completion of the execution of the first set of transactions.

Example 13 includes the subject matter of Example 9, wherein the first set of transactions comprise patch transactions relating to one or more patch processes including compute shading of graphics contents, wherein the second set of transactions comprise submit transactions relating
25 to one or more submit processes including three-dimensional (3D) rendering of graphics contents.

Example 14 includes the subject matter of Example 9 or 10, further comprising maintaining the second command buffer to remain open for one or more API calls to be subsequently issued by the application, wherein the second command buffer is used for
30 executing additional transactions relating to the one or more API calls.

Example 15 includes the subject matter of Example 9, further comprising detecting a barrier API call, wherein the barrier API call to requisite graphics processing unit (GPU) resources for other processes.

35 Example 16 includes the subject matter of Example 15, wherein in response to the barrier API call, terminating the second command buffer to relinquish the GPU resources for the other

processes.

Example 17 includes at least one machine-readable medium comprising a plurality of instructions, when executed on a computing device, to implement or perform a method or realize an apparatus as claimed in any preceding claims.

5 Example 18 includes at least one non-transitory or tangible machine-readable medium comprising a plurality of instructions, when executed on a computing device, to implement or perform a method or realize an apparatus as claimed in any preceding claims.

Example 19 includes a system comprising a mechanism to implement or perform a method or realize an apparatus as claimed in any preceding claims.

10 Example 20 includes an apparatus comprising means to perform a method as claimed in any preceding claims.

Example 21 includes a computing device arranged to implement or perform a method or realize an apparatus as claimed in any preceding claims.

15 Example 22 includes a communications device arranged to implement or perform a method or realize an apparatus as claimed in any preceding claims.

Some embodiments pertain to Example 23 includes a system comprising a storage device having instructions, and a processor to execute the instructions to facilitate a mechanism to perform one or more operations comprising: detecting an application programming interface (API) call to perform a plurality of transactions, wherein the API call is issued by an application
20 at a first command buffer, wherein the plurality of transactions include a first set of transactions and a second set of transactions; creating a second command buffer and appending the second command buffer to the first command buffer, wherein creating further includes separating the first set transactions from the second set of transactions; and executing, via the second command buffer, the first set of transactions, prior to executing the first set of transactions.

25 Example 24 includes the subject matter of Example 23, wherein one or more operations further comprise issuing a patch execution command to execute the first set of transactions, prior to executing the first set of transactions.

Example 25 includes the subject matter of Example 23, wherein one or more operations further comprise issuing a submit execution command to execute the second set of transactions
30 upon completion of the execution of the first set of transactions.

Example 26 includes the subject matter of Example 23 or 25, wherein one or more operations further comprise executing the second set of transactions, upon completion of the execution of the first set of transactions.

Example 27 includes the subject matter of Example 23, wherein the first set of transactions
35 comprise patch transactions relating to one or more patch processes including compute shading

of graphics contents, wherein the second set of transactions comprise submit transactions relating to one or more submit processes including three-dimensional (3D) rendering of graphics contents.

Example 28 includes the subject matter of Example 23 or 24, wherein one or more operations further comprise maintaining the second command buffer to remain open for one or more API calls to be subsequently issued by the application, wherein the second command buffer is used for executing additional transactions relating to the one or more API calls.

Example 29 includes the subject matter of Example 23, wherein one or more operations further comprise detecting a barrier API call, wherein the barrier API call to requisite graphics processing unit (GPU) resources for other processes.

Example 30 includes the subject matter of Example 29, wherein in response to the barrier API call, terminating the second command buffer to relinquish the GPU resources for the other processes.

Some embodiments pertain to Example 31 includes an apparatus comprising: means for detecting an application programming interface (API) call to perform a plurality of transactions, wherein the API call is issued by an application at a first command buffer, wherein the plurality of transactions include a first set of transactions and a second set of transactions; means for creating a second command buffer and appending the second command buffer to the first command buffer, wherein creating further includes separating the first set transactions from the second set of transactions; and means for executing, via the second command buffer, the first set of transactions, prior to executing the first set of transactions.

Example 32 includes the subject matter of Example 31, wherein one or more operations further comprise means for issuing a patch execution command to execute the first set of transactions, prior to executing the first set of transactions.

Example 33 includes the subject matter of Example 31, wherein one or more operations further comprise means for issuing a submit execution command to execute the second set of transactions upon completion of the execution of the first set of transactions.

Example 34 includes the subject matter of Example 31 or 33, wherein one or more operations further comprise means for executing the second set of transactions, upon completion of the execution of the first set of transactions.

Example 35 includes the subject matter of Example 31, wherein the first set of transactions comprise patch transactions relating to one or more patch processes including compute shading of graphics contents, wherein the second set of transactions comprise submit transactions relating to one or more submit processes including three-dimensional (3D) rendering of graphics contents.

Example 36 includes the subject matter of Example 31 or 32, wherein one or more operations further comprise maintaining the second command buffer to remain open for one or more API calls to be subsequently issued by the application, wherein the second command buffer is used for executing additional transactions relating to the one or more API calls.

5 Example 37 includes the subject matter of Example 31, wherein one or more operations further comprise detecting a barrier API call, wherein the barrier API call to requisite graphics processing unit (GPU) resources for other processes.

10 Example 38 includes the subject matter of Example 31, wherein in response to the barrier API call, terminating the second command buffer to relinquish the GPU resources for the other processes.

Example 39 includes at least one non-transitory or tangible machine-readable medium comprising a plurality of instructions, when executed on a computing device, to implement or perform a method as claimed in any of claims or examples 9-16.

15 Example 40 includes at least one machine-readable medium comprising a plurality of instructions, when executed on a computing device, to implement or perform a method as claimed in any of claims or examples 9-16.

Example 41 includes a system comprising a mechanism to implement or perform a method as claimed in any of claims or examples 9-16.

20 Example 42 includes an apparatus comprising means for performing a method as claimed in any of claims or examples 9-16.

Example 43 includes a computing device arranged to implement or perform a method as claimed in any of claims or examples 9-16.

Example 44 includes a communications device arranged to implement or perform a method as claimed in any of claims or examples 9-16.

25 The drawings and the forgoing description give examples of embodiments. Those skilled in the art will appreciate that one or more of the described elements may well be combined into a single functional element. Alternatively, certain elements may be split into multiple functional elements. Elements from one embodiment may be added to another embodiment. For example, orders of processes described herein may be changed and are not limited to the manner described
30 herein. Moreover, the actions any flow diagram need not be implemented in the order shown; nor do all of the acts necessarily need to be performed. Also, those acts that are not dependent on other acts may be performed in parallel with the other acts. The scope of embodiments is by no means limited by these specific examples. Numerous variations, whether explicitly given in the specification or not, such as differences in structure, dimension, and use of material, are
35 possible. The scope of embodiments is at least as broad as given by the following claims.

CLAIMS

What is claimed is:

1. An apparatus to facilitate efficient graphics command generation and execution for improved graphics performance, comprising:

5 reception/detection logic to detect an application programming interface (API) call to perform a plurality of transactions, wherein the API call is issued by an application at a first command buffer, wherein the plurality of transactions include a first set of transactions and a second set of transactions;

10 patch phase logic to create a second command buffer and append the second command buffer to the first command buffer, wherein the patch phase logic is further to separate the first set transactions from the second set of transactions; and

application/execution logic to execute, via the second command buffer, the first set of transactions, prior to executing the first set of transactions.

15 2. The apparatus of claim 1, wherein the patch phase logic is further to issue a patch execution command to facilitate the application/execution logic to execute the first set of transactions, prior to executing the first set of transactions.

3. The apparatus of claim 1, further comprising submit phase logic to issue a submit execution command to facilitate the application/execution logic to execute the second set of transactions upon completion of the execution of the first set of transactions.

20 4. The apparatus of claim 1 or 3, wherein the application/execution logic to execute the second set of transactions, upon completion of the execution of the first set of transactions.

25 5. The apparatus of claim 1 or 2, wherein the first set of transactions comprise patch transactions relating to one or more patch processes including compute shading of graphics contents, wherein the second set of transactions comprise submit transactions relating to one or more submit processes including three-dimensional (3D) rendering of graphics contents.

6. The apparatus of claim 1, wherein the patch phase logic is further to maintain the second command buffer to remain open for one or more API calls to be subsequently issued by the application, wherein the second command buffer is used for executing additional transactions relating to the one or more API calls.

30 7. The apparatus of claim 1, wherein the reception/detection logic is further to detect a barrier API call, wherein the barrier API call to requisite graphics processing unit (GPU) resources for other processes.

35 8. The apparatus of claim 7, wherein in response to the barrier API call, the application/execution logic is further to terminate the second command buffer to relinquish the GPU resources for the other processes.

9. A method for facilitating efficient graphics command generation and execution for improved graphics performance, comprising:

detecting an application programming interface (API) call to perform a plurality of transactions, wherein the API call is issued by an application at a first command buffer, wherein
5 the plurality of transactions include a first set of transactions and a second set of transactions;

creating a second command buffer and appending the second command buffer to the first command buffer, wherein creating further includes separating the first set transactions from the second set of transactions; and

10 executing, via the second command buffer, the first set of transactions, prior to executing the first set of transactions.

10. The method of claim 9, further comprising issuing a patch execution command to execute the first set of transactions, prior to executing the first set of transactions.

11. The method of claim 9, further comprising issuing a submit execution command to execute the second set of transactions upon completion of the execution of the first set of
15 transactions.

12. The method of claim 9 or 11, further comprising executing the second set of transactions, upon completion of the execution of the first set of transactions.

13. The method of claim 9 or 10, wherein the first set of transactions comprise patch transactions relating to one or more patch processes including compute shading of graphics
20 contents, wherein the second set of transactions comprise submit transactions relating to one or more submit processes including three-dimensional (3D) rendering of graphics contents.

14. The method of claim 9, further comprising maintaining the second command buffer to remain open for one or more API calls to be subsequently issued by the application, wherein the second command buffer is used for executing additional transactions relating to the
25 one or more API calls.

15. The method of claim 9, further comprising detecting a barrier API call, wherein the barrier API call to requisite graphics processing unit (GPU) resources for other processes.

16. The method of claim 15, wherein in response to the barrier API call, terminating the second command buffer to relinquish the GPU resources for the other processes.

30 17. At least one machine-readable medium comprising a plurality of instructions, when executed on a computing device, to implement or perform a method as claimed in any of claims 9-16.

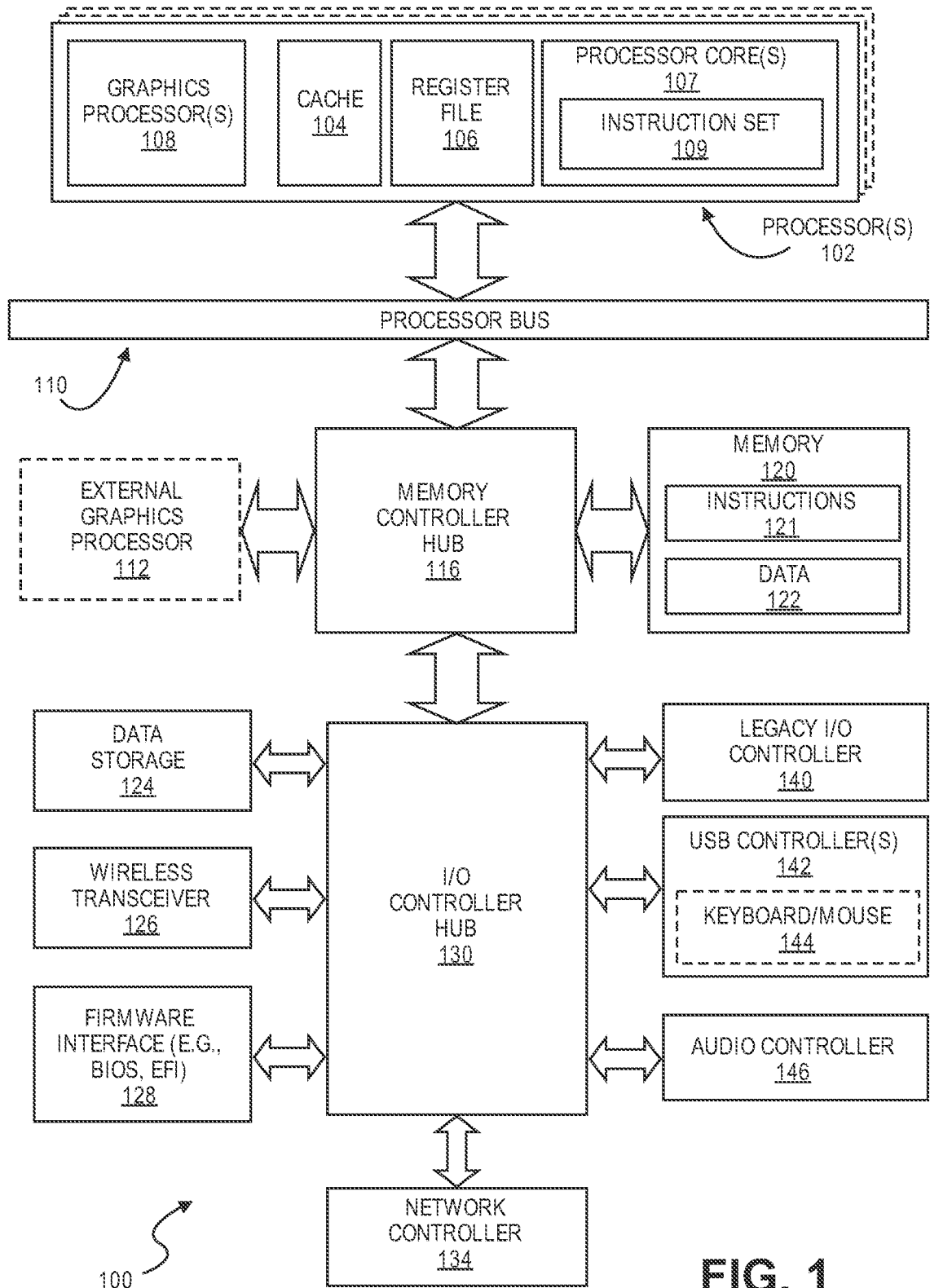
18. A system comprising a mechanism to implement or perform a method as claimed in any of claims 9-16.

19. An apparatus comprising means for performing a method as claimed in any of claims 9-16.

20. A computing device arranged to implement or perform a method as claimed in any of claims 9-16.

5 21. A communications device arranged to implement or perform a method as claimed in any of claims 9-16.

1/21

**FIG. 1**

PROCESSOR
200

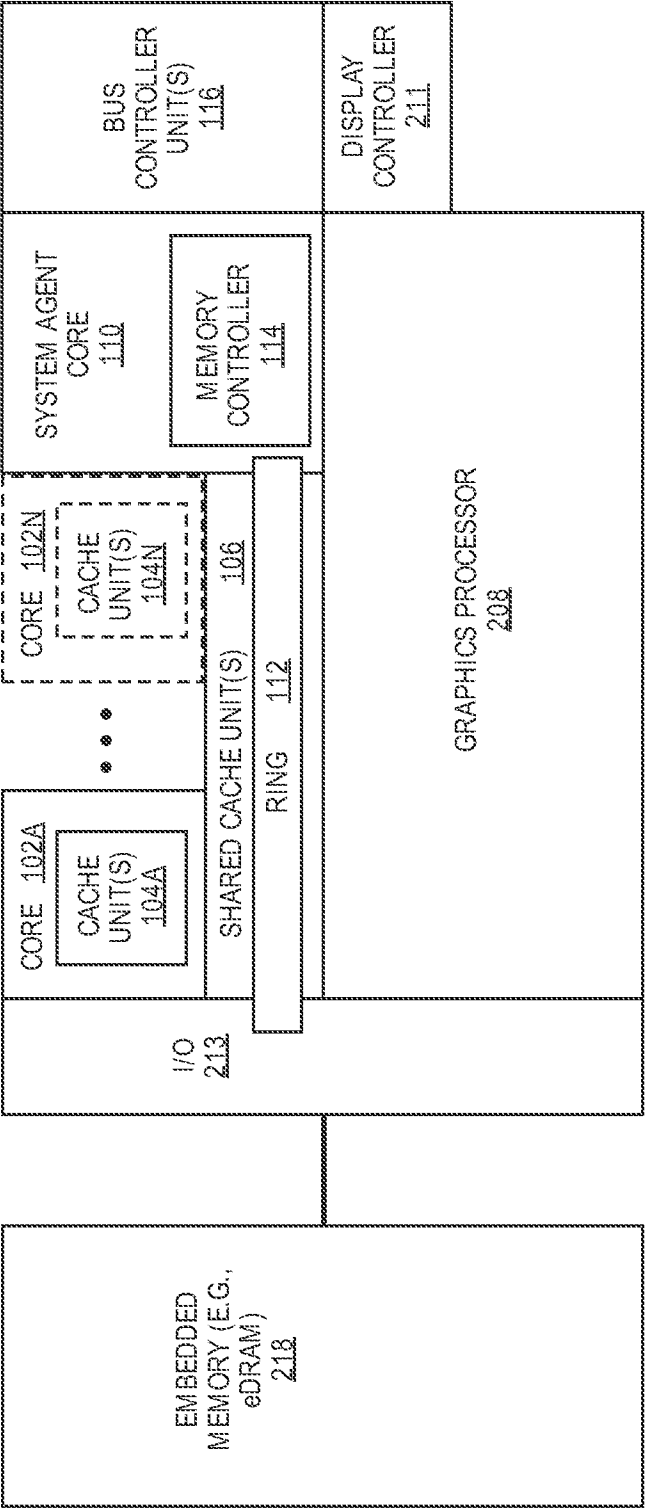
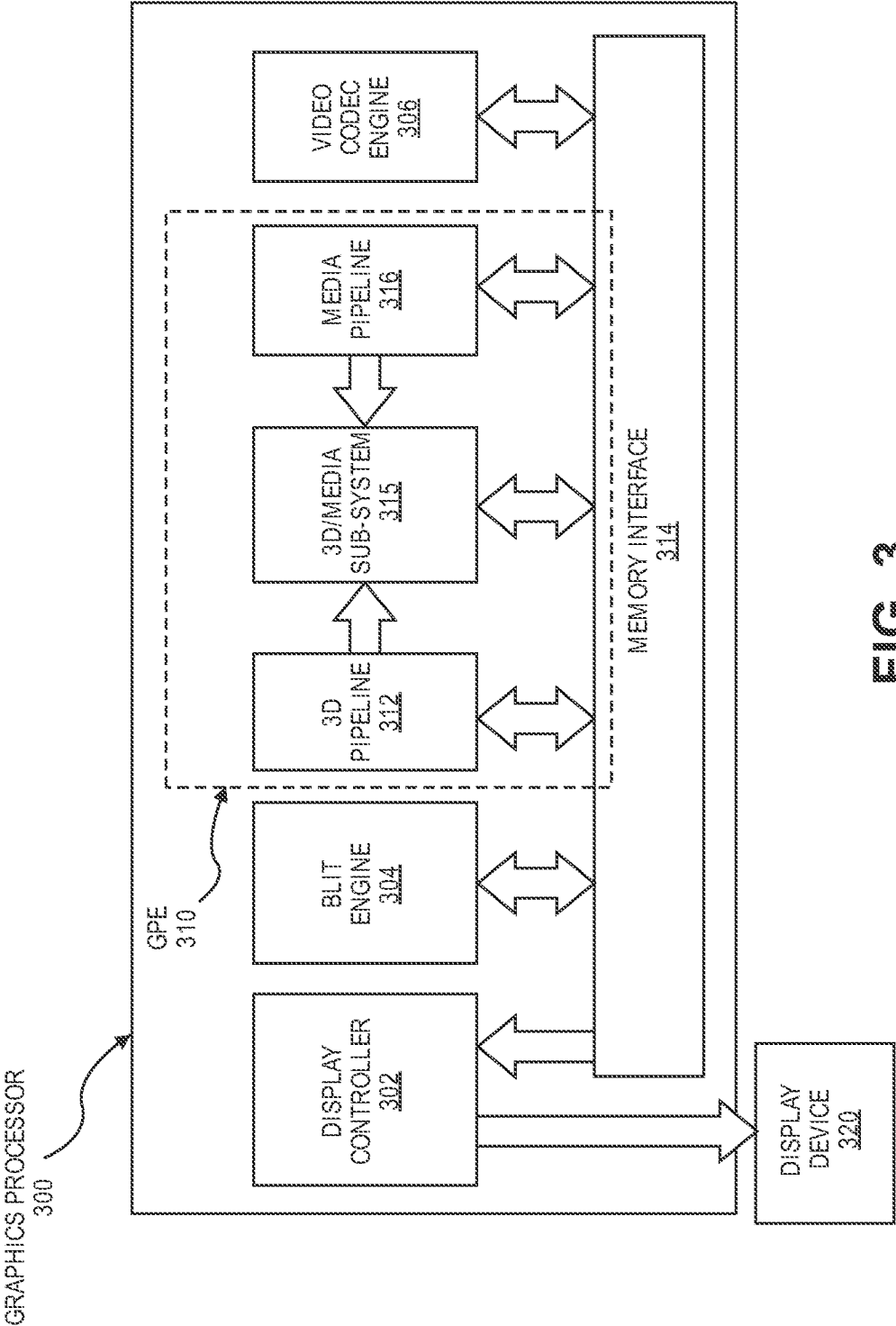


FIG. 2



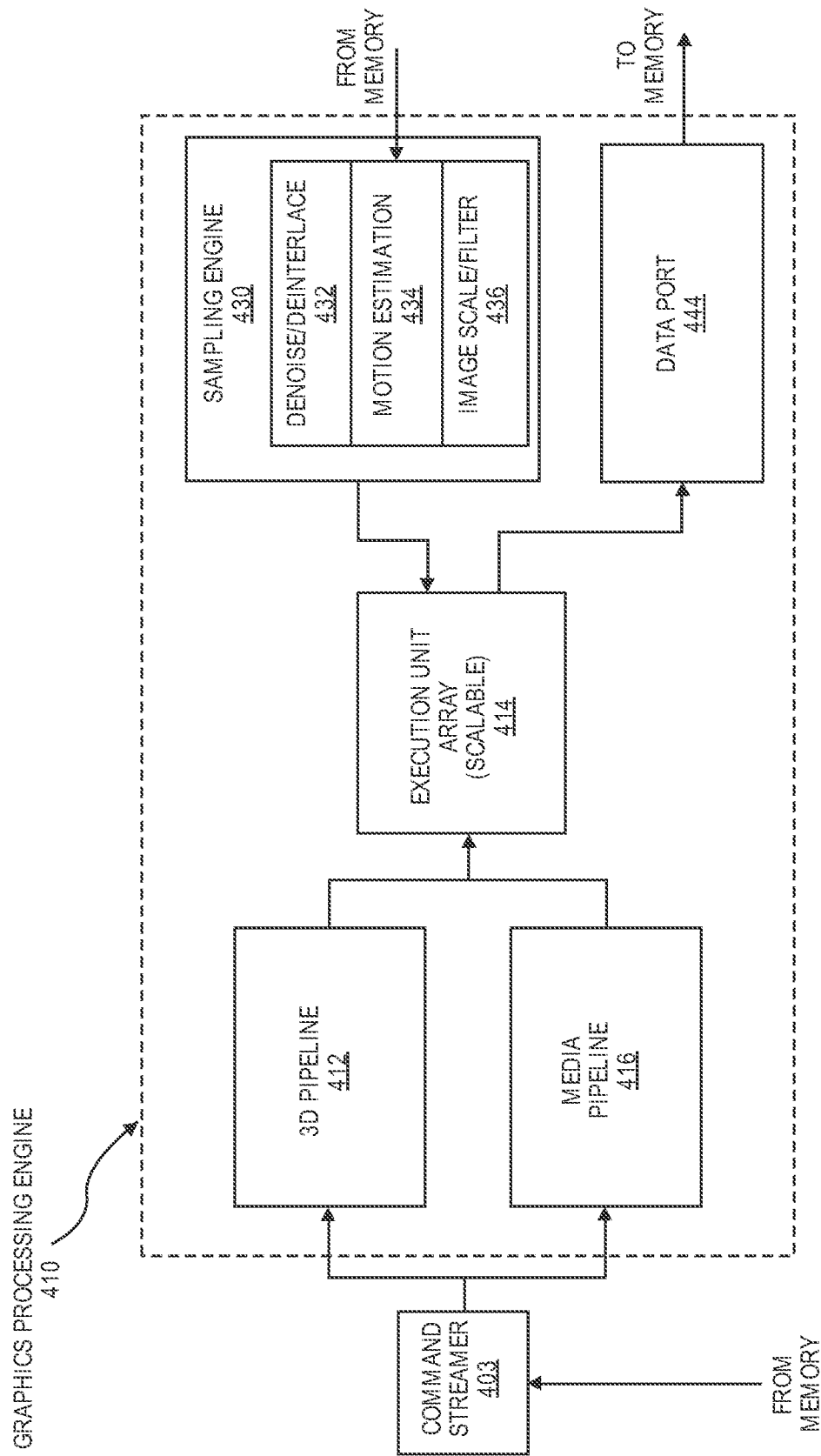


FIG. 4

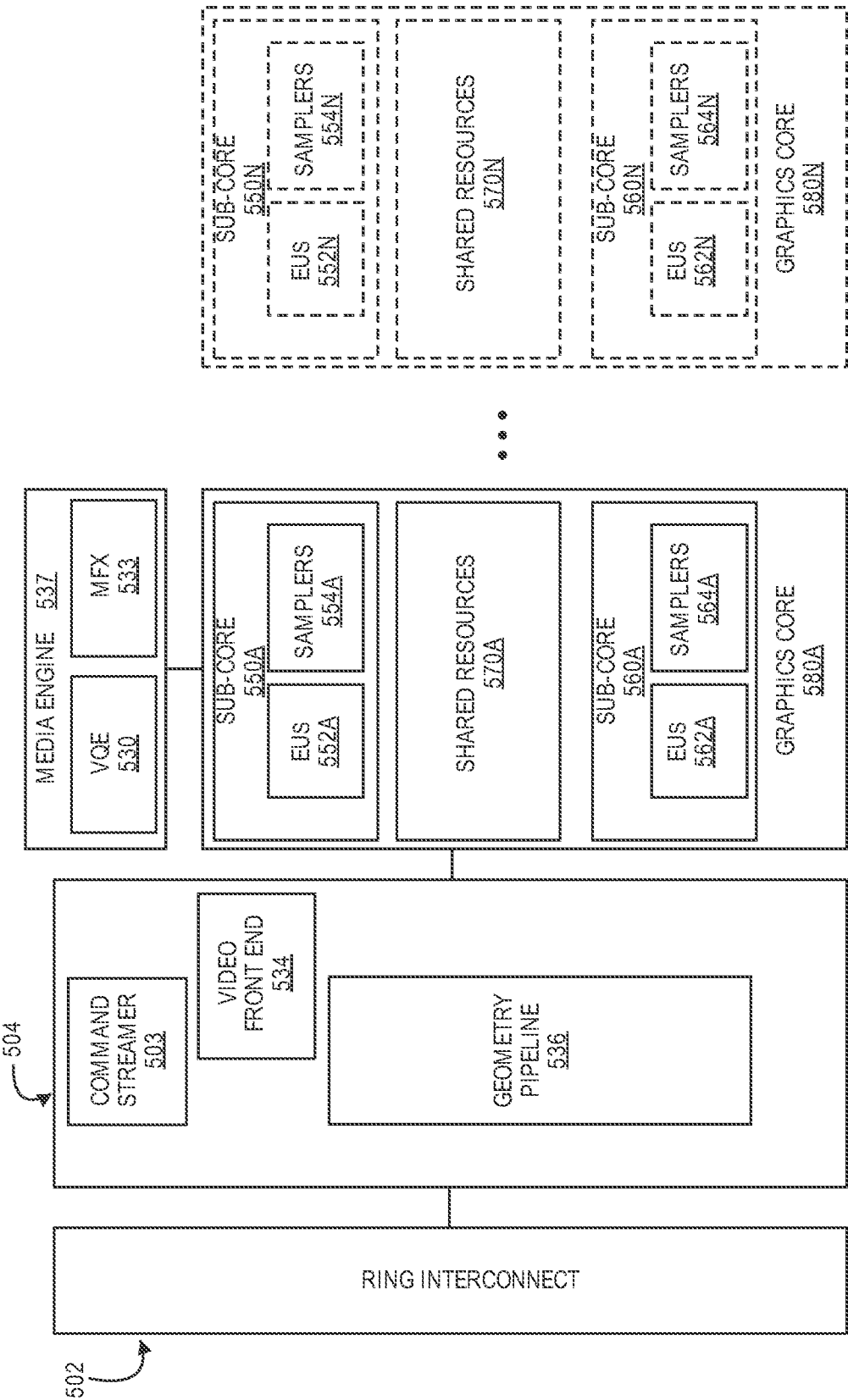


FIG. 5

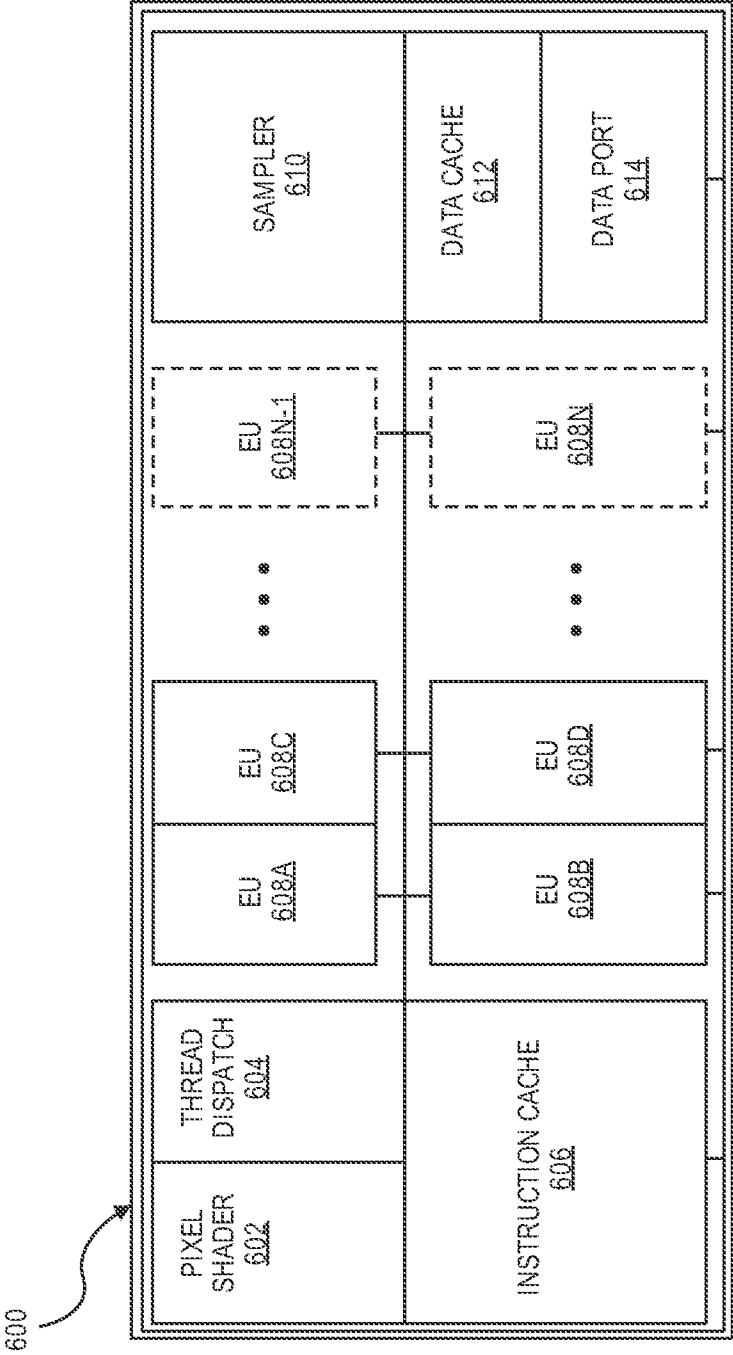
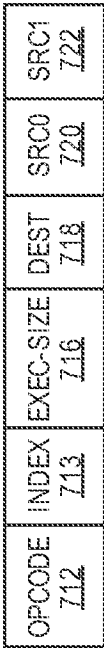
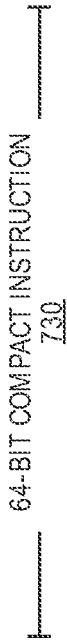
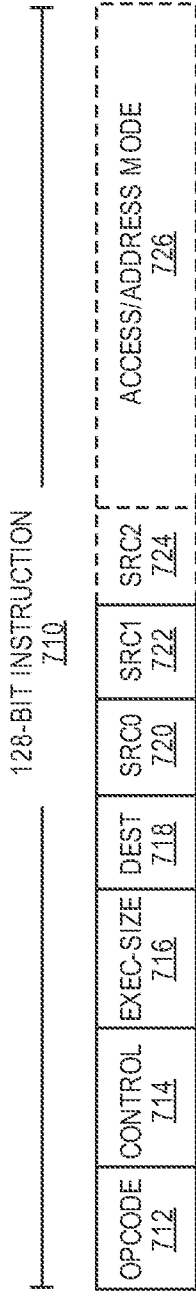


FIG. 6

GRAPHICS CORE INSTRUCTION FORMATS

700



OPCODE DECODE

740

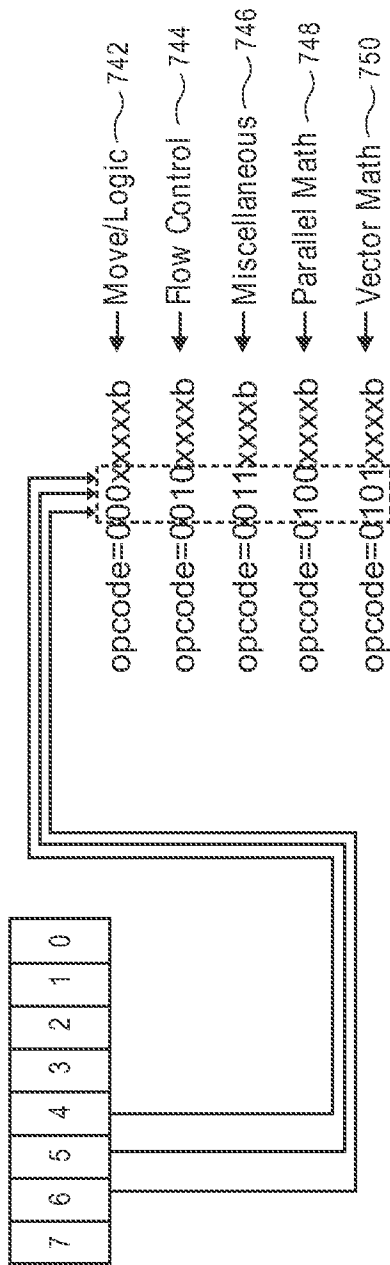


FIG. 7

8/21

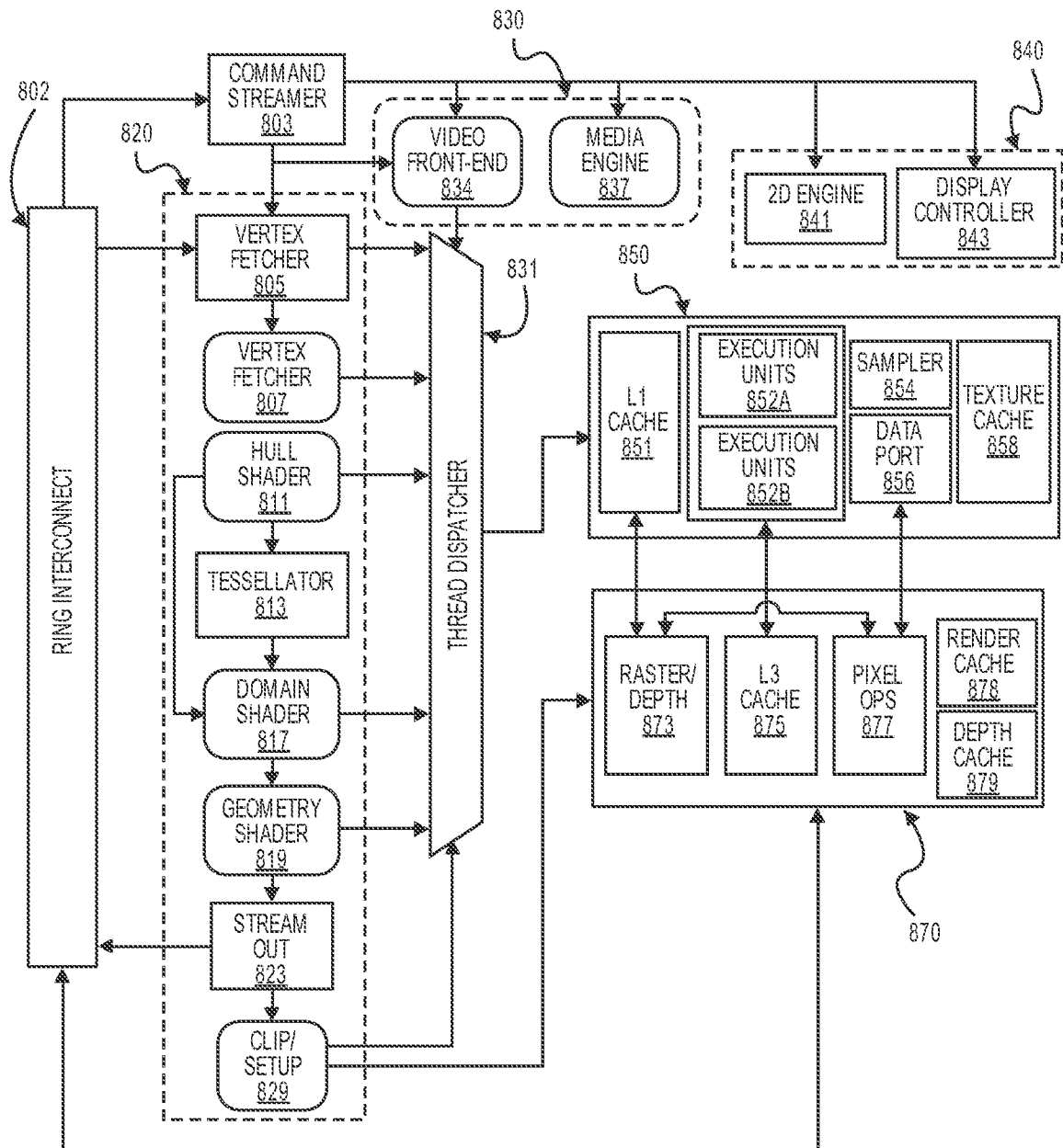


FIG. 8

FIG. 9A

SAMPLE COMMAND FORMAT
900

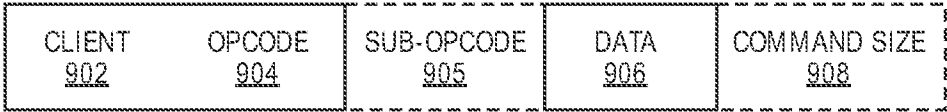
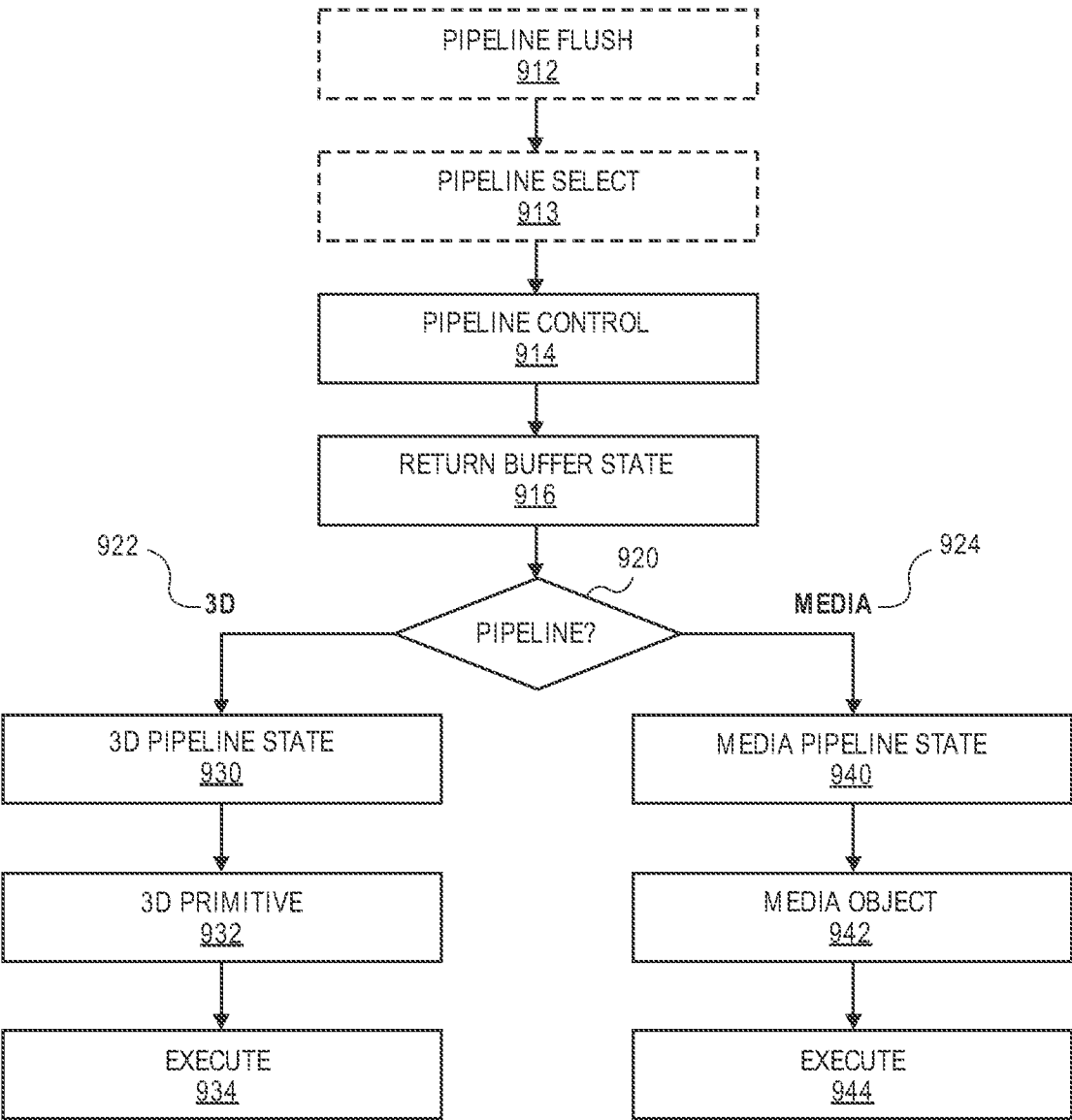
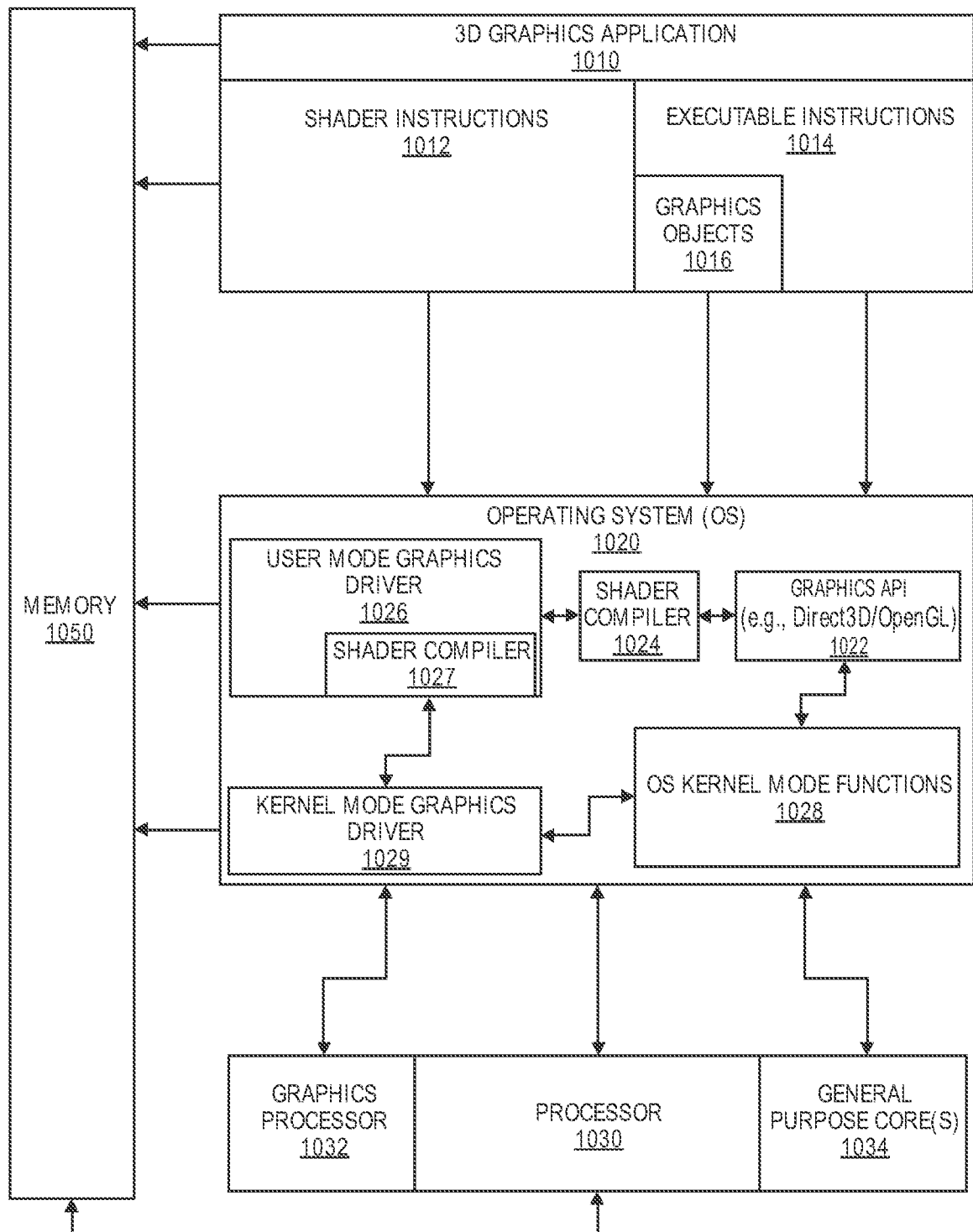


FIG. 9B

SAMPLE COMMAND SEQUENCE
910



10/21

**FIG. 10**

11/21

IP CORE DEVELOPMENT - 1100

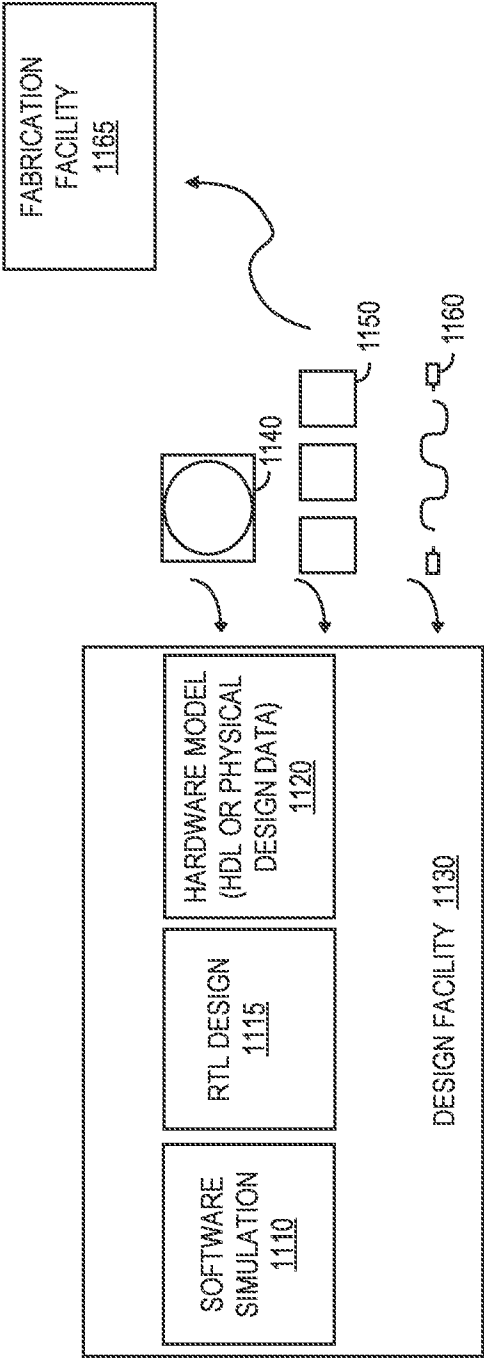
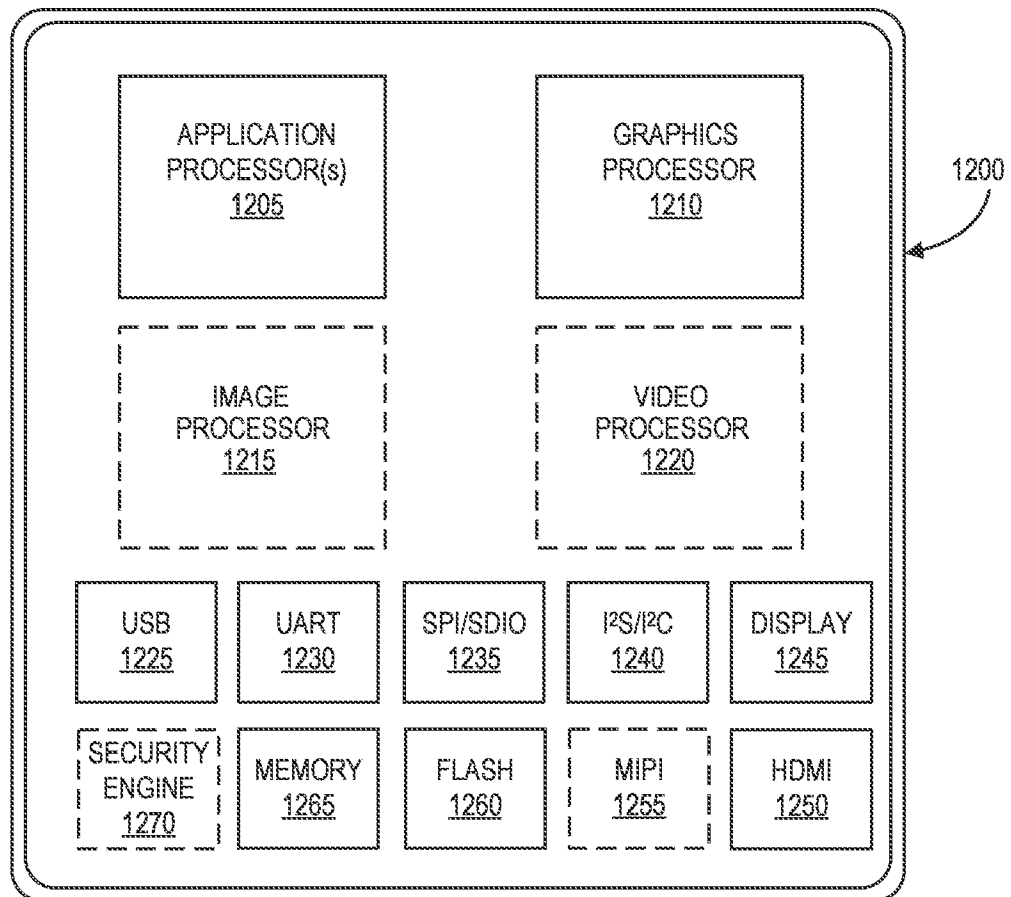


FIG. 11

12/21

**FIG. 12**

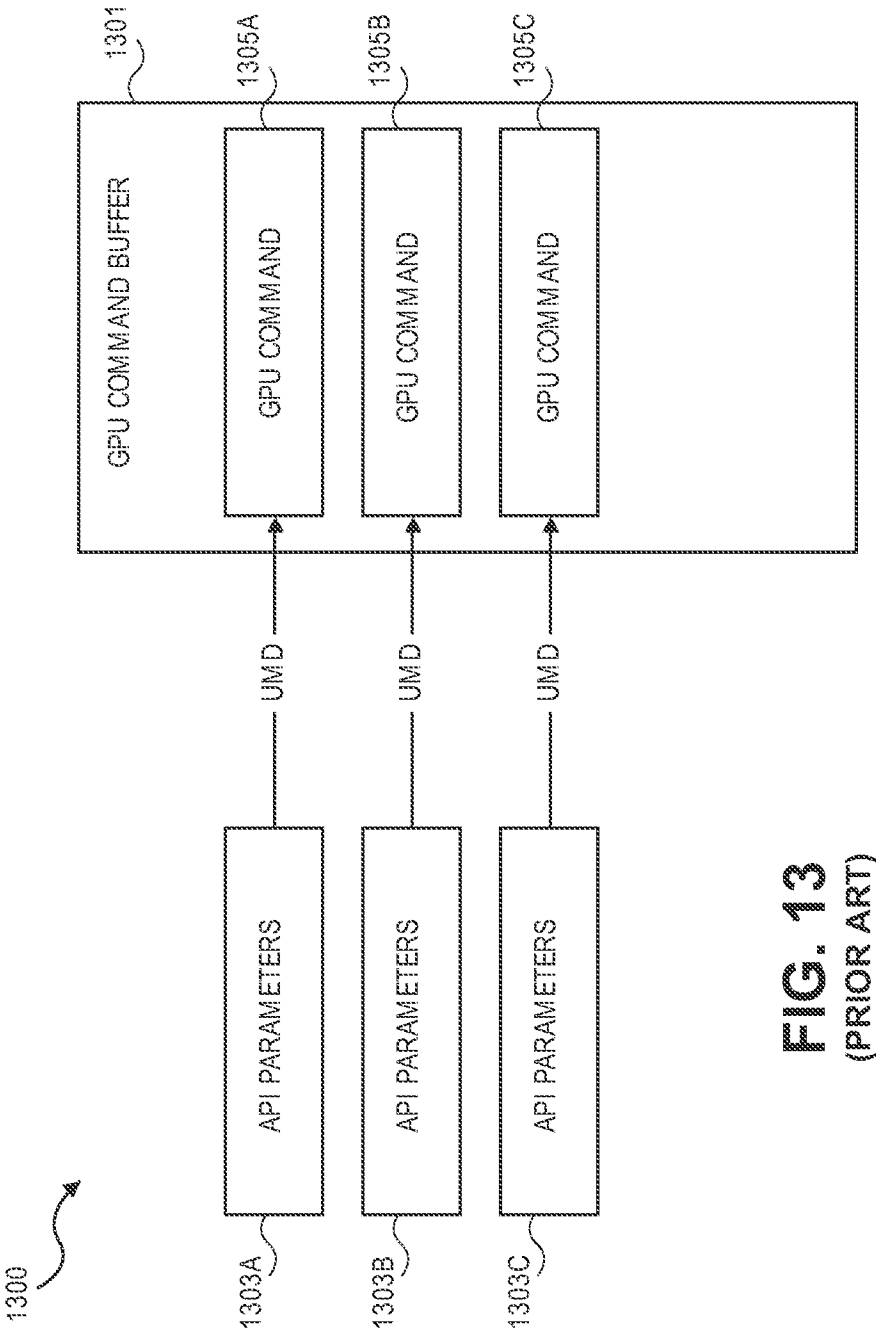


FIG. 13
(PRIOR ART)

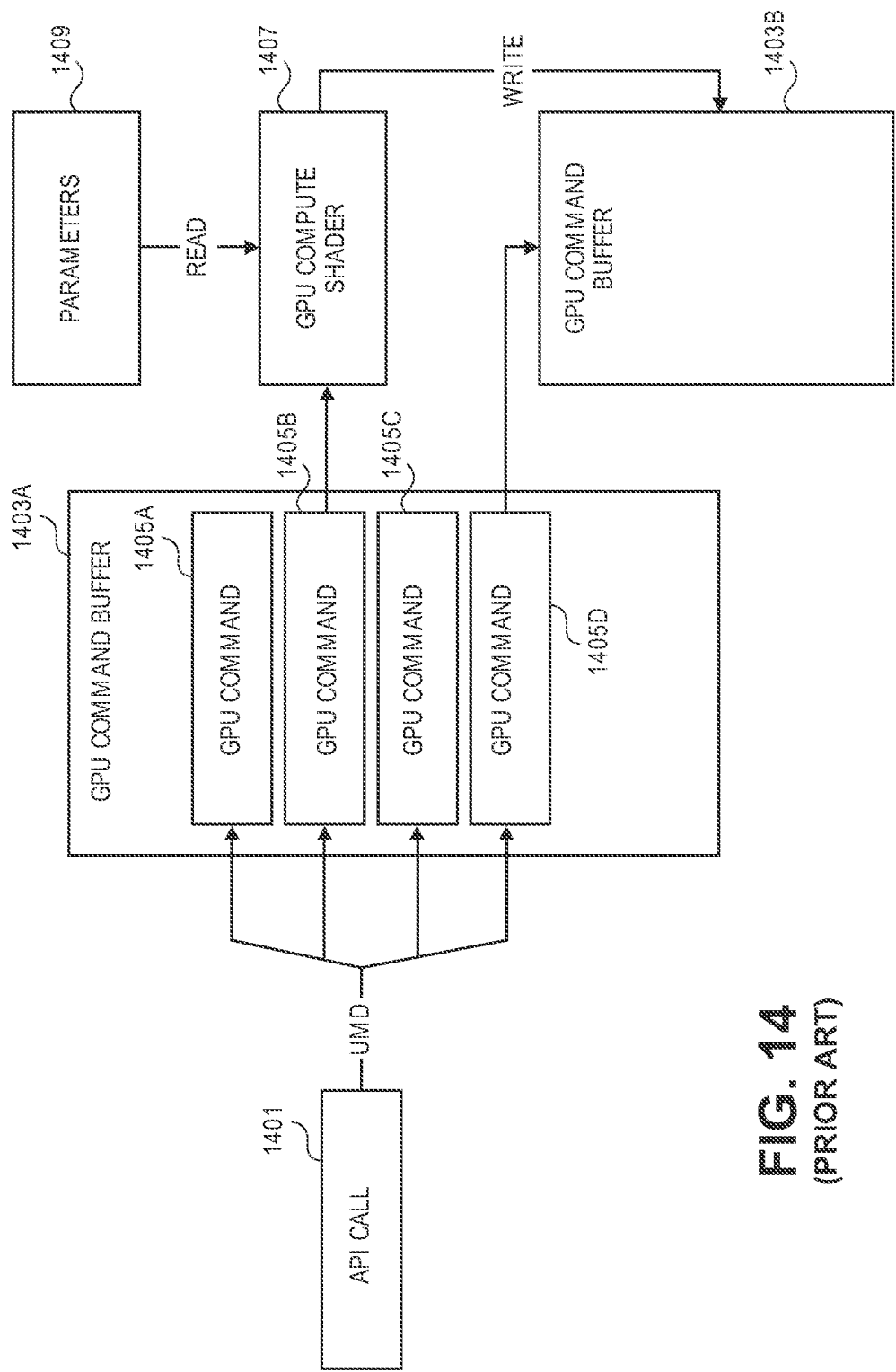


FIG. 14
(PRIOR ART)

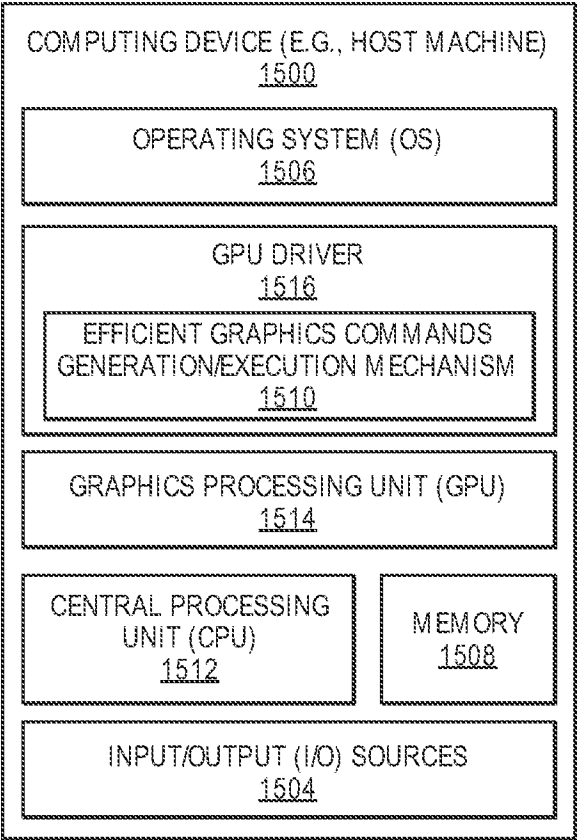
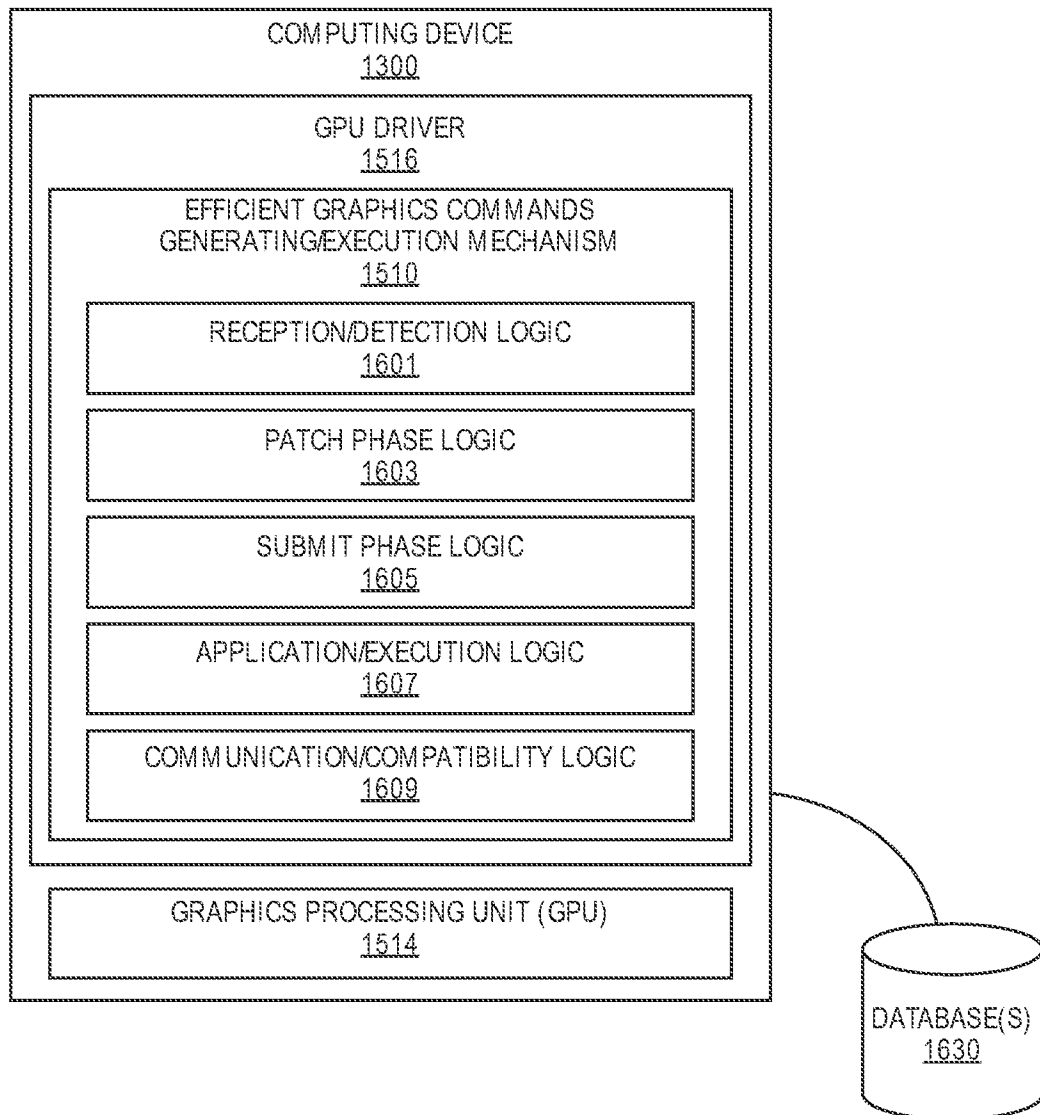


FIG. 15

16/21

**FIG. 16**

17/21

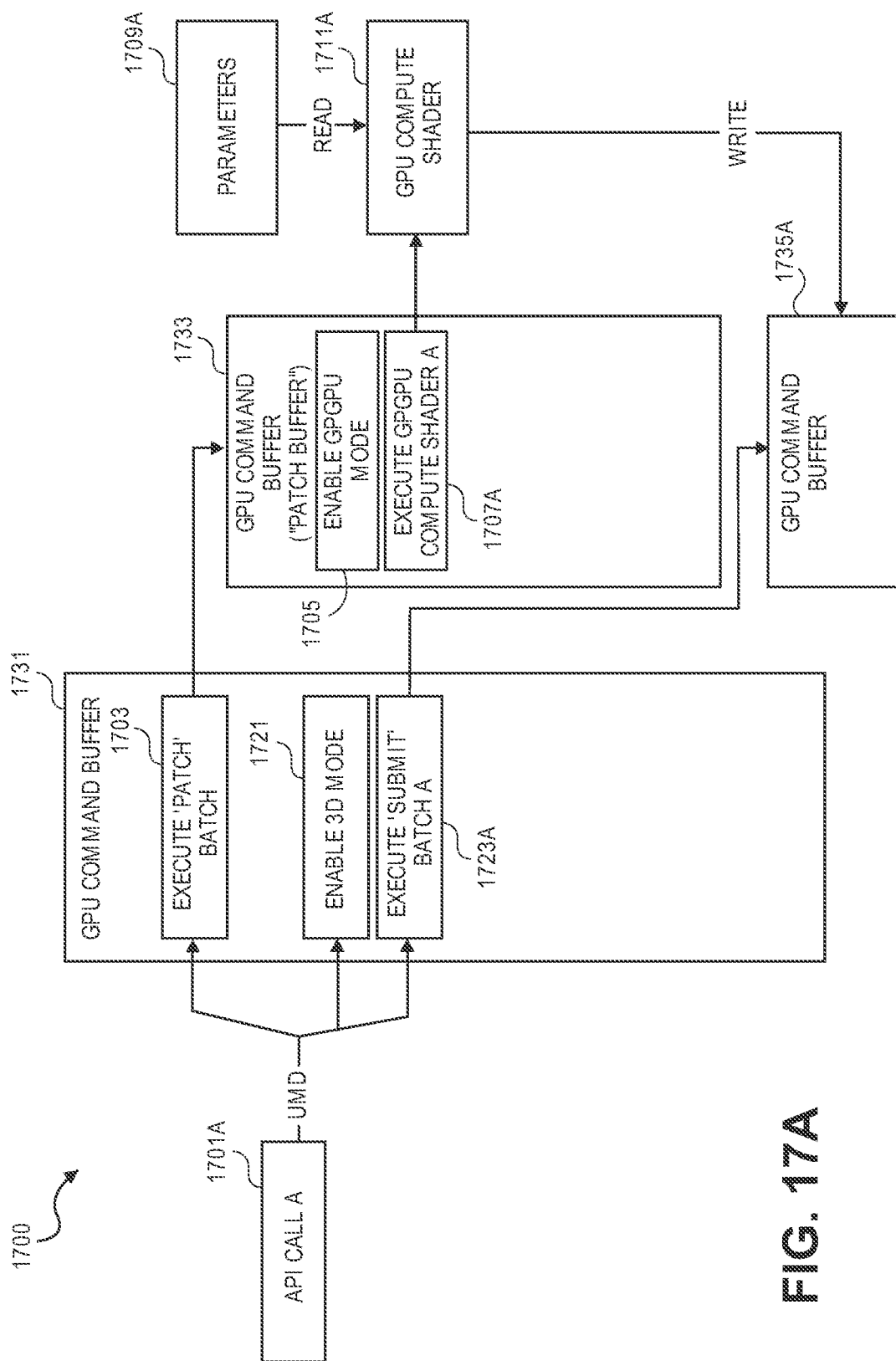


FIG. 17A

18/21

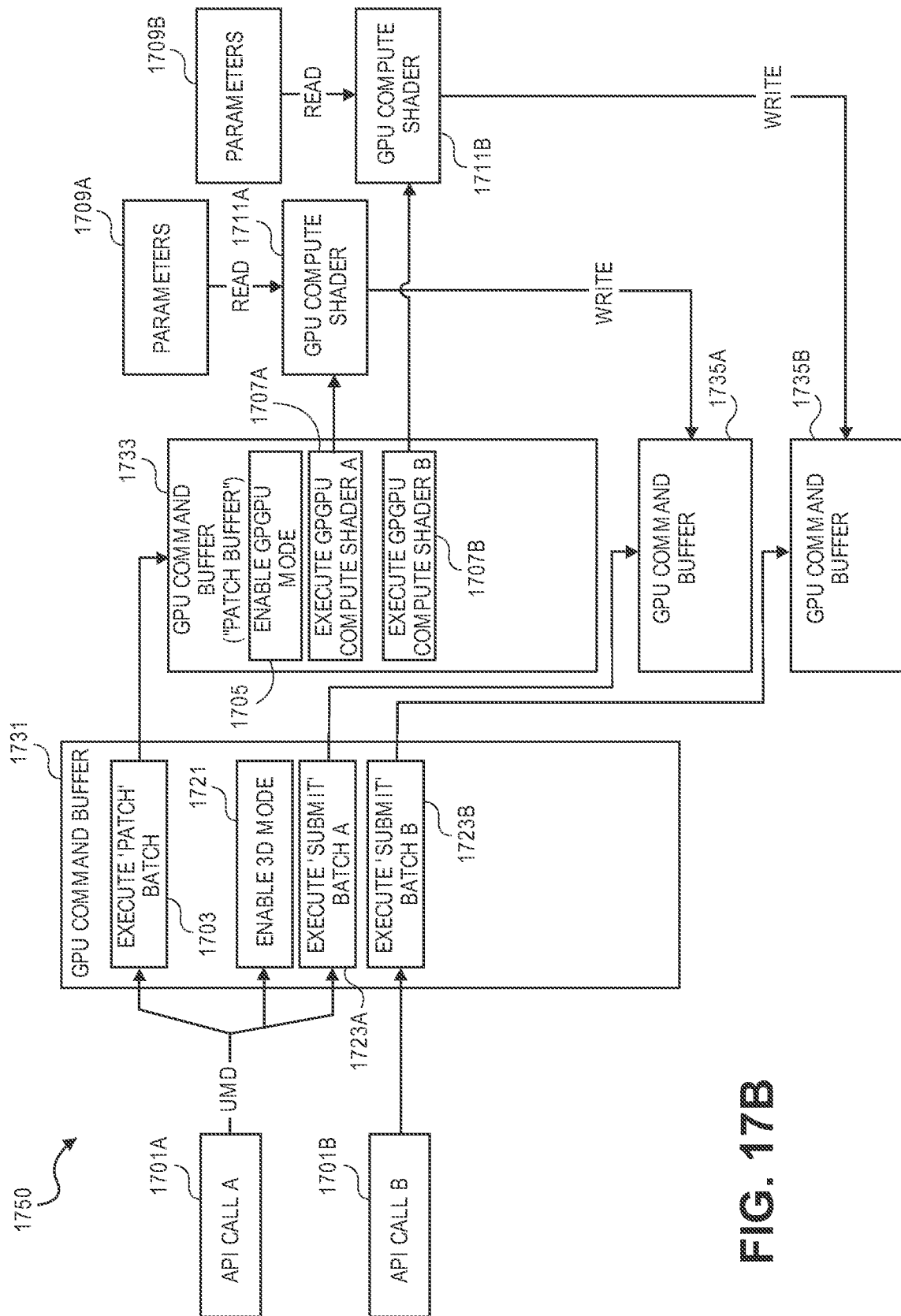


FIG. 17B

19/21

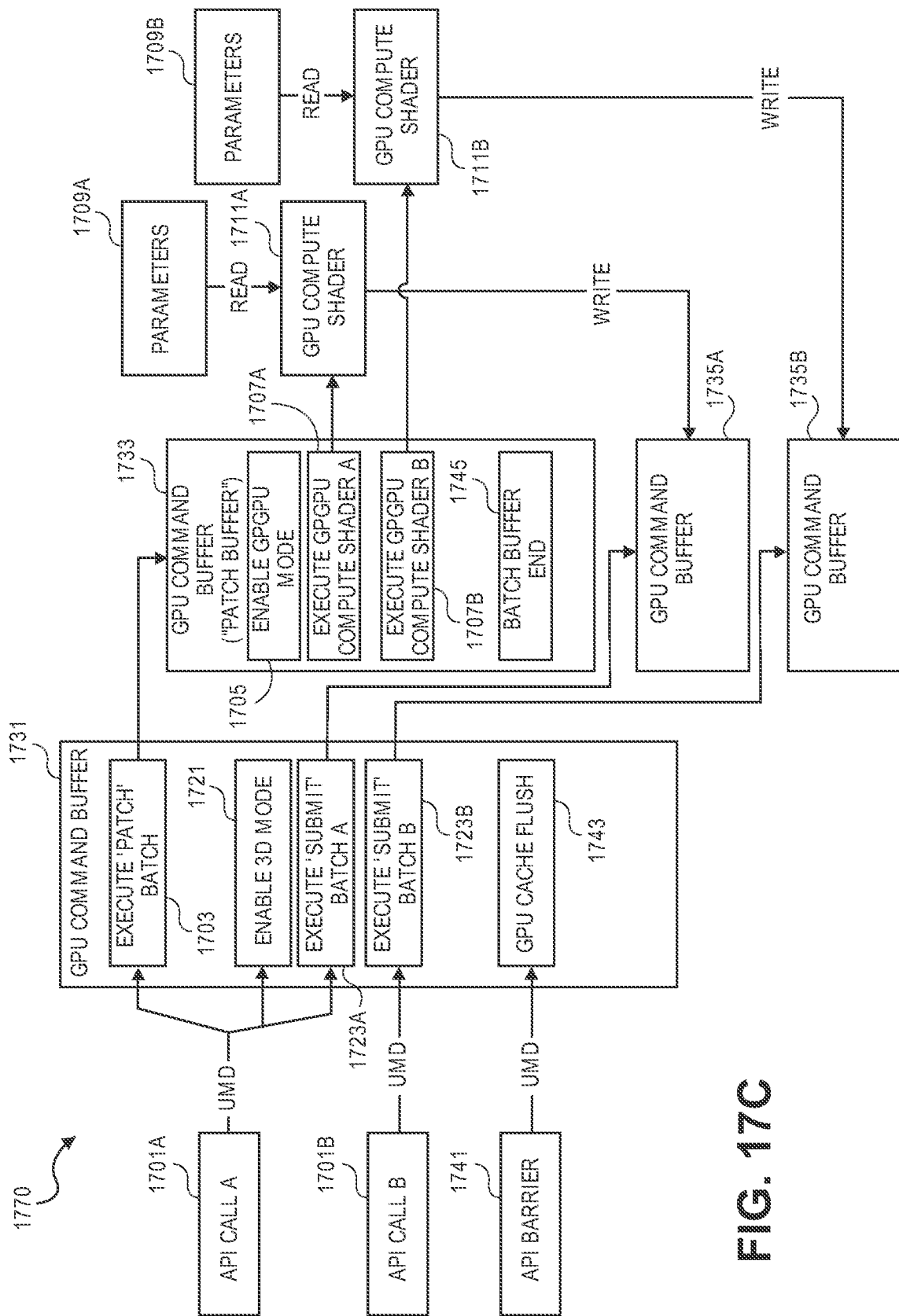


FIG. 17C

20/21

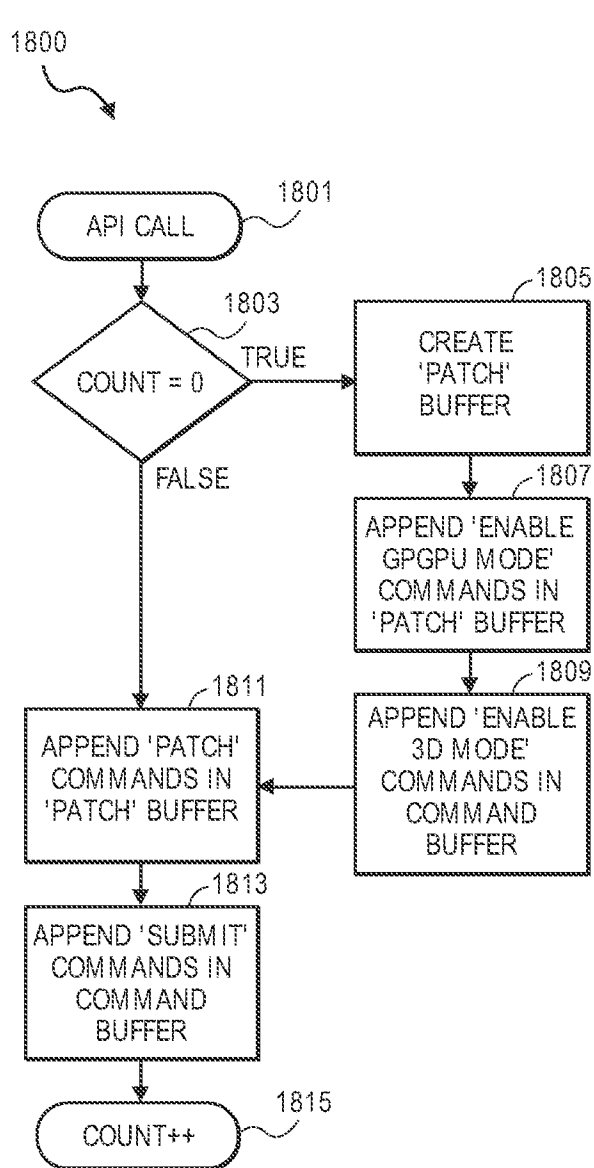


FIG. 18A

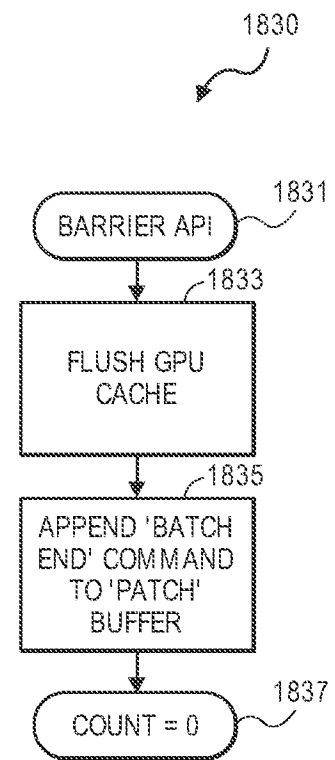
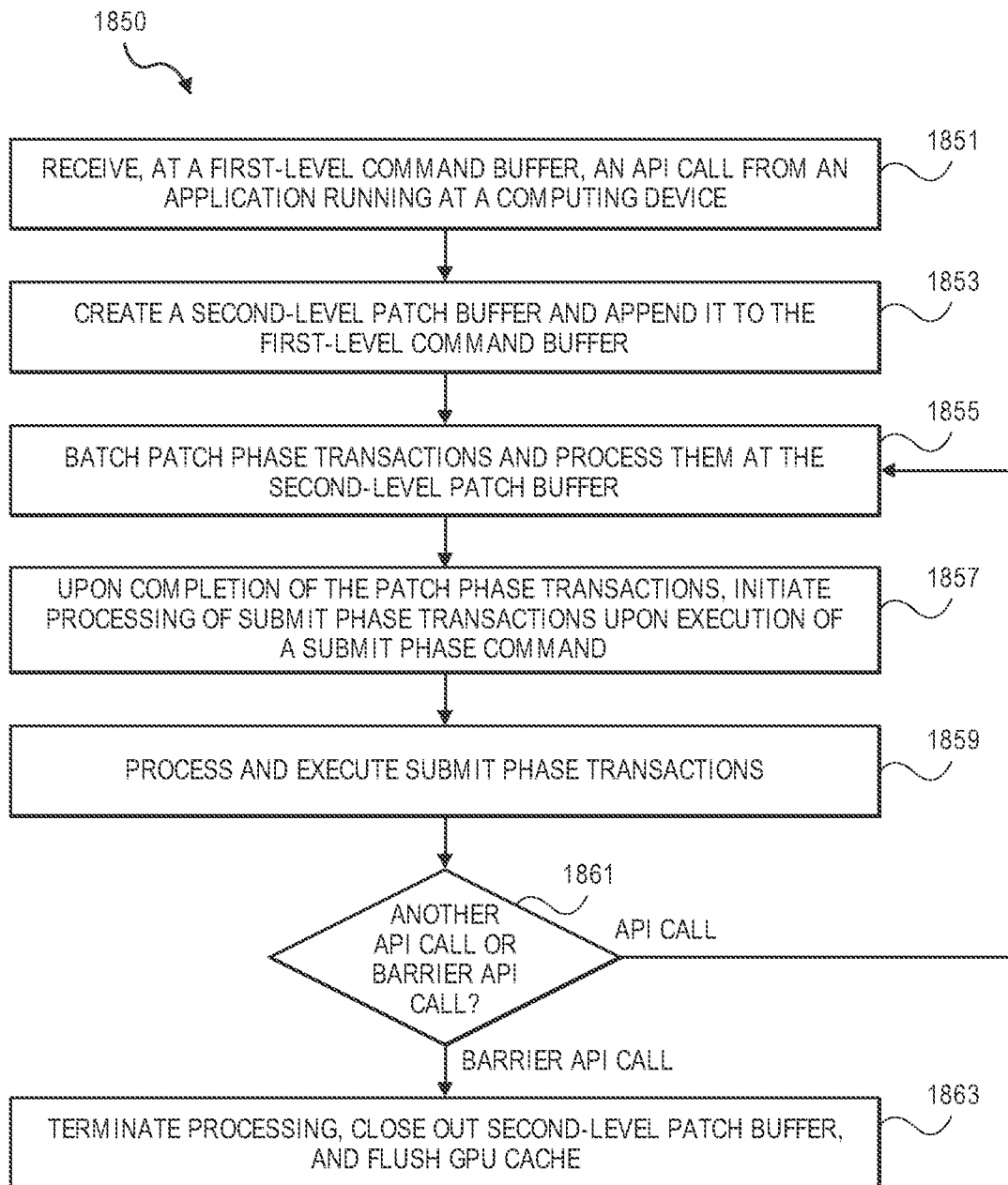


FIG. 18B

21/21

**FIG. 18C**

INTERNATIONAL SEARCH REPORT

International application No.
PCT/US2016/031872**A. CLASSIFICATION OF SUBJECT MATTER****G06F 9/38(2006.01)I, G06F 9/54(2006.01)I, G06T 1/20(2006.01)I**

According to International Patent Classification (IPC) or to both national classification and IPC

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

G06F 9/38; G06F 15/16; G06F 3/00; G09G 5/39; G06T 1/20; G06T 1/60; G09G 5/377; G06F 9/54

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Korean utility models and applications for utility models

Japanese utility models and applications for utility models

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

eKOMPASS(KIPO internal) & Keywords: graphic processing unit (GPU), application programming interface (API), second command, command buffer, another buffer**C. DOCUMENTS CONSIDERED TO BE RELEVANT**

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
X	US 8941671 B2 (MEHER PRASAD MALAKAPALLI et al.) 27 January 2015	1-6,9-14,17-21
Y	See column 2, line 19 - column 12, line 18; claims 1, 5; and figure 1.	7-8,15-16
Y	US 2011-0063313 A1 (JEFFREY A. BOLZ et al.) 17 March 2011 See paragraph [0047].	7-8,15-16
A	US 2015-0035860 A1 (APPLE INC.) 05 February 2015 See paragraphs [0034]-[0053]; and figures 3-4.	1-21
A	US 2007-0088856 A1 (GUOFENG ZHANG) 19 April 2007 See paragraphs [0019]-[0021]; and figures 3-4A.	1-21
A	US 2014-0313214 A1 (APPLE INC.) 23 October 2014 See paragraphs [0040]-[0043]; and figure 4.	1-21



Further documents are listed in the continuation of Box C.



See patent family annex.

* Special categories of cited documents:

"A" document defining the general state of the art which is not considered to be of particular relevance

"E" earlier application or patent but published on or after the international filing date

"L" document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified)

"O" document referring to an oral disclosure, use, exhibition or other means

"P" document published prior to the international filing date but later than the priority date claimed

"T" later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention

"X" document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone

"Y" document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art

"&" document member of the same patent family

Date of the actual completion of the international search

29 August 2016 (29.08.2016)

Date of mailing of the international search report

30 August 2016 (30.08.2016)

Name and mailing address of the ISA/KR

International Application Division

Korean Intellectual Property Office

189 Cheongsa-ro, Seo-gu, Daejeon, 35208, Republic of Korea

Facsimile No. +82-42-481-8578

Authorized officer

CHIN, Sang Bum

Telephone No. +82-42-481-8398



INTERNATIONAL SEARCH REPORT

Information on patent family members

International application No.

PCT/US2016/031872

Patent document cited in search report	Publication date	Patent family member(s)	Publication date
US 8941671 B2	27/01/2015	CN 104040494 A EP 2802982 A1 EP 2802982 A4 US 2013-0181999 A1 WO 2013-106491 A1	10/09/2014 19/11/2014 09/09/2015 18/07/2013 18/07/2013
US 2011-0063313 A1	17/03/2011	US 2011-0063294 A1 US 2011-0063296 A1 US 2011-0063318 A1 US 9024946 B2 US 9245371 B2 US 9324175 B2	17/03/2011 17/03/2011 17/03/2011 05/05/2015 26/01/2016 26/04/2016
US 2015-0035860 A1	05/02/2015	EP 2250558 A1 US 2009-225093 A1 US 2013-335443 A1 US 8477143 B2 US 8842133 B2 WO 2009-111045 A1	17/11/2010 10/09/2009 19/12/2013 02/07/2013 23/09/2014 11/09/2009
US 2007-0088856 A1	19/04/2007	CN 100489820 C CN 100495435 C CN 100568277 C CN 101025717 A CN 1991903 A CN 1991903 B CN 1991904 A CN 1991906 A TW 200821978 A TW 200821979 A TW I322354 B TW I328779 B TW I329845 B TW I354240 B US 2007-091097 A1 US 2007-091098 A1 US 2007-091099 A1 US 7812849 B2 US 7889202 B2 US 7903120 B2	20/05/2009 03/06/2009 09/12/2009 29/08/2007 04/07/2007 12/05/2010 04/07/2007 04/07/2007 16/05/2008 16/05/2008 21/03/2010 11/08/2010 01/09/2010 11/12/2011 26/04/2007 26/04/2007 26/04/2007 12/10/2010 15/02/2011 08/03/2011
US 2014-0313214 A1	23/10/2014	AU 2011-256745 A1 AU 2011-256745 B2 CA 2795365 A1 CA 2795365 C CN 102870096 A CN 102870096 B DE 112011101725 T5 EP 2572276 A1	01/11/2012 23/01/2014 24/11/2011 08/12/2015 09/01/2013 13/01/2016 25/04/2013 27/03/2013

INTERNATIONAL SEARCH REPORT

Information on patent family members

International application No.

PCT/US2016/031872Patent document
cited in search reportPublication
datePatent family
member(s)Publication
date

GB 201108084 D0	29/06/2011
GB 2480536 A	23/11/2011
GB 2480536 B	02/01/2013
JP 2013-528861 A	11/07/2013
JP 2015-007982 A	15/01/2015
JP 5583844 B2	03/09/2014
KR 10-1477882 B1	30/12/2014
KR 10-1558831 B1	08/10/2015
KR 10-2013-0004351 A	09/01/2013
KR 10-2014-0117689 A	07/10/2014
MX 2012012534 A	17/12/2012
US 2011-285729 A1	24/11/2011
US 2015-187322 A1	02/07/2015
US 8723877 B2	13/05/2014
US 8957906 B2	17/02/2015
WO 2011-146197 A1	24/11/2011